

Logique élémentaire

Jacques Zahnd

Logique élémentaire

Cours de base
pour informaticiens

Dans la même collection:

Algorithmes et structures de données avec Ada, C++ et Java

A. GUERID, P. BREGUET, H. RÖTHLISBERGER

VHDL du langage à la modélisation

R. AIRIAU, J.-M. BERGÉ, V. OLIVE, J. ROUILLARD

Programmation séquentielle avec ADA 95

PIERRE BREGUET, LUIGI ZAFFALON

Programmation concurrente et temps réel avec ADA 95

LUIGI ZAFFALON, PIERRE BREGUET

Programmation concurrente

ANDRÉ SCHIPER

Informatique industrielle

ROGER HERSCH

Programmation orientée objets en c++

MARYLÈNE MICHELOUD

MEDARD RIEDER

Les Presses polytechniques et universitaires romandes sont une fondation scientifique dont le but est principalement la diffusion des travaux de l'Ecole polytechnique fédérale de Lausanne ainsi que d'autres universités et écoles d'ingénieurs francophones.

Le catalogue de leurs publications peut être obtenu par courrier aux Presses polytechniques et universitaires romandes : EPFL – Centre Midi, CH-1015 Lausanne, par E-Mail à ppur@epfl.ch, par téléphone au (0)21 693 41 40, ou par fax au (0)21 693 40 27.

www.ppur.org

Deuxième édition

ISBN 2-88074-360-5

© 1998, **2003, réimpression corrigée**

Presses polytechniques et universitaires romandes,
CH-1015 Lausanne.

Tous droits réservés.

Reproduction, même partielle, sous quelque forme
ou sur quelque support que ce soit, interdite sans l'accord écrit de l'éditeur.
Imprimé en Italie.

TABLE DES MATIÈRES

1	Introduction	1
1.1	Buts généraux	1
1.2	La logique	3
1.3	Plan de l'ouvrage	11
2	Langages formels	15
2.1	Expressions d'un langage	15
2.2	Syntaxe et sémantique	17
2.3	Termes et propositions	22
2.4	Langage et métalangage	24
3	Symboles d'un langage du premier ordre	27
3.1	Exemple mathématique	27
3.2	Exemple informatique	31
3.3	Variables et constantes individuelles	33
3.4	Symboles fonctionnels	37
3.5	Symboles relationnels	40
3.6	Connecteurs logiques	42
3.7	Quantificateurs	48
3.8	Donnée d'un langage du premier ordre	50
4	Syntaxe d'un langage du premier ordre	53
4.1	Diagrammes et règles syntaxiques	53
4.2	Constructions génératrices	54
4.3	Arbres des termes et propositions	58
4.4	Variables libres et variables liées	61
4.5	Substitutions	66
4.6	Termes librement substituables	71
4.7	Exercices	75
5	Théories ou systèmes de déduction	79
5.1	Déductions	79
5.2	Hypothèses et jugements	84
5.3	Théories	88
5.4	Démonstrations et déductions	93
5.5	Théorèmes	97

5.6	Propositions vraies, propositions fausses	98
5.7	Extensions d'une théorie	100
6	Logique propositionnelle	103
6.1	Déductions élémentaires	103
6.2	Quelques démonstrations	110
6.3	Démonstrations naturelles	116
6.4	La loi du tiers exclu et la logique intuitionniste	126
6.5	Déductions structurelles	131
7	Déductions dérivées de logique propositionnelle	139
7.1	Proposition fausse et contradictions	139
7.2	Schémas d'implication	141
7.3	Disjonctions de cas contraires	147
7.4	Propriétés de l'équivalence	150
7.5	Substitutivité de l'équivalence	155
7.6	Schémas de déduction booléens	158
7.7	La règle de dualité	166
7.8	Formes booléennes de l'implication et de l'équivalence	170
7.9	Simplifications conditionnelles	177
7.10	Conditions nécessaires et suffisantes	177
7.11	Exercices	182
8	Introduction à la théorie des ensembles	183
8.1	Le rôle de la théorie des ensembles	183
8.2	L'univers de la théorie des ensembles	185
9	Logique des prédicats	191
9.1	Déductions élémentaires	191
9.2	Convention d'un langage universel	198
9.3	Particularisations	200
9.4	Généralisations	204
9.5	Rôle des déductions structurelles	208
9.6	Introduction de quantificateurs existentiels	212
9.7	Élimination de quantificateurs existentiels	220
10	Déductions dérivées de logique des prédicats	229
10.1	Substitutivité de l'équivalence	229
10.2	Dualité en logique des prédicats	232
10.3	Permutation distribution de quantificateurs	237
10.4	Changements de variables liées	245
10.5	Particularisations simultanées	253
10.6	Exemplifications	259
10.7	Application à la théorie des ensembles	263
10.8	Formes prénexes	274

10.9	Quantificateurs typés	278
11	Logique des prédicats avec égalité	281
11.1	Déductions élémentaires	281
11.2	Schémas de déduction dérivés	286
11.3	Application à la théorie des ensembles	294
11.4	Quantificateurs d'unicité	300
11.5	Exercices	310
12	Extensions définitionnelles	313
12.1	Axiomes et schémas définitionnels	313
12.2	Couples et graphes	320
13	Langages du premier ordre à opérateurs généraux	329
13.1	Opérateurs généraux	329
13.2	Sélecteurs	335
13.3	Termes conditionnels	346
14	Opérateurs de réunion et de collection de la théorie des ensembles	351
14.1	Opérateurs de réunion	351
14.2	Opérateurs de collection simples	359
14.3	Opérateurs de collection généraux	370
15	Fonctions	381
15.1	Fonctions et applications	381
15.2	Opérateurs d'abstraction	384
15.3	Propriétés de fonctions	391
16	Annexe	397
	Démonstrations de théorie des ensembles	397
	Bibliographie	423
	Index	427

REMERCIEMENTS

L'auteur remercie les collaborateurs des Presses polytechniques et universitaires romandes qui ont contribué à la réalisation de ce livre, en particulier, pour leurs conseils rédactionnels, Mme Claire-Lise Delacrausaz et M. Fabio Tonasso.

Il remercie le directeur du Laboratoire de systèmes logiques de l'EPFL, le professeur Daniel Mange, pour son soutien de très longue durée.

1

Introduction

1.1 Buts généraux

Résumé. La faculté de raisonner sur des objets concrets ou abstraits est une aptitude fondamentale dans toutes les disciplines qui traitent d'objets dont les propriétés sont bien définies, telles que les mathématiques, l'informatique, les sciences exactes en général. L'informaticien, en particulier, doit être capable de raisonner sur des objets tels que programmes, structures de données, systèmes, processus, circuits logiques, etc. Ce livre enseigne les techniques de base pour effectuer ces raisonnements, non pas de la manière informelle et incertaine usuelle, mais de façon formelle, c'est-à-dire en appliquant des règles précises de manipulation de symboles qui ne laissent aucune place à l'ambiguïté et à l'erreur. Ce sont les règles de la logique.

Cet ouvrage a été conçu comme un cours de base pour les informaticiens qui désirent maîtriser ce qu'on appelle les méthodes formelles en informatique. Celles-ci sont un sujet dont nous allons parler dans cette introduction. En bref, ce sont des méthodes mathématiques, mais de celles qui appartiennent aux *bases* des mathématiques, non aux mathématiques avancées. La maîtrise de ces méthodes nécessite une étude sérieuse de ces bases, une connaissance précise de ce qu'est la démarche mathématique elle-même. On peut dire que les méthodes formelles sont simplement l'application de *la méthode mathématique* aux objets de l'informatique. Une formation dans ce domaine n'est autre qu'une formation de base à la méthode mathématique. Elle présente donc un intérêt qui va bien au-delà des méthodes formelles en informatique. Elle est utile dans l'étude de toutes les branches des mathématiques et des sciences exactes qui les appliquent.

Méthodes formelles en informatique. Celui qui participe au développement de logiciels complexes est confronté d'abord à un problème de langage ou de notation. Il lui faut, pour travailler avec sûreté, une description précise des problèmes qu'il doit résoudre et une description précise des systèmes auxquels

il a affaire et dont il doit appréhender la structure et le fonctionnement afin de pouvoir développer les composants informatiques adaptés à leur gestion.

Ces descriptions peuvent être « dans sa tête » après un certain temps passé à étudier le problème, mais leur notation sous une forme appropriée est un support pratiquement indispensable dans le cas de systèmes complexes. Il est reconnu depuis longtemps que des notations précises dans ce domaine sont un outil extrêmement utile. Le génie logiciel en connaît de nombreuses.

Les *langages de spécification* sont des systèmes de notation « universels », conçus pour formuler avec précision n'importe quelle tâche de développement de logiciel ou de matériel. Ils parviennent à cette universalité grâce à des notations qui correspondent aux formes et aux opérations de base de la pensée, lorsque nous analysons la structure et le comportement d'un système quelconque, et ce sont aussi les formes de base de la pensée mathématique. Car les concepts de la *logique* et de la *théorie des ensembles* sont les abstractions fondamentales suivant lesquelles notre cerveau élabore une représentation d'un système d'objets. Il y a des ensembles d'objets, des ensembles d'états, des relations entre les éléments de ces ensembles, des correspondances, des fonctions, etc. D'autre part, les énoncés qui décrivent les propriétés de ces systèmes font intervenir le formalisme de la logique des prédicats.

Les mathématiques ont découvert depuis longtemps le rôle fondamental de ces concepts. Les informaticiens le découvrent à leur tour, à propos de systèmes dont les composants ne sont pas a priori de nature mathématique, par exemple les composants d'une base de donnée industrielle, mais dont la structure est telle que nous la décrivons naturellement au moyen de ces abstractions.

Les méthodes formelles ne se limitent pas à l'emploi de notations précises pour formuler un problème. Elles incluent des méthodes de développement qui doivent aboutir à un programme ou à un plan de circuit corrects par rapport à la tâche formulée. Elles rendent possible la preuve formelle que le produit final satisfait à la spécification initiale donnée.

Une formation préalable. Il existe une littérature assez étendue sur les méthodes formelles en informatique et sur les langages de spécification (voir Bibliographie). Mais la mise en pratique de ces méthodes nécessite une formation préalable qui fait défaut à la plupart des informaticiens, à savoir la familiarité avec les formes d'expression du langage mathématique que l'on retrouve dans les langages de spécification et l'entraînement au raisonnement formel. Nous pensons que la méconnaissance de ces sujets constitue une « barrière culturelle » dont on sous-estime l'importance. L'enseignement mathématique traditionnel, où l'on s'intéresse aux « faits » mathématiques bien plus qu'au langage utilisé, et où le raisonnement est presque toujours sémantique et non formel, ne suffit pas. On y prend peu conscience de la *forme* de la pensée, car cette forme n'est pas le sujet du discours.

Or c'est une autre vision des mathématiques qu'il faut avoir pour étudier et appliquer les méthodes formelles. Cette vision est celle d'un *langage* obéissant

à des règles de syntaxe précises. La déduction logique formelle dans ce langage est un genre de *calcul* dans lequel on applique des règles de manipulation de symboles et d'expressions. Ce livre propose un programme d'apprentissage et d'entraînement assez poussé à cette forme de calcul. Il s'adresse à des étudiants de première année de l'Université qui devraient en bénéficier non seulement dans l'étude de méthodes formelles en informatique mais encore dans l'ensemble des branches mathématiques.

1.2 La logique

L'étude des formes du raisonnement. Il n'est guère possible de présenter une science en quelques mots sans en omettre des aspects importants. Ce que nous pourrions dire de la logique dans ce paragraphe ne sera donc qu'une description sommaire, visant seulement à esquisser le sujet de notre étude.

Pour partir d'une définition très succincte, nous dirons que la logique est *l'étude des formes du raisonnement*. Cette définition n'a de sens évidemment que pour celui qui sait plus ou moins ce qu'est le raisonnement. Nous n'allons pas tenter de le définir ici, car cela peut être considéré justement comme l'une des tâches de la logique. Nous supposons seulement que le lecteur sait *par expérience* quelle espèce d'activité mentale on désigne par ce terme et par quel genre de textes — les démonstrations — elle se traduit. Nous supposons qu'il en a donc une notion *pratique*, telle qu'elle s'acquiert à l'école secondaire dans les branches mathématiques.

Le sujet de notre étude est évidemment fondamental. Le raisonnement, la déduction, sont l'une des formes caractéristiques de la pensée scientifique. Cela suffirait d'un point de vue culturel à justifier l'étude de la logique. Mais celle-ci ne présente pas qu'un intérêt culturel. Celui qui désire maîtriser les aspects théoriques d'une discipline scientifique quelconque en tirera toujours un profit considérable. Car la logique nous fournit réellement la grammaire de n'importe quel discours rigoureux. Elle seule permet d'acquérir en particulier la maîtrise du langage mathématique, commun à toutes les sciences exactes. De ce fait, sa portée est universelle et nous ne voyons pas quelle science pourrait être plus fondamentale.

La logique va d'ailleurs plus loin que la simple étude des règles du raisonnement. Elle décrit la forme générale de ces édifices de la pensée qu'on appelle des *théories*. Elle précise la notion même de théorie. Elle analyse les relations qui peuvent exister entre différentes théories, précisant en particulier la notion de *modèle*, si fondamentale dans toutes les sciences exactes. C'est dire que son sujet, au delà du raisonnement, est de façon générale l'étude de certaines formes caractéristiques d'un discours scientifique.

Validité formelle d'un raisonnement. Le raisonnement est une activité mentale qui se manifeste par un certain *genre de discours*. C'est de ce point de vue linguistique que nous le considérerons. La logique ne s'intéresse pas

directement à l'activité mentale qu'il représente, à ses aspects psychologiques ou physiologiques. Elle ne considère que l'expression de la pensée par des mots, des phrases, des textes.

Considérant le raisonnement de ce point de vue, nous pouvons le comparer à d'autres genres de discours tels que le récit, le poème, le texte de loi, le cahier des charges, l'inventaire, etc. Chacun de ces genres possède, à des degrés divers, des qualités qui lui sont propres. Le raisonnement possède une qualité principale qui n'a que deux degrés possibles: il est *correct* — on dit aussi *valide* — ou il ne l'est pas.

Le premier but de la logique est de caractériser les raisonnements corrects, du point de vue de leur *forme* seulement. Comme toute espèce de discours, un raisonnement présente deux aspects: celui de sa forme — on parle aussi de ses propriétés *syntactiques* — et de sa signification (propriétés *sémantiques*). La logique — tout au moins la logique élémentaire — ne s'intéresse qu'à la *forme*, et fait abstraction de toute signification. Son hypothèse de travail est que la validité d'un raisonnement, pour peu qu'il soit formulé dans un langage adéquat, est une propriété purement syntaxique, c'est-à-dire ne dépend que de la forme des phrases et des enchaînements de phrases. Deux raisonnements de même forme, sur des objets totalement différents, sont simultanément valides ou invalides. Considérons un exemple simple, pour illustrer ce propos:

*Deux hommes qui parlent la même langue peuvent communiquer;
Monsieur Marius et Monsieur Brun parlent la même langue;
donc Monsieur Marius et Monsieur Brun peuvent communiquer.*

Les petits exemples comme celui-ci, que l'on rencontre au début des livres de logique, ont toujours quelque chose de simpliste et le lecteur se demande parfois s'il vaut bien la peine de discuter de pareilles banalités. C'est un fait que les opérations de déduction élémentaires, dont l'exemple ci-dessus est une illustration, sont si simples qu'on s'abstient souvent en pratique de les exprimer complètement. C'est pourtant bien la première tâche de la logique que de mettre ces opérations en évidence. Mais elle ne se limite pas à cela, de même que la théorie du jeu d'échecs ne se borne pas à enseigner les lois du déplacement des pièces. Tout en partant des opérations logiques les plus élémentaires, il s'agit d'aboutir à des raisonnements complexes tels qu'on les rencontre dans les mathématiques.

Le raisonnement suivant porte sur des êtres mathématiques, mais on peut dire, sans rien connaître aux nombres complexes, qu'il est à des détails près de la même forme que le précédent:

*Deux nombres complexes ayant la même amplitude ont un rapport réel;
 z_1 et z_2 ont la même amplitude;
donc le rapport z_1/z_2 est réel.*

Les ambiguïtés de la langue naturelle. Ces deux exemples de déduction sont clairement perçus comme étant *en gros* de la même forme, et comme étant valides en vertu de celle-ci seulement. Mais nous prenons bien la précaution de souligner « en gros », car il y a des degrés dans la similitude de textes, et le premier problème de la logique se situe là comme nous allons le voir. L'idée que la validité d'un raisonnement ne doit dépendre que de sa forme paraît assez claire au premier abord. Encore faut-il savoir *reconnaître* la forme d'un raisonnement et savoir reconnaître, en particulier, que deux raisonnements sont de même forme.

Or la similitude de forme peut être trompeuse, et l'on s'en rend compte facilement, lorsque la langue du raisonnement est une langue naturelle — le français par exemple. Ces langues sont trop riches et subtiles, elles admettent trop de tournures variées, trop de formes grammaticales, et leurs termes sont chargés de trop de sous-entendus pour qu'il soit possible de juger de la validité d'un raisonnement uniquement d'après la forme apparente des phrases. En voici quelques exemples.

Lorsque nous entendons, dans un pays de religion chrétienne, la déduction suivante :

*Les cloches sonnent toujours le dimanche;
demain est dimanche;
donc demain les cloches sonneront,*

celle-ci nous paraît trivialement correcte, car nous interprétons naturellement sa première phrase comme un rappel du fait habituel que les cloches des églises sonnent *tous* les dimanches. Par contre la déduction suivante d'un citoyen suisse, de forme identique à la précédente, sera immédiatement douteuse pour ses concitoyens, car ceux-ci n'interpréteront pas le mot « toujours » de la même manière que dans l'exemple précédent :

*Les votations populaires ont toujours lieu le dimanche;
demain est dimanche;
donc demain il y aura des votations populaires.*

Lorsqu'une personne tient le discours suivant, nous trouvons sa déduction correcte — que nous partagions ou non sa conception de la culture — car nous pensons naturellement que par « un homme cultivé » elle veut dire « tout homme cultivé » :

*Un homme cultivé a lu les aventures de Tintin;
Monsieur X est un homme cultivé;
donc Monsieur X a lu les aventures de Tintin.*

Par contre, on trouvera bien douteux le raisonnement suivant — de même forme — commis par cette personne lorsque, voyant un parapluie oublié sur un banc public, elle déclarera :

*Un homme distrait a oublié ce paraphrasis;
Monsieur Tournesol est un homme distrait;
donc Monsieur Tournesol a oublié ce paraphrasis.*

L'exemple que voici est emprunté à Church [4]¹. Cet auteur propose d'abord un raisonnement correct :

*J'ai vu un portrait de John Wilkes Booth;
John Wilkes Booth a assassiné Abraham Lincoln;
donc j'ai vu le portrait d'un assassin.*

Puis un raisonnement incorrect, de même forme apparente que le premier :

*J'ai vu le portrait de quelqu'un;
quelqu'un a inventé la roue;
donc j'ai vu le portrait d'un inventeur.*

Il serait facile évidemment de corriger la forme de ces raisonnements — sans changer le sens des phrases — de façon à éliminer toute concordance de forme entre les exemples corrects et incorrects. En outre, dans le troisième exemple, on peut fort bien objecter que les deux raisonnements ne *sont pas* de la même forme; que le mot « quelqu'un » ne peut pas s'employer de la même manière qu'un nom propre tel que « John Wilkes Booth ». Objection valable, mais qui montre immédiatement que l'on sera confronté à toute la complexité et à toute l'imprécision de la grammaire de la langue naturelle, dès que l'on voudra donner des *règles* permettant de reconnaître deux textes différents de cette langue comme étant de la même forme — ces règles étant telles que deux raisonnements de même forme soient toujours simultanément corrects ou simultanément incorrects.

La vérification mécanique du raisonnement. De toute évidence, la formulation de telles règles ne pourra se faire qu'au prix de restrictions du langage, d'une limitation des formes et tournures utilisables.

Jusqu'où faut-il aller dans ce sens? La réponse dépend clairement de ce que l'on attend de ces règles. Or on attend de celles-ci qu'il n'y ait aucun désaccord possible entre lecteurs d'un raisonnement sur la validité de celui-ci, aucune place à la subjectivité dans son évaluation. Pour cela, on exige que cette évaluation ne fasse appel à « aucune intelligence ». Elle doit pouvoir s'effectuer de façon purement *mécanique* — autant dire par une machine.

Cette exigence est d'une importance fondamentale. Elle conditionne entièrement le développement de la logique, et l'on peut dire qu'elle fait partie de son essence même. Il convient donc de s'arrêter brièvement à l'adjectif « mécanique » employé ci-dessus.

¹ Les numéros entre crochets se rapportent à la bibliographie en fin du livre.

Ce terme doit être interprété dans le sens suivant : il qualifie une opération qui pourrait être effectuée par une *machine*, ou tout au moins qui en *donne l'impression* — et nous laisserons volontairement subsister cet élément subjectif dans notre définition. Toute tentative de définition objective nous entraînerait trop loin. En ce qui concerne les opérations dont nous parlons ici, à savoir l'analyse de la forme d'un raisonnement et son évaluation comme correct ou incorrect, le lecteur informaticien verra lui-même très vite qu'un *programme* adéquat pourrait les effectuer. C'est dans ce sens qu'il convient de prendre l'adjectif « mécanique ». Effectuée par un être humain, une opération qui mérite ce qualificatif ne demande que des facultés mentales élémentaires — telles que de savoir lire et compter — considérées comme ne faisant appel à aucune intelligence imaginative et créatrice. Pour employer le terme propre, c'est une opération pour laquelle on dispose d'un *algorithme*.

Dans un cours de logique élémentaire, il n'y a pas lieu de préciser davantage le sens du terme « mécanique ». Mais nous devons signaler que les découvertes les plus remarquables de ce siècle en logique (les théorèmes de Gödel) ont été le fruit d'une étude plus poussée de cette notion, et des limitations qu'elle impose au raisonnement formel.

Nécessité d'un langage formel. Pour satisfaire à l'exigence fondamentale discutée ci-dessus — la possibilité de vérifier mécaniquement la validité d'un raisonnement — les limitations du langage doivent être telles qu'on ne peut plus parler d'une simple restriction du langage naturel, mais bien de la création d'un *langage artificiel*, spécialement conçu dans ce but. Un tel langage, soumis comme un langage de programmation à des règles de syntaxe absolument strictes, sera dit *formel* ou *formalisé*, par opposition au langage naturel que l'on dit aussi informel.

Comment se présente un langage formel et que peut-on bien exprimer dans celui-ci, telles sont sans doute les premières questions qui viennent à l'esprit, à la lecture de ces indications. À la première de ces questions il serait trop long de répondre ici. Plusieurs des chapitres qui suivent sont entièrement consacrés à la description de tels langages. Quant à la seconde, la meilleure réponse que nous puissions lui donner est la suivante : les logiciens du dix-neuvième siècle et du début du vingtième siècle ont décrit des langages formels dans lesquels on peut traduire — en principe — les oeuvres complètes des mathématiciens parues à ce jour. C'est une manière de dire que leur pouvoir d'expression ne laisse rien à désirer.

Si nous nous référons ici aux mathématiques, c'est évidemment parce que le raisonnement y occupe une place plus grande, et surtout plus explicite, que dans toute autre science. Il n'y a guère que dans les ouvrages de mathématique, ou dans les parties mathématiques des ouvrages de science exacte, que l'on rencontre des paragraphes intitulés « Démonstration » et qui méritent ce titre aux yeux d'un logicien. Il n'est pas exagéré de considérer cet aspect comme étant à tel point caractéristique de la pensée et de l'écriture mathématique

qu'on puisse le prendre comme *définition* même du terme « mathématique ». Nous reviendrons plus bas sur ce sujet.

La mathématique formelle. Le lecteur qui aborde ici pour la première fois l'étude de la logique peut être surpris lorsque nous parlons de la possibilité de *traduire* les oeuvres des mathématiciens dans un langage formel. Lorsqu'on ne sait pas encore comment se présente un tel langage — au sens des logiciens — on peut penser que les ouvrages des mathématiciens sont *déjà* et *toujours* rédigés par leurs auteurs dans un langage « formel ». L'adjectif « formel » n'évoque d'abord que l'emploi de « formules » et d'une multitude de symboles spéciaux. Or s'il y a bien de cela dans un langage formel, ce n'en est pas l'essentiel. L'important est qu'il satisfasse à l'exigence fondamentale de la logique que nous avons posée plus haut et que nous rappelons ici : il doit être possible à une *machine* de vérifier un texte écrit dans ce langage, c'est-à-dire de vérifier tous les raisonnements — les démonstrations — de l'auteur. Cette vérification doit donc pouvoir s'effectuer d'après la seule forme des phrases et ne doit demander *aucune compréhension* de leur sens. Pour une machine, un texte formalisé n'est qu'une suite de symboles dépourvus de toute signification, dans laquelle seul importe le respect de certaines règles d'assemblage et d'enchaînement. Sa vérification peut être comparée à celle d'une addition ou celle du respect des règles du jeu dans la description symbolique d'une partie d'échecs.

Ce sont là des exigences extrêmement fortes qui nécessitent un degré de formalisation très supérieur à celui des textes mathématiques courants. Ces derniers ne sont pas rédigés pour être lus et vérifiés par une machine, mais par des lecteurs intelligents auxquels ils veulent surtout communiquer des connaissances. Les mathématiciens ne se servent que très partiellement d'un langage formalisé. Ils combinent celui-ci avec des expressions et des phrases entières du langage naturel, et surtout, ils omettent un nombre considérable de maillons dans les chaînes d'assertions qui constituent les démonstrations. Ils supposent que l'intelligence des lecteurs suppléera d'elle-même à ces lacunes. Quiconque a déjà peiné comme lecteur à suivre les déductions d'un texte mathématique — et cela arrive même avec ceux qui se veulent didactiques — sait bien que ce travail requiert une intelligence que l'on ne saurait reconnaître à une machine.

Les mathématiques réelles. S'il est relativement facile de savoir comment travaille une machine, il en va tout autrement du travail mental d'un être humain. Il n'est pas dans notre propos de nous étendre sur ce sujet. Nous nous bornons ici à revenir sur la distinction faite entre *syntaxe* et *sémantique*. La compréhension d'un texte mathématique par un lecteur humain s'appuie surtout sur le *sens* des mots et des phrases, et nous dirons que les textes mathématiques usuels sont faits pour être appréhendés de manière sémantique. On ne peut donner évidemment que de très vagues indications sur ce qu'est de manière générale le « sens » des énoncés mathématiques. Mais on ne saurait nier en tout cas l'existence d'un tel sens, pour ceux qui lisent un texte

mathématique dans un domaine qui leur est familier. Ce texte leur parle d'objets qui sont sans doute des abstractions mais dont ils ont une « image » mentale. Il ne s'agit pas nécessairement d'une image au sens visuel car les êtres mathématiques ne sont pas tous des concepts géométriques, mais d'une représentation mentale qui rend plus ou moins évidentes les propriétés de ces objets. Cette représentation joue un rôle considérable dans la compréhension des démonstrations mathématiques usuelles et cette manière d'appréhender un texte mathématique est fort éloignée du contrôle purement *syntactique*, ignorant le sens des mots, qu'une machine peut faire d'un texte formalisé.

Les mathématiciens d'aujourd'hui utilisent donc, comme leurs prédécesseurs des siècles passés un langage en bonne partie informel. Ils le font pour de bonnes raisons, malgré le risque de déductions erronées que cela comporte et que même les meilleurs d'entre eux commettent parfois.

La mathématique formelle comme idéal de rigueur. Depuis la fin du dix-neuvième siècle toutefois, les mathématiciens disposent de langages formels auquel ils *peuvent* recourir ou dont ils *peuvent* se rapprocher lorsqu'ils le jugent nécessaire. Ils savent ainsi comment mettre en forme leurs démonstrations d'une manière qui élimine tout appel à l'intuition du lecteur, d'une manière qui les rend plus indiscutables — ce qui veut dire *mécaniquement vérifiables*. Cette mise en forme s'effectue chaque fois que le risque d'une erreur ou celui d'une interprétation erronée de la part des lecteurs est ressenti.

On peut dire que le rôle de la logique, vis-à-vis des mathématiques, est de décrire ce qu'est la forme de textes dont la rigueur est poussée à son point extrême, celui-ci étant caractérisé par la possibilité d'une vérification mécanique, ignorant toute signification. Ainsi donc, si la rigueur est une vertu cardinale des textes mathématiques, la logique décrit ce qui est la perfection de ce point de vue.

Dans la pratique, le mathématicien ne cherche jamais à atteindre cette extrême rigueur. Il se contente de développer ses démonstrations jusqu'au point où il aperçoit clairement la possibilité de les mettre en forme complètement. En rédigeant ses textes de cette manière, il se comporte un peu comme un informaticien qui se contenterait d'esquisser ses programmes juste assez pour qu'il puisse se rendre compte que leur mise en forme, c'est-à-dire leur rédaction dans un strict langage de programmation, ne serait plus qu'un exercice assez évident.

Ce que signifie « mathématique ». Nous revenons ici à la question: de quoi peut-on parler dans un langage formel au sens de la logique, c'est-à-dire un langage formel permettant la vérification mécanique des raisonnements? Nous avons répondu précédemment que l'on pouvait exposer dans un tel langage toutes les théories mathématiques. Quant à savoir si l'on peut parler d'autre chose que des « êtres mathématiques », la réponse dépend de ce que l'on définit comme étant « mathématique ».

Mais d'une part la meilleure réponse est *non*, et d'autre part cette réponse

n'est pas limitative. Car la meilleure définition que l'on puisse donner actuellement du terme « mathématique » est qu'il s'applique justement à tous les sujets susceptibles d'être traités dans un langage formalisé au sens de la logique. Dès lors qu'on tient un discours de type déductif traduisible dans un tel langage, on « fait » par définition de la mathématique.

À une époque antérieure, où la mathématique ne traitait que de nombres et de figures géométriques, on pouvait encore définir celle-ci comme la « science des grandeurs ». Mais cette caractérisation est périmée depuis longtemps. Les mathématiques traitent d'une multitude toujours croissante d'autres concepts. Il se trouve que le trait le plus caractéristique de cette science est son *langage*, la *manière* dont elle parle des choses, bien plus que les « êtres » dont elle parle — dont la liste est ouverte.

Il faut bien avoir à l'esprit la largeur de cette définition lorsque nous parlons des « mathématiques ». Est mathématique au sens large toute activité intellectuelle essentiellement soumise aux lois de la logique. Il en est ainsi de la programmation, même si dans la pratique actuelle le rôle de l'intuition dans ce domaine est encore bien plus dominant que dans les mathématiques réelles et le recours à la preuve formelle bien plus rare. Comme nous l'avons dit précédemment, il faut maîtriser la logique et la mathématique formelle pour exploiter la nature mathématique de la programmation et bénéficier dans ce domaine de la sûreté que peut procurer la rigueur.

Logique élémentaire et logique mathématique. Le lecteur qui aurait déjà consulté le catalogue d'une bibliothèque sous la rubrique « Logique » aura pu remarquer que bon nombre de titres dans ce domaine parlent de *Logique mathématique*. Il n'est pas certain que le sens de l'adjectif « mathématique », traditionnellement appliqué au substantif « Logique », soit le même pour tous les auteurs. Nous tenons, pour terminer cette introduction, à préciser *notre* interprétation de ce titre, et à préciser du même coup le sens que nous donnons à *Logique élémentaire*, qui en est l'*opposé* en quelque sorte. Il se peut, malgré nos explications, qu'un lecteur débutant ne voie pas clairement la différence. Cela ne présentera pour lui aucun inconvénient dans la suite. Ce que nous en disons s'adresse plutôt aux spécialistes.

La logique mathématique et la logique élémentaire sont deux manières d'étudier un même sujet — qui sommairement désigné est le raisonnement. Une science exacte peut toujours étudier son objet au travers d'un *modèle mathématique* de cet objet et il n'en va pas autrement de la logique. Le modèle mathématique que l'on peut donner d'un langage formel est une certaine espèce de structure algébrique abstraite — comme l'est un espace vectoriel — et un raisonnement est représenté par une suite finie d'éléments de cet espace. Avec ce modèle, on peut raisonner mathématiquement *sur le raisonnement mathématique* en démontrant les propriétés de ces suites. Pour être suivie avec succès et en toute sûreté, une telle approche nécessite *déjà* la maîtrise du langage mathématique et du raisonnement ainsi que de solides connaissances

en algèbre et en théorie des ensembles. Le but d'une telle étude ne saurait donc être l'acquisition de ce savoir. Il est de découvrir des propriétés générales non évidentes du langage, des raisonnements et des théories. De même que le physicien peut prédire les résultats d'expériences grâce aux lois de ses modèles de la nature, le logicien mathématicien peut caractériser par avance le résultat de certains raisonnements ou les propriétés de certaines théories, grâce à son modèle de la démarche mathématique. Dans une telle approche, les mathématiques sont appliquées à l'étude de leur propre démarche.

On peut dire que la logique mathématique est l'étude de la logique avec des *moyens mathématiques*. L'envergure de ces moyens peut d'ailleurs varier beaucoup d'un exposé à l'autre. Mais ils comprennent toujours un minimum de théorie des ensembles (jusqu'au lemme de Zorn) et d'algèbre universelle (algèbres libres, homomorphismes).

À l'opposé de cette approche, la *logique élémentaire* se veut une étude qui *précède* le développement de toute mathématique, en tant que *fondement* de celle-ci. Elle ne peut donc utiliser *aucune* connaissance mathématique. Cela ne veut pas dire qu'un exposé de logique élémentaire puisse s'adresser à de jeunes esprits n'ayant pas encore commencé l'apprentissage des mathématiques — possibilité pédagogique que personne ne recommanderait. Il faut avoir déjà pratiqué un peu les mathématiques pour comprendre les motivations de la logique. Mais la culture mathématique du lecteur ne servira que de « toile de fond » à notre exposé. Le lecteur peut se considérer, vis-à-vis de ce cours, dans la même position qu'une personne qui, tout en sachant plus ou moins bien parler sa langue maternelle, n'en a jamais étudié la grammaire et le fait pour la première fois.

La plupart des ouvrages de logique ne font pas une distinction aussi nette entre logique élémentaire et logique mathématique. La plupart d'entre eux sont assez mathématiques, même lorsqu'ils sont considérés comme élémentaires par leurs auteurs. Nous pensons qu'il convient de tracer une frontière tout à fait nette entre ces « étapes » de la logique, quels que soient les noms qu'on leur donne: celle que l'on parcourt avant de construire les mathématiques, comme fondement de celles-ci, et les autres, celles qui utilisent les mathématiques. Pour le lecteur déjà familiarisé avec le sujet, il doit être clair par exemple que le théorème de complétude de la logique des prédicats de Gödel n'appartient pas à la logique élémentaire.

1.3 Plan de l'ouvrage

Langages du premier ordre simples (Chap. 2–4). La première partie de l'ouvrage se compose des chapitres 2 à 4. Elle est consacrée à la description des langages du premier ordre simples, dont les seuls opérateurs sont des symboles fonctionnels et des symboles relationnels ainsi que les connecteurs logiques et quantificateurs usuels. Les opérateurs plus généraux, indispensables à la

formalisation du langage mathématique, ne sont introduits qu'au chapitre 13. Les chapitres 2 et 3 sont une introduction aux langages du premier ordre. Les définitions syntaxiques rigoureuses, notamment sur les variables libres et les substitutions, viennent au chapitre 4.

Théories ou systèmes de déductions (Chap. 5). La logique classique est un système de déductions parmi d'autres. Le concept général de système de déductions, que l'on appelle parfois «système formel», auquel se rattachent les notions de démonstration et de théorème, fait l'objet de ce chapitre. Une déduction dans un langage formel $\mathcal{L}$ est définie comme une séquence de jugements dans ce langage et un jugement se compose de deux parties: une liste d'hypothèses et une assertion.

Logique propositionnelle (Chap. 6–7). Le premier système de déductions étudié, dans un langage $\mathcal{L}$, est celui de la logique propositionnelle, fixant l'usage des connecteurs logiques *et*, *ou*, *non*, $\Rightarrow$, $\Leftrightarrow$. On part d'un ensemble de déductions élémentaires défini par des schémas de déduction qui sont, sous une autre forme, ceux de la «déduction naturelle» de Gentzen [9]. Le chapitre 6 présente ce système de déductions, définit ce qu'est une démonstration de forme naturelle, et introduit la notation utilisée dans cet ouvrage pour ces démonstrations. Cette notation, déjà employée par d'autres auteurs [10], est celle d'un texte muni d'une structure de blocs imbriqués qui correspondent aux états d'une pile d'hypothèses. Ce chapitre contient une brève orientation sur la logique intuitionniste et les mathématiques constructives.

Le chapitre 7 est entièrement consacré à l'élaboration d'un répertoire de schémas de déduction dérivés couvrant la plupart des déductions propositionnelles usuelles en mathématique.

Logique des prédicats (Chap. 8–10). Le système de déductions de la logique des prédicats dans un langage $\mathcal{L}$ est présenté dans le chapitre 9. Les quatre schémas de déduction élémentaires d'introduction et d'élimination des quantificateurs y sont discutés et illustrés par de nombreux exemples.

Le chapitre 8 est une brève introduction à la théorie des ensembles pour les besoins des chapitres suivants, où cette théorie est prise comme domaine d'application de la logique. Il contient d'abord une brève description du rôle de la théorie des ensembles en mathématique et en informatique (langages de spécification) puis il prépare le lecteur à aborder une théorie *formelle* des ensembles en décrivant «l'univers de l'interprétation» d'une telle théorie.

Dans le chapitre 10, un répertoire suffisant de schémas de déduction dérivés est élaboré comme pour la logique propositionnelle (Chap. 7). Ces schémas sont appliqués à la théorie des ensembles, où ils permettent de démontrer le corpus de théorèmes standard sur l'inclusion, la réunion, l'intersection, la complémentation d'ensembles.

Logique des prédicats avec égalité (Chap. 11). L'emploi du symbole d'égalité en mathématique est régi par les schémas de déduction introduits dans ce

chapitre. Ils sont appliqués en théorie des ensembles dans la démonstration des théorèmes sur les ensembles à un ou deux éléments.

Une première extension des langages du premier ordre a lieu dans ce chapitre, avec l'introduction des quantificateurs d'unicité «il existe au plus un», «il existe un et un seul» et de schémas de déduction y relatifs.

Extensions définitionnelles (Chap. 12). Une étude approfondie du mécanisme d'extension d'une théorie au moyen de définitions et de la notion plus générale d'extension conservative sort du cadre de cet ouvrage. Ce chapitre donne seulement la définition de ces notions, sans démontrer qu'une extension définitionnelle est conservative. Le lecteur doit seulement prendre conscience des propriétés logiques de ce genre d'extension. Cette notion est illustrée en théorie des ensembles par la définition des couples et des graphes (ensembles de couples).

Langages du premier ordre à opérateurs généraux (Chap. 13). Les langages du premier ordre simples (Chap. 3–4) ne suffisent pas — pratiquement — comme formalisation du langage mathématique courant. Il leur manque entre autres des opérateurs tels que ceux de sommation ou d'intégration, permettant de construire des termes contenant des variables locales (liées). Ce chapitre décrit sommairement la syntaxe d'un langage du premier ordre contenant de tels opérateurs.

Nous introduisons notamment les opérateurs logiques de sélection (ou choix) de Hilbert [11] qui permettent de mettre sous la forme d'une égalité toutes les définitions mathématiques dans lesquelles un terme est défini comme désignant «l'objet x tel que $P(x)$ ». L'introduction de ces opérateurs permet de traiter formellement l'emploi du pseudo-prédicat «existe» dans le langage mathématique informel, par exemple dans l'expression «la limite de la suite (x_n) existe» en analyse, ou dans l'expression «l'ensemble $\{x \mid x \notin x\}$ n'existe pas» en théorie des ensembles. Nous introduisons aussi l'opérateur logique de condition « Si », permettant de construire des *termes conditionnels* de la forme «Si A alors T_1 sinon T_2 », présents en mathématique dans les définitions par cas et dans les langages de programmation.

Opérateurs de collection et de réunion (Chap. 14). Ce chapitre et le suivant sont entièrement consacrés à la théorie des ensembles. Le premier contient notamment les schémas de déduction spécifiques de la théorie: le schéma de réunion et celui de séparation. Ils sont appliqués au traitement des produits cartésiens, des projections d'un graphe, du transformé d'un ensemble par un graphe, du graphe réciproque d'un graphe, du graphe identique d'un ensemble et du composé de deux graphes.

Nous appelons *collections de forme générale* les expressions de la forme «l'ensemble des $E(x_1, \dots, x_n)$ tels que $P(x_1, \dots, x_n)$ », où certaines des variables $x_1, \dots, x_n$ sont locales (liées). Ces expressions sont d'un emploi courant en mathématique, mais leur traitement logique rigoureux est généralement omis dans les textes d'introduction à la théorie des ensembles. Ceux-ci se bornent

aux collections *simples* où $E(x_1, \dots, x_n)$ se réduit à une variable x_i . L'une des raisons de leur traitement dans ce livre est que ces expressions appartiennent à certains langages de spécification en informatique [13].

Fonctions (Chap. 15). À part un exposé de la matière standard — définition des fonctions et applications (nous faisons une distinction), image d'un élément, égalité de fonctions, composition de fonctions, applications injectives, surjectives, bijectives — ce chapitre contient un traitement rigoureux des *opérateurs d'abstraction fonctionnelle*, pour lesquels la notation du « lambda-calcul » de Church [29] est utilisée. La motivation est d'une part de bien mettre en évidence les abus de langage usuels en mathématique dans ce domaine et d'autre part de présenter une notation qui est à la base des langages de programmation dits fonctionnels [30, 31].

Annexe (Chap. 16). Cette annexe contient les démonstrations formelles de tous les théorèmes de théorie des ensembles énoncés sans démonstration dans les chapitres précédents. Cette démonstration est considérée comme un *exercice* dont la solution est fournie dans cette annexe.

2

Langages formels

2.1 Expressions d'un langage

Avant d'aborder la description des langages formels du premier ordre, nous allons parler brièvement de la notion de langage en général, afin de permettre au lecteur de se familiariser avec quelques concepts généraux : les expressions d'un langage, leur syntaxe, leur sémantique, et les deux espèces d'expressions principales — termes et propositions — des langages déclaratifs.

Le sens général du mot « langage » est très large, puisqu'on l'utilise à propos de phénomènes aussi divers que celui des langues naturelles parlées ou écrites (le français, le chinois, etc.) et celui des langages de programmation. On l'emploie aussi dans des sens imagés : le langage de la musique, le langage des gestes, etc. Cependant toutes ces significations du mot « langage » ont des traits communs, lorsqu'on se borne à considérer l'aspect extérieur de ces phénomènes, leur pure forme.

Lorsqu'on cherche à décrire un langage, considéré de l'extérieur, celui-ci apparaît toujours comme une certaine classe d'*expressions*, qui sont des suites de signaux visuels ou sonores, ou des suites de symboles enregistrés sous forme électromagnétique dans la mémoire d'un ordinateur.

Les expressions d'un langage sont toujours construites par assemblage de signes ou signaux élémentaires, pris dans un répertoire déterminé, qui constituent les « atomes » du langage. Ce sont par exemple les lettres de l'alphabet d'une langue naturelle écrite, ou les phonèmes (sons élémentaires) de la langue parlée correspondante, ou encore les caractères d'un langage de programmation. Nous appellerons *symboles* d'un langage ces signaux élémentaires, quelle que soit leur nature physique et leur forme.

Dans le cas d'un phénomène aussi complexe que celui d'une langue naturelle, il y a plusieurs façons de décomposer les expressions en éléments, c'est-à-dire que la décomposition peut être plus ou moins fine. Par exemple, on peut considérer les mots entiers, et non les lettres, comme étant les atomes de la langue, et alors ses expressions sont des suites de mots. Dans ce cas, ce

sont les mots eux-mêmes qui constituent les « symboles » du langage. Il faut donc distinguer la notion de langage en tant que phénomène et la description qu'on en donne, description qui inclut le choix des atomes (symboles) que l'on considère. Dans ce qui suit, le mot « langage » est pris dans le sens de langage *décrit*, c'est-à-dire que ce sens inclut le choix des symboles pris comme atomes.

La donnée d'un langage commencera donc toujours par celle de son répertoire de symboles. Elle se poursuivra par la donnée des règles à observer lorsqu'on assemble des symboles pour former des expressions. Car ce qui est typique d'un langage, c'est qu'il n'autorise que certaines combinaisons de symboles et ce sont ces combinaisons que l'on appelle les expressions du langage.

Nous ferons ici une hypothèse générale, à savoir que les expressions d'un langage sont toujours formées de symboles placés les uns *après* les autres, dans un ordre linéaire (en ligne), dans le temps ou dans l'espace, et qu'une expression est analysée (lue, écoutée) en parcourant la suite de ses symboles dans cet ordre. C'est à cet ordre des symboles que l'on fait allusion en disant qu'une expression est une *suite* de symboles. On peut parler plus précisément d'une suite *finie*, car une expression comporte toujours un nombre fini de symboles. Le terme de *séquence* de symboles est aussi utilisé. Il a le même sens que celui de « suite finie ». Nous admettons donc que les expressions d'un langage sont toujours des séquences de symboles et nous utiliserons la notation

$$s_1 s_2 \cdots s_n$$

pour désigner une expression composée de n symboles ($n = 1, 2, 3, \dots$).

Cette hypothèse d'une forme séquentielle des expressions peut paraître trop restrictive à première vue. On peut penser en effet que la notion générale de langage devrait inclure toutes les formes de communication de l'information, notamment celles qui utilisent des images, par exemple des schémas techniques et autres diagrammes dont les éléments ne sont pas disposés dans un ordre linéaire. Cependant les mots perdent leur utilité lorsqu'on élargit trop leur sens. En ce qui concerne celui du mot « langage », nous pensons qu'il est préférable de le limiter aux formes d'expression séquentielles. Cette restriction se justifie particulièrement dans le domaine de l'informatique, où l'on utilise des langages de description de schémas, de formes, d'images, qui décrivent ces données graphiques au moyen d'expressions séquentielles. Les données graphiques elles-mêmes ne sont pas perçues comme des objets de nature linguistique.

Notons que dans une expression $s_1 s_2 \cdots s_n$ un même symbole peut figurer plusieurs fois, à savoir que les symboles s_i et s_j qui figurent en i -ième et j -ième position, avec $i \neq j$, peuvent être identiques. On parle alors de deux *occurrences* distinctes du même symbole dans l'expression, l'une en i -ième position, l'autre en j -ième position. Par exemple, la lettre « e » possède deux occurrences dans le mot « expression » : une occurrence en première position et une occurrence en cinquième position. Il faut donc distinguer, en principe, les symboles d'une expression, de leurs occurrences dans cette expression. Par exemple, on dit couramment que le mot « expression » est un mot de dix lettres. Mais ce nombre

est celui des *occurrences* de lettres, car seules huit lettres figurent dans ce mot: les lettres e, x, p, r, s, i, o, n. Lorsque nous désignons une expression par $s_1s_2 \cdots s_n$, il est clair que n est le nombre d'*occurrences* de symboles dans celle-ci, alors que si l'on compte seulement le nombre de symboles distincts qui figurent dans l'expression, on obtient un nombre m qui est en général inférieur à n . Le lecteur peut prendre note en passant du vocabulaire que nous utilisons ici et qui reviendra constamment dans la suite. En particulier, nous disons qu'un symbole s *figure* dans une expression, s'il y a au moins une occurrence de s dans celle-ci.

2.2 Syntaxe et sémantique

Les expressions d'un langage sont des séquences de symboles de ce langage. Mais ce ne sont pas n'importe quelles séquences. Par exemple, on sait bien que n'importe quelle suite de mots français n'est pas une phrase admissible, de même que n'importe quelle suite de caractères d'un langage de programmation ne constitue pas un programme. La définition d'un langage comporte donc des règles qui déterminent quelles séquences de symboles sont des expressions du langage. Les séquences de symboles qui ne sont pas conformes à ces règles ne font pas partie des expressions du langage.

Pour qu'un langage puisse être dit *formel*, il faut que ces règles soient non seulement absolument précises, mais encore qu'il existe un procédé mécanique (un algorithme) permettant de déterminer si une séquence quelconque de symboles est une expression du langage ou non, c'est-à-dire si elle est construite conformément aux règles de ce langage. C'est le cas des langages de programmation.

Les langues naturelles ne sont évidemment pas des langages formels. Les règles que contiennent les livres de grammaire ne décrivent ces langues qu'approximativement et elles ne sont pas suffisamment précises pour qu'on puisse en tirer un procédé mécanique de reconnaissance des expressions. Le développement de langages formels qui soient de bonnes approximations des langues naturelles est un domaine de recherche en informatique: celui du traitement automatique du langage naturel — traduction automatique de documents, dialogue homme-machine en langage naturel, etc.

L'idée principale que nous voulons introduire dans ce paragraphe est la distinction entre deux aspects d'un langage: sa *syntaxe* d'une part et sa *sémantique* d'autre part. Cette distinction correspond à celle que l'on peut faire entre la *forme* et le *sens* (ou *signification*) des expressions.

Cette distinction est bien claire, même s'il peut être difficile de définir en toute généralité les concepts de forme et de sens. Elle s'impose d'elle-même à notre esprit sitôt que deux expressions différentes ont le même sens — en particulier lorsqu'elles désignent le même objet. Car on distingue alors entre elles une différence qui est de forme seulement.

« L'Angleterre », « le Royaume-Uni », « la patrie de Shakespeare et des Beatles » et pour certains « la perfide Albion », sont autant d'expressions utilisées pour désigner le même pays. De même, deux expressions arithmétiques de formes différentes, par exemple $a(x + y)$ et $ax + ay$, peuvent représenter le même nombre, et avoir ainsi le même sens.

Un langage possède une *syntaxe* bien définie lorsqu'on peut caractériser ses expressions par des règles de construction qui ne font pas intervenir le sens de ces expressions. C'est le cas des langages formels et notamment des langages de programmation. Les règles qui déterminent si une suite de caractères d'un tel langage est admise comme étant un programme ne demandent aucune compréhension du sens de ce texte. Elles peuvent être très complexes, mais dans leur essence elles ne sont pas différentes d'une règle qui admet la séquence de caractères « $(x + y)$ » comme une expression arithmétique, mais qui n'admet pas la séquence « $(x +)y$ », à cause d'une parenthèse mal placée. Les règles générales sur le parenthésage des expressions arithmétiques peuvent être formulées d'une manière qui ignore complètement le sens de ces expressions.

Pour les langues naturelles, on ne connaît pas de système de règles de syntaxe indépendantes de la signification des expressions. Par exemple, en français, on connaît bien les règles sur la terminaison des adjectifs au pluriel et sur celle des verbes à la troisième personne du pluriel, et l'on peut reconnaître que les phrases suivantes respectent ces règles :

Les bouchers sales la tranchent.

Les bouchers salent la tranche.

Encore faut-il reconnaître un adjectif dans le mot « sales » de la première phrase et un verbe dans le mot « salent » de la seconde. Pour cela, il faut comprendre le *sens* de ces phrases.

La définition d'un langage doit donc comporter, en général, une partie syntaxique et une partie sémantique. Nous ne pouvons pas entrer dans une discussion approfondie de ces sujets et notamment des différentes manières de décrire la syntaxe d'un langage formel d'une part, et sa sémantique d'autre part. Nous nous limiterons à quelques brèves indications.

Syntaxe d'un langage formel. En ce qui concerne la syntaxe des langages formels, nous nous bornerons à rappeler au lecteur la manière dont est définie celle des langages de programmation, en admettant que cette manière lui est plus ou moins familière. C'est de cette manière aussi que nous définirons plus loin la syntaxe — beaucoup plus simple, d'ailleurs — des langages du premier ordre.

Comme nous l'avons mentionné plus haut, la définition d'un langage formel commence par la donnée de ses symboles. Ceux des langages de programmation correspondent plus ou moins aux caractères qui se trouvent sur les touches d'un clavier « standard ». On définit ensuite des classes de symboles, auxquelles on donne un nom, en énumérant explicitement les symboles qu'elles contiennent. Par exemple, on définit les *chiffres* (en anglais *digits*) comme étant les symboles

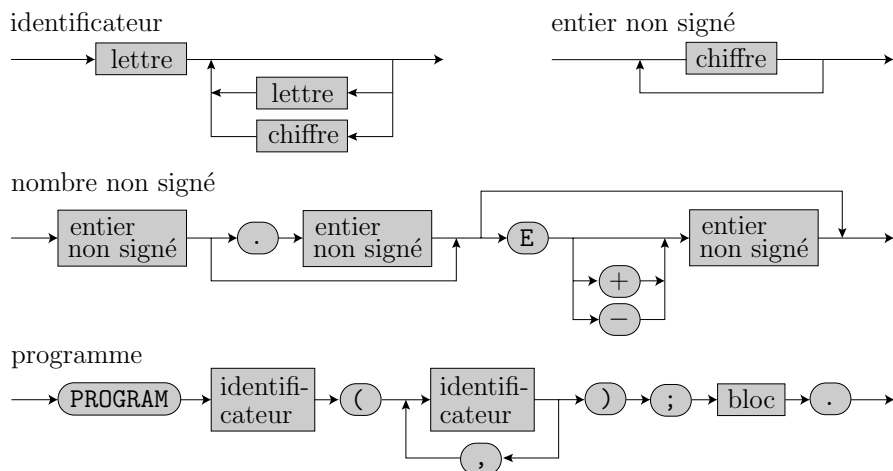


Figure 2.1 Les expressions définies par un diagramme syntaxique sont celles que l'on obtient lorsque, en parcourant le diagramme depuis son entrée à gauche jusqu'à sa sortie à droite, en suivant une route du diagramme (suivre les flèches), on écrit successivement un symbole ou une expression de chacune des espèces nommées dans les boîtes rectangulaires que l'on traverse, ainsi qu'une occurrence de chacun des symboles rencontrés dans les boîtes rondes.

0, 1, 2, ..., 9 et les *lettres* comme étant les symboles A, B, ..., Z et a, b, ..., z. On définit ensuite, de proche en proche, différentes espèces ou classes d'expressions. La première est souvent celle des *identificateurs*, définis par exemple comme étant les expressions dont chaque symbole est une lettre ou un chiffre et dont le premier symbole est une lettre. Cette classe d'expressions est représentée par le premier diagramme de la figure 2.1, que l'on appelle un *diagramme syntaxique*. Nous admettons que le lecteur sait interpréter ce genre de diagramme (une brève explication est donnée avec la figure). Le deuxième et le troisième diagramme définissent deux autres catégories d'expressions, celle des *entiers non signés* et celle des *nombres sans signe*. Le langage considéré est le langage PASCAL. Le quatrième diagramme définit l'espèce d'expression *principale*, celle des *programmes*. Toutes les autres espèces sont des espèces *auxiliaires*. La définition des programmes se réfère à une autre espèce, celle des *blocs*, et la définition des blocs nécessite à son tour celle d'une douzaine d'autres espèces auxiliaires.

Ce style de définition syntaxique d'un langage est ce qu'on appelle une *grammaire indépendante du contexte*. C'est le genre de grammaire le plus utilisé pour les langages de programmation, en raison de sa simplicité. Mais en général il ne permet pas de définir complètement la syntaxe de ces langages, et l'on doit compléter la grammaire par des règles imposant des restrictions supplémentaires. Il existe par ailleurs d'autres genres de grammaires formelles.

Sémantique d'un langage formel. Considérons maintenant le problème de définir la sémantique d'un langage, c'est-à-dire le sens de ses expressions. Nous commencerons aussi par le cas d'une langue naturelle — disons le français. S'il s'agit d'apprendre cette langue à une personne d'une autre langue — par exemple polonaise — la solution est évidente: il suffit de traduire la première langue dans la seconde, c'est-à-dire de donner les instruments qui permettent cette traduction, en particulier un dictionnaire bilingue. On peut dire que la solution est triviale dans ce cas, même si ce n'est pas une petite affaire que d'apprendre une langue étrangère.

Plus intéressant est le problème de la définition de la sémantique d'une langue naturelle pour les usagers de cette même langue. Il existe aussi, pour ces derniers, des dictionnaires qui définissent le sens des mots, mais ce sont des dictionnaires monolingues et il est clair qu'il faut déjà comprendre un peu cette langue pour pouvoir les utiliser. À côté des dictionnaires, il faut mentionner d'ailleurs tous les ouvrages et documents qui définissent un vocabulaire spécialisé — scientifique et technique par exemple. Ce qu'il faut remarquer ici surtout, c'est qu'à partir d'un certain vocabulaire de base, il est possible de définir le sens d'une langue naturelle *dans cette langue elle-même*. Il s'agit aussi d'une forme de traduction. Seulement la langue n'est pas traduite ici dans une langue «étrangère», mais dans un «sous-ensemble» d'elle-même. Ce sous-ensemble s'accroît d'ailleurs pour l'utilisateur, au fur et à mesure des nouvelles connaissances qu'il acquiert. Ce mécanisme est particulièrement visible dans les ouvrages scientifiques, dont le vocabulaire spécialisé est défini soigneusement, de proche en proche, chaque terme nouveau étant défini au moyen de termes définis précédemment.

Nous laissons de côté la question de savoir comment l'individu parvient à ce stade de compréhension de la langue naturelle à partir duquel celle-ci peut être utilisée pour sa propre définition. Ceci est un sujet qui sort du cadre de cet exposé — et aussi de notre compétence.

La sémantique des langages formels que nous utiliserons sera donnée de la même manière. Il y aura toujours certaines expressions de base dont le sens sera supposé compris intuitivement et le sens de toutes les autres expressions sera défini par un mécanisme de traduction dans le sous-langage composé des seules expressions de base.

C'est aussi l'approche des mathématiques, dont le langage utilise un grand nombre de symboles — non seulement les signes spéciaux tels que 0 , $+$, $=$, $\sqrt{}$, $<$, mais encore tout le vocabulaire de mots spécifiques: fonction, matrice, limite, graphe, rationnel, etc. Le sens de presque tous ces symboles et de ce vocabulaire est donné par une succession de définitions qui permettent de les traduire tous au moyen d'un très petit nombre de symboles de base dont le sens est compris intuitivement. Ces symboles de base seront par exemple le symbole d'égalité ($=$) et le symbole d'appartenance ($\in$) de la théorie des ensembles. Le sens des expressions $x = y$ et $x \in y$ ne sera pas défini formellement, c'est-à-

dire par une traduction de celles-ci en d'autres expressions mathématiques. Il sera simplement expliqué, si besoin est, dans la langue naturelle. Notamment, pour l'expression $x \in y$, l'explication se basera sur les notions d'ensemble et d'éléments d'un ensemble. Ces notions sont indéfinissables formellement, c'est-à-dire qu'on ne peut pas les ramener à d'autres notions mathématiques. On ne peut qu'en donner l'*intuition* au moyen d'exemples et de commentaires.

Par contre, à titre d'exemple, le sens du symbole d'inclusion ($\subset$) d'une proposition de la forme $x \subset y$ (x est un sous-ensemble de l'ensemble y) sera défini *formellement*, c'est-à-dire par traduction de cette expression au moyen du symbole de base $\in$. On dira en effet que la proposition $x \subset y$ signifie exactement ceci: tout élément de x est élément de y . Or cela peut s'exprimer au moyen de symboles logiques et du symbole d'appartenance $\in$ en disant:

$$\text{quel que soit } a, \text{ si } a \in x \text{ alors } a \in y, \quad (1)$$

ou encore $\forall a(a \in x \Rightarrow a \in y)$, en utilisant les notations que nous introduirons plus loin. L'expression (1) donne le sens de la proposition $x \subset y$.

Cas des langages de programmation. Les langages de programmation sont des langages formels. Leur syntaxe est définie avec une parfaite précision et il existe des algorithmes permettant de vérifier de façon mécanique si une expression donnée est conforme à cette syntaxe. Ce sont les algorithmes d'analyse lexicale et syntaxique utilisés par les compilateurs.

Pour la plupart des utilisateurs, la sémantique des langages de programmation n'est pas définie de manière formelle. Le sens des expressions d'un langage leur est expliqué de manière informelle. Une partie des notations de ces langages est proche des notations mathématiques usuelles et il n'est pas jugé nécessaire de définir ce sens, supposé connu. Mais à côté de ces notations, les langages de programmation introduisent des notations qui sortent du langage mathématique usuel. Il y a des opérations d'affectation (assignation) qui changent la valeur des expressions tout en se mêlant aux opérations arithmétiques traditionnelles; l'évaluation des expressions peut provoquer des « effets de bord »; il y a des instructions dites de contrôle (répétitions, branchements, itérations, etc.).

Le sens de ces notations est expliqué dans les manuels de programmation, mais en général de manière informelle, en recourant surtout à des exemples. Cependant, en voyant la précision de la définition syntaxique de ces langages, on peut se demander si le sens de ces notations ne pourrait pas faire aussi l'objet de définitions précises, qui les ramèneraient toutes à des notations de base connues, comme cela se fait en mathématique.

L'étude de cette question constitue le sujet de la *sémantique formelle* des langages de programmation, sur lequel il existe une abondante littérature. La question s'est révélée beaucoup plus complexe qu'on aurait pu le penser, et ce sujet demeure encore, actuellement, un sujet de recherche en informatique théorique. On connaît différentes manières de définir formellement la

sémantique des constructions usuelles des langages de programmation, mais la définition formelle complète d'un langage « réel » (Ada, Lisp, C++, etc.) reste une tâche formidable. Certains y voient un défaut de ces langages et attendent de cette étude des progrès importants dans la conception même des langages de programmation.

2.3 Termes et propositions

Les expressions des langages formels que nous allons étudier seront divisées en deux classes: une classe d'expressions appelées *termes* et une classe d'expressions appelées *propositions*.

On peut déjà distinguer ces deux types d'expressions dans le langage naturel ou parmi les expressions mathématiques. Nous allons le faire ici, afin de permettre au lecteur de se familiariser avec cette distinction. Pour le moment, celle-ci se basera sur le sens des expressions. Il s'agira donc d'une distinction sémantique. Par contre plus loin, dans le cadre des langages formels, les deux espèces d'expressions pourront être distinguées d'après leur forme seulement. La distinction sera alors syntaxique.

Termes. Les *termes* sont les expressions qui *désignent des objets*, ce dernier mot étant pris dans un sens très large: objets concrets ou abstraits, animés ou inanimés, etc. Par exemple, les six expressions suivantes, occupant chacune une ligne, constituent six termes:

la Terre,
 $\sqrt{2}$,
 $x + 3$,
le Président des États-Unis,
le plus petit entier k tel que $2k > 7$,
l'ensemble des nombres complexes de module inférieur à 1.

On peut dire qu'un terme est une expression qui peut être prise comme *nom* d'un objet, même si cet objet n'est pas déterminé, comme dans le cas du terme $x + 3$ ci-dessus.

Propositions. Les *propositions* sont les expressions qui *énoncent un fait* — cet énoncé pouvant être vrai ou faux. Par exemple, chacune des six lignes suivantes contient une proposition:

la Terre tourne,
 $\sqrt{2} < 1$,
 $x + 3$ est divisible par 2,
si n est pair alors $n + 1$ est impair,
tout nombre pair supérieur à 3 est la somme de deux nombres premiers,
il existe un homme qui est infallible.

Le fait que l'on puisse qualifier une proposition de vraie ou de fausse — sans nécessairement pouvoir décider lequel des deux convient — est ce qui distingue essentiellement les propositions des termes, du point de vue sémantique. Pour des expressions telles que « *la Terre* » ou « $x + 3$ », prises dans leur sens usuel, on ne se posera jamais la question de savoir s'il s'agit d'expressions vraies ou fausses, car cette question n'aurait pas de sens. Par contre, la question est tout à fait sensée à propos des expressions « *la Terre tourne* » et « $x + 3$ est divisible par 2 ».

Dans les langages formels que nous considérerons, tout comme dans la langue naturelle, les propositions seront l'espèce d'expressions principale. Les termes sont une espèce auxiliaire, dont on se sert pour construire des propositions. On le voit bien dans l'exemple des propositions « *la Terre tourne* » et « $x + 3$ est divisible par 2 », dont la construction utilise les termes « *la Terre* », « $x + 3$ », et « 2 ».

Nous appelons *langage déclaratif* tout langage possédant une espèce principale d'expressions auxquelles peuvent s'appliquer les qualificatifs « vrai » et « faux » et que nous appelons *propositions*. D'autres appellations possibles seraient : *énoncés*, *assertions*, *jugements*, *déclarations*, etc. Dans ce sens, les langages formels que nous étudierons seront des langages déclaratifs. Dans ce sens également, la plupart des langages de programmation ne sont pas des langages déclaratifs. Leurs expressions de l'espèce principale, à savoir les programmes, ne sont pas des affirmations dont il s'agit de savoir si elles sont vraies ou fausses, mais sont des ordres à exécuter, ou des termes à évaluer.

Vérité et contexte. Si la vérité et la fausseté sont des propriétés caractéristiques (du point de vue sémantique) de l'espèce d'expression que nous appelons *propositions*, il faut bien voir qu'en général la valeur de vérité d'une proposition dépend du contexte dans lequel elle est lue. Cela est particulièrement évident dans le cas d'une proposition telle que « $x + 3$ est divisible par 2 », contenant une inconnue x . Cette proposition sera vraie dans un contexte où x désigne un nombre impair, et fausse dans un contexte où x désigne un nombre pair.

Lorsque nous lisons la phrase « *la Terre tourne* », nous admettons naturellement que le terme « *la Terre* » désigne notre planète et nous considérons cette phrase comme vraie. Mais la phrase en question pourrait être interprétée différemment, par exemple par un programmeur qui aurait appelé « *la Terre* » un programme qu'il est en train de développer et qui justement ne « tourne » pas encore. Dans un texte où il serait question de notre planète et des œuvres du programmeur en question, le sens d'une occurrence du terme « *la Terre* » ou du verbe « tourner » dépendra du contexte de cette occurrence, lequel devra indiquer clairement de quelle Terre il est question à cet endroit.

Ce sont là des remarques banales, mais il sera important de les avoir présentes à l'esprit lorsque nous aborderons ces notions dans le cadre de la logique formelle.

2.4 Langage et métalangage

Nous avons parlé déjà plusieurs fois des langages formels que nous « allons étudier ». Lorsqu'un auteur parle ainsi « d'étudier » un sujet, il veut dire qu'il va faire un *exposé* sur ce sujet, autrement dit tenir un *discours* sur ce sujet. Dans ce sens du mot « étudier », on peut dire qu'un langage est un sujet d'étude qui présente une particularité unique, parmi tous les sujets d'étude possibles. En effet, l'étude d'un sujet quelconque nécessite elle-même un langage, à savoir celui que l'on utilise pour parler du sujet étudié, pour en décrire les propriétés. Lorsque ce sujet est lui-même un langage, on se trouve en présence de deux langages qu'il ne faut pas confondre : le langage $\mathcal{L}$ que l'on *étudie*, autrement dit le langage *dont on parle*, pour en décrire par exemple la syntaxe et ses propriétés, et le langage que l'on *utilise* pour effectuer cette description et que l'on appelle le *métalangage*.

Une situation dans laquelle cette distinction est particulièrement claire est celle de l'enseignement d'une langue naturelle, disons l'allemand, aux usagers d'une autre langue, disons le français. Un manuel d'allemand pour écoliers francophones utilise le français comme métalangage pour décrire à ses lecteurs la structure de l'allemand.

La distinction est beaucoup moins claire dans le cas d'un livre de grammaire d'une langue naturelle écrit dans cette langue elle-même, par exemple un livre en français sur la grammaire du français, à l'usage de ceux qui comprennent déjà cette langue. Le métalangage s'identifie ici au langage étudié. Cependant le même langage joue ici deux *rôles* différents et l'on peut toujours dire dans lequel de ces rôles apparaissent les expressions du livre. Par exemple si l'on y rencontre les phrases

manger est un verbe transitif et intransitif (1)

manger est un acte quotidien (2)

il sera clair que la première appartient au métalangage — c'est une affirmation sur le mot « manger » lui-même, donc sur le langage — alors que la seconde est une affirmation sur l'acte et non sur le mot — c'est une phrase du langage sur un sujet qui n'est pas le langage lui-même.

Ces remarques ont pour but d'inviter le lecteur à bien distinguer, plus loin, les expressions des langages formels dont nous décrirons la syntaxe et les propriétés, des expressions du métalangage que nous utiliserons pour cette description.

La raison de cette mise en garde est que le risque de confusion existe. Il se trouve en effet qu'un langage formel est un « objet » dont les propriétés sont rigoureusement définies, comme celles d'un objet mathématique, à tel point que l'on peut formuler des « théorèmes » sur cet objet. Or les langages formels dont il est question dans un ouvrage de logique servent *eux-mêmes* à formuler des théorèmes sur les objets de telle ou telle espèce, par exemple sur les nombres réels, sur les ensembles, sur les programmes. On se trouve donc en présence

de deux niveaux de théorèmes: les théorèmes formulés *dans* le langage formel, et les théorèmes *sur le langage formel*, formulés dans le métalangage. On peut parler pour ces derniers de *métathéorèmes*. Une comparaison peut être faite avec les phrases (1) et (2) ci-dessus: la première peut être considérée comme un « métathéorème » *sur* le langage qu'est le français; la seconde comme un « théorème » *en* français, donc *dans* le langage, sur la vie des animaux.

3

Symboles d'un langage du premier ordre

3.1 Exemple mathématique

La syntaxe des langages du premier ordre peut être définie de manière très succincte et cela sera fait au chapitre 4. Mais pour se familiariser avec ces langages, il est utile d'en voir d'abord des exemples et d'apprendre à connaître les différentes *espèces de symboles* qu'ils utilisent. Il y en a six. Nous allons les présenter au moyen de deux exemples. Dans ce paragraphe, nous considérons une proposition exprimant une propriété bien connue des nombres réels, et dans le suivant, certaines propositions décrivant quelques propriétés d'une base de données.

Formalisation d'une proposition. La proposition sur les nombres réels que nous allons formaliser est la suivante:

Si le produit de deux nombres réels quelconques est nul, alors l'un de ces nombres est nul. (1)

La manière courante d'exprimer que le produit de deux nombres x , y est nul est évidemment d'écrire

$$x \cdot y = 0. \quad (2)$$

Pour affirmer que l'un des deux nombres x , y est nul, nous écrivons

$$x = 0 \text{ ou } y = 0, \quad (3)$$

en utilisant le symbole logique « ou ». En logique, et en particulier en mathématique, le sens du mot « ou » est *inclusif*, c'est-à-dire que l'assertion (3) signifie que l'un *au moins* des deux nombres x , y est nul et cela inclut le cas où les deux sont nuls. L'assertion (1) combine les propositions (2) et (3) sous la forme:

$$\text{Si } (x \cdot y = 0) \text{ alors } (x = 0 \text{ ou } y = 0). \quad (4)$$

Au lieu des mots *si ... alors ...*, nous utiliserons le symbole $\Rightarrow$, en écrivant

$$(x \cdot y = 0) \Rightarrow (x = 0 \text{ ou } y = 0). \quad (5)$$

Le symbole $\Rightarrow$ est appelé le *symbole d'implication*. Il se lit « implique », ou « entraîne », mais on peut aussi lire (5) en utilisant les mots *si ... alors ...* comme dans (4). Une proposition de la forme (5), qui combine deux propositions au moyen du symbole d'implication est appelée elle-même une *implication*. Le symbole d'implication est son symbole principal.

Il reste à formaliser un mot essentiel de l'assertion (1), à savoir le mot « quelconques ». Ce mot exprime que l'implication (5) n'a pas lieu seulement pour deux nombres x et y particuliers, mais pour n'importe quels nombres x et y . On formalise ceci en écrivant

$$\forall x \forall y [(x \cdot y = 0) \Rightarrow (x = 0 \text{ ou } y = 0)]. \quad (6)$$

Le signe $\forall$ se lit « pour tout ». Il s'agit d'un A renversé, la première lettre du mot anglais *All*. Donc la proposition (6) se lit :

pour tout x pour tout y [...].

Les combinaisons de symboles $\forall x$, $\forall y$ sont appelées des *quantificateurs universels*. Nous verrons plus loin les quantificateurs existentiels $\exists x$, $\exists y$ — il existe un x , il existe un y .

Nous présentons la proposition (6) comme une traduction de la proposition (1) dans un langage formel du premier ordre — mais bien entendu, le lecteur qui n'a pas l'habitude des notations logiques n'est nullement censé approuver cette traduction, c'est-à-dire reconnaître que (1) et (6) ont le même sens. Notre but n'est pas de l'en convaincre immédiatement, mais seulement de présenter un symbolisme. Le sens des symboles logiques s'acquiert par l'usage et par l'étude de la logique. C'est un langage qu'il faut « apprendre à parler ».

Commentaire. Il ne faut pas s'attendre d'ailleurs à ce que l'on puisse *prouver* l'équivalence sémantique des expressions (1) et (6) — le fait qu'elles aient le même sens. Il n'existe pas de document officiel donnant les règles de traduction de la langue informelle des mathématiques, à laquelle appartient la phrase (1), dans un langage formel. De manière générale, la traduction de textes d'une langue naturelle ou informelle dans un langage formel quelconque est une opération dont on ne peut pas prouver qu'elle soit correcte. Le génie logiciel moderne, par exemple, recommande l'utilisation de langages de spécification formelle pour exprimer les tâches et les buts des systèmes à développer, afin de ne pas s'engager dans de grands travaux de programmation avec des objectifs flous. Mais un projet de développement de logiciel commence toujours par un document informel, en langage naturel augmenté éventuellement de schémas divers, décrivant les objectifs à atteindre. Le cycle de développement doit donc commencer par une formalisation de ce document. Il en résulte un document de spécification formelle, auquel le logiciel développé par la suite peut être comparé selon des méthodes rigoureuses. Mais il n'est pas possible

de prouver que la spécification formelle soit conforme au document informel initial. Cette conformité ne peut être que *reconnue* par ceux qui comprennent les deux langages et les deux documents en question.

Il en va de même des expressions (1) et (6), et plus généralement du passage de la langue mathématique usuelle à sa traduction dans un langage formel. Seules l'habitude et une certaine maîtrise des deux langages permettent de reconnaître la justesse de cette traduction. Le lecteur de cet ouvrage n'est pas censé posséder cette maîtrise, au stade présent de sa lecture. L'auteur veut croire, cependant, qu'il en ira différemment à la fin de celle-ci.

Par ailleurs, lorsqu'on maîtrise le rapport entre le langage mathématique usuel et un langage formel, ce rapport s'inverse en quelque sorte. Ce n'est plus le langage formel qui est considéré comme une « traduction » du langage usuel, mais bien le *contraire*. Le langage usuel est perçu comme une traduction — souvent assez libre — du langage formel. Celui-ci devient la référence.

Les six espèces de symboles. Considérée comme proposition d'un langage du premier ordre, la proposition (6) contient les six espèces de symboles d'un tel langage, à savoir :

x, y	Variables individuelles
0	Constante individuelle
$\cdot$	Symbole fonctionnel
$=$	Symbole relationnel
ou, $\Rightarrow$	Connecteurs logiques
$\forall x, \forall y$	Quantificateurs.

L'expression (6) contient encore des parenthèses, mais ce sont des symboles auxiliaires que nous mettons à part. Nous verrons qu'il y a une manière d'écrire les expressions en se passant de parenthèses. Le mot « ou » est considéré comme un seul symbole. Au lieu de ce mot, on emploie aussi le signe « $\vee$ ». Il y aura trois autres connecteurs logiques : le symbole de négation « non », pour lequel on utilise aussi le signe « $\neg$ » ; le symbole de conjonction « et » (signe $\wedge$) ; le symbole d'équivalence ou double implication « $\Leftrightarrow$ », pour lequel on utilise aussi différentes tournures de mots, par exemple « si et seulement si ».

Un langage du premier ordre comporte donc six espèces de symboles : des variables individuelles, des constantes individuelles, des symboles fonctionnels, des symboles relationnels, des connecteurs logiques, des quantificateurs. Le lecteur est invité à apprendre rapidement cette terminologie, en se basant sur l'exemple qui précède et sur les quelques commentaires d'ordre sémantique qui suivent.

Sémantique. Un langage du premier ordre sert à parler d'une certaine classe d'objets constituant ce que l'on appelle le *domaine du discours*, ou *l'univers du discours*, ou encore *le domaine ou univers de l'interprétation*. Dans l'exemple qui précède, ces objets sont les nombres réels. Les *variables individuelles* sont des symboles utilisés pour désigner des objets arbitraires du domaine considéré.

Ce sont des noms d'objets, mais qui ne sont pas associés à des objets fixés. Par contre, les *constantes individuelles*, comme le symbole 0 de l'exemple, désignent des objets bien déterminés, et leur sens ne varie pas.

Les *symboles fonctionnels*, comme le signe de multiplication de l'exemple, représentent des *opérations* de construction d'objets du domaine, au moyen d'autres objets. Par exemple le terme $x \cdot y$ désigne le nombre construit en prenant deux nombres x et y et en leur appliquant l'opération de multiplication. Le symbole de la multiplication « $\cdot$ » est appelé un symbole fonctionnel *bin aire*, parce qu'il s'applique à deux arguments (x et y dans l'exemple). Il peut y avoir en général des symboles fonctionnels unaires (un seul argument), binaires (deux arguments), ternaires (trois arguments), et plus généralement n -aires (n arguments).

Les *symboles relationnels*, comme le symbole d'égalité, s'appliquent aussi à un nombre déterminé d'arguments, et il peut y avoir des symboles relationnels unaires, binaires, ternaires, etc. Les symboles relationnels binaires servent à exprimer des propriétés déterminées de certains couples d'objets. Par exemple, l'expression $x = y$ exprime une propriété bien claire (du moins nous voulons l'admettre) des objets désignés par x et y : celle d'être égaux. D'autres symboles relationnels binaires utilisés couramment dans un langage parlant des nombres réels sont les symboles d'inégalité $<$, $\leq$. Nous parlerons plus loin de symboles relationnels n -aires, ou comme on dit, *d'arité n* autre que 2. « L'arité » d'un symbole fonctionnel ou d'un symbole relationnel est le nombre fixe d'arguments auxquels il s'applique dans les expressions. Un symbole n -aire est un symbole d'arité n .

Les connecteurs logiques et les quantificateurs servent à former des propositions complexes en combinant des propositions simples. Leur sens sera discuté plus loin. Ces symboles appartiennent à *tous* les langages du premier ordre, alors que les constantes individuelles, les symboles fonctionnels, et les symboles relationnels peuvent changer d'un langage à un autre. Nous admettrons aussi que les variables individuelles, dont nous parlerons plus loin, sont les mêmes dans tous les langages du premier ordre que nous considérons. Ce sont donc seulement les trois espèces de symboles suivantes qui font la « personnalité » d'un langage du premier ordre: ses constantes individuelles, ses symboles fonctionnels, et ses symboles relationnels.

Si nous pouvons appeler l'expression (6) une *proposition* (§2.3), c'est en nous basant sur sa signification: elle est la traduction formelle d'une *assertion* (1) et il est naturel de penser qu'elle possède une valeur de vérité (vrai, faux), même si l'on ne connaît pas cette valeur. La syntaxe exacte des langages du premier ordre sera précisée au chapitre suivant, au moyen de diagrammes syntaxiques. Nous pourrions alors reconnaître dans l'expression (6) une proposition d'après sa forme seulement. Pour le moment, en nous basant toujours sur le sens des expressions, nous pouvons distinguer dans l'expression (6) diverses parties ou sous-expressions qui sont elles-mêmes des propositions, et d'autres

qui sont des termes. Nous distinguons ainsi quatre termes :

$$0, \quad x, \quad y, \quad x \cdot y,$$

et sept propositions :

$$\begin{aligned} x &= 0, \\ y &= 0, \\ x \cdot y &= 0, \\ (x = 0 \text{ ou } y = 0), \\ (x \cdot y = 0) &\Rightarrow (x = 0 \text{ ou } y = 0), \\ \forall y [(x \cdot y = 0) &\Rightarrow (x = 0 \text{ ou } y = 0)], \\ \forall x \forall y [(x \cdot y = 0) &\Rightarrow (x = 0 \text{ ou } y = 0)]. \end{aligned}$$

Le lecteur remarquera que les variables et constantes individuelles sont elles-mêmes des termes (réduits à un seul symbole). Les trois premières propositions de cette liste (trois égalités) sont appelées des propositions *atomiques*, alors que les suivantes sont des propositions *composées*. Une proposition atomique est toujours formée au moyen d'un symbole relationnel et de termes, alors qu'une proposition composée est formée au moyen d'un symbole logique (connecteur ou quantificateur) et d'autres propositions. Nous avons inclus dans cette liste, en dernier lieu, la proposition (6) elle-même, en tant que sous-proposition d'elle-même.

3.2 Exemple informatique

Un langage du premier ordre peut être utilisé pour décrire le contenu d'une base de données. Il ne s'agit pas d'exprimer une à une toutes les informations qu'elle contient mais d'exprimer des généralités sur ce contenu, notamment les conditions que doivent respecter les informations introduites dans cette base. Cette description fait partie de la conception d'une base de données. Les « requêtes » de l'utilisateur, c'est-à-dire les questions qu'il peut poser au système pour en extraire des informations peuvent être aussi formulées dans ce langage.

Supposons qu'il s'agisse de développer un système d'informations utile à l'administration d'une université. On veut y faire figurer toutes les informations nécessaires sur : les étudiants, le personnel, les bâtiments et les locaux, les départements, les cours, les examens, les projets de recherche, etc.

Un tel système d'informations se décrit de manière très naturelle au moyen d'un langage du premier ordre. Dans la pratique on utilise souvent des notations graphiques (schémas), mais il y a une correspondance directe entre ces notations et les expressions d'un langage du premier ordre.

Les informations concerneront diverses catégories « d'objets » : les étudiants, les professeurs, les départements, etc. Elles exprimeront certaines

associations entre objets, par exemple en associant à chaque professeur le département dont il est membre, les cours qu'il donne, etc.

Une base de données est en bonne partie une représentation d'une certaine réalité. Une description de cette base de données est en même temps une description de la réalité. La base de données peut contenir par exemple, sous une forme quelconque, l'information qu'un nommé Euclide est professeur. D'autre part, qu'il y a un cours intitulé Géométrie, et troisièmement que ce cours est donné par Euclide. Un langage du premier ordre permettant d'exprimer ces informations utilisera par exemple les noms `euclide` et `géométrie`. Ces identificateurs seront des constantes individuelles du langage. On utilisera deux symboles relationnels unaires, par exemple `est-un-prof` et `est-un-cours`, permettant d'écrire les propositions

`est-un-prof(euclide), est-un-cours(géométrie).`

Un symbole relationnel binaire, `enseigne`, permettra de former la proposition

`enseigne(euclide,géométrie).`

Nous utilisons ici la manière d'écrire qui place toujours les symboles relationnels avant leurs arguments, manière qui sera présentée plus loin.

Des propositions telles que celles-ci, formées d'un symbole relationnel et de constantes individuelles sont appelées parfois des « faits ». Une base de données contient en général un très grand nombre de faits de ce genre, mais sous la forme de tableaux, et il n'y a pas de raison d'écrire tous ces faits de la manière ci-dessus.

Cependant l'ensemble des faits est soumis à ce qu'on appelle des *contraintes*, c'est-à-dire des conditions générales qui méritent d'être exprimées formellement au stade de la conception. On voudra par exemple qu'à chaque cours nommé dans la base, celle-ci associe un enseignant qui soit professeur. Cette condition s'exprimera dans le langage du premier ordre considéré de la manière suivante:

$$\forall x [\text{est-un-cours}(x) \Rightarrow \exists p (\text{est-un-prof}(p) \text{ et } \text{enseigne}(p, x))].$$

On utilise ici des variables individuelles, x , p , et des quantificateurs.

De nombreuses contraintes sont ce qu'on appelle des contraintes de type. Par exemple, si la base contient l'information `enseigne(beethoven,musique)`, elle doit contenir aussi les informations `est-un-prof(beethoven)` et `est-un-cours(musique)`, car le symbole relationnel `enseigne` ne doit être appliqué qu'à un professeur et à un cours et non à d'autres types d'objets. On peut exprimer cela par la proposition suivante, que devra satisfaire la base de données:

$$\forall x \forall y [\text{enseigne}(x, y) \Rightarrow (\text{est-un-prof}(x) \text{ et } \text{est-un-cours}(y))].$$

Une base de données est faite pour être « interrogée ». On doit pouvoir lui demander des renseignements. On le fait en lui adressant des questions que l'on appelle des requêtes. Par exemple, le directeur de l'établissement, soupçonnant

l'oisiveté de certains membres du corps enseignant, demande qu'on lui fournisse le nom des professeurs qui ne donnent aucun cours. Avec les symboles que nous venons d'introduire, la requête s'exprime comme suit : trouver tous les p tels que

$$\text{est-un-prof}(p) \text{ et } \text{non} \exists x(\text{enseigne}(p, x)).$$

Il est possible que le langage d'interrogation de la base de données réelle ne soit pas exactement celui-ci, qu'il soit plus proche du langage naturel de l'utilisateur par exemple. Mais conceptuellement c'est bien par rapport à un langage du premier ordre que l'on peut caractériser les possibilités d'interrogation d'une base de données.

L'exemple de langage que nous venons de voir n'utilise que cinq des six espèces de symboles des langages du premier ordre : variables individuelles ($x, y, p, \dots$); constantes individuelles (**beethoven**, **géométrie**, $\dots$); symboles relationnels (**est-un-cours**, **enseigne**, $\dots$); connecteurs logiques ($\Rightarrow$, **non**); quantificateurs.

Ce langage n'utilise pas de symboles fonctionnels et c'est souvent le cas dans la description des bases de données. Par exemple, si notre base de données contient des informations sur le salaire du personnel, on pourrait en parler au moyen du symbole fonctionnel **le-salaire**, permettant par exemple de construire le terme **le-salaire**(**beethoven**). Le symbole relationnel $<$ permettrait alors de demander s'il existe un professeur mieux payé qu'**euclide**:

$$\exists y \left(\begin{array}{l} \text{est-un-prof}(y) \\ \text{et} \\ \text{le-salaire}(\text{euclide}) < \text{le-salaire}(y) \end{array} \right) ?$$

Toutefois ces informations peuvent être décrites sans utiliser de symbole fonctionnel, mais au moyen d'un symbole *relationnel* binaire, **gagne**, permettant d'écrire par exemple la proposition **gagne**(**euclide**, 55'000). Cependant la requête précédente est plus difficile à formuler. Il faut demander s'il existe un professeur y et s'il existe deux nombres a et b tels que l'on ait $a < b$, **gagne**(**euclide**, a), et **gagne**(y , b):

$$\exists y \exists a \exists b \left(\begin{array}{l} \text{est-un-prof}(y) \text{ et} \\ \text{gagne}(\text{euclide}, a) \text{ et } \text{gagne}(y, b) \\ \text{et } a < b \end{array} \right).$$

3.3 Variables et constantes individuelles

Un langage du premier ordre $\mathcal{L}$ possède toujours une collection de symboles appelés les *variables individuelles* de $\mathcal{L}$, ou simplement les *variables*. Dans les exemples qui précèdent, nous avons utilisé dans ce rôle les lettres a, b, x, y, p . Mais on peut choisir comme variables d'un langage n'importe quels symboles. L'essentiel est qu'on sache les reconnaître, c'est-à-dire les distinguer des autres

symboles du langage, et ceci de façon mécanique. Dans un langage destiné à un traitement automatique, on prendra naturellement comme variables une collection d'identificateurs, reconnaissables à telle ou telle particularité, par exemple tous les identificateurs qui commencent par une lettre majuscule.

Sémantique des variables. Le rôle sémantique des variables est de désigner des éléments arbitraires de la classe d'objets constituant ce qu'on appelle le domaine du discours, ou le domaine de l'interprétation, ou encore l'univers du discours. Dans le premier des exemples précédents, ces objets étaient les nombres réels; dans le second, certaines entités d'une université (cours, professeurs, étudiants, etc.). Les objets du domaine du discours sont appelés aussi les *individus* de ce domaine, ou de cet univers, d'où l'adjectif « individuelles » utilisé pour les variables en question.

D'autres espèces de langages, notamment les langages du deuxième ordre et ceux d'ordres supérieurs, se réfèrent aussi, sémantiquement, à un univers d'objets déterminé, mais ils possèdent non seulement des variables individuelles, mais encore d'autres catégories de variables. Ces variables peuvent servir par exemple à désigner non pas des individus, mais des ensembles arbitraires d'individus, ou des opérations arbitraires sur les individus, etc. Comme nous ne considérons dans cet ouvrage que les langages du premier ordre, l'adjectif « individuelles » est superflu, mais il est traditionnel.

Infinité du répertoire de variables. La collection des variables individuelles d'un langage doit toujours être *infinie*. La raison de ceci est que toute limitation du nombre des variables à un nombre fini, quelconque mais fixé, entraîne des impossibilités d'expression dont on ne voit pas la raison. Par exemple, supposons que le symbole $<$ soit un symbole relationnel binaire d'un langage $\mathcal{L}$. Si l'on ne dispose que de deux variables individuelles, on ne peut pas énoncer une propriété de transitivité telle que

$$\forall x \forall y \forall z ((x < y \text{ et } y < z) \Rightarrow x < z),$$

qu'elle soit vraie ou fausse. Si l'on ne dispose que de trois variables, on ne peut pas formuler

$$\forall x \forall y \forall z \forall w ((x < y \text{ et } y < z \text{ et } z < w) \Rightarrow x < w),$$

et ainsi de suite. Il est peu probable que l'on ait jamais à écrire de proposition contenant mille variables distinctes. Cependant, on ne voit pas de bonne raison de se fixer arbitrairement une limite quelconque. Ce n'est d'ailleurs pas un problème que de se donner un répertoire illimité de symboles. On peut choisir par exemple la famille de lettres accentuées A, A', A'', A''', A'''' , etc. ou celle des identificateurs (sans limitation de longueur) d'un langage de programmation.

Répertoire de variables standard. Dans cet ouvrage, nous utiliserons toujours comme variables individuelles les symboles donnés dans la figure 3.1. Le lecteur est invité à bien prendre note de ce choix: ce sont les lettres *latines italiques* $a, \dots, z$ et $A, \dots, Z$ ainsi que toutes leurs variantes munies d'accents

a', a'', a''' , etc. C'est un répertoire infini, comme il se doit, car on peut mettre à une lettre autant d'accents que l'on veut. On veillera à ne pas confondre ces symboles avec d'autres familles de caractères que nous utiliserons, notamment les lettres droites telles que x, y, A, B que nous emploierons systématiquement dans le métalangage (§2.4).

$$\begin{array}{l} a, b, c, \dots, z, A, B, C, \dots, Z \\ a', b', c', \dots, z', A', B', C', \dots, Z' \\ a'', b'', c'', \dots, z'', A'', B'', C'', \dots, Z'' \\ \dots \end{array}$$

Figure 3.1 Répertoire de variables individuelles pour l'ensemble de l'ouvrage.

Ainsi tous les langages du premier ordre considérés dans ce livre posséderont le répertoire de variables individuelles déclaré dans la figure 3.1. Ils ne se distingueront que par leurs répertoires de constantes individuelles, ou de symboles fonctionnels, ou de symboles relationnels, puisque les symboles logiques (connecteurs, quantificateurs) sont les mêmes dans tous les langages du premier ordre.

Le langage mathématique traditionnel n'est pas standardisé. Notamment il ne possède pas de répertoire déclaré de variables individuelles. À côté des lettres latines, on emploie couramment des lettres grecques et parfois d'autres alphabets encore. Mais on emploie aussi les mêmes symboles comme constantes individuelles, par exemple la lettre e lorsqu'on parle du nombre d'Euler. C'est le contexte qui détermine en général le rôle d'un symbole, alors que dans un langage formel tous les symboles ont un rôle — variable ou constante individuelle, symbole fonctionnel, symbole relationnel — fixé une fois pour toutes. En particulier, la lettre e appartient au répertoire de variables individuelles que nous venons de fixer. Donc, si nous voulons une constante individuelle pour désigner le nombre d'Euler, nous prendrons peut-être la lettre $\mathbf{e}$, choisie dans une autre police de caractères comme le lecteur peut le remarquer.

Ce qu'il peut remarquer aussi, c'est que la lecture d'un livre de logique demande que l'on porte une certaine attention à la typographie — par exemple que l'on ne confonde pas les lettres e et $\mathbf{e}$.

Variables et quantificateurs. Une particularité des variables individuelles, parmi les symboles d'un langage du premier ordre, est qu'à chacune d'elles, par exemple x , sont associés deux quantificateurs, $\forall x$ et $\exists x$. Ces quantificateurs sont employés pour construire des propositions de la forme $\forall x(\dots)$, ou de la forme $\exists x(\dots)$, dans lesquelles $(\dots)$ peut être une proposition quelconque. Ceci nous amène à discuter la notion de variable, en relation avec les quantificateurs, afin de mettre en garde le lecteur contre une interprétation du mot « variable » qui est encore assez répandue, mais qu'il doit bannir de sa pensée.

Au paragraphe 3.1, nous avons formalisé l'énoncé suivant :

$$\textit{Si le produit de deux nombres réels quelconques est nul, alors l'un de ces nombres est nul} \quad (1)$$

en écrivant :

$$\forall x \forall y [(x \cdot y = 0) \Rightarrow (x = 0 \text{ ou } y = 0)]. \quad (2)$$

Plus d'un lecteur, sans doute, se serait contenté d'écrire

$$(x \cdot y = 0) \Rightarrow (x = 0 \text{ ou } y = 0), \quad (3)$$

sans quantificateurs. Cela provient, nous semble-t-il, d'une « ancienne version » de la notion de variable, encore assez répandue. Selon cette interprétation, les variables x et y de la proposition (3) désignent — par définition du mot « variable » — des nombres *quelconques*, donc l'adjectif « quelconques » qui figure dans l'énoncé (1) est déjà traduit par l'usage de *variables* x et y dans (3). L'adjonction de quantificateurs, comme dans (2), ne paraît certainement pas fausse, mais en tout cas *superflue*.

C'est ce sens du mot « variables » que le lecteur de cet ouvrage doit bannir de sa pensée. Le sens qui subsistera est un autre sens, que l'on rencontre *aussi*, et même plus fréquemment que le premier, dans les textes mathématiques. Considérons par exemple les phrases suivantes, que l'on pourrait rencontrer dans un manuel de géométrie :

Soient x et y les longueurs des côtés du triangle rectangle ABC,
et z la longueur de l'hypothénuse. On a

$$z^2 = x^2 + y^2. \quad (4)$$

Il est clair que dans ce contexte, les variables x , y , z de l'équation (4) n'ont rien de « variable ». Elles désignent des nombres *particuliers*, que l'on ne connaît peut-être pas, mais particuliers tout de même. D'ailleurs l'adjonction de quantificateurs universels $\forall x \forall y \forall z$ à (4) n'entre pas en considération. Elle ne serait pas superflue, mais simplement fausse.

On voit qu'il y a dans le langage mathématique courant une catégorie de symboles tels que x , y , que l'on interprète tantôt comme des variables, en ce sens que la proposition dans laquelle ils figurent est vraie quelle que soit la valeur qu'on leur attribue, tantôt comme des noms propres désignant des objets particuliers indiqués dans le contexte.

La première de ces interprétations n'aura jamais lieu dans cet ouvrage. En l'absence de quantificateurs universels $\forall x$, $\forall y$, ... judicieusement placés, les variables x , y , ... désignent toujours des objets *particuliers*. Ces objets dépendent du contexte, et peuvent changer d'un contexte à un autre. C'est seulement à cause de cette dépendance du contexte que l'on a conservé la dénomination de « variables ».

En lisant une proposition sans quantificateurs telle que (3), le lecteur doit donc s'habituer à considérer les variables qu'elle contient comme désignant

des objets *particuliers*, même s'il ne possède aucune autre information sur ces objets. Le sens des propositions (2) et (3) n'est donc pas le même et les quantificateurs de la première sont *indispensables* pour traduire (1).

Constantes individuelles. Les constantes individuelles d'un langage servent à désigner des objets fixes, particuliers, de l'univers de l'interprétation. Nous avons vu le symbole 0 dans l'exemple du paragraphe 3.1, et les identificateurs *beethoven* et *euclide* dans l'exemple du paragraphe 3.2. Contrairement aux variables individuelles, les constantes individuelles désignent des objets qui ne dépendent pas du contexte d'une expression. C'est pourquoi nous parlons d'objets fixes.

Dans le langage mathématique courant, il n'y a que peu de symboles dont la signification soit universellement fixée. Ainsi le symbole 0 n'est pas utilisé seulement pour désigner le nombre zéro, mais aussi par exemple le vecteur nul d'un espace vectoriel. Mais on peut admettre que ce sont deux 0 « différents », que l'on distingue mentalement, donc qu'il s'agit de constantes individuelles différentes. D'autres exemples de constantes individuelles du langage mathématique sont les symboles $1, 2, \dots, \infty, \emptyset, \mathbb{R}, \mathbb{N}$.

Dans un langage du premier ordre, le répertoire des constantes individuelles est fixé et ne contient aucun symbole appartenant aussi à une autre catégorie de symboles du langage (variables individuelles, symboles fonctionnels, etc.). Ces catégories n'ont aucun élément commun. Contrairement au répertoire des variables individuelles qui est infini, le répertoire des constantes individuelles d'un langage peut être de taille quelconque. Il peut être fini, vide en particulier, ou infini.

3.4 Symboles fonctionnels

Les expressions mathématiques font constamment intervenir des *symboles d'opérations*, comme par exemple le symbole $+$ dans l'expression $x + 1$. Le terme $x + 1$ est un *terme composé*, résultant de l'assemblage des termes simples x (variable individuelle) et 1 (constante individuelle) au moyen du symbole d'opération $+$.

Dans un langage du premier ordre, les symboles d'opération sont appelés généralement les *symboles fonctionnels* du langage. C'est une terminologie qui n'est pas très heureuse, mais elle est courante et nous l'emploierons aussi. Nous conseillons cependant au lecteur de ne pas chercher à mettre en relation l'adjectif « fonctionnel » avec la notion de fonction qu'il peut avoir, car cela pourrait l'induire en erreur. Nous reviendrons plus loin sur ce sujet.

Notation préfixée des symboles fonctionnels. Le symbole fonctionnel $+$ de l'opération usuelle d'addition de nombres se combine toujours avec deux termes, qui peuvent être deux fois le même terme comme dans l'expression $x + x$. On dit qu'il s'agit d'un symbole fonctionnel *binaire*. Dans l'expression $x + 1$, le symbole fonctionnel (symbole d'opération) est placé entre les deux

termes avec lesquels il forme un terme composé. Une autre notation consiste à écrire

$$+ x 1$$

en mettant le symbole fonctionnel à gauche des deux termes. Il forme ainsi un préfixe de l'expression et l'on dit qu'il est *préfixé*. On parle encore de la *notation préfixée* ou de la *notation polonaise*¹. Une notation analogue, avec des parenthèses en plus, est utilisée dans le langage de programmation LISP, dans lequel on écrit $(+ x 1)$. Une troisième notation analogue est la notation $+(x, 1)$ dite aussi « fonctionnelle ». La notation

$$x1+,$$

symétrique de la notation préfixée, est dite *postfixée* ou *polonaise inverse*. Quant à la notation habituelle $x + 1$, elle est dite *infixée*.

Autres notations d'opérations. Dans le langage mathématique courant les opérations ne sont pas toujours notées au moyen d'un symbole fonctionnel préfixé, infixé, ou postfixé. Par exemple, l'opération d'exponentiation a^b et l'opération de multiplication ab sont notées sans signes d'opération, la division $\frac{a}{b}$ est notée verticalement, etc. Mais il est toujours possible d'introduire pour ces opérations des symboles fonctionnels binaires explicites. Ceux-ci peuvent être préfixés comme dans le langage LISP où l'on écrit $(\text{EXPT } a \ b)$ pour l'exponentiation, $(* \ a \ b)$ pour la multiplication, et $(/ \ a \ b)$ pour la division. On peut choisir aussi la notation infixée: $a \ \text{EXPT} \ b$, $a * b$, a/b .

Opérations non apparentes. Certaines opérations binaires en mathématique ne sont pas apparentes dans la notation traditionnelle et il faut effectuer un certain travail de traduction pour mettre ces notations en relation avec celles d'un langage du premier ordre. Nous voulons citer deux exemples. Le premier est la notation $\{a, b\}$ désignant l'ensemble dont les éléments sont deux objets désignés par a et b . Nous avons ici une opération binaire de *construction d'ensemble*: à partir de deux objets a et b , on construit l'ensemble formé par ces deux objets. Ce sont les trois signes « $\{ , \}$ » qui constituent le symbole de cette opération. On peut dire qu'il s'agit d'un symbole fonctionnel binaire composé de trois parties, l'une préfixée, l'autre infixée, et la troisième postfixée. Au lieu de cette notation, on pourrait utiliser un seul symbole fonctionnel spécifique, par exemple Δ , et écrire Δab (notation préfixée) ou $a \ \Delta \ b$ (notation infixée), au lieu de la notation usuelle $\{a, b\}$.

L'opération d'application d'une fonction à un argument. L'autre exemple d'une opération non apparente que nous voulons citer réside dans la notation traditionnelle $f(x)$ pour désigner la valeur d'une fonction f correspondant à un argument x . Pour fixer les idées, supposons que f soit une application de l'ensemble $\mathbb{R}$ dans lui-même. Il faut distinguer bien entendu le *symbole* f de la *fonction* dont ce symbole est le nom. Du point de vue de la classification

¹ Notation introduite par le logicien polonais Łukasiewicz, en 1929.

des symboles, la lettre f appartient *comme* la lettre x à la classe des *variables individuelles*. L'une désigne ici une fonction, l'autre un nombre réel. Nous soulignons que f *n'est pas un symbole fonctionnel*. C'est la raison de notre mise en garde plus haut contre une mauvaise interprétation de la dénomination « symbole fonctionnel ». Ce n'est pas parce qu'un symbole désigne une fonction qu'il est nécessairement un symbole fonctionnel. On peut se servir de la variable f pour désigner tout ce qu'on veut, elle reste une variable quoi qu'elle désigne.

Nous rappelons maintenant au lecteur notre interprétation des variables en l'absence de quantificateurs universels : dans un contexte donné, une variable représente un objet particulier, même si l'on ne sait rien de cet objet. Donc, dans le terme $f(x)$, la variable f désigne une fonction particulière, la variable x désigne un nombre *particulier*, et le terme $f(x)$ désigne le nombre que l'on appelle la valeur de la fonction f correspondant au nombre x , ou encore le résultat de l'application de la fonction f au nombre x , ou encore — jargon des informaticiens — la valeur retournée par la fonction f pour la donnée x . Nous sommes ici en présence d'une opération *binaire*, car le terme $f(x)$ représente un objet qui est construit à partir de deux objets : la fonction f d'une part, le nombre x d'autre part. Il s'agit de l'opération « appliquer une fonction donnée à un argument donné ». Adoptons comme symbole de cette opération le symbole $@$, en tant que symbole fonctionnel binaire infixé. Au lieu du terme $f(x)$ nous écrirons alors $f @ x$, ce qui peut toujours se lire « f de x ».

Autres arités. On rencontre aussi, dans le langage mathématique courant, des symboles fonctionnels *unaires*, c'est-à-dire qui s'appliquent à un seul terme, comme le symbole « $-$ » dans l'expression $-x$. La notation est préfixée dans cet exemple. Dans celui de la factorielle $x!$ le symbole fonctionnel « $!$ » est postfixé. Dans d'autres exemples, la notation n'est ni pré- ni postfixée : $\sqrt{x}$, $|x|$, $\bar{x}$, etc. Mais on peut introduire pour chacune de ces opérations un symbole fonctionnel spécifique préfixé, si l'on désire se conformer à la syntaxe des langages du premier ordre. En LISP, par exemple, on écrit ainsi (SQRT x) pour la racine carrée, (ABS x) pour la valeur absolue, (NOT x) pour la complémentarité booléenne, etc.

Les symboles fonctionnels *ternaires*, ou d'arité plus élevée, c'est-à-dire s'appliquant à trois arguments ou plus, sont assez rares dans le langage mathématique courant — nous parlons de symboles ayant une signification largement connue, et dont l'usage est répandu. Mais chaque auteur est libre de créer son propre langage. Par exemple, l'auteur d'un manuel de géométrie pourrait introduire la notation $\triangle ABC$ pour désigner le triangle dont les sommets sont les points A , B , C . Le symbole $\triangle$ serait dans ce langage un symbole fonctionnel ternaire.

Symboles fonctionnels du langage naturel. De nombreux termes du langage naturel peuvent être considérés comme étant composés à partir d'autres termes au moyen d'un symbole fonctionnel. Par exemple, le terme

Le père d'Archimède

peut être considéré comme étant construit à partir du terme « Archimède » au moyen du symbole fonctionnel « le père de ». On pourrait rencontrer des symboles proches de ces mots dans un langage du premier ordre décrivant une base de données généalogiques : constante individuelle **archimède**, symbole fonctionnel **père**. L'expression **père**(**archimède**), ou (**père archimède**), ou simplement **père archimède**, suivant la notation que l'on adopte pour les termes composés, serait alors un terme de ce langage. Un autre symbole fonctionnel unaire pourrait être « l'âge de », comme dans le terme « l'âge du capitaine ». Un terme tel que **âge**(**père**(**archimède**)), où **âge** et **père** sont deux symboles fonctionnels unaires, est doublement composé. Il est construit avec le symbole fonctionnel unaire **âge**, appliqué au terme composé **père**(**archimède**).

Certains termes du langage naturel peuvent être considérés comme étant composés au moyen d'un symbole fonctionnel *binaire*. Par exemple :

la distance de Paris à Berlin

est un terme composé qui pourrait s'écrire **distance**(**paris**, **berlin**) dans un langage dans lequel **distance** est un symbole fonctionnel binaire et **paris**, **berlin** sont deux constantes individuelles.

Résumé. Un langage du premier ordre comporte une collection de symboles appelés les symboles fonctionnels du langage, à chacun desquels un nombre fixe est associé. C'est le nombre de termes avec lesquels le symbole se combine toujours, dans une expression. Ce nombre est appelé *l'arité* d'un symbole fonctionnel — d'après les adjectifs unaire, binaire, ternaire, etc. Un symbole fonctionnel d'arité 1 est un symbole fonctionnel unaire, un symbole d'arité 2 est binaire, etc. Pour une arité n (nombre entier 1, 2, 3, etc. non précisé), on parle d'un symbole n -aire. Un symbole fonctionnel se combine toujours avec des termes et l'expression ainsi formée est un terme composé. La collection des symboles fonctionnels d'un langage, tout comme celle des constantes individuelles, peut être finie — en particulier vide — ou infinie.

3.5 Symboles relationnels

Un langage du premier ordre possède une collection de symboles dits *relationnels*, servant à construire des propositions à partir de *termes*. Une proposition d'égalité telle que $x = a + 1$, par exemple, est construite au moyen du symbole relationnel $=$ à partir des deux termes x et $a + 1$. D'autres symboles relationnels du langage mathématique courant sont les symboles $\leq$, $<$, $\in$, $\subset$. Ce sont tous des symboles relationnels *binaires*, puisqu'ils relient toujours *deux* termes. Ceux-ci peuvent être d'ailleurs deux occurrences du même terme, comme dans l'égalité $a + 1 = a + 1$.

On peut avoir dans un langage des symboles relationnels d'arité n quelconque ($n = 1, 2, 3$, etc.). Dans le langage mathématique courant, les symboles relationnels d'arité autre que 2 sont constitués en général par des mots ou des

groupes de mots empruntés au langage naturel et pris dans un sens spécial. Lorsqu'on écrit par exemple

f est une fonction,

le groupe de mots «est une fonction» doit être considéré comme *un seul symbole*, du type relationnel unaire, postfixé. Ce symbole est appliqué ici au terme f , lequel est une variable individuelle, et la combinaison des deux est une proposition. On pourrait utiliser à la place des mots «est une fonction» un symbole spécial plus court, par exemple la lettre $\mathfrak{F}$ ou l'abréviation **Fct**, que l'on mettrait plutôt en préfixe, en écrivant $\mathfrak{F}(f)$ ou **Fct**(f), ou simplement $\mathfrak{F}f$ ou **Fct** f , au lieu de « f est une fonction». Cependant les symboles relationnels de ce genre, utilisant les mots

«est un ...» ou «est une ...»,

sont si nombreux qu'il serait difficile de trouver pour chacun d'eux un symbole spécial facile à mémoriser. Il y a en effet autant de symboles de ce genre qu'il y a de concepts, ou d'espèces d'objets auxquels on peut penser: un nombre réel, une fonction, un graphe, un espace vectoriel, une cône, un programme, un mammifère, etc. Pour chacun de ces concepts on a l'occasion d'employer le symbole relationnel «est un ...» correspondant. Dans l'exemple de langage du paragraphe 3.2, nous avons introduit les symboles relationnels unaires **est-un-cours**, **est-un-prof**, correspondant à deux espèces d'objets représentés dans une certaine base de données.

De même lorsque, parlant d'une fonction, on écrit « f est continue», le groupe de mots «est continue» constitue un symbole relationnel unaire. De nombreuses propriétés d'objets s'expriment de cette manière, au moyen de symboles relationnels unaires formés avec des adjectifs de la langue naturelle.

C'est aussi avec des mots de la langue naturelle, pris dans un sens spécial, que sont construits de nombreux symboles relationnels ternaires du langage mathématique. Par exemple, dans la phrase

f est une application de A dans B ,

il y a trois termes: les variables individuelles f , A , B . Ces trois termes sont reliés par les mots

... est une application de ... dans ...

et ces mots constituent un symbole relationnel ternaire. On pourrait employer une abréviation telle que **Appl**(f, A, B), ou simplement **Appl** fAB .

Résumé. Un langage du premier ordre comporte une collection de symboles appelés les *symboles relationnels* du langage, à chacun desquels un nombre fixe est associé. Ce nombre entier n ($n = 1, 2, 3$, etc.) est appelé *l'arité* du symbole. Le symbole doit être combiné avec n termes et la combinaison est une *proposition* du langage. Les propositions de cette forme sont appelées les propositions *atomiques* du langage. Les autres propositions du langage sont construites à

partir de propositions atomiques au moyen de connecteurs logiques (et, ou non, $\Rightarrow$, $\Leftrightarrow$) et de quantificateurs. Ces propositions sont dites *composées*.

Le nombre de symboles relationnels d'un langage peut être quelconque, mais il faut qu'il y en ait au moins un pour que le langage présente un intérêt. Les propositions sont en effet l'espèce d'expression principale d'un langage du premier ordre, et sans symbole relationnel on ne peut pas construire de proposition.

3.6 Connecteurs logiques

Les connecteurs logiques sont, avec les quantificateurs, les symboles qui permettent de construire les propositions composées d'un langage du premier ordre. Ils ne sont pas spécifiques à un langage particulier, mais communs à *tous* les langages du premier ordre. Nous avons déjà utilisé la plupart d'entre eux dans les exemples des paragraphes 3.1 et 3.2: les symboles de négation (non), de disjonction (ou), de conjonction (et), d'implication ($\Rightarrow$). Nous avons mentionné aussi le symbole de double implication ($\Leftrightarrow$). Chacun de ces symboles possède une arité. Le symbole de négation est unaire. Les quatre autres sont binaires.

Nous ajouterons encore à ce répertoire le symbole spécial $\perp$, que nous considérons comme un connecteur logique nullaire (arité zéro). Ce symbole est à lui seul une proposition, qui sera toujours fausse. Nous l'appellerons *la proposition fausse par excellence*.

La négation. Dans toutes les langues naturelles il y a une manière standard de former une phrase exprimant le contraire d'une phrase donnée. C'est ce qu'on appelle prendre la négation de cette phrase. Dans nos langages formels, la négation d'une proposition A sera l'expression

$$\neg A.$$

Le symbole $\neg$ est le connecteur logique de négation, préfixé. L'expression $\neg A$ se lit « non A » et nous l'écrirons aussi en général de cette manière, c'est-à-dire en utilisant le mot « non » à la place du symbole $\neg$.

Note typographique. Le lecteur remarquera ici l'usage qui est fait de la lettre droite A , comme symbole du métalangage, ou métasympbole. Lorsque nous annonçons que nous noterons $\text{non}A$ la négation d'une proposition A , ce n'est pas le symbole A lui-même qui est cette proposition. La lettre A sert ici à désigner une proposition quelconque, indéterminée, d'un langage du premier ordre. On pourrait se passer d'un métasympbole en disant simplement: la négation d'une proposition quelconque sera construite en lui ajoutant le préfixe $\neg$. On peut dire que la lettre A est employée comme *métavariable* ou variable du métalangage — le langage que nous utilisons en ce moment pour décrire et commenter la syntaxe des langages du premier ordre. Nous emploierons régulièrement les lettres majuscules droites A , B , C , R comme métavariables pour désigner des propositions quelconques d'un langage formel.

Il ne faut pas les confondre avec les lettres italiques A , B , C , R , qui sont des *symboles* des langages formels que nous considérons, à savoir des variables individuelles de ces langages. Elles appartiennent en effet au répertoire de variables individuelles que nous nous sommes fixé (§3.3).

La barre de négation. La négation d'une proposition atomique construite avec un symbole relationnel binaire est souvent notée en « barrant » ce symbole. L'exemple le plus connu est celui d'une proposition d'égalité $x = y$. Au lieu d'écrire $\neg(x = y)$, ou $\text{non}(x = y)$, on écrit généralement $x \neq y$, en barrant le symbole relationnel $=$. La même notation s'emploie avec les symboles relationnels binaires $\in$, $\subset$, $<$, $\leq$, etc. Le lecteur se rappellera que $a \neq b$, $a \notin b$, $a \not\subset b$, etc. sont des *abréviations* qui représentent les propositions $\text{non}(a = b)$, $\text{non}(a \in b)$, $\text{non}(a \subset b)$, etc.

Négations répétées. Dans un langage du premier ordre, il est possible de prendre la négation de *n'importe quelle proposition*. Par conséquent, étant donné une proposition quelconque A , on peut former la proposition $\neg A$, puis la proposition $\neg\neg A$, puis la proposition $\neg\neg\neg A$, et ainsi de suite. Toutes ces propositions sont *distinctes*. Chacune d'elles est la négation de la précédente. En particulier, nous attirons l'attention du lecteur sur ceci: la négation d'une proposition de la forme $\neg A$ n'est pas la proposition A , mais bien la proposition $\neg\neg A$, qui est munie d'un symbole de négation supplémentaire et qui donc, syntaxiquement, est bien distincte de A .

Sémantiquement, il est assez naturel d'attribuer à une proposition $\neg\neg A$, qui est la double négation d'une proposition A , le même sens qu'à A . Les règles de la logique classique correspondent en effet à cette interprétation et permettent donc de substituer A à $\neg\neg A$ dans n'importe quel contexte. Nous voulons cependant inviter le lecteur à ne pas identifier sommairement A et $\neg\neg A$ avant d'avoir étudié les règles en question. D'ailleurs, bien que cet ouvrage traite de la logique classique, nous ferons quelques allusions à la logique intuitionniste et à l'approche dite *constructive* des mathématiques — fondée sur une critique des mathématiques classiques qu'il n'est plus possible d'ignorer actuellement. En logique intuitionniste, la double négation $\neg\neg A$ d'une proposition A n'est pas équivalente à A . Le lecteur doit noter d'ailleurs comme règle générale que dans un ouvrage de logique *aucune* manipulation de symboles, si intuitive qu'elle puisse paraître, ne « va de soi ». Toute manipulation doit être autorisée par une convention explicite.

Dans les langues naturelles un peu connues de l'auteur, la manière de prendre la négation d'une phrase ne peut pas s'appliquer deux ou plusieurs fois comme dans les langages formels ($\neg\neg A$). Une telle construction est impraticable. Par exemple, en français, en allemand, et en anglais, la négation de la phrase « je comprends, ich verstehe, I understand », c'est la phrase: je ne comprends pas, ich verstehe nicht, I do not understand. Mais on ne peut pas la nier à son tour en itérant simplement la construction, ce qui donnerait: je ne ne comprends pas pas, ich verstehe nicht nicht, I do not do not understand!

La disjonction et la conjonction. Avec deux propositions A et B , pouvant être la même proposition, on peut former une proposition qui est appelée la *disjonction* de A et de B . Dans la langue française, cette construction utilise le mot « ou », comme dans la phrase

Arthur a oublié notre rendez-vous *ou* il a eu un accident, (1)

qui est la disjonction des phrases « Arthur a oublié notre rendez-vous », « il a eu un accident ». Cependant, contrairement à ce qui se passe dans les langues naturelles, le « ou » des langages formels a toujours le sens que l'on dit *inclusif*. Par exemple, dans la situation de personnes qui se demandent pourquoi un nommé Arthur n'est pas venu à un certain rendez-vous, le « ou » de la phrase (1) a le sens *inclusif* s'il veut dire que l'un *au moins* des deux faits suivants s'est produit :

Arthur a oublié le rendez-vous

Arthur a eu un accident.

Cela *inclut* la possibilité que les *deux* faits se soient produits. Comme le « ou » de la langue naturelle n'a pas toujours ce sens inclusif, celui qui voudra être sûr de bien se faire comprendre dans ce sens ajoutera à la phrase (1) des mots tels que « ou les deux à la fois ».

Dans le langage mathématique usuel, le « ou » a toujours, normalement, le sens inclusif. Par exemple la proposition

$$x = 0 \text{ ou } y = 0 \quad (2)$$

signifie que l'un au moins des nombres désignés par x et y est nul, et l'on inclut dans ce sens la possibilité que les deux soient nuls.

Nous adopterons le connecteur logique binaire $\vee$ comme symbole de la disjonction dans les langages formels. Cependant, de même que nous écrirons généralement « non A » au lieu de $\neg A$, nous emploierons pratiquement toujours le mot « ou » à la place du symbole de disjonction $\vee$. Il en ira de même pour le symbole de conjonction, $\wedge$, à la place duquel nous emploierons pratiquement le mot « et ». C'est une question de lisibilité. Mélangés avec d'autres symboles mathématiques, les symboles $\neg, \vee, \wedge$ se distinguent moins bien que les mots non, ou, et.

Notation préfixée et notation usuelle. La syntaxe des termes et des propositions d'un langage du premier ordre sera précisée au chapitre 4, et cela sera fait en utilisant pour tous les opérateurs la notation préfixée. Par *opérateurs*, nous entendons tous les symboles qui servent à construire des expressions au moyen d'autres expressions, donc les symboles fonctionnels, les symboles relationnels, les connecteurs logiques et les quantificateurs. Par exemple, selon cette syntaxe, la proposition $x = 0$, doit s'écrire

$$= x 0,$$

avec le symbole relationnel binaire $=$ en position préfixée. La disjonction de

deux propositions A , B doit s'écrire $\vee AB$. Donc la proposition (2) s'écrit

$$\vee = x0 = y0. \quad (3)$$

Nous verrons plus loin l'intérêt de cette notation « polonaise ». Nous mentionnons seulement ici le fait qu'elle permet de se passer complètement de parenthèses, même dans les expressions les plus compliquées. Elle n'utilise donc que les constituants essentiels des expressions.

Évidemment, cette notation est peu lisible et contraire aux habitudes. C'est pourquoi nous emploierons *pratiquement* la forme traditionnelle (2) au lieu de la forme (3). Mais celle-ci sera bien la forme « réglementaire » et nous admettrons que l'expression (2) n'appartient pas elle-même au langage formel, mais qu'elle ne fait que *représenter* dans le dialogue l'expression formelle (3). Nous emploierons ainsi continuellement des expressions qui ne seront pas elles-mêmes les véritables expressions formelles, mais qui serviront à désigner ou représenter ces dernières.

La *conjonction* de deux propositions, dans un langage formel, correspond à la combinaison de deux phrases de la langue naturelle avec le mot « et », comme dans la phrase

Arthur a oublié notre rendez-vous *et* il a eu un accident.

Le sens du mot « et » dans une telle combinaison est toujours le même et il n'est pas nécessaire de le préciser comme nous avons dû le faire pour le « ou » d'une disjonction. Dans un langage formel, suivant la syntaxe définie plus loin, la conjonction de deux propositions A , B sera l'expression $\wedge AB$, construite avec le connecteur logique binaire $\wedge$. Pratiquement, nous la représenterons par « A et B ». Ainsi, l'expression

$$x = 0 \text{ et } y = 0$$

désignera la proposition formelle $\wedge = x0 = y0$.

L'implication. Avec deux propositions A et B de la langue naturelle, ou du langage mathématique usuel, on peut construire la proposition *si A alors B* . Par exemple, parlant d'un nombre réel x , on pourra énoncer la proposition vraie

$$\text{si } x \geq 2 \text{ alors } x^2 \geq 4. \quad (4)$$

Les mots *si ... alors ...* constituent un connecteur logique binaire, reliant ici les propositions $x \geq 2$ et $x^2 \geq 4$. Dans un langage formel, le connecteur correspondant sera le symbole $\Rightarrow$. Avec deux propositions A et B , on pourra donc former la proposition $\Rightarrow AB$, que nous représenterons pratiquement par

$$A \Rightarrow B,$$

et parfois aussi par « si A , alors B » comme dans (4). Une proposition de cette forme est appelée une *implication*. L'expression $A \Rightarrow B$ se lit aussi « A implique B ».

En utilisant le mot **Exp** comme symbole fonctionnel binaire pour l'exponentiation, le terme x^2 s'écrit **Exp** x 2 en notation préfixée et la proposition (4) devient dans cette notation :

$$\Rightarrow \geq x\ 2 \geq \text{Exp } x\ 2\ 4.$$

Le symbole d'implication que nous utilisons dans cet ouvrage, $\Rightarrow$, est celui du langage mathématique courant. Dans les ouvrages de logique, on rencontre souvent d'autres symboles dans ce rôle. En particulier, au lieu de $A \Rightarrow B$, on trouve souvent la notation $A \supset B$, ou encore $A \rightarrow B$.

Sémantique de l'implication. Nous admettons que le sens (sémantique) de l'implication, c'est-à-dire le sens des mots « si ... alors ... » dans une proposition telle que (4) est plus ou moins le même pour tous les lecteurs. Une erreur d'interprétation possible serait de voir toujours dans l'implication l'expression d'une causalité. On rencontre dans certains ouvrages de logique des exemples qui pourraient facilement susciter cette interprétation, comme :

S'il pleut, alors les routes sont mouillées.

Dans cet exemple il se trouve que le fait « il pleut », s'il se produit, est bien la *cause* physique du fait que les routes soient mouillées, lequel se produit alors inévitablement. Mais du point de vue de la logique, ce n'est pas cette causalité qu'expriment les mots *si ... alors ...*. C'est uniquement la concomitance des deux faits, à savoir que l'on n'observe jamais le premier sans le second. Nous verrons d'ailleurs plus loin qu'en logique classique une proposition de la forme $A \Rightarrow B$ est équivalente à la proposition non(A et non B).

Comme exemple d'implication, dans la langue naturelle, qui ne suggère pas l'idée de causalité nous pouvons imaginer à nouveau la situation d'une personne qui ne voit pas venir son ami Arthur au rendez-vous qu'ils se sont fixé et qui déclare :

Si Arthur n'a pas oublié notre rendez-vous, alors il a eu un accident. (5)

Dans l'esprit de celui qui s'exprime ainsi, il n'y a pas de rapport de cause à effet. Il pourrait exprimer la même pensée en disant : ou bien Arthur a oublié notre rendez-vous, ou bien il a eu un accident. Et en effet, en logique classique, une implication $A \Rightarrow B$ est encore équivalente à $((\text{non}A) \text{ ou } B)$ comme nous le verrons. La phrase (5) a le même sens que la phrase (1).

Une autre erreur d'interprétation assez fréquente du symbole $\Rightarrow$ consiste à prendre celui-ci pour un symbole du *métalangage*, c'est-à-dire à interpréter une expression de la forme $A \Rightarrow B$ comme une affirmation *sur les propositions* A , B . On risque de glisser facilement vers cette interprétation lorsqu'on lit $A \Rightarrow B$ en disant « A implique B ». Car on peut prendre cette expression pour une phrase avec le *verbe* « implique », le *sujet* A , et le *complément* B , donc bien comme une phrase du métalangage sur A et B . Un pas de plus dans cette

direction est d'ajouter le mot « la proposition » devant A et devant B, en disant

La proposition A implique la proposition B.

Une pareille expression est clairement une assertion dans le métalangage sur les propositions A et B elles-mêmes. Elle exprime une relation entre ces propositions. Or une implication telle que

$$x \geq 2 \Rightarrow x^2 \geq 4$$

n'est pas du tout une assertion sur la proposition $x \geq 2$ et sur la proposition $x^2 \geq 4$. C'est une assertion sur *les objets* nommés $x, 2, 4$. Il n'y a pas de doute là-dessus lorsqu'on lit cette proposition en disant « si $x \geq 2$ alors $x^2 \geq 4$ ». C'est pourquoi il est préférable, au début tout au moins, de lire le symbole $\Rightarrow$ en utilisant la tournure « si ... alors ... » plutôt que le verbe « implique ».

L'équivalence. Avec le connecteur logique binaire $\Leftrightarrow$ et deux propositions A, B on peut former la proposition $\Leftrightarrow AB$ (en notation préfixée), que nous représenterons par $A \Leftrightarrow B$. Nous lirons cette expression en disant « A équivaut à B » ou « A si et seulement si B ». Par exemple, en analyse, la proposition suivante est vraie :

$$(x^2 \geq 4) \Leftrightarrow (x \leq -2 \text{ ou } x \geq 2). \quad (6)$$

Le symbole $\Leftrightarrow$ ne doit pas être pris non plus pour un symbole du métalangage. Une proposition de la forme $A \Leftrightarrow B$ n'est pas une assertion *sur les propositions* A et B, mais sur les objets dont il est question dans ces propositions. Par exemple, (6) n'est pas une assertion sur la proposition $(x^2 \geq 4)$ et sur la proposition $(x \leq -2 \text{ ou } x \geq 2)$. C'est une assertion sur les objets nommés $x, 2, 4$. Il n'y a pas de doute à ce sujet lorsqu'on lit (6) en disant

$$x^2 \geq 4 \text{ si et seulement si } (x \leq -2 \text{ ou } x \geq 2)$$

plutôt qu'avec le verbe « équivaut » : $x^2 \geq 4$ équivaut à $(x \leq -2 \text{ ou } x \geq 2)$.

Dans le langage mathématique usuel, à la place de « si et seulement si », on se sert encore d'autres tournures, utilisant les mots « il faut et il suffit » ou l'expression « condition nécessaire et suffisante ». Nous discuterons ces tournures plus loin (§7.10).

Le symbole $\Leftrightarrow$ est appelé *symbole d'équivalence logique* et une proposition de la forme $A \Leftrightarrow B$ est appelée une équivalence. On rencontre d'autres symboles dans ce rôle dans la littérature sur la logique. Le plus fréquent d'entre eux est le symbole $\equiv$, avec lequel on écrit $A \equiv B$ au lieu de $A \Leftrightarrow B$. On rencontre aussi les notations $A \leftrightarrow B$, $A \supset \subset B$.

La proposition fausse $\perp$. Nous avons introduit le symbole $\perp$ au début de ce paragraphe comme étant un connecteur logique nullaire, ou d'arité zéro. De façon générale lorsqu'on assemble un connecteur logique n -aire avec n propositions on obtient une proposition. Donc le connecteur nullaire $\perp$ forme à lui seul, c'est-à-dire combiné avec zéro propositions, une proposition. Celle-ci aura la propriété d'être toujours fausse et nous l'appellerons *la proposition fausse*.

Mais il y a toujours, dans un contexte donné, une infinité d'autres propositions fausses, en particulier toutes celles qui sont la négation d'une proposition vraie. La proposition $\perp$ peut être appelée la proposition fausse « par excellence ».

3.7 Quantificateurs

À chaque variable individuelle x d'un langage du premier ordre sont associés deux *quantificateurs* : le quantificateur *universel* $\forall x$ (pour tout x) et le quantificateur *existentiel* $\exists x$ (il existe un x). Avec ceux-ci et une proposition A quelconque, on peut former les propositions $\forall xA$ et $\exists xA$. D'autres auteurs appellent quantificateurs les symboles $\forall$ et $\exists$ eux-mêmes, sans variable adjointe. Nous soulignons donc que dans le présent ouvrage ce sont les combinaisons de symboles de la forme $\forall x$, $\exists x$, où x est une variable individuelle, qui sont appelées quantificateurs. Ces combinaisons sont considérées chacune comme un seul symbole. Il y en a une infinité, puisqu'il y a une infinité de variables individuelles.

Remarque : métavariabiles. Nous attirons encore l'attention du lecteur sur l'usage que nous faisons des lettres droites x et A ci-dessus. Ce sont toujours des *métavariabiles* (variables du métalangage). Nous utilisons systématiquement les lettres A, B, C pour désigner des propositions arbitraires, et les lettres x, y, z pour désigner des *variables individuelles arbitraires* d'un langage du premier ordre. Nous rappelons que nous avons choisi et fixé (§3.3) comme répertoire de variables individuelles, pour tous nos langages du premier ordre, les lettres *italiques*

$$a, b, \dots, z, A, B, \dots, Z$$

et leurs variantes munies d'accents $a', \dots, Z', a'', \dots, Z'', a''', \dots$. Lorsque nous parlons d'une proposition de la forme $\forall xA$, nous utilisons la métavARIABLE x pour désigner une variable quelconque, indéterminée, de ce répertoire. Lorsque nous disons qu'à chaque variable individuelle x d'un langage du premier ordre est associé le quantificateur universel $\forall x$, cela veut dire qu'à la variable a est associé le quantificateur $\forall a$, qu'à la variable b est associé le quantificateur $\forall b$, et ainsi de suite. C'est pour éviter cette énumération que nous nous servons d'une métavARIABLE x pouvant représenter n'importe laquelle de ces variables. En particulier, il ne faut pas confondre les métavariabiles x, y, z avec les variables particulières x, y, z du répertoire standard fixé.

Deux métavariabiles distinctes peuvent désigner le même symbole ou la même expression d'un langage formel, si on ne précise pas le contraire. Par exemple, lorsque nous parlons d'une proposition de la forme $(A \text{ ou } B)$, les propositions désignées par A et par B peuvent être la même proposition. De même, les métavariabiles x et y peuvent désigner la même variable du répertoire standard, si on ne précise pas le contraire. Par exemple, dans une proposition de la forme $\forall xA$, la proposition A peut elle-même commencer par un quantificateur, c'est-à-dire être elle-même de la forme $\forall yB$, où B est une proposition.

On obtient ainsi une proposition de la forme $\forall x \forall y B$. Les propositions de cette forme sont toutes les propositions qui commencent par deux quantificateurs universels (au moins, car B peut en contenir d'autres). Mais, comme il n'est pas précisé que x et y désignent deux variables distinctes, ces deux quantificateurs peuvent être *en particulier* deux occurrences du même quantificateur.

Expressions usuelles pour les quantificateurs. Dans le langage mathématique usuel, de nombreuses tournures tiennent lieu de quantificateurs. Le quantificateur $\forall x$ s'exprime normalement par

pour tout x ,
pour chaque x ,
quel que soit x .

Ces expressions viennent normalement avant la proposition à laquelle elles se rapportent, mais elles sont souvent placées après celle-ci. Au lieu de

$$\forall x (x^2 \geq 0), \tag{1}$$

on écrit par exemple: $x^2 \geq 0$ quel que soit x , ou encore $x^2 \geq 0$ ($\forall x$).

Mais il y a d'autres tournures, dans lesquelles le quantificateur universel se trouve plus « dissimulé ». Pour la proposition (1) ce sont par exemple, par ordre de dissimulation croissante:

le carré de tout nombre x est supérieur à 0,
le carré d'un nombre x est toujours supérieur à 0,
le carré d'un nombre est toujours supérieur à 0,
le carré d'un nombre x est supérieur à 0,
le carré d'un nombre est supérieur à 0.

Les deux dernières sont ambiguës: il faut comprendre que « un nombre » veut dire « tout nombre », et pas « un nombre particulier ».

Pour les quantificateurs existentiels, les tournures de langage employées en mathématique sont aussi nombreuses. L'expression normale pour une proposition de la forme $\exists x A$ est

il existe un x tel que A . (2)

Lorsqu'il y a deux quantificateurs existentiels au début d'une proposition, $\exists x \exists y A$, on devrait donc dire « il existe un x tel que (il existe un y tel que A) ». Au lieu de cela, on dit souvent: « il existe un x et un y tels que A ».

La tournure favorite de certains pédagogues pour une proposition de la forme $\exists x A$ est

on peut trouver un x tel que A ,

quand bien même il y a une certaine différence entre affirmer l'existence d'une chose et prétendre qu'on peut la trouver. Lorsqu'un crime est commis, il existe un coupable, mais cela ne veut pas dire qu'on peut le trouver, la police en sait quelque chose! L'interprétation de « il existe » est d'ailleurs le point sur lequel l'école des mathématiques dites *constructives* est en désaccord avec celle

des mathématiques classiques. Les constructivistes ne reconnaissent l'existence d'une chose précisément que lorsqu'on dispose d'un procédé permettant de la trouver *effectivement*, alors que les mathématiques classiques admettent d'autres preuves d'existence et démontrent ainsi l'existence de choses que personne n'a jamais « pu trouver ».

Les quantificateurs existentiels se « cachent », encore plus que les quantificateurs universels, dans beaucoup d'expressions usuelles et il faut apprendre à les déceler. Nous en donnons ici un exemple. Considérons un nombre réel b et un intervalle I de la droite réelle. Pour la proposition

$$\exists x(x \in I \text{ et } b = x^2), \quad (3)$$

on pourra rencontrer les expressions suivantes:

il existe un $x \in I$ tel que $x^2 = b$,
 I contient un nombre dont le carré est égal à b ,
 on a $b = x^2$ pour un $x \in I$,
 b est de la forme x^2 avec $x \in I$.

On est obligé de constater que la langue mathématique éprouve le même besoin esthétique de varier ses tournures que la langue littéraire. Elle n'y gagne pas toujours en clarté. Lorsqu'un étudiant n'a eu affaire qu'à des enseignants qui au lieu de (3) disent toujours « on a $b = x^2$ pour un $x \in I$ », il est normal qu'il éprouve quelques difficultés lorsqu'il est confronté à la définition d'une fonction continue, par exemple, avec ses vrais quantificateurs « pour tout x , pour tout ε , il existe un δ , pour tout y , etc. ».

Quantificateurs typés. Le langage mathématique utilise rarement des quantificateurs « purs » tels que « pour tout x » et « il existe un x tel que ». Le x du quantificateur est presque toujours qualifié comme étant d'un certain *type*. On dit par exemple:

il existe un nombre réel x tel que,
 il existe un x dans l'intervalle $[a, b]$ tel que.

et de même avec « pour tout ». Il n'y a pas de quantificateurs de ce genre dans les langages du premier ordre que nous traiterons. Mais on peut s'en passer. Par exemple, les propositions

il existe un nombre réel x tel que $x^2 = 2$,
 pour tout nombre réel x , $x^2 \geq 0$

peuvent se traduire respectivement par:

$$\begin{aligned} \exists x(x \in \mathbb{R} \text{ et } x^2 = 2), \\ \forall x(x \in \mathbb{R} \Rightarrow x^2 \geq 0). \end{aligned}$$

3.8 Donnée d'un langage du premier ordre

Pour conclure ce chapitre sur les symboles d'un langage du premier ordre, nous voulons résumer en quoi consiste la *donnée* d'un tel langage. Rappelons que

les six catégories de symboles d'un tel langage sont :

- (a) les variables individuelles
- (b) les *constantes individuelles*
- (c) les *symboles fonctionnels*
- (d) les *symboles relationnels*
- (e) les connecteurs logiques
- (f) les quantificateurs.

D'après nos conventions, les variables individuelles sont les mêmes pour tous les langages du premier ordre. Ce répertoire de symboles a été fixé une fois pour toutes (§3.3). Il en va de même des connecteurs logiques et des quantificateurs. Ces symboles ne sont pas spécifiques à un langage du premier ordre particulier. Ce qui caractérise un langage du premier ordre particulier, ce sont ses trois répertoires de *constantes individuelles*, de *symboles fonctionnels*, et de *symboles relationnels*. La donnée d'un langage du premier ordre consiste donc dans la donnée ou déclaration des symboles de ces trois classes. Le nombre de symboles de chacune d'elles peut être quelconque, y compris zéro ou une infinité. Nous rappelons aussi que chaque symbole fonctionnel et chaque symbole relationnel doit être donné avec son *arité*: unaire, binaire, ternaire, etc.

Ces données ne concernent que la syntaxe d'un langage et ne disent rien de sa sémantique. En principe on peut ne s'intéresser qu'aux aspects syntaxiques mais en général un langage est introduit avec une certaine sémantique en vue, qu'on appelle précisément la *sémantique prévue*, ou *l'interprétation prévue*. Il est donc rare qu'un langage soit donné sans quelques indications sur cette interprétation.

Exemple. Les déclarations de symboles suivantes définissent un langage du premier ordre que nous appellerons le *langage de l'arithmétique du premier ordre*:

- 0 constante individuelle;
- σ symbole fonctionnel unaire;
- $+$ symbole fonctionnel binaire;
- $\cdot$ symbole fonctionnel binaire;
- $=$ symbole relationnel binaire.

La sémantique prévue est la suivante: les objets, ou individus, de l'univers du discours sont les entiers naturels, autrement dit les nombres $0, 1, 2, \dots$. Les symboles $0, +, \cdot, =$ ont leur signification usuelle: zéro, addition, multiplication, égalité. Le symbole fonctionnel unaire σ (lettre grecque sigma) représente l'opération «prendre le successeur» d'un nombre, à savoir le nombre qui lui succède dans l'ordre normal des entiers. Autrement dit, le successeur σx d'un nombre x est le nombre $x + 1$. Mais le symbole 1 n'appartient pas au langage introduit ici. Au lieu de $x + 1$ on écrit justement σx .

Ce langage permet d'écrire les propositions suivantes, que l'on prend comme axiomes d'une théorie appelée *Arithmétique du premier ordre*. La notion de théorie et les notions connexes d'axiome, de théorème, de démonstration, de déduction, seront longuement présentées au chapitre 5.

$$\begin{aligned} & \forall x(\sigma x \neq 0) \\ & \forall x \forall y(\sigma x = \sigma y \Rightarrow x = y) \\ & \forall x(x + 0 = x) \\ & \forall x(x \cdot 0 = 0) \\ & \forall x \forall y(x + \sigma y = \sigma(x + y)) \\ & \forall x \forall y(x \cdot \sigma y = x \cdot y + x). \end{aligned}$$

Ces axiomes sont présentés ici seulement comme des exemples de propositions du langage considéré et nous n'allons pas les discuter, ni en déduire de théorèmes. Le lecteur familiarisé avec le raisonnement par récurrence pourra essayer de démontrer par exemple le théorème de commutativité de l'addition, $\forall x \forall y(x + y = y + x)$, qui ne fait pas partie de ces axiomes. Le principe de récurrence est une méthode de déduction — nous parlerons d'un schéma de déduction — qui s'ajoute à ces axiomes, comme nous le verrons.

Suivant la syntaxe des langages du premier ordre qui sera définie au paragraphe 4.1, les symboles fonctionnels et relationnels doivent être préfixés et il n'y a pas de parenthèses. Ainsi l'expression $\forall x(\sigma x \neq 0)$ donnée ci-dessus comme premier axiome représente la proposition

$$\forall x \neg = \sigma x 0.$$

4

Syntaxe d'un langage du premier ordre

4.1 Diagrammes et règles syntaxiques

La syntaxe des termes et des propositions d'un langage du premier ordre $\mathcal{L}$ est définie par les diagrammes syntaxiques de la figure 4.1. Ces diagrammes sont une manière de représenter les règles suivantes.

Règles de construction des termes. (a) Si x est une variable individuelle de $\mathcal{L}$, l'expression qui se compose du seul symbole x est un terme de $\mathcal{L}$. (b) Si s est une constante individuelle de $\mathcal{L}$, l'expression qui se compose du seul symbole s est un terme de $\mathcal{L}$. (c) Si f est un symbole fonctionnel de $\mathcal{L}$ d'arité n ($n = 1, 2, \dots$), et si $T_1, \dots, T_n$ sont des termes de $\mathcal{L}$ au sens des présentes règles, alors l'expression

$$fT_1 \dots T_n$$

est un terme de $\mathcal{L}$. On désigne ainsi l'expression obtenue en écrivant successivement : le symbole f , puis les termes T_1 à T_n . Ces termes ne doivent pas être nécessairement tous distincts. Plusieurs d'entre eux peuvent être des occurrences d'un même terme. (d) Une séquence de symboles de $\mathcal{L}$ n'est un terme de $\mathcal{L}$ que si elle est construite suivant les règles (a)–(c).

Règles de construction des propositions. (a) Si r est un symbole relationnel de $\mathcal{L}$ d'arité n ($n = 1, 2, \dots$), et si $T_1, \dots, T_n$ sont des termes de $\mathcal{L}$ au sens des règles qui précèdent, l'expression

$$rT_1 \dots T_n,$$

obtenue en écrivant successivement le symbole r , puis les termes T_1 à T_n , est une proposition. Les propositions de cette forme sont appelées les propositions *atomiques* de $\mathcal{L}$. (b) L'expression réduite au seul symbole $\perp$ est une proposition de $\mathcal{L}$, appelée la *proposition fausse par excellence* de $\mathcal{L}$. (c) Si A est une proposition de $\mathcal{L}$ au sens des présentes règles, l'expression $\neg A$ est une proposition de $\mathcal{L}$, appelée la *négation* de A . (d) Si A et B sont des propositions de

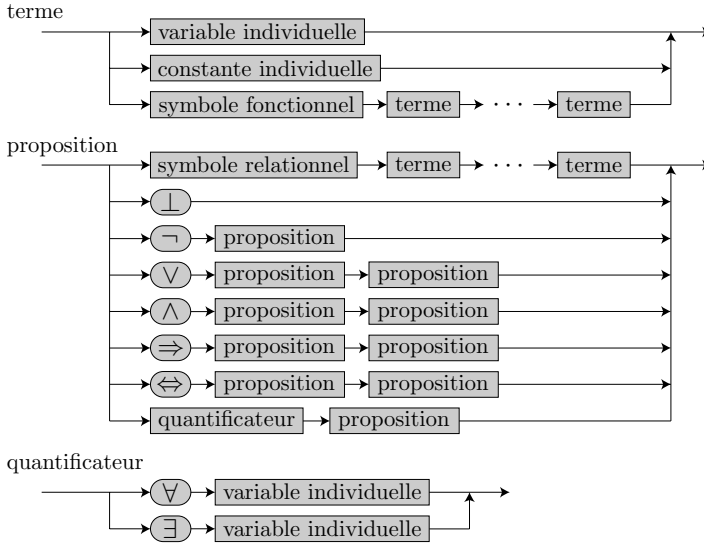


Figure 4.1 Syntaxe des termes et des propositions d'un langage du premier ordre. Il est sous-entendu que le nombre de termes combinés avec un symbole fonctionnel ou relationnel est égal à l'arité de ce symbole.

$\mathcal{L}$ au sens des présentes règles, les expressions

$$\vee AB, \quad \wedge AB, \quad \Rightarrow AB, \quad \Leftrightarrow AB$$

sont des propositions de $\mathcal{L}$. Les deux premières sont appelées respectivement la *disjonction* et la *conjonction* de A et de B . (e) Si A est une proposition de $\mathcal{L}$ au sens des présentes règles, et si x est une variable individuelle de $\mathcal{L}$, les expressions

$$\forall xA, \quad \exists xA$$

sont des propositions de $\mathcal{L}$. (f) Une séquence de symboles de $\mathcal{L}$ n'est une proposition de $\mathcal{L}$ que si elle est construite suivant les règles (a)–(e).

4.2 Constructions génératrices

Le diagramme syntaxique des termes de la figure 4.1 décrit la manière de construire un terme en utilisant d'autres termes, et en les combinant avec un symbole fonctionnel. Les termes utilisés dans cette construction doivent obéir eux-mêmes à cette syntaxe. C'est ce qu'on entend dans l'article (c) des règles de construction des termes en parlant de termes $T_1, \dots, T_n$ au sens des présentes règles. Donc ces termes doivent être eux-mêmes : soit des variables ou des constantes individuelles, soit des termes construits avec un symbole fonctionnel et d'autres termes encore. Et ces autres termes, à leur tour, doivent

être construits suivant ces règles. Mais cette décomposition ne peut pas se poursuivre indéfiniment, car le nombre de symboles d'un terme est fini.

L'idée que représentent le diagramme syntaxique ou les règles de construction des termes est qu'un terme doit être *engendré* à partir des termes de base que sont les variables et les constantes individuelles, au moyen d'un nombre *fini* d'opérations de combinaison avec des symboles fonctionnels. La notion de *construction génératrice* que nous introduisons dans ce paragraphe précise cette idée. Nous en donnons d'abord un exemple.

Exemple 1. Considérons un langage $\mathcal{L}$ dans lequel il y a les symboles suivants (il peut y en avoir d'autres, mais ils n'interviennent pas dans l'exemple):

- 0 constante individuelle
- symbole fonctionnel unaire
- $+$, $*$ symboles fonctionnels binaires
- $=$ symbole relationnel binaire.

Nous voulons montrer en quoi l'expression

$$+ * y - x 0 \tag{1}$$

est un terme au sens défini par les règles de construction des termes (a)–(c) (§4.1) ou par le diagramme syntaxique de la figure 4.1. Cette expression correspond à la suivante en notation infixée usuelle:

$$(y * (-x)) + 0, \tag{2}$$

mais le passage de l'une à l'autre n'est pas ce qui nous intéresse ici. En vertu des règles mentionnées, les symboles x et y sont des termes de $\mathcal{L}$, parce que ce sont des variables individuelles de $\mathcal{L}$ (§3.3, Fig. 3.1). L'expression $-x$ est donc un terme de $\mathcal{L}$, d'après la règle (c), puisqu'elle se compose d'un symbole fonctionnel unaire de $\mathcal{L}$ suivi d'un terme de $\mathcal{L}$. L'expression $*y - x$ est un terme de $\mathcal{L}$ d'après la même règle, puisqu'elle se compose d'un symbole fonctionnel binaire de $\mathcal{L}$ suivi de deux termes de $\mathcal{L}$ que nous venons de mentionner: y et $-x$. Le symbole 0 est un terme $\mathcal{L}$ d'après la règle (b), puisque c'est une constante individuelle de $\mathcal{L}$. Finalement, l'expression (1) toute entière est un terme de $\mathcal{L}$, puisqu'elle se compose d'un symbole fonctionnel binaire suivi de deux termes que nous venons de mentionner: $*y - x$ et 0 . Cette analyse peut être résumée par la séquence d'expressions T_1 à T_6 suivante:

$$\begin{array}{ll} T_1: & x \\ T_2: & y \\ T_3: & -x \\ T_4: & *y - x \\ T_5: & 0 \\ T_6: & + *y - x 0 \end{array} \tag{3}$$

Chacune des expressions de cette suite est un terme du langage $\mathcal{L}$ considéré, en vertu de la manière dont cette suite est construite, c'est-à-dire de la propriété

suivante: chaque expression T_k de la suite qui ne se réduit pas à une variable ou une constante individuelle de $\mathcal{L}$ est composée avec un symbole fonctionnel de $\mathcal{L}$ et des expressions qui *figurent avant* T_k dans la suite.

Par exemple, T_4 est l'expression $*T_2T_3$, et T_6 est l'expression $+T_4T_5$.

Définition. Une séquence d'expressions $T_1, \dots, T_p$ composées de symboles d'un langage du premier ordre $\mathcal{L}$ est appelée une *construction génératrice de termes* de $\mathcal{L}$ si, pour chacune des expressions T_k de cette suite, l'une des conditions suivantes est vérifiée:

(a) T_k est une variable individuelle de $\mathcal{L}$.

(b) T_k est une constante individuelle de $\mathcal{L}$.

(c) T_k se compose d'un symbole fonctionnel de $\mathcal{L}$ suivi de n expressions qui figurent avant T_k dans la suite, n étant l'arité du symbole fonctionnel.

En nous servant de ce concept, nous pouvons donner un sens précis à la règle (d) (§4.1) de construction des termes. Elle signifie qu'une expression T n'est un terme de $\mathcal{L}$ que si l'on peut écrire une construction génératrice de termes $T_1, \dots, T_p$ de $\mathcal{L}$ telle que T soit identique à T_p . C'est ce que nous avons fait dans l'exemple qui précède pour l'expression (1). Celle-ci est identique à la dernière expression de la construction génératrice (3).

Finalement, les règles de construction de termes du paragraphe 4.1 doivent être considérées comme des règles sur la manière d'écrire les constructions génératrices de termes de $\mathcal{L}$, et les termes de $\mathcal{L}$ sont *par définition* les expressions qui figurent dans ces constructions.

Exemple 2. Soit $\mathcal{L}$ le langage de l'exemple 1. Il s'agit cette fois-ci de vérifier que l'expression suivante est une proposition de $\mathcal{L}$:

$$\forall x \forall y \Rightarrow = * x y 0 \vee = x 0 = y 0. \quad (4)$$

L'expression correspondante, dans la notation infixée usuelle, est

$$\forall x \forall y (x * y = 0 \Rightarrow (x = 0 \vee y = 0)). \quad (5)$$

Le procédé d'analyse est analogue à celui de l'exemple 1. On écrit une *construction génératrice de propositions* de $\mathcal{L}$, dans laquelle on forme des sous-expressions de tailles croissantes de l'expression (4), qui sont toutes des propositions de $\mathcal{L}$, et dont la dernière est l'expression (4) elle-même:

$$\begin{aligned} R_1: & \quad = x 0 \\ R_2: & \quad = y 0 \\ R_3: & \quad \vee = x 0 = y 0 \\ R_4: & \quad = * x y 0 \\ R_5: & \quad \Rightarrow = * x y 0 \vee = x 0 = y 0 \\ R_6: & \quad \forall y \Rightarrow = * x y 0 \vee = x 0 = y 0 \\ R_7: & \quad \forall x \forall y \Rightarrow = * x y 0 \vee = x 0 = y 0. \end{aligned}$$

Chacune des expressions R_k de cette suite est une proposition de $\mathcal{L}$, d'après l'une des règles de construction de propositions (a)–(e) (§4.1): les expressions

R_1 et R_2 sont des propositions de $\mathcal{L}$ d'après la règle (a), parce que x , y , 0 sont des termes de $\mathcal{L}$ et le symbole $=$ est un symbole relationnel binaire de $\mathcal{L}$; l'expression R_3 est une proposition de $\mathcal{L}$ d'après la règle (d) parce qu'elle est la disjonction $\vee R_1 R_2$; l'expression R_4 est une proposition de $\mathcal{L}$ d'après la règle (a), car elle se compose du symbole relationnel binaire $=$, du terme $*xy$ et du terme 0 ; l'expression R_5 est une proposition de $\mathcal{L}$ d'après la règle (d) parce qu'elle est identique à $\Rightarrow R_4 R_3$; l'expression R_6 est une proposition de $\mathcal{L}$ d'après la règle (e) parce qu'elle est identique à $\forall y R_5$; et finalement l'expression R_7 , qui est l'expression (4), est une proposition de $\mathcal{L}$ d'après la règle (e), puisqu'elle est identique à $\forall x R_6$.

On voit que la propriété essentielle de cette suite d'expressions est que chacune des expressions R_k qui n'est pas une proposition atomique de $\mathcal{L}$ est composée au moyen d'un connecteur logique ou d'un quantificateur et d'une ou de deux expressions qui *figurent avant* R_k dans la suite. C'est ce qui permet de montrer de proche en proche que chacune des expressions R_k est une proposition de $\mathcal{L}$ d'après l'une ou l'autre des règles de construction de propositions (a)–(e).

Définition. Une séquence d'expressions $R_1, \dots, R_p$ composées de symboles d'un langage du premier ordre $\mathcal{L}$ est appelée une *construction génératrice de propositions* de $\mathcal{L}$ si, pour chacune des expressions R_k de cette suite, l'une des conditions suivantes est vérifiée:

- (a) R_k est une proposition atomique de $\mathcal{L}$, c'est-à-dire se compose d'un symbole relationnel de $\mathcal{L}$ suivi de n termes de $\mathcal{L}$, n étant l'arité du symbole relationnel.
- (b) R_k se compose du seul symbole $\perp$.
- (c) R_k se compose du connecteur logique $\neg$ suivi d'une expression qui figure avant R_k dans la suite.
- (d) R_k se compose d'un des quatre connecteurs logiques $\vee$, $\wedge$, $\Rightarrow$, $\Leftrightarrow$, suivi de deux expressions qui figurent avant R_k dans la suite.
- (e) R_k se compose d'un quantificateur suivi d'une expression qui figure avant R_k dans la suite.

Nous pouvons répéter pour les propositions ce que nous avons dit pour les termes. La règle (f) (§4.1) de construction des propositions signifie qu'une expression R n'est une proposition de $\mathcal{L}$ que si l'on peut écrire une construction génératrice de propositions $R_1, \dots, R_p$ de $\mathcal{L}$ telle que R soit identique à R_p . C'est ce que nous avons fait dans l'exemple 2 pour l'expression (4).

Finalement, les règles de construction de propositions du paragraphe 4.1 doivent être considérées comme des règles sur la manière d'écrire les constructions génératrices de propositions de $\mathcal{L}$, et les propositions de $\mathcal{L}$ sont *par définition* les expressions qui figurent dans ces constructions.

Remarque sur les notations du métalangage. Le lecteur aura remarqué la manière dont nous exprimons qu'une méta-expression désigne une expression

particulière, ou que deux méta-expressions désignent la même expression d'un langage formel. Par exemple, à propos de la construction génératrice de termes de l'exemple 1, nous disons:

$$\begin{aligned} T_4 &\text{ est l'expression } *y - x; \\ T_6 &\text{ est l'expression } +T_4T_5. \end{aligned}$$

À propos de la construction génératrice de propositions de l'exemple 2, nous disons:

$$\begin{aligned} R_3 &\text{ est la disjonction } \vee R_1R_2; \\ R_5 &\text{ est identique à } \Rightarrow R_4R_3. \end{aligned}$$

Au lieu des mots « est l'expression », ou « est identique à », on est naturellement tenté d'utiliser un métasymbole auquel on donnerait ce sens. C'est bien entendu possible, mais il serait maladroit de se servir pour cela du symbole d'égalité ($=$). Car ce symbole est constamment utilisé dans les langages formels, et l'on ne devrait pas prendre comme symbole du métalangage un symbole du langage formel. Par exemple, au lieu de « R_4 est l'expression $= *xy0$ », il serait bien maladroit d'écrire $R_4 = *xy0$. Certains auteurs utilisent le symbole $\equiv$. Par exemple, $R_5 \equiv \Rightarrow R_4R_3$. Mais ce symbole est aussi fréquemment utilisé dans les langages formels. Nous utilisons parfois le signe « : » pour associer un métasymbole à une expression particulière, comme dans les constructions génératrices des exemples 1 et 2. Par exemple $R_4 : \quad = *xy0$.

4.3 Arbres des termes et propositions

En dehors des expressions réduites à un seul symbole, c'est-à-dire les variables, les constantes individuelles et la proposition $\perp$, toutes les expressions d'un langage du premier ordre sont de la forme générale

$$sA_1 \cdots A_n, \tag{1}$$

avec les possibilités suivantes: s est un symbole fonctionnel ou un symbole relationnel n -aire et $A_1, \dots, A_n$ sont des termes; s est le connecteur logique unaire $\neg$ ou un quantificateur, n est égal à 1, et A_1 est une proposition; s est un des connecteurs logiques binaires ($\vee, \wedge, \Rightarrow, \Leftrightarrow$), n est égal à 2, et A_1, A_2 sont des propositions.

Dans tous les cas, nous dirons que le symbole s est le *symbole principal* de l'expression (1) et que les expressions $A_1, \dots, A_n$ lui sont *subordonnées* dans l'expression (1). Chacune des sous-expressions A_i qui n'est pas réduite à un seul symbole se compose elle-même d'un symbole principal auquel sont subordonnées un certain nombre d'expressions. Nous voyons ainsi dans une expression générale une structure hiérarchique et pour mettre cette hiérarchie en évidence nous représenterons parfois les expressions par des *arbres*.

Exemples. Nous considérons le même langage $\mathcal{L}$ que dans les exemples du

paragraphe §4.2, et trois expressions. Premièrement, l'expression

$$+ x 0,$$

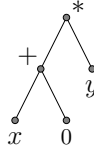
est un terme de $\mathcal{L}$ et s'écrit $x + 0$ en notation infixée. L'arbre correspondant est le diagramme suivant :



Le symbole fonctionnel $+$ est le symbole principal de ce terme et les termes x et 0 lui sont subordonnés, ce qui est représenté par leur position inférieure dans l'arbre. Deuxièmement, le terme

$$* + x 0 y,$$

en notation infixée $(x + 0) * y$, est représenté par l'arbre



Troisièmement, la proposition

$$\forall x \forall y \Rightarrow = * x y 0 \vee = x 0 = y 0,$$

en notation usuelle $\forall x \forall y (x * y = 0 \Rightarrow (x = 0 \text{ ou } y = 0))$, est représentée par l'arbre de la figure 4.2.

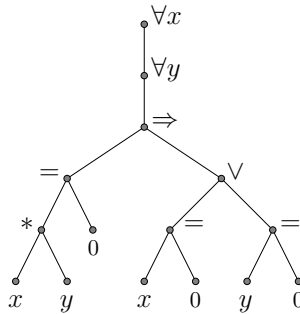


Figure 4.2 Arbre de la proposition $\forall x \forall y (x * y = 0 \Rightarrow (x = 0 \text{ ou } y = 0))$.

Définition. Un arbre est un diagramme composé de *nœuds*, qui sont des « points » ou des petits cercles, et d'*arcs*, qui sont des traits reliant chacun deux points. Chaque nœud de l'arbre d'une expression (terme ou proposition) d'un langage du premier ordre « porte » un symbole de l'expression, qui est noté près du nœud. Il ne faut pas confondre un nœud avec le symbole qu'il porte, car deux nœuds distincts peuvent porter le même symbole. Par exemple,

dans l'arbre de la figure 4.2, il y a trois nœuds portant le symbole d'égalité, et deux nœuds portant la variable x .

Les arcs des arbres que nous considérons sont dirigés de haut en bas par rapport à la page. Ils descendent. On dit qu'un nœud K d'un arbre est un *successeur* d'un nœud N s'il y a un arc qui descend de N à K . Inversement, on dit que N est *prédécesseur* de K . La structure d'arbre est caractérisée par les propriétés suivantes:

- (a) Il y a un nœud sans prédécesseur et un seul, appelé *racine* de l'arbre.
- (b) Les autres nœuds possèdent chacun un prédécesseur et un seul.
- (c) En partant de n'importe quel nœud autre que la racine, si l'on répète l'opération « remonter au prédécesseur », on aboutit toujours après un nombre fini de pas à la racine de l'arbre.

Les arbres que nous considérons sont *orientés*, non seulement de haut en bas mais encore de gauche à droite dans le sens suivant: les successeurs de chaque nœud sont numérotés de 1 à n , où n est le nombre de successeurs du nœud. Il y a donc un ordre parmi ces successeurs: le premier, le deuxième, etc. Ces numéros peuvent être notés à côté des arcs correspondants, mais nous les omettons en dessinant toujours des arcs descendants qui ne se croisent pas, et en convenant que la numérotation va de gauche à droite.

Les nœuds qui n'ont pas de successeur sont appelés les *feuilles* de l'arbre. Dans l'arbre d'une expression, ces nœuds portent toujours des variables ou des constantes individuelles, et ces deux catégories de symboles ne se trouvent que sur les feuilles de l'arbre.

On appelle *branche* d'un arbre toute suite de nœuds $N_1, \dots, N_p$ où N_1 est la racine de l'arbre, N_p est une feuille, et pour $k = 2, \dots, p$, N_k est un successeur de N_{k-1} . Chaque feuille d'un arbre appartient à une branche et une seule. Par exemple, dans l'arbre de la figure 4.2, la branche de la seconde occurrence de x porte les symboles $\forall x, \forall y, \Rightarrow, \vee, =, x$.

La manière générale de construire l'arbre d'une expression

$$sA_1 \cdots A_n$$

est représentée dans la figure 4.3. La racine de l'arbre porte le symbole principal s . Il y a n arcs qui partent de la racine et vont vers les racines de n sous-arbres, représentant les expressions subordonnées A_1 à A_n . Chacun de ces n sous-arbres est construit à son tour selon le même principe.

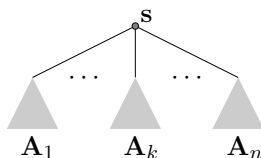


Figure 4.3 Arbre d'une expression $sA_1 \cdots A_n$.

4.4 Variables libres et variables liées

Introduction. Dans de nombreuses notations mathématiques il y a des variables qu'un informaticien qualifierait de *locales*, et d'autres qui sont *globales*. Par exemple, dans le terme

$$\sum_{n=1}^p n^2, \quad (1)$$

qui représente la somme des carrés des p premiers nombres entiers non nuls, la variable n est locale, alors que la variable p est globale. Le caractère local de n se manifeste de deux manières. Premièrement, on peut remplacer toutes les occurrences de cette variable par n'importe quelle autre variable (sauf p) sans changer la valeur de l'expression. Par exemple, le terme

$$\sum_{k=1}^p k^2 \quad (2)$$

représente le même nombre que le terme (1). Secondement, si l'on utilise le terme (1) comme partie d'une expression dans laquelle il y a d'autres occurrences de n , ces occurrences n'ont rien à voir avec celles de n dans (1). Considérons par exemple la proposition suivante :

$$n \leq p \Rightarrow n \leq \sum_{k=1}^p k^2. \quad (3)$$

Puisque nous admettons que les termes (1) et (2) désignent le même nombre, nous pouvons remplacer le second par le premier dans (3), ce qui donne la proposition

$$n \leq p \Rightarrow n \leq \sum_{n=1}^p n^2. \quad (4)$$

Dans cette dernière proposition, les deux occurrences de n en dehors de la somme n'ont rien à voir avec les occurrences de n dans la somme, puisqu'on peut remplacer celles-ci par k . Lorsqu'on lit la proposition (4), on distingue mentalement « deux » variables n : l'une qui est globale et l'autre qui est locale dans la somme. La proposition (3) est sans doute plus lisible que (4) et l'on évite en général, en mathématique, d'utiliser la même variable à la fois comme variable globale et comme variable locale d'une expression. Mais cela n'a rien d'illégal.

Nous allons voir que dans une proposition d'un langage du premier ordre qui commence par un quantificateur la variable du quantificateur est locale à la proposition. Supposons par exemple que a désigne un ensemble. La proposition

$$\exists x(x \in a) \quad (5)$$

signifie que cet ensemble possède au moins un élément. Il est clair que les variables a et x ont, sémantiquement, des statuts très différents dans cette proposition. On peut remplacer la variable x , dans le quantificateur et dans la

proposition $x \in a$ par une autre variable, disons y , sans changer le sens de la proposition :

$$\exists y(y \in a). \quad (6)$$

Mais si l'on remplace a par b , cela peut changer complètement le sens de la proposition, puisque b peut désigner un autre ensemble que a . La variable x est locale dans (5), alors que a est globale. Si l'on combine (5) avec une autre proposition qui contient des occurrences de x , celles-ci n'ont rien à voir avec celles de (5). Par exemple, dans la proposition

$$x = \emptyset \text{ et } \exists x(x \in a), \quad (7)$$

il est affirmé qu'un certain ensemble, désigné par x , est vide, et que l'ensemble désigné par a possède au contraire un élément. Le x de $x = \emptyset$ n'a rien à voir avec celui de $\exists x(x \in a)$. Dans ce cas aussi, on préférera changer de variable locale et écrire « $x = \emptyset$ et $\exists y(y \in a)$ » au lieu de (7), mais la proposition (7) est parfaitement légale syntaxiquement.

Un quantificateur $\exists x$ ou $\forall x$, en plus de la signification « il existe » ou « pour tout », peut être considéré comme une déclaration que x est variable locale dans la proposition qui suit le quantificateur. Mais nous devons préciser les choses car il peut y avoir, dans une même proposition, plusieurs déclarations de ce genre, imbriquées, pour la même variable x . Par exemple, on peut écrire une proposition de la forme

$$\exists x((A \text{ et } \exists xB) \text{ ou } C). \quad (8)$$

Dans une telle proposition, le second quantificateur « masque » le premier dans le sens suivant : supposons qu'il n'y ait pas d'autres quantificateurs $\exists x$ ou $\forall x$ dans les propositions A , B , et C , ce qui compliquerait la situation; alors les occurrences de x dans la partie A et dans la partie C de (8) sont des occurrences de la variable déclarée par le premier quantificateur, alors que les occurrences de x dans la partie B sont des occurrences de la variable du second quantificateur. Nous dirons que les occurrences de x dans A et dans C sont *liées* au premier quantificateur, alors que celles de x dans B sont liées au second quantificateur.

En raison de ces liens, on utilise en logique les termes de *variables liées* pour les variables locales d'une expression, et de *variables libres* pour les variables globales. Nous allons maintenant définir exactement ces concepts.

Liens d'une proposition. Pour cela, nous représenterons graphiquement au moyen de traits les liens entre les occurrences liées des variables et les quantificateurs correspondants. Exemple : la proposition $\exists x(x \in a)$ sera notée

$$\exists x \overbrace{(x \in a)}. \quad (9)$$

Le trait qui relie l'occurrence de x dans $x \in a$ et le quantificateur $\exists x$ indique précisément : 1° que cette occurrence de x est liée; et 2° à quel quantificateur

elle est liée. Un exemple plus compliqué :

$$\exists x \left(x \in a \text{ et } \forall x \left(x \in a \Rightarrow x \in b \right) \right). \quad (10)$$

Règles définissant les liens d'une proposition. (a) Dans une proposition, chaque occurrence d'une variable x est soit *libre*, soit *liée*, l'un excluant l'autre; lorsqu'elle est liée, elle est liée à une occurrence particulière du quantificateur $\forall x$ ou du quantificateur $\exists x$, et nous dirons qu'il y a un lien entre les deux. (b) Dans une proposition atomique, toutes les occurrences de variables sont libres. (c) Lorsqu'on forme une proposition composée en ajoutant à une proposition A un quantificateur $\forall x$ (resp. $\exists x$), les occurrences de x qui étaient libres dans A deviennent liées à ce préfixe dans la proposition composée $\forall x A$ (resp. $\exists x A$). Nous disons que l'adjonction du quantificateur *crée* les liens correspondants. (d) Les autres opérations de construction de propositions, au moyen de connecteurs logiques, ne créent pas de liens.

Ces règles déterminent les liens d'une proposition lorsqu'on suit les étapes d'une construction génératrice de celle-ci. Par exemple, pour la proposition (10), une construction génératrice est la suite de propositions :

$$\begin{aligned} R_1: & \quad x \in a \\ R_2: & \quad x \in b \\ R_3: & \quad x \in a \Rightarrow x \in b \\ R_4: & \quad \forall x (x \in a \Rightarrow x \in b) \\ R_5: & \quad x \in a \text{ et } \forall x (x \in a \Rightarrow x \in b) \\ R_6: & \quad \exists x (x \in a \text{ et } \forall x (x \in a \Rightarrow x \in b)). \end{aligned}$$

Les créations de liens ont lieu en passant de R_3 à R_4 et de R_5 à R_6 .

Liens dans l'arbre d'une proposition. En considérant l'arbre d'une proposition, on peut définir les occurrences libres et les occurrences liées des variables, ainsi que leurs liens, par deux règles très simples. Celles-ci sont illustrées par la figure 4.4 contenant l'arbre de la proposition (10). On rappelle d'abord que les occurrences d'une variable, dans l'arbre d'une proposition, se trouvent toujours sur des feuilles de l'arbre — nous ne comptons pas comme occurrences les variables qui figurent dans les quantificateurs. Les règles sont les suivantes :

(a) Une occurrence d'une variable x dans une proposition A est *libre* s'il n'y a pas de quantificateur avec la variable x ($\forall x$ ou $\exists x$) sur la branche de l'arbre qui contient cette occurrence de x .

(b) S'il y a un ou plusieurs quantificateurs avec la variable x sur la branche d'une occurrence de x , alors cette occurrence est liée au premier de ces quantificateurs que l'on rencontre en remontant la branche.

Variables libres d'une proposition et d'un terme. On dit qu'une variable x *figure librement* dans une proposition A s'il y a dans cette proposition au

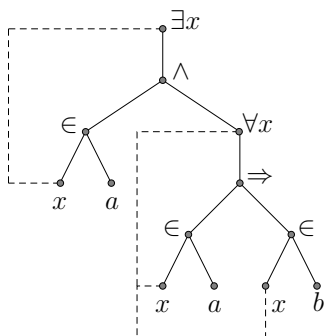


Figure 4.4 Liens dans l'arbre d'une proposition.

moins une occurrence libre de x . Nous aurons souvent affaire à la négation de cette assertion, à savoir la phrase « x ne figure pas librement dans A ». Cela veut dire qu'il n'y a pas d'occurrence libre de x dans A . Il est clair que cette condition est vérifiée, en particulier, si A est une proposition dans laquelle il n'y a aucune occurrence de x , ni libre, ni liée. Par exemple la variable y ne figure pas librement dans les propositions

$$\exists y(y + x = 0), \quad z + x = 0.$$

Une variable x peut avoir des occurrences libres et des occurrences liées dans une même proposition A . Un exemple est fourni par la proposition R_5 de la construction génératrice donnée plus haut. La première occurrence de x dans cette proposition est libre; la deuxième et la troisième sont liées. Nous appelons *variables libres* d'une proposition A toutes les variables qui figurent librement dans A . Ce sont donc toutes les variables qui possèdent au moins une occurrence libre dans A .

Nous dirons qu'une variable x *figure librement* dans un terme T d'un langage $\mathcal{L}$ si elle figure simplement dans T . Il est clair qu'il ne peut pas y avoir d'occurrence liée de x dans T , puisqu'un terme ne contient pas de quantificateur. Il peut donc sembler inutile de dire que x figure *librement* dans T lorsque x figure dans T . Mais plus loin (chap. 13) nous introduirons des langages du premier ordre plus généraux, dans lesquels il y aura d'autres opérateurs qui lient les variables comme les quantificateurs, et de tels opérateurs figureront dans certains termes. En prévision de ce développement, nous employons dès maintenant la formule « x figure librement dans T ». Pour le moment, elle est simplement synonyme de « x figure dans T ».

Règles sur les variables libres. Étant donné une variable individuelle x d'un langage du premier ordre $\mathcal{L}$, les règles suivantes permettent de déterminer *mécaniquement*, pour chaque terme ou proposition A de $\mathcal{L}$, si x figure librement dans A ou non. Un programmeur qui maîtrise la programmation dite récursive peut facilement écrire, d'après ces règles, un programme de test de la condition « x figure librement dans A ».

- (a) x figure librement dans le terme de $\mathcal{L}$ réduit au seul symbole x .
- (b) Si y est une variable individuelle de $\mathcal{L}$ distincte de x , x ne figure pas librement dans le terme de $\mathcal{L}$ réduit au seul symbole y .
- (c) Si c est une constante individuelle de $\mathcal{L}$, x ne figure pas librement dans le terme de $\mathcal{L}$ réduit au seul symbole c .
- (d) Si f est un symbole fonctionnel n -aire de $\mathcal{L}$, et si $T_1, \dots, T_n$ sont des termes de $\mathcal{L}$, alors x figure librement dans le terme composé $fT_1 \dots T_n$ de $\mathcal{L}$ si et seulement si x figure librement dans l'un au moins des termes T_i .
- (e) Si r est un symbole relationnel n -aire de $\mathcal{L}$, et si $T_1, \dots, T_n$ sont des termes de $\mathcal{L}$, alors x figure librement dans la proposition atomique $rT_1 \dots T_n$ de $\mathcal{L}$ si et seulement si x figure librement dans l'un au moins des termes T_i .
- (f) Si A est une proposition de $\mathcal{L}$, x figure librement dans la proposition $\neg A$ de $\mathcal{L}$ si et seulement si x figure librement dans A .
- (g) Si A et B sont des propositions de $\mathcal{L}$, x figure librement dans la proposition $A \vee B$ (resp. $A \wedge B$, $A \Rightarrow B$, $A \Leftrightarrow B$) si et seulement si x figure librement dans l'une au moins des propositions A , B .
- (h) Si A est une proposition de $\mathcal{L}$, x ne figure pas librement dans les propositions $\forall xA$, $\exists xA$.
- (i) Si A est une proposition de $\mathcal{L}$ et si y est une variable de $\mathcal{L}$ distincte de x , x figure librement dans la proposition $\forall yA$ (resp. $\exists yA$) de $\mathcal{L}$ si et seulement si x figure librement dans A .

Variables libres et sémantique. Du point de vue sémantique, une proposition dans laquelle il y a des variables libres est une assertion sur les objets désignés par ces variables. Il est possible en général d'exprimer cette assertion dans un autre langage, en particulier dans le langage naturel, sans utiliser de variables liées. Considérons par exemple la proposition suivante, dans le langage de la théorie des ensembles :

$$\exists x(x \in a).$$

Cette proposition est une assertion à propos d'un certain ensemble a , disant : il existe un objet x tel que x est élément de a . On peut exprimer la même propriété de l'ensemble a sans utiliser de variable liée, en disant simplement : l'ensemble a possède un élément.

Prenons, dans le même langage, une proposition plus compliquée :

$$(a \subset x \text{ et } \exists x(x \in a)) \Rightarrow \exists a(a \in x). \quad (11)$$

Les variables a et x possèdent chacune deux occurrences libres et une occurrence liée dans cette proposition. Sémantiquement, ce dont il est question dans cette proposition, ce sont deux ensembles, désignés par les variables *libres* a et x . On affirme que si a est un sous-ensemble de x , et si a possède un élément, alors x possède un élément. En utilisant les tournures « a possède un élément » et « x possède un élément », nous évitons l'emploi des variables liées x et a des propositions $\exists x(x \in a)$ et $\exists a(a \in x)$.

4.5 Substitutions

De nombreuses déductions logiques s'accompagnent d'opérations de *substitution*. Par exemple, lorsqu'on déduit du théorème général

$$(a + b)^2 = a^2 + 2ab + b^2 \quad (1)$$

le théorème particulier

$$((x - 5) + 8)^2 = (x - 5)^2 + 2(x - 5)8 + 8^2, \quad (2)$$

on substitue le terme $(x - 5)$ à la variable a , et le terme 8 à la variable b dans la proposition (1).

Nous allons définir exactement, dans ce paragraphe, ce que l'on entend par *substituer un terme T à une variable individuelle x dans une expression A*, celle-ci pouvant être un terme ou une proposition. La définition précise sera donnée sous la forme d'un ensemble de règles qui déterminent exactement l'expression résultant de cette opération, étant donné une expression A, une variable x , et un terme T quelconques.

Mais le principe de cette opération est simple: elle consiste à remplacer *chaque occurrence libre* de la variable x dans A par une copie du terme T. L'expression obtenue de cette manière sera désignée par la notation

$$(T|x)A,$$

que nous lisons « T remplace x dans A ». Nous soulignons que seules les occurrences *libres* de x dans A sont remplacées par T, et nous convenons encore que s'il n'y a pas d'occurrences libres de x dans A, autrement dit si x ne figure pas librement dans A, l'expression désignée par $(T|x)A$ sera l'expression A elle même.

Exemple 1. Soit A la proposition

$$a + 0 = a,$$

d'un langage où les symboles 0, +, = sont respectivement une constante individuelle, un symbole fonctionnel binaire, et un symbole relationnel binaire. Si nous désignons par T le terme $x + y$ de ce langage, alors $(T|a)A$ désigne la proposition

$$(x + y) + 0 = x + y,$$

obtenue en remplaçant les deux occurrences de la variable a dans A par le terme considéré. Ces deux occurrences de a dans A sont en effet libres (il n'y a pas d'occurrence liée de a dans A). Les parenthèses ne sont pas nécessaires si nous écrivons A et T en notation préfixée:

$$\begin{array}{ll} A: & = + a 0 a \\ T: & + x y \\ (T|a)A: & = + + x y 0 + x y. \end{array}$$

Les règles de substitution que nous donnerons plus bas se rapportent à la notation préfixée, et ne feront donc pas mention de parenthèses.

Une particularité de cet exemple est que le terme T , à savoir $x + y$, ne contient pas lui-même d'occurrence de la variable à laquelle il est substitué (variable a). Pour cette raison, si l'on part d'une expression A dans laquelle il y a n occurrences libres de la variable a , on peut obtenir l'expression $(T|a)A$ en répétant n fois l'opération: remplacer *une* occurrence libre quelconque de a dans l'expression en cours de transformation. Plus précisément, cela consiste à écrire une suite d'expressions $A_0, A_1, \dots, A_n$ dans laquelle A_0 est identique à A , et dans laquelle chaque expression A_k (pour $k > 0$) est construite en remplaçant une occurrence libre quelconque de a dans l'expression précédente A_{k-1} .

Il est clair cependant que cette méthode ne convient pas lorsque le terme T contient une ou plusieurs occurrences de la variable à laquelle il doit être substitué. Par exemple, si nous désignons par T_1 le terme $x + a$, et si A désigne encore la proposition $a + 0 = a$, alors $(T_1|a)A$ est la proposition

$$(x + a) + 0 = x + a.$$

Toutefois, si l'on remplace d'abord l'une des deux occurrences de a dans A par T_1 , par exemple la seconde, en écrivant la proposition intermédiaire

$$a + 0 = x + a,$$

il faut évidemment se rappeler *laquelle* des deux occurrences de a dans cette proposition intermédiaire doit être encore remplacée par T_1 pour qu'on obtienne la proposition $(T_1|a)A$.

Pour cette raison, on précise souvent le sens de la notation $(T|x)A$ en disant qu'elle désigne l'expression obtenue en remplaçant *simultanément* toutes les occurrences libres de x dans A par T .

Exemple 2. Cet exemple est pris dans le langage de la théorie des ensembles, dans lequel les symboles $\in$ et $\subset$ sont des symboles relationnels binaires, et le symbole $\cup$ est un symbole fonctionnel binaire. Désignons par A la proposition suivante:

$$(a \subset x \text{ et } \exists x \left(\begin{array}{|c} \square \\ x \end{array} (x \in a) \right)) \Rightarrow \exists y \left(\begin{array}{|c} \square \\ y \end{array} (y \in x) \right).$$

La variable x possède deux occurrences libres et une occurrence liée dans cette proposition A , dont l'arbre est donné dans la figure 4.5. Désignons par T le terme $X \cup Y$. Alors $(T|x)A$ est la proposition

$$(a \subset (X \cup Y) \text{ et } \exists x(x \in a)) \Rightarrow \exists y(y \in (X \cup Y)),$$

dont l'arbre est donné aussi dans la figure 4.5. Cet arbre s'obtient en prenant celui de A et en y remplaçant les deux feuilles des occurrences libres de x par l'arbre du terme $X \cup Y$.

Règles de substitution. Étant donné un terme T d'un langage du premier ordre $\mathcal{L}$ et une variable individuelle x , les règles suivantes déterminent exactement, pour toute expression A de $\mathcal{L}$ (terme ou proposition), quelle est l'expression désignée par $(T|x)A$.

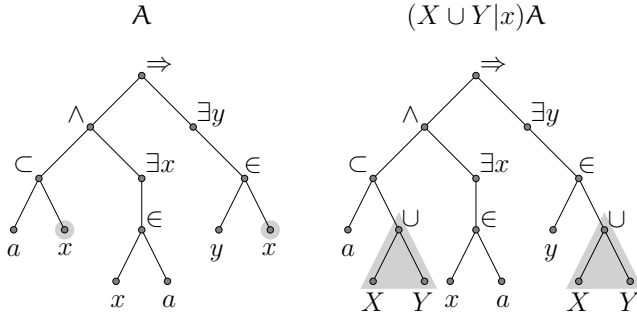


Figure 4.5 La proposition $(X \cup Y|x)A$ est obtenue en remplaçant toutes les occurrences libres de x dans A par $X \cup Y$.

(a) $(T|x)x$ est le terme T . Commentaire: l'expression x se compose d'une seule occurrence de la variable x . Celle-ci est libre et elle est remplacée par T .

(b) Si a est une variable individuelle de $\mathcal{L}$ *distincte* de x , ou si a est une constante individuelle de $\mathcal{L}$, $(T|x)a$ est identique à a . Commentaire: il n'y a pas d'occurrence libre de x dans l'expression réduite au seul symbole a .

(c) Si A est une expression composée de la forme $sA_1 \cdots A_n$, où s est un symbole fonctionnel ou un symbole relationnel n -aire de $\mathcal{L}$ et $A_1, \dots, A_n$ sont des termes de $\mathcal{L}$, alors $(T|x)A$ est l'expression

$$sA'_1 \cdots A'_n,$$

dans laquelle A'_i , pour $i = 1, \dots, n$, est l'expression $(T|x)A_i$. Commentaire: la substitution de T à x dans $sA_1 \cdots A_n$ revient à effectuer cette opération dans chacune des expressions subordonnées A_i .

(d) $(T|x)\perp$ est identique à $\perp$. Commentaire: il n'y a pas d'occurrence libre de x dans l'expression réduite au seul symbole $\perp$.

(e) Si A est une proposition de $\mathcal{L}$, l'expression $(T|x)(\neg A)$ est identique à $\neg A'$, où A' est l'expression $(T|x)A$. Commentaire: la substitution de T à x dans $\neg A$ revient à effectuer cette opération dans la sous-expression A .

(f) Si A et B sont des propositions de $\mathcal{L}$, les expressions $(T|x)(A \vee B)$, $(T|x)(A \wedge B)$, $(T|x)(A \Rightarrow B)$, $(T|x)(A \Leftrightarrow B)$ sont identiques respectivement à $A' \vee B'$, $A' \wedge B'$, $A' \Rightarrow B'$, $A' \Leftrightarrow B'$, où A' et B' sont les expressions $(T|x)A$ et $(T|x)B$. Commentaire: la substitution de T à x dans $A \vee B$ revient à effectuer cette opération dans chacune des sous-expressions A , B . De même pour $A \wedge B$, $A \Rightarrow B$, $A \Leftrightarrow B$.

(g) Si A est une proposition de $\mathcal{L}$, les expressions $(T|x)(\forall xA)$ et $(T|x)(\exists xA)$ sont identiques respectivement aux propositions $\forall xA$ et $\exists xA$. Commentaire: x ne figure pas librement dans les propositions $\forall xA$ et $\exists xA$.

(h) Si A est une proposition de $\mathcal{L}$, et si y est une variable individuelle de $\mathcal{L}$ *distincte* de x , les expressions $(T|x)(\forall yA)$ et $(T|x)(\exists yA)$ sont identiques respectivement à $\forall yA'$ et $\exists yA'$, où A' est l'expression $(T|x)A$. Commentaire: la substitution de T à x dans $\forall yA$ (resp. $\exists yA$) revient à effectuer cette opération dans la sous-expression A .

Exemple 3. Pour illustrer ces règles, nous considérons la proposition A suivante du langage de la théorie des ensembles (cf. Exemple 2), dans lequel le symbole $\cap$ est aussi un symbole fonctionnel binaire :

$$(a \subset (x \cap y) \text{ et } \exists x \overline{(x \in a)}) \Rightarrow \exists y \overline{(y \in x)}. \quad (A)$$

Pour un terme T quelconque de ce langage, l'expression $(T|x)A$ sera la proposition

$$(a \subset ((T) \cap y) \text{ et } \exists x(x \in a)) \Rightarrow \exists y(y \in (T)).$$

obtenue en remplaçant les deux occurrences libres de x par T . Les parenthèses autour de T ne sont nécessaires que dans la notation infixée. La figure 4.6 montre comment cette expression est déterminée par les règles de substitution qui précèdent. Elle contient d'abord une construction génératrice de termes et de propositions $A_1, \dots, A_{11}$, dans laquelle A_{11} est l'expression A . La deuxième partie de la figure contient les expressions $(T|x)A_i$ correspondantes, désignées par A'_i , telles qu'elles sont déterminées par les règles de substitution mentionnées à droite. Par exemple, A_5 étant la proposition $\subset A_1 A_4$, A'_5 est la proposition $\subset A'_1 A'_4$ d'après la règle (c).

Il est souligné dans la figure que la proposition A'_7 est identique à A_7 , c'est-à-dire à $\exists x A_6$, d'après la règle (g). À cause de cela, l'expression A'_6 n'est pas utilisée dans la construction des expressions suivantes, A'_7 à A'_{11} .

Propriété de la substitution. Étant donné une variable individuelle x et un terme T d'un langage du premier ordre $\mathcal{L}$, les règles de substitution qui précèdent définissent l'expression $(T|x)A$ pour toute expression A de $\mathcal{L}$, que ce soit un terme ou une proposition. Il est assez évident que si A est un terme de $\mathcal{L}$, alors $(T|x)A$ est un terme de $\mathcal{L}$, et si A est une proposition de $\mathcal{L}$, alors $(T|x)A$ est une proposition de $\mathcal{L}$. Néanmoins cela mérite d'être examiné.

Par définition, un terme A de $\mathcal{L}$ est une expression qui figure dans une construction génératrice de termes de $\mathcal{L}$ (§4.2), et une proposition A de $\mathcal{L}$ est une expression qui figure dans une construction génératrice de propositions de $\mathcal{L}$. Comme une proposition contient des termes, pour vérifier qu'une expression A est une proposition de $\mathcal{L}$, il faut écrire une construction génératrice de termes et une construction génératrice de propositions. On peut les combiner en une seule construction génératrice de termes et de propositions, comme nous l'avons fait dans la figure 4.6. Dans cet exemple, où T est un terme indéterminé, on peut vérifier que pour chacune des expressions A_i de la construction génératrice $A_1, \dots, A_{11}$, l'expression $(T|x)A_i$, désignée par A'_i est un terme si A_i est un terme, et une proposition si A_i est une proposition.

Cette vérification est effectuée dans la figure 4.7, qui résume le contenu de la figure 4.6. Chacune des expressions A_4 à A_{11} est construite au moyen d'expressions A_i qui la précèdent, de la manière indiquée dans la deuxième colonne de ce tableau. Dans la colonne intitulée « classe », on a noté la classe de chacune des expressions A_i , en désignant par t la classe des termes de $\mathcal{L}$ et par p celle des propositions. La notation $\cap t t$ désigne la classe des expressions

A_1	a		
A_2	x		
A_3	y		
A_4	$x \cap y$	$\cap A_2 A_3$	
A_5	$a \subset (x \cap y)$	$\subset A_1 A_4$	
A_6	$x \in a$	$\in A_2 A_1$	
A_7	$\exists x(x \in a)$	$\exists x A_6$	
A_8	$a \subset (x \cap y) \text{ et } \exists x(x \in a)$	$\wedge A_5 A_7$	
A_9	$y \in x$	$\in A_3 A_2$	
A_{10}	$\exists y(y \in x)$	$\exists y A_9$	
A_{11}	$(a \subset (x \cap y) \text{ et } \exists x(x \in a)) \Rightarrow \exists y(y \in x)$	$\Rightarrow A_8 A_{10}$	
A'_1	a		(b)
A'_2	$\top$		(a)
A'_3	y		(b)
A'_4	$(\top) \cap y$	$\cap A'_2 A'_3$	(c)
A'_5	$a \subset ((\top) \cap y)$	$\subset A'_1 A'_4$	(c)
A'_6	$(\top) \in a$	$\in A'_2 A'_1$	(c)
A'_7	$\exists x(x \in a)$	$\exists x A'_6$	(g)
A'_8	$a \subset ((\top) \cap y) \text{ et } \exists x(x \in a)$	$\wedge A'_5 A'_7$	(f)
A'_9	$y \in (\top)$	$\in A'_3 A'_2$	(c)
A'_{10}	$\exists y(y \in (\top))$	$\exists y A'_9$	(h)
A'_{11}	$(a \subset ((\top) \cap y) \text{ et } \exists x(x \in a)) \Rightarrow \exists y(y \in (\top))$	$\Rightarrow A'_8 A'_{10}$	(f)

Figure 4.6 Construction génératrice $A_1, \dots, A_{11}$ de termes et de propositions. Séquence des expressions $(\top|x)A_i$ désignées par A'_i et déterminées par les règles de substitution (a)–(h).

de $\mathcal{L}$ qui se composent du symbole fonctionnel binaire $\cap$ suivi de deux termes. Les expressions de cette forme sont des termes de $\mathcal{L}$ d'après les règles de construction des termes du paragraphe 4.1. L'observation principale à faire est qu'à partir de la quatrième ligne du tableau le contenu de la troisième colonne est identique à celui de la deuxième colonne, aux accents près. Comme A_1, A_2, A_3 et A'_1, A'_2, A'_3 appartiennent à la classe des termes ($\mathbf{t}$), chacune des expressions A'_i suivantes ($i = 4, \dots, 11$) appartient à la même classe, $\mathbf{t}$ ou $\mathbf{p}$, que l'expression A_i . Ainsi A_4 et A'_4 sont toutes deux de la classe $\cap \mathbf{t} \mathbf{t}$, donc de la classe $\mathbf{t}$. Par suite, A_5 et A'_5 sont toutes deux de la classe $\subset \mathbf{t} \mathbf{t}$, donc de la classe $\mathbf{p}$; et ainsi de suite.

Il est clair que ces propriétés ne sont pas particulières à cet exemple. La similitude de contenu de la deuxième et de la troisième colonne à partir de la quatrième ligne provient des règles de substitution mentionnées dans la quatrième colonne qui déterminent le contenu de la troisième (A'_i) en fonction de celui de la deuxième (A_i). Pour la ligne 6 par exemple, c'est la règle (c) qui dit que $(\top|x)(\in A_2 A_1)$ est identique à $\in((\top|x)A_2)((\top|x)A_1)$. Pour $i = 1, \dots, 3$,

i	A_i	A'_i	règle		classe
1	a	a	(b)		t
2	x	T	(a)		t
3	y	y	(b)		t
4	$\cap A_1 A_2$	$\cap A'_1 A'_2$	(c)	$\cap t t$	t
5	$\subset A_1 A_4$	$\subset A'_1 A'_4$	(c)	$\subset t t$	p
6	$\in A_2 A_1$	$\in A'_2 A'_1$	(c)	$\in t t$	p
7	$\exists x A_6$	$\exists x A_6$	(g)	$\exists x p$	p
8	$\wedge A_5 A_7$	$\wedge A'_5 A'_7$	(f)	$\wedge p p$	p
9	$\in A_3 A_2$	$\in A'_3 A'_2$	(c)	$\in t t$	p
10	$\exists y A_9$	$\exists y A'_9$	(h)	$\exists y p$	p
11	$\Rightarrow A_8 A_{10}$	$\Rightarrow A'_8 A'_{10}$	(f)	$\Rightarrow p p$	p

Figure 4.7 Résumé de la figure 4.6. On désigne par t la classe des termes de $\mathcal{L}$ et par p celle des propositions.

la similitude de classe de A_i et A'_i provient aussi des règles de substitutions (a) et (b).

4.6 Termes librement substituables

Pour illustrer la notion introduite dans ce paragraphe, nous allons employer un langage du premier ordre $\mathcal{L}$ dont le domaine d'interprétation prévu est l'ensemble des êtres humains, et dont les symboles spécifiques seront: le mot «Socrate» comme constante individuelle; le mot composé «est-fille-de» comme symbole relationnel binaire (infixé), par exemple dans la proposition « y est-fille-de Socrate»; le mot composé «le-père-de» comme symbole fonctionnel unaire, par exemple dans le terme «le-père-de Socrate».

Considérons la proposition A suivante de ce langage $\mathcal{L}$:

$$\exists y(y \text{ est-fille-de } x). \quad (1)$$

Le sens de cette proposition est clair: l'être humain désigné par x a une fille. La question que nous voulons examiner est ce que devient ce sens lorsqu'on substitue un terme T du langage $\mathcal{L}$ à la variable x dans cette proposition. Autrement dit, quel est le sens de la proposition $(T|x)A$. Un terme T du langage $\mathcal{L}$ désigne toujours un être humain — c'est l'interprétation choisie. La question est de savoir si pour n'importe quel terme T la proposition $(T|x)A$, c'est-à-dire

$$\exists y(y \text{ est-fille-de } T)$$

signifie que l'être humain désigné par T a une fille. Nous allons voir que c'est toujours le cas si y ne figure pas librement dans le terme T, mais que ce n'est pas toujours le cas lorsque y figure librement dans T.

Dans les exemples suivants, le terme T substitué à x dans A ne contient pas

d'occurrence de y :

$$\begin{aligned} & \exists y(y \text{ est-fille-de Socrate}) \\ & \exists y(y \text{ est-fille-de (le-père-de Socrate)}) \\ & \exists y(y \text{ est-fille-de (le-père-de (le-père-de } x))) \end{aligned}$$

Dans les trois cas, la proposition signifie « T a une fille »: Socrate a une fille, le père de Socrate a une fille, le grand-père de x a une fille.

Par contre, dans les exemples suivants, le terme T substitué à x dans A contient une occurrence de la variable y :

$$\begin{aligned} & \exists y(y \text{ est-fille-de } y) \\ & \exists y(y \text{ est-fille-de (le-père-de } y)) \\ & \exists y(y \text{ est-fille-de (le-père-de (le-père-de } y))) \end{aligned} \tag{2}$$

Dans les trois cas, la proposition ne signifie pas du tout que l'être humain désigné par T a une fille. Dans les trois cas, il n'est même pas question d'un être humain désigné par T, ni même d'un être humain particulier. La première proposition signifie qu'il existe une femme qui est sa propre fille; la seconde qu'il existe une femme qui est la fille de son père; la troisième qu'il existe une femme qui est la fille de son grand-père. Nous laissons au lecteur le soin de débattre de la véracité de ces propositions. Ce qui nous importe ici c'est que le sens de ces propositions n'est pas « T a une fille ».

La cause de ces changements de signification réside dans le fait que la variable y du terme T dans les trois propositions (2) est *liée* au quantificateur $\exists y$. Nous dirons que dans les trois exemples (2), la substitution du terme T à la variable x dans A (1) a *créé un lien*. Nous pouvons le mettre en évidence en récrivant ces propositions avec des liens explicites:

$$\begin{aligned} & \exists y \left(\overbrace{y}^{\quad} \text{ est-fille-de } x \right) \\ & \exists y \left(\overbrace{y}^{\quad} \text{ est-fille-de } \overbrace{y}^{\quad} \right) \\ & \exists y \left(\overbrace{y}^{\quad} \text{ est-fille-de (le-père-de } \overbrace{y}^{\quad})} \right) \\ & \exists y \left(\overbrace{y}^{\quad} \text{ est-fille-de (le-père-de (le-père-de } \overbrace{y}^{\quad}))} \right) \end{aligned} \tag{A}$$

Les déductions logiques s'accompagnent souvent de substitutions. Par exemple, s'il est vrai que tout homme a une fille, il est vrai que Socrate a une fille, et plus généralement, si T est un terme quelconque, il est vrai que « T a une fille ». En formalisant « tout homme a une fille » par la proposition

$$\forall x \exists y(y \text{ est-fille-de } x),$$

on voit que si cette proposition est vraie, il en est de même de la proposition

$$\exists y(y \text{ est-fille-de T}),$$

pour autant que T soit un terme dans lequel la variable y ne figure pas, c'est-à-dire tel que la substitution de T à x dans la proposition $\exists y(y \text{ est-fille-de } x)$

ne crée pas de lien.

Lorsque la substitution d'un terme T à une variable x dans une proposition A ne crée pas de lien, on dit que T est *librement substituable* à x dans A . Le but de ce paragraphe est de définir cette notion syntaxique de manière précise.

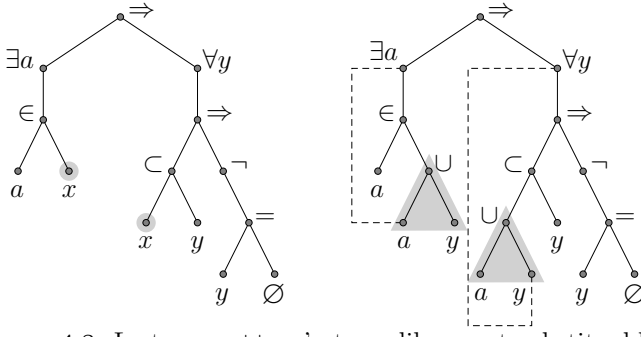


Figure 4.8 Le terme $a \cup y$ n'est pas librement substituable à x dans la proposition $\exists a(a \in x) \Rightarrow \forall y(x \subset y \Rightarrow y \neq \emptyset)$. La substitution crée deux liens.

Exemple 2. Examinons d'abord un exemple un peu plus complexe, en désignant par A la proposition suivante du langage de la théorie des ensembles :

$$\exists a(a \in x) \Rightarrow \forall y(x \subset y \Rightarrow y \neq \emptyset). \quad (3)$$

Cette proposition affirme que si l'ensemble x possède un élément, alors tout ensemble qui contient x est non vide. Cette proposition est sans doute vraie, mais ce n'est pas ce qui importe ici.

Il y a dans cette proposition deux occurrences de la variable x , toutes deux libres. La première se situe dans le champ d'action du quantificateur $\exists a$, et la seconde dans celui du quantificateur $\forall y$. Si l'on remplace la première par un terme qui contient une occurrence de la variable a , celle-ci sera liée au quantificateur $\exists a$, et la substitution créera ainsi un lien. De même, si l'on remplace la seconde occurrence de x par un terme contenant une occurrence de la variable y , celle-ci sera liée au quantificateur $\forall y$. Un terme *librement substituable* à x dans cette proposition A est un terme T pour lequel il n'y a pas de création de lien lorsqu'on le substitue à *toutes* les occurrences libres de x dans A . C'est donc un terme T dans lequel a et y ne figurent pas librement.

L'arbre de la proposition A (3) est représenté dans la figure 4.8, et il permet de préciser ce qu'on entend par *champ d'action* d'un quantificateur dans une proposition. Il s'agit de la partie de la proposition qui correspond au sous-arbre issu du nœud qui porte ce quantificateur. Par exemple le champ d'action du quantificateur $\forall y$ de la proposition (3) correspond au sous-arbre dont la racine est le nœud qui porte ce quantificateur, et qui représente la proposition $\forall y(x \subset y \Rightarrow y \neq \emptyset)$. La seconde occurrence libre de x dans A se situe donc dans le champ d'action du quantificateur $\forall y$. Nous dirons aussi qu'elle est

dominée par ce quantificateur, parce que celui-ci se trouve sur la branche de l'arbre qui porte cette occurrence de x à son extrémité. Quant à la première occurrence libre de x dans A , elle est dominée par le quantificateur $\exists a$.

Si l'on remplace la première occurrence libre de x dans A par un terme T , la feuille correspondante de l'arbre de A est remplacée par une copie de l'arbre de T , et si ce terme contient des occurrences de variables, celles-ci sont toutes dominées par le quantificateur $\exists a$. S'il y a des occurrences de a parmi les variables de T , celles-ci seront alors *liées* à ce quantificateur (§4.4). L'arbre de droite de la figure est celui de la proposition $(T|x)A$ où T est le terme $a \cup y$, et la figure montre les liens créés par cette substitution.

Définition. Un terme T d'un langage $\mathcal{L}$ est dit *librement substituable* à une variable x dans une proposition A de $\mathcal{L}$ lorsqu'aucune occurrence libre de x dans A n'est dominée par un quantificateur $\exists y$ ou $\forall y$ dont la variable y figure librement dans T .

Il faut prendre garde de ne pas interpréter cette terminologie de manière erronée. Dire que T est librement substituable à x dans A ne signifie pas qu'il y ait des occurrences libres de x dans A . Cela veut seulement dire que *s'il y en a*, elles ne sont pas dominées par des quantificateurs susceptibles de lier les variables de T . En particulier, s'il n'y a aucune occurrence libre de x dans A , n'importe que terme T est *par définition* librement substituable à x dans A . Nous rappelons que dans ce cas la proposition $(T|x)A$ est définie comme étant identique à A (§4.5). Un terme T est toujours « substituable » à une variable x dans une proposition A , en ce sens que la proposition $(T|x)A$ est toujours définie. Il n'est *librement* substituable à x que dans certaines propositions, en particulier celles dans lesquelles x ne figure pas librement.

Règles sur les termes librement substituables. Étant donné un terme T d'un langage du premier ordre $\mathcal{L}$ et une variable individuelle x , les règles suivantes déterminent exactement, pour toute proposition A de $\mathcal{L}$, si T est librement substituable à x dans A ou non.

(a) Si A est une proposition atomique de $\mathcal{L}$, T est librement substituable à x dans A .

(b) T est librement substituable à x dans la proposition $\perp$.

(c) Si A est une proposition de $\mathcal{L}$, T est librement substituable à x dans la proposition $\neg A$ si et seulement si T est librement substituable à x dans A .

(d) Si A et B sont des propositions de $\mathcal{L}$, T est librement substituable à x dans la proposition $A \vee B$ (resp. $A \wedge B$, $A \Rightarrow B$, $A \Leftrightarrow B$) si et seulement si T est librement substituable à x dans chacune des propositions A , B .

(e) Si A est une proposition de $\mathcal{L}$ et si y est une variable, T est librement substituable à x dans la proposition $\forall y A$ (resp. $\exists y A$) si et seulement si l'une au moins des deux conditions suivantes est vérifiée:

– x ne figure pas librement dans $\forall y A$ (resp. $\exists y A$);

– T est librement substituable à x dans A et y ne figure pas librement dans T .

4.7 Exercices

Exercice 1. On considère un langage du premier ordre $\mathcal{L}$ dans lequel il y a les symboles suivants:

- symbole fonctionnel unaire;
- $+, *$ symboles fonctionnels binaires;
- $0, 1$ constantes individuelles.

(a) Donner sous forme d'arbre, et en notation préfixée, le terme de $\mathcal{L}$ représenté par l'expression suivante utilisant la notation infixée usuelle:

$$-(((-(a + (b * c))) * (1 * (-(a + d)))). \quad (1)$$

(b) Donner une méthode systématique ou algorithmique pour écrire un terme de $\mathcal{L}$, en notation préfixée, à partir de l'arbre correspondant. On suppose que l'exécutant comprend uniquement les instructions simples de parcours de l'arbre en suivant les arcs et de lecture et copie des symboles portés par les nœuds. Il sait: se déplacer d'un nœud à ses successeurs; distinguer le successeur gauche du successeur droit lorsqu'un nœud a deux successeurs; reconnaître qu'un nœud est une feuille; passer d'un nœud à son prédécesseur; reconnaître la racine de l'arbre; écrire le symbole porté par un nœud; remplacer le symbole porté par un nœud par un symbole spécial servant à marquer ce nœud lorsqu'il est déjà traité.

(c) Donner une méthode systématique pour l'opération inverse, c'est-à-dire pour construire l'arbre d'un terme de $\mathcal{L}$ donné en notation préfixée, et appliquer cette méthode au terme

$$* + - * a b * 1 c - + * a c - 1. \quad (2)$$

(d) Écrire une construction génératrice du terme suivant sans construire son arbre:

$$- + 0 * - + a b + 1 * c + 1 d. \quad (3)$$

On demande que les termes $T_1, T_2, \dots, T_p$ de la construction soient dans un ordre de hauteur croissante, en appelant *hauteur* d'un terme le nombre entier défini comme suit: la hauteur d'un terme réduit à une seule variable ou une seule constante individuelle est le nombre 0; la hauteur d'un terme composé $sT_1 \cdots T_n$, où s est un symbole fonctionnel n -aire, est égale à la hauteur du ou des termes T_i de hauteur maximale parmi $T_1, \dots, T_n$ augmentée de 1. Par exemple:

- a : hauteur 0
- $-a$: hauteur 1
- $+ab$: hauteur 1
- $+ - ab$: hauteur 2
- $+ - a - b$: hauteur 2.

Exercice 2. On considère un langage ayant le même répertoire de symboles que celui de l'exercice précédent, mais dans lequel on veut adopter, pour les termes composés, une notation infixée des symboles $+$ et $*$. Pour cela, on ajoute aux symboles de $\mathcal{L}$ les symboles de parenthèse gauche « $($ » et droite « $)$ ». Définir les termes de façon précise par un diagramme syntaxique, et donner la définition correspondante des constructions génératrices. Les parenthèses, dans une notation infixée, ont pour but de garantir l'unique décomposabilité des termes: le symbole principal d'un terme composé et les termes qui lui sont subordonnés sont déterminés de manière unique grâce aux parenthèses. Pour cela on peut faire un usage plus ou moins abondant des parenthèses. Deux solutions sont demandées (diagrammes syntaxiques), dont l'une utilise le nombre minimum de parenthèses.

Écrire les termes de l'exercice 1 dans l'une des notations proposées, puis effacer dans chacun de ces termes toutes les occurrences du symbole de parenthèse droite et vérifier qu'il n'y a qu'une manière de les rétablir. Autrement dit, vérifier que l'on peut conserver la propriété d'unique décomposabilité en n'utilisant que le symbole de parenthèse gauche. Cette notation n'est jamais utilisée en pratique, mais l'idée est appliquée parfois de la manière suivante: on supprime (à choix) certaines des occurrences de la parenthèse droite, et l'on remplace les occurrences correspondantes de la parenthèse gauche, devenues «*veuves*», par un point. Par exemple, on écrit $(a*.b+c)$ au lieu de $(a*(b+c))$. Mettre sous une telle forme les termes de l'exercice 1, en leur donnant la meilleure lisibilité possible. Donner un diagramme syntaxique des termes qui permette ce choix.

Exercice 3. Soit $\mathcal{L}$ un langage du premier ordre dans lequel il y a les symboles suivants:

- $0, 1$ constantes individuelles;
- $+, \cdot$ symboles fonctionnels binaires;
- $=, <, \leq$ symboles relationnels binaires.

Le domaine de l'interprétation est celui des nombres entiers positifs ou négatifs $0, \pm 1, \pm 2, \dots$, mais le symbole « $-$ » n'appartient pas au langage $\mathcal{L}$. On demande de formaliser dans ce langage les propositions qui suivent, en convenant que les symboles de $\mathcal{L}$ ont leur signification usuelle. On donnera chacune de ces propositions d'abord en notation infixée, puis en notation préfixée et avec l'arbre correspondant.

(a) Tout nombre possède un opposé, c'est-à-dire un nombre qui, ajouté au premier, donne 0.

(b) Un nombre quelconque possède *au plus* un opposé, autrement dit ne possède jamais deux opposés distincts.

(c) Le nombre x est un carré, c'est-à-dire l'un des nombres: 0, 1, 4, 9, 16, 25, etc.

(d) x est divisible par 2.

(e) Tout nombre qui est un carré est divisible par deux — proposition fausse.

(f) z est un carré supérieur à n .

(g) z est le plus petit carré supérieur à n . Indication: z est un carré supérieur à n et tout carré supérieur à n est supérieur à z .

(h) y est un diviseur de x .

(i) z est un nombre premier, c'est-à-dire z est différent de 1 et n'a pas d'autres diviseurs que 1 et lui-même.

(j) Conjecture de Goldbach: tout nombre pair supérieur à 3 est la somme de deux nombres premiers.

Exercice 4. Écrire chacune des propositions (a)–(g) ci-dessous en traçant explicitement les liens entre quantificateurs et occurrences liées de variables. Dessiner les arbres de ces propositions et y tracer également les liens. Dire quelles sont les variables qui figurent librement dans chacune de ces propositions, souligner leurs occurrences libres, et dire si ces variables ont aussi des occurrences liées dans la proposition.

Les symboles du langage considéré ont leur type et leur arité usuelle: $=$, $\subset$, $\in$, $\leq$, $<$ sont des symboles relationnels binaires; $+$, $-$, $\cdot$ sont des symboles fonctionnels binaires. Les deux barres de valeur absolue $||$ représentent un symbole fonctionnel unaire **Abs** (utiliser celui-ci dans l'arbre); $f(x)$ est une abréviation du terme $f @ x$, composé des variables f , x et du symbole fonctionnel binaire $@$, symbole de l'opération d'application d'une fonction à un argument.

(a) $\forall x \exists y (x + y = z)$.

(b) $\exists y (x = y \cdot y) \Rightarrow \exists y (x = y + y)$.

(c) $\forall y (x \leq y \Leftrightarrow \exists a (x + a = y))$.

(d) $\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y))$.

(e) $(x \subset y \text{ et } \exists y (y \in x)) \Rightarrow \exists x (x \in y)$.

(f) $(x \leq z \text{ et } \exists y (z = y \cdot y)) \text{ et } \forall y ((x \leq y \text{ et } \exists z (y = z \cdot z)) \Rightarrow z \leq y)$.

(g) $\forall b (b > 0 \Rightarrow \exists a (a > 0 \text{ et } \forall y (|y - x| < a \Rightarrow |f(y) - f(x)| < b)))$.

Exercice 5. On désigne par $A_1, \dots, A_7$ les propositions de l'exercice 4, et par T un terme indéterminé du langage considéré. Pour chacune des propositions A_i , écrire les trois propositions $(T|x)A_i$, $(T|y)A_i$, $(T|z)A_i$, et dire à quelles conditions T est librement substituable à x (resp. à y , resp. à z) dans A_i .

Exercice 6. On désigne par A la proposition particulière $\exists b (a + b = c)$, et l'on demande de trouver toutes les variables individuelles y telles que la proposition $(a|y)((y|a)A)$ soit différente de A . Les variables y telles que $(a|y)((y|a)A)$ soit identique à A sont dites réversiblement substituables à a dans A . En se basant sur cet exemple, on demande de trouver des conditions générales, nécessaires et suffisantes, pour qu'une variable y soit réversiblement substituable à une variable x dans une proposition A d'un langage $\mathcal{L}$, c'est-à-dire pour que la proposition $(x|y)((y|x)A)$ soit identique à A .

Exercice 7. Si A désigne une proposition d'un langage du premier ordre $\mathcal{L}$, $T_1, \dots, T_n$ des termes de $\mathcal{L}$, et $x_1, \dots, x_n$ des variables individuelles *distinctes*,

on désigne par

$$(T_1|x_1, \dots, T_n|x_n)A \quad (4)$$

la proposition obtenue en remplaçant *simultanément*, ou *en parallèle*, dans A : chaque occurrence libre de x_1 par T_1 , chaque occurrence libre de x_2 par T_2 , etc. On demande de trouver des conditions suffisantes pour que la proposition (4) obtenue de cette manière soit identique à la proposition

$$(T_n|x_n)(\dots(T_2|x_2)((T_1|x_1)A)\dots) \quad (5)$$

obtenue en substituant successivement T_1 à x_1 dans A , puis T_2 à x_2 dans la proposition obtenue par la première substitution, et ainsi de suite.

5

Théories ou systèmes de déduction

5.1 Dédutions

Au chapitre 2 (§2.3) nous avons défini un langage *déclaratif* comme un langage qui possède une espèce principale d'expressions auxquelles peuvent s'appliquer les qualificatifs « vrai » et « faux », et que nous appelons *propositions*. Cette classe d'expressions est considérée non seulement en logique mais encore en grammaire et elle a reçu bien des noms différents, parmi lesquels nous avons déjà mentionné : assertions, énoncés, phrases, déclarations, formules.

Nous avons remarqué aussi, dès qu'il a été question de la valeur de vérité d'une proposition (§2.3), que cette valeur dépend toujours du *contexte* dans lequel la proposition est considérée. La même proposition peut être vraie dans un contexte, fausse dans un autre.

Le but de ce chapitre est de donner un sens précis à ces notions de vérité et de contexte. Nous n'avons pas la prétention de définir ce qu'est la vérité dans le sens le plus large du terme, philosophique, qui n'est guère définissable avec la précision que l'on attend d'une science exacte. Nous nous bornons à définir le *critère de vérité* qui est le seul reconnu par la logique formelle et qui est le suivant : une proposition est acceptée comme vraie lorsque ceci a été *démontré* ou *prouvé*. C'est un point de vue volontairement limité, dont le principal avantage est qu'il peut être défini avec une grande précision, et de ce fait, ne laisse place à aucune subjectivité dans l'appréciation de la vérité d'une proposition.

Le sens du mot vérité se ramène donc à celui de *démonstration*. Mais une définition précise du concept de démonstration nécessite l'introduction d'autres concepts fondamentaux de la logique. Une démonstration est une suite finie ou séquence de *jugements*. Elle se déroule dans un cadre logique général appelé une *théorie*. Les jugements de la séquence sont reliés les uns aux autres par des *dédutions*. Ce bref résumé contient, soulignés, les quatre mots-clés de ce

chapitre — jugements, théories, déductions, démonstrations — auxquels il faut ajouter encore : théorèmes.

La définition exacte de ces termes peut se formuler de manière très concise, et c'est à cela que nous aboutirons dans ce chapitre. Mais présentée directement sous cette forme cette définition serait trop abstraite pour le lecteur qui n'est pas déjà familiarisé avec la logique formelle. Nous allons donc introduire ces concepts progressivement. Malgré cela, à la première lecture, les définitions de ce chapitre paraîtront plutôt abstraites. Mais tout le reste de l'ouvrage en fournira une longue illustration, grâce à laquelle ces notions prendront un sens très concret. Il est donc recommandé au lecteur de relire plus tard le présent chapitre en relation avec les exemples fournis par les chapitres suivants.

Le lecteur de cet ouvrage doit avoir acquis normalement à l'école secondaire, dans les branches mathématiques, une certaine idée générale de ce qu'est une démonstration. Démontrer un théorème, dans le contexte d'une théorie particulière, par exemple la géométrie euclidienne, c'est le *déduire* de théorèmes déjà démontrés ou de propositions *admisses comme vraies*, appelées les *axiomes* de la théorie. Peut-être, l'idée du lecteur n'est-elle même pas si précise. Une démonstration est un raisonnement dans lequel on utilise des théorèmes déjà démontrés ou des axiomes pour aboutir, par des déductions, au théorème que l'on veut démontrer. Cette idée générale, si vague soit-elle, contient déjà deux notions importantes. Premièrement, qu'il y a en général, dans un contexte logique, des propositions admises comme vraies ou axiomes. La donnée de ces propositions fait partie du contexte logique d'une démonstration. Nous en discuterons plus en détail au paragraphe 5.3, avec des exemples. L'autre notion fondamentale qui intervient dans une démonstration est celle de déduction, l'opération logique par excellence. Le but principal de ce paragraphe est de préciser en quoi elle consiste.

Une déduction met en jeu un ou plusieurs *jugements*. Le terme de jugement, dans le sens le plus général, ne sera défini complètement qu'au paragraphe 5.2. Nous nous limiterons pour le moment à une forme particulière de jugements, et cette notion particulière de jugement est presque identique à celle de proposition. Le lecteur peut donc se contenter, à ce stade, de considérer un jugement comme étant simplement une proposition.

Que signifie déduire une proposition d'autres propositions? C'est simplement reconnaître que si l'on écrit cette proposition à la suite de celles dont on la déduit, on obtient une séquence de propositions qui est d'une certaine forme. C'est cette forme qui est le sujet de la logique et non ce qui se passe dans le cerveau de celui qui fait la déduction. Nous ne nous intéressons qu'à la déduction *écrite*, de sorte que nous identifions le concept de déduction à celui d'une séquence de propositions d'une certaine forme. En utilisant le terme de jugement plutôt que celui de proposition, nous dirons donc ceci :

Premièrement, une déduction est une séquence de jugements. Deuxièmement, dans un contexte donné, seules certaines séquences de jugements sont

admises comme déductions. Troisièmement, les critères suivants lesquels une séquence de jugements est admise comme déduction sont des critères de *forme* seulement.

Il n'y a pas d'autre manière de définir l'opération de déduction qu'en se référant à un catalogue ou répertoire de formes de déductions admises. Ce répertoire dépend du contexte, ainsi que nous venons de l'indiquer. Plus exactement, c'est ce répertoire lui-même qui, avec les axiomes, *constitue* le contexte logique d'une démonstration. Ceci sera précisé et illustré au paragraphe 5.3. Nous donnons ici quelques exemples de déductions de même forme, afin de montrer comment se présente un *schéma de déduction*, c'est-à-dire la description d'une certaine forme de déduction.

Exemple. Considérons la célèbre séquence de jugements

$$\begin{array}{l} \vdash \text{tout homme est mortel} \\ \vdash \text{Socrate est un homme} \\ \hline \vdash \text{Socrate est mortel.} \end{array} \quad (1)$$

Le signe $\vdash$ est un signe de ponctuation qui n'a pas d'importance pour le moment et que le lecteur peut ignorer. Ce signe sera présent cependant dans tous les jugements, et pas toujours au début de ceux-ci (§5.2). Dans notre exemple, à ce signe de ponctuation près, les trois jugements sont simplement trois propositions. Les deux premiers jugements sont appelés les *prémisses* de la déduction et le troisième est la *conclusion* de la déduction, déduite des prémisses. La conclusion est séparée des prémisses par une barre, que nous appellerons *barre de déduction*. Elle marque l'opération de déduction, comme la barre d'addition sous une liste de nombres que l'on additionne sépare ces nombres du résultat de l'addition. Mais cette barre ne joue qu'un rôle optique. C'est la séquence des trois jugements qui constitue une déduction, et ce qui distingue la conclusion des prémisses, c'est uniquement le fait qu'elle est le dernier jugement de la séquence. La barre de déduction n'ajoute rien. Au lieu de cette barre, on insère souvent le mot « donc » entre les prémisses et la conclusion, en écrivant :

$$\begin{array}{l} \vdash \text{tout homme est mortel} \\ \vdash \text{Socrate est un homme} \\ \text{donc } \vdash \text{Socrate est mortel.} \end{array}$$

La disposition verticale des jugements d'une déduction n'est pas essentielle non plus. Dans la suite, les prémisses seront souvent écrites sur une même ligne, de gauche à droite :

$$\begin{array}{l} \vdash \text{tout homme est mortel} \quad \vdash \text{Socrate est un homme} \\ \hline \vdash \text{Socrate est mortel.} \end{array}$$

Il va sans dire que cette déduction est admise par la logique classique, ou simplement par le bon sens. Mais l'ordre des trois jugements a évidemment son importance. La séquence de jugements suivante n'est certainement pas

admissible comme déduction :

$$\frac{\begin{array}{l} \vdash \text{Socrate est un homme} \\ \vdash \text{Socrate est mortel} \end{array}}{\vdash \text{tout homme est mortel.}} \quad (2)$$

La conclusion « tout homme est mortel » est bien vraie — pauvres de nous. Mais cela ne se déduit pas logiquement du fait qu'un individu particulier, nommé Socrate, soit mortel. Il suffit de remplacer « est mortel » par un attribut moins universel, par exemple « est un philosophe grec », pour s'apercevoir que la forme de la séquence de jugements (2) n'est pas de celles qui garantissent la vérité de la conclusion, étant donné celle des prémisses.

Ceci nous amène à caractériser la *forme* de la déduction (1), afin de pouvoir reconnaître toutes les déductions de cette forme. Aristote avait déjà constitué un catalogue de déductions de ce genre, appelées syllogismes. Mais il le faisait dans un langage formel limité qui ne se prête pas à la description d'autres formes de déductions appliquées en mathématique. Nous allons donc exprimer d'abord la déduction de cet exemple dans un langage formel du premier ordre, en l'écrivant

$$\frac{\begin{array}{l} \vdash \forall x(x \text{ est un homme} \Rightarrow x \text{ est mortel}) \\ \vdash \text{Socrate est un homme} \end{array}}{\vdash \text{Socrate est mortel.}} \quad (3)$$

Dans ce langage, les mots « est un homme » forment un symbole relationnel unaire postfixé, de même que les mots « est mortel ». Le mot « Socrate » est une constante individuelle.

Pour identifier la forme de cette déduction, désignons par les métasymboles A et B les propositions qui figurent à gauche et à droite du symbole $\Rightarrow$ de la première prémisses, à savoir

$$\begin{array}{ll} A : & x \text{ est un homme} \\ B : & x \text{ est mortel.} \end{array}$$

Pour caractériser la forme générale de cette déduction, il est important de reconnaître que la deuxième prémisses et la conclusion de la déduction sont les propositions que l'on obtient lorsqu'on substitue le terme « Socrate » à la variable x dans la proposition A, respectivement dans la proposition B :

$$\begin{array}{ll} \text{Socrate est un homme} & (\text{Socrate}|x)A \\ \text{Socrate est mortel} & (\text{Socrate}|x)B. \end{array}$$

Avec ces notations, nous pouvons écrire la déduction (3) de la manière suivante, qui met en évidence sa forme générale :

$$\frac{\begin{array}{l} \vdash \forall x(A \Rightarrow B) \\ \vdash (\text{Socrate}|x)A \end{array}}{\vdash (\text{Socrate}|x)B.} \quad (4)$$

Schéma de déduction. Nous verrons que suivant la logique classique, semblable déduction peut se faire avec n'importe quelles propositions A et B,

n'importe quelle variable x au lieu de la variable particulière x , et n'importe quel terme T , au lieu du terme «Socrate», pour autant que le terme T soit librement substituable à la variable x dans chacune des propositions A et B . La forme de toutes ces déductions est la suivante:

$$\frac{\begin{array}{l} \vdash \forall x(A \Rightarrow B) \\ \vdash (T|x)A \end{array}}{\vdash (T|x)B}, \quad (5)$$

le terme T devant satisfaire à la condition mentionnée.

La séquence de méta-expressions (5), jointe à la condition mentionnée sur T , constitue ce qu'on appelle un *schéma de déduction*. Ce schéma représente toutes les déductions que l'on peut faire en prenant pour A , B , x , et T deux propositions, une variable et un terme particuliers satisfaisant à la condition. C'est dire que dans n'importe quel langage du premier ordre digne d'intérêt il y a une infinité de déductions possibles obéissant à ce schéma.

Prenons par exemple pour A la proposition $a \in \mathbb{R}$; pour B la proposition $a^2 \geq 0$; pour x la variable a ; pour T le terme $-\pi/2$. On obtient la déduction suivante, de schéma (5):

$$\frac{\begin{array}{l} \vdash \forall a(a \in \mathbb{R} \Rightarrow a^2 \geq 0) \\ \vdash -\pi/2 \in \mathbb{R} \end{array}}{\vdash (-\pi/2)^2 \geq 0}. \quad (6)$$

Il faut faire attention ici à un point important. Un schéma de déduction tel que (5) n'exige pas que les deux premières propositions soient vraies. Dans l'exemple (6), nous avons pris deux prémisses notoirement vraies — dans le contexte de la théorie des nombres réels. Le schéma de déduction garantit dans ce cas que la conclusion est vraie. Mais une déduction de la forme (5) est toujours *correcte*, quelle que soit la valeur de vérité de ses prémisses. Un lecteur des déductions (3) et (6) peut très bien ne pas connaître la valeur de vérité des prémisses de ces déductions. Mais il est capable de reconnaître que les déductions sont correctes, en ce sens qu'elles suivent un schéma admis.

Le choix des formes de déduction admises dans un contexte est une convention, comme l'est celui des axiomes, c'est-à-dire des propositions admises comme vraies. Ensemble, les deux choix déterminent toutes les propositions vraies d'un contexte, puisque ce sont par définition celles que l'on peut déduire des axiomes en effectuant des déductions admises.

Dédutions sans prémisses. Le nombre de prémisses d'une déduction n'est jamais très élevé. Il est rarement supérieur à trois. D'autre part, il peut être nul. Il y a en effet de nombreuses déductions qui ne comportent qu'une conclusion et point de prémisses. Il est un peu abusif de parler dans ce cas de « déductions », puisque leur conclusion ne se déduit de rien du tout — elle se tire du néant si l'on peut dire. Mais on simplifie le vocabulaire en considérant cela comme le cas limite d'une déduction à n prémisses. Une telle déduction sera notée avec une barre de déduction au-dessus de laquelle il n'y a rien.

Par exemple, la logique classique admet, pour n'importe quelle proposition A , la déduction sans prémisses

$$\frac{}{\vdash A \text{ ou } \text{non}A.} \quad (7)$$

La méta-expression (7) est un schéma de déduction sans prémisses de la logique classique, et comme exemples de déductions conformes à ce schéma, on peut écrire :

$$\frac{\vdash x = 0 \text{ ou } x \neq 0}{\vdash \text{Socrate est mortel ou Socrate est immortel}} \\ \frac{}{\vdash \exists a(a \in x) \text{ ou } \text{non}\exists a(a \in x).}$$

Le schéma de déduction (7) est appelé traditionnellement règle, ou loi, ou principe, du *tiers exclu*. Nous lui consacrerons un paragraphe au chapitre 6, car il est au centre d'une critique de la logique classique de la part d'une école de logiciens et de mathématiciens qui s'abstient d'en faire usage.

5.2 Hypothèses et jugements

L'examen de n'importe quel ouvrage de mathématique, sans même chercher à en comprendre le détail, permet de constater dans les raisonnements la présence de fréquentes *hypothèses*. Celles-ci sont des propositions introduites par des déclarations telles que : « supposons que ... » ou « soit ... ».

Exemple 1. Dans un texte mathématique où il est question de nombres réels, on pourra rencontrer des déclarations telles que :

- supposons que x appartienne à l'intervalle $[a, b]$;
- soit $x \in [a, b]$;
- soit f une application continue de $[a, b]$ dans $\mathbb{R}$.

Les deux premières déclarations sont équivalentes. Elles donnent le statut d'hypothèse à la proposition $x \in [a, b]$. La troisième déclare comme hypothèse la proposition « f est une application continue de $[a, b]$ dans $\mathbb{R}$ ».

Faire une hypothèse est une opération qui ressemble, du point de vue logique, à celle qui consiste à déclarer une proposition comme un axiome. C'est admettre la vérité d'une certaine proposition. Un contexte logique est déterminé par ces deux groupes de propositions dont la vérité est admise : les axiomes (§5.3) et les hypothèses.

La raison pour laquelle on distingue *deux* catégories de propositions admises comme vraies dans un contexte est que les hypothèses, contrairement aux axiomes, ont toujours une durée de validité assez limitée. Par exemple, lorsque l'auteur d'un livre d'analyse écrit quelque part « supposons que $x \geq 0$ », il fait une hypothèse qui n'est certainement pas valable pour tout le reste de

son ouvrage, mais seulement dans une portion bien déterminée et limitée de celui-ci. Les hypothèses changent constamment d'un endroit à l'autre dans un texte mathématique. Par contre les axiomes, par exemple les axiomes sur les nombres réels dans le cas de l'analyse, peuvent être valables pour l'ensemble d'un traité en plusieurs volumes. Les deux catégories de propositions admises comme vraies se distinguent donc par leur *portée*. Celle des hypothèses est toujours *locale*, alors que les axiomes ont une portée globale. Cette différence de statut se reflète dans les schémas de déduction de la logique des prédicats, comme nous le verrons au chapitre 9.

Il y a toujours un endroit où une hypothèse cesse d'être en vigueur, même si cet endroit n'est pas explicitement signalé par un auteur. Ce sera d'ailleurs l'une des différences importantes entre les démonstrations informelles des mathématiques usuelles et les démonstrations formelles, comme elles se présenteront dans cet ouvrage. Dans celles-ci, la fin de validité d'une hypothèse sera toujours explicitement signalée. Chaque entrée en vigueur d'une nouvelle hypothèse, chaque cessation de validité d'une hypothèse, sont des variations du contexte logique. Ces notions se préciseront naturellement plus loin et nous ne cherchons pour le moment qu'à en donner une première idée.

Exemple 2. Souvent, des hypothèses sont faites au début de l'énoncé d'un théorème. Par exemple, dans le même cadre mathématique que celui de l'exemple 1, à savoir celui de l'analyse réelle, on pourra rencontrer un énoncé tel que le suivant :

Théorème¹. *Soient a, b deux nombres réels tels que $a \leq b$. Soient f, g deux fonctions continues sur $[a, b]$ et supposons que $f(x) \leq g(x)$ pour tout $x \in [a, b]$. Alors*

$$\int_a^b f \leq \int_a^b g. \quad (1)$$

Cet énoncé commence par déclarer comme hypothèses les six propositions

$$\begin{aligned} a &\in \mathbb{R} \\ b &\in \mathbb{R} \\ a &\leq b \\ f &\text{ est une application continue de } [a, b] \text{ dans } \mathbb{R} \\ g &\text{ est une application continue de } [a, b] \text{ dans } \mathbb{R} \\ \forall x (x \in [a, b] \Rightarrow f(x) &\leq g(x)). \end{aligned} \quad (2)$$

Puis il énonce la proposition (1) qui est vraie *dans le contexte ainsi défini*, c'est-à-dire celui de la théorie des nombres réels avec ses axiomes et ses schémas de déduction et *sous les hypothèses* (2). Avec les axiomes de la théorie, celles-ci constituent les propositions admises comme vraies pour démontrer (1).

¹ Serge Lang, *Undergraduate Analysis*, Springer-Verlag 1983, p. 95.

Jugements. L'énoncé de théorème qui précède est un exemple de ce que nous appellerons désormais un *jugement*, à savoir la combinaison d'une liste d'hypothèses et d'une proposition considérée sous ces hypothèses. Formellement, si A désigne la proposition (1), et si Γ (lettre grecque gamma majuscule) désigne la liste de propositions (2), prises comme hypothèses, nous désignerons par

$$\Gamma \vdash A$$

le jugement qui est la combinaison des deux, autrement dit l'énoncé complet intitulé « théorème ».

Le signe $\vdash$ n'a pas d'interprétation particulière. Il n'est qu'un signe de ponctuation servant à séparer la liste des hypothèses Γ d'un jugement de la proposition A qui est considérée par rapport à ces hypothèses et que nous appellerons *l'assertion* du jugement.

Syntaxiquement, un jugement n'est donc rien d'autre qu'une proposition A combinée avec une liste de propositions, et le signe $\vdash$ indique la combinaison des deux.

Il se trouve souvent que la liste d'hypothèses d'un jugement est *vide*, autrement dit qu'aucune hypothèse n'est faite dans ce jugement. Dans ce cas, il n'y a rien à gauche du signe $\vdash$ dans la notation du jugement, qui s'écrit simplement $\vdash A$, si A est son assertion. Ce sont des jugements de ce genre que nous avons vus dans les exemples du paragraphe 5.1.

Le mot « jugement » a donc dans le présent ouvrage un sens syntaxique très précis. Lorsque nous parlons d'un jugement $\Gamma \vdash A$ d'un langage $\mathcal{L}$, ou « dans » un langage $\mathcal{L}$, il s'agit d'un jugement dont l'assertion A est une proposition de $\mathcal{L}$ et dont toutes les hypothèses, c'est-à-dire toutes les propositions de la liste Γ sont des propositions de $\mathcal{L}$. Si cette liste est vide, il s'agit simplement d'un jugement $\vdash A$ dont l'assertion A est une proposition de $\mathcal{L}$.

Le vocabulaire de la logique est loin d'être unifié. Il varie beaucoup dans la littérature. Notamment le terme de « jugement » n'est pas utilisé par tous les auteurs, et lorsqu'il est utilisé, il n'a pas toujours le sens exact que nous lui donnons. Cependant il désigne assez généralement les « unités de texte » dont se compose une déduction, dans le sens où une déduction peut être décrite comme une séquence de telles unités. Nous incluons dans chacune de ces unités une liste d'hypothèses. On peut présenter la logique différemment en ne gardant dans les jugements que leur assertion et en traitant les hypothèses séparément.

Notations. Nous aurons souvent à représenter des listes d'hypothèses ainsi que des opérations sur celles-ci, par exemple l'opération d'adjonction d'une hypothèse à une liste ou la suppression d'une hypothèse dans une liste. Nous introduisons ici quelques notations à cet usage.

Nous désignerons toujours par le métasymbole Γ (gamma) une liste d'hypothèses quelconque. Ce peut être par exemple la liste des hypothèses du théorème de l'exemple 2. Il s'agit donc simplement d'une liste de propositions d'un langage $\mathcal{L}$. Une liste d'hypothèses peut être *vide* et il faut toujours penser

à cette possibilité lorsque nous parlons d'une liste d'hypothèses Γ *quelconque*. Une liste de n propositions $A_1, A_2, \dots, A_n$ prises comme hypothèses sera notée

$$A_1; A_2; \dots; A_n. \quad (3)$$

Si Γ désigne une liste d'hypothèses quelconque et si A désigne une proposition, la notation

$$\Gamma; A \quad (4)$$

désignera la liste d'hypothèses obtenue en adjoignant à Γ l'hypothèse supplémentaire A . Le passage d'une liste d'hypothèses Γ à une liste augmentée $\Gamma; A$ a lieu, dans un document déductif, chaque fois que l'on déclare une nouvelle hypothèse — en disant « supposons que ... » ou « soit ... », car ces déclarations reviennent toujours à admettre comme vraie une certaine proposition A . Une liste de n hypothèses (3) est donc obtenue à partir de la liste vide par n opérations d'adjonction (4). Nous dirons qu'une liste d'hypothèses Γ' est une *extension* d'une liste Γ si Γ' est identique à Γ ou si Γ' est obtenue en adjoignant à Γ une ou plusieurs hypothèses :

$$\underbrace{\Gamma; A_1; \dots; A_n}_{\Gamma'}.$$

Ces notations sont utilisées systématiquement dans les schémas de déduction et nous en donnons ici un exemple.

Exemple 3. Lorsqu'on veut démontrer qu'une proposition de la forme $A \Rightarrow B$ est vraie dans un certain contexte, la manière la plus courante de procéder est de *faire l'hypothèse* A et de démontrer que B est vraie sous cette hypothèse. Par exemple, ayant à démontrer dans le contexte de la théorie des nombres réels la proposition

$$(0 \leq a \text{ et } a \leq b) \Rightarrow a^2 \leq b^2, \quad (5)$$

on dira ceci : supposons que $0 \leq a$ et $a \leq b$, puis, admettant cette proposition comme vraie, on démontrera que $a^2 \leq b^2$. Cette démarche est si importante qu'elle peut être prise comme définition même de la vérité d'une implication : une proposition de la forme $A \Rightarrow B$ est vraie *par définition* si B est vraie lorsqu'on fait l'hypothèse A .

Mais décrivons cette démarche de manière plus précise. Premièrement, il nous faut préciser que le théorème exact à démontrer est le jugement sans hypothèses

$$\vdash (0 \leq a \text{ et } a \leq b) \Rightarrow a^2 \leq b^2.$$

Désignons ce jugement par $\Gamma \vdash A \Rightarrow B$. Le métasymbole Γ représente dans ce cas la liste d'hypothèses vide, mais nous l'utilisons quand même, car la méthode de déduction s'applique pour une liste d'hypothèses Γ quelconque. Cette méthode consiste à démontrer que B est vraie sous la liste d'hypothèses augmentée $\Gamma; A$, autrement dit à démontrer le jugement $\Gamma; A \vdash B$. Lorsqu'on

y est parvenu, on en *conclut* que $A \Rightarrow B$ est vraie sous les hypothèses Γ . Autrement dit, la démonstration s'achève par la déduction

$$\frac{\Gamma; A \vdash B}{\Gamma \vdash A \Rightarrow B.} \quad (6)$$

Le schéma de déduction (6) est l'un des plus universellement admis qui soient. Il appartient bien entendu à la logique classique et sera introduit en tant que tel au chapitre 6, mais il appartient encore à toutes les autres espèces de logiques connues de l'auteur.

5.3 Théories

La vérité d'une proposition d'un langage $\mathcal{L}$ n'est déterminée que par rapport à un contexte, et du point de vue de la logique elle n'est reconnue que lorsqu'elle a fait l'objet d'une démonstration. Dans une démonstration, la vérité d'une proposition est établie par des déductions à partir de propositions dont la vérité a déjà été démontrée, ou de propositions dont la vérité est admise sans démonstration et que l'on appelle des axiomes. Du point de vue de la logique, le contexte par rapport auquel on établit la vérité d'une proposition est donc constitué par les axiomes et les déductions permises.

Ce ne sont pas les seules composantes d'un contexte logique. Comme nous l'avons vu, la vérité d'une proposition est souvent considérée par rapport à des hypothèses. Celles-ci font donc partie du contexte. Mais dans ce paragraphe, nous ne nous occupons que des axiomes et des déductions permises.

Les axiomes sont les propositions admises comme base de ce qu'on appelle une *théorie*. Nous avons vu au chapitre 2 qu'un langage du premier ordre $\mathcal{L}$ sert à parler d'une certaine classe d'objets qu'on appelle le domaine de l'interprétation prévue de $\mathcal{L}$. Les propositions de $\mathcal{L}$ sont des assertions sur ces objets et l'on s'intéresse naturellement à celles de ces assertions qui sont vraies. Si l'on prend le mot « vrai » dans le sens opérationnel de « démontrable », on doit toujours admettre comme vraies, sans démonstration, un nombre limité de propositions de $\mathcal{L}$ exprimant des propriétés fondamentales des objets du domaine considéré. Ce sont les axiomes. Toutes les autres propriétés de ces objets seront déduites logiquement de ces axiomes. La déduction devra se faire suivant des schémas fixés, dont on admet qu'ils conviennent, tout comme les axiomes, au domaine d'interprétation considéré. La donnée conjointe d'un langage $\mathcal{L}$, d'axiomes, et de schémas de déduction *définit* une théorie.

Ce que nous appelons *schéma de déduction* est une description de toutes les déductions d'une certaine forme. Nous avons vu par exemple (§5.1) que les déductions particulières

$$\begin{array}{l|l} \vdash \forall x(x \text{ est un homme} \Rightarrow x \text{ est mortel}) & \vdash \forall a(a \in \mathbb{R} \Rightarrow a^2 \geq 0) \\ \vdash \text{Socrate est un homme} & \vdash -\pi/2 \in \mathbb{R} \\ \hline \vdash \text{Socrate est mortel} & \vdash (-\pi/2)^2 \geq 0 \end{array}$$

sont toutes deux de la forme

$$\frac{\begin{array}{l} \vdash \forall x(A \Rightarrow B) \\ \vdash (T|x)A \end{array}}{\vdash (T|x)B}, \quad (1)$$

où A et B sont deux propositions, x est une variable et T est un terme librement substituable à x dans A et dans B . Une déduction est une séquence de jugements et un schéma de déduction décrit une certaine classe de séquences de jugements.

Ce qui importe, dans un schéma de déduction, c'est la classe de déductions qu'il décrit, et non la manière dont il décrit cette classe. Car la même classe de déductions peut-être décrite de différentes manières. Par exemple, le choix des métasymboles A , B , x , T du schéma (1) est déjà arbitraire. D'autres métasymboles feraient tout aussi bien l'affaire. En principe on pourrait même décrire cette classe de déductions en se passant de métasymboles. On pourrait dire par exemple qu'il s'agit de toutes les déductions à deux prémisses qui satisfont aux conditions suivantes: les deux prémisses et la conclusion ont la même liste d'hypothèses (vide); l'assertion de la première prémisse est une implication précédée d'un quantificateur universel; les assertions de la deuxième prémisse et de la conclusion sont les propositions obtenues en substituant le même terme à la variable du quantificateur dans le membre de gauche, respectivement dans le membre de droite de l'implication de la première prémisse, ce terme étant librement substituable à cette variable dans ces deux propositions. Cette description est évidemment moins précise que le schéma (1), mais pour certains schémas plus simples une description précise sans métavariabes est parfaitement possible.

La donnée d'une théorie comporte ainsi celle d'une *classe de déductions*, que nous dirons *permises* ou *admises* dans cette théorie. Pratiquement, cette classe est décrite au moyen de plusieurs schémas de déduction. Mais comme nous venons de le voir, ces schémas peuvent être formulés de diverses manières, et ce qui importe c'est la classe de déductions qu'ils décrivent. Chacun d'eux définit une sous-classe de la classe des déductions admises dans la théorie.

Nous pouvons donc résumer de la manière suivante en quoi consiste une théorie, au sens de la logique. Une *théorie* $\mathcal{T}$ est définie par la donnée de:

- (a) un langage $\mathcal{L}$, appelé le *langage de* $\mathcal{T}$;
- (b) une liste de propositions de $\mathcal{L}$ appelées les *axiomes* de $\mathcal{T}$;
- (c) une classe de déductions dans le langage $\mathcal{L}$ appelées les *déductions élémentaires* de $\mathcal{T}$.

Ce que nous appelons ici une déduction *dans le langage* $\mathcal{L}$ est naturellement une séquence de jugements dans ce langage (§5.2).

Ces données constituent un cadre logique dans lequel, partant de la vérité admise des axiomes, on démontre celle d'autres propositions, éventuellement soumises à des hypothèses. On élargit ainsi par le mécanisme de la déduction, la

classe des propositions reconnues comme vraies, qui se compose au départ des axiomes uniquement. Mais on élargit aussi, comme nous le verrons, la classe des déductions admises dans une théorie. Cette classe ne se compose initialement que des déductions dites *élémentaires*, faisant partie de la donnée de la théorie. Nous verrons que les déductions élémentaires peuvent se combiner entre elles pour former des déductions « moins élémentaires », qui *dérivent* comme on dit des déductions élémentaires.

Théories logiques. Dans les mathématiques courantes on n'a affaire qu'à des théories dont les déductions élémentaires comprennent toutes celles qui appartiennent à ce qu'on appelle la logique classique. Ce sont ces déductions, et celles qui en dérivent, qui constituent le sujet principal de cet ouvrage.

Nous parlerons dans ce sens de théories *logiques*. C'est le cas notamment de la seule théorie à laquelle nous consacrerons d'assez longs développements et dont nous démontrerons un assez grand nombre de théorèmes, à savoir la théorie des ensembles.

En plus des déductions qui appartiennent au répertoire de la logique, une théorie logique $\mathcal{T}$ peut admettre d'autres déductions qui n'appartiennent pas à ce répertoire. On parlera de celles-ci comme étant des déductions *spécifiques* à $\mathcal{T}$. Ces déductions seront décrites par des schémas de déduction que l'on dira aussi spécifiques à $\mathcal{T}$. Nous en donnons ci-dessous un exemple.

Exemple de théorie: l'arithmétique du premier ordre. Le langage $\mathcal{L}$ et les axiomes de cette théorie ont déjà été présentés au paragraphe 3.8. Le domaine de l'interprétation prévue de $\mathcal{L}$ est celui des entiers naturels 0, 1, 2, etc. muni des opérations d'addition et de multiplication et de l'opération « successeur » : passage d'un nombre x à son successeur $x+1$, noté σx . Le langage $\mathcal{L}$ a pour symboles spécifiques: la constante individuelle 0, le symbole fonctionnel unaire σ , les symboles fonctionnels binaires $+$ et $\cdot$, et le symbole relationnel binaire $=$. Par un choix purement arbitraire, les symboles 1, 2, 3, etc. n'appartiennent pas à ce langage — mais les nombres correspondants peuvent être représentés si besoin est par les termes $\sigma 0$, $\sigma \sigma 0$, $\sigma \sigma \sigma 0$, etc. Nous reproduisons ici les six propositions de ce langage prises comme axiomes :

$$\begin{aligned} &\forall x(\sigma x \neq 0) \\ &\forall x \forall y(\sigma x = \sigma y \Rightarrow x = y) \\ &\forall x(x + 0 = x) \\ &\forall x(x \cdot 0 = 0) \\ &\forall x \forall y(x + \sigma y = \sigma(x + y)) \\ &\forall x \forall y(x \cdot \sigma y = x \cdot y + x). \end{aligned}$$

Les déductions élémentaires de cette théorie sont d'une part toutes celles de la logique classique présentées plus loin — il s'agit donc d'une théorie logique au sens ci-dessus — et d'autre part les déductions conformes à un schéma de déduction spécifique à cette théorie, appelé *schéma de récurrence* ou *principe*

de récurrence. Il s'agit des déductions de la forme suivante, où A est une proposition quelconque du langage $\mathcal{L}$ et x une variable individuelle quelconque :

$$\frac{\begin{array}{l} \Gamma \vdash (0|x)A \\ \Gamma \vdash \forall x(A \Rightarrow (\sigma x|x)A) \end{array}}{\Gamma \vdash \forall xA.} \quad (2)$$

Les notations $(0|x)A$ et $(\sigma x|x)A$ représentent les propositions obtenues par substitution de 0, respectivement de σx , à x dans A (§4.5).

Ce schéma de déduction est celui qui permet les démonstrations par *récurrence*, dont le lecteur a sans doute déjà vu des exemples. Rappelons que cette méthode de démonstration est la suivante : pour démontrer — dans un contexte dans lequel une liste d'hypothèses Γ est en vigueur — que tout entier naturel possède une certaine propriété, ce qui s'exprime par une proposition de la forme $\forall xA$ où A est la proposition exprimant que le nombre désigné par x possède la propriété en question, il suffit de démontrer que le nombre 0 possède cette propriété, ce qui s'exprime par la proposition $(0|x)A$, et que pour tout nombre x possédant cette propriété, le successeur de x possède également cette propriété, et cela s'exprime par $\forall x(A \Rightarrow (\sigma x|x)A)$. Suivant le schéma (2), on peut en conclure alors que $\forall xA$ est vraie dans le contexte.

Les déductions de la forme (2) sont des déductions élémentaires de cette théorie $\mathcal{T}$. Toutes ses autres déductions élémentaires appartiennent au répertoire de la logique classique. Cela comprend par exemple toutes les déductions de la forme suivante, où A et B sont deux propositions de $\mathcal{L}$:

$$\frac{\begin{array}{l} \Gamma \vdash A \Rightarrow B \\ \Gamma \vdash B \Rightarrow A \end{array}}{\Gamma \vdash A \Leftrightarrow B.}$$

Théorie d'une généalogie. Nous allons donner ici un exemple de ce que nous appelons une théorie « ad hoc », c'est-à-dire une théorie construite spécialement pour traiter d'un sujet bien particulier — par opposition au sujet très général des grandes théories mathématiques. Celle que nous présentons ici ne sert qu'à illustrer les définitions qui précèdent et n'a pas d'autre utilité. En informatique, l'application de méthodes formelles comprend parfois le développement de théories ad hoc sur les données particulières traitées dans un programme.

Une théorie est l'expression formelle d'une certaine *image* que l'on se fait d'une partie de la réalité. Celle qui suit traduit la vision de l'ensemble des êtres humains descendant tous d'un premier homme et d'une première femme. Il n'est pas question du temps dans cette vision. Les êtres humains ne sont considérés que du point de vue de leurs relations de parenté. Pour nommer cette théorie, nous l'appellerons la *théorie de la genèse*, ou $\mathcal{T}_G$. Le langage $\mathcal{L}$ de cette théorie sera le langage du premier ordre ayant pour symboles spécifiques les symboles suivants :

- a, e : constantes individuelles;
- F, H : symboles relationnels unaires;

= : symbole relationnel binaire;
 G : symbole relationnel ternaire.

Interprétation. Les constantes individuelles **a** et **e** désignent le premier homme (Adam) et la première femme (Eve). Les propositions atomiques Fx , Hx signifient que l'être humain désigné par x est une femme, respectivement un homme. La proposition atomique $Gxyz$ signifie que les êtres humains désignés par x et y sont respectivement le père et la mère (les géniteurs) de l'être humain désigné par z . La proposition $x = y$ s'interprète comme le fait que x et y désignent le même être humain.

Comme *axiomes* de $\mathcal{T}_G$, nous prenons les propositions suivantes:

Ha et Fe;
 non $\exists x \exists y Gxya$;
 non $\exists x \exists y Gxye$;
 non $\exists x (Hx \text{ et } Fx)$;
 $\forall x \forall y \forall z (Gxyz \Rightarrow (Hx \text{ et } Fy))$;
 $\forall x \forall y \forall z (Gxyz \Rightarrow (z \neq x \text{ et } z \neq y))$;
 $\forall x \forall y \forall z (Gxyz \Rightarrow (Hz \text{ ou } Fz))$;
 $\forall x \forall y \forall x' \forall y' \forall z ((Gxyz \text{ et } Gx'y'z) \Rightarrow (x = x' \text{ et } y = y'))$.

La signification de ces axiomes est claire: **a** est un homme; **e** est une femme; **a** n'a pas de parents; **e** non plus; personne n'est à la fois homme et femme; le père et la mère d'un être humain sont respectivement un homme et une femme; un être humain est différent de ses parents; l'être engendré par deux êtres humains est un homme ou une femme; un être humain n'a pas plus d'un père ni d'une mère.

Les déductions élémentaires de la théorie $\mathcal{T}_G$ sont celles de la logique classique, ainsi que les déductions de la forme suivante, où **R** est une proposition du langage $\mathcal{L}$ et **x**, **y**, **z** sont trois variables *distinctes*, satisfaisant à la condition que **y** et **z** soient librement substituables à **x** dans **R**:

$$\frac{\begin{array}{l} \Gamma \vdash ((a|x)R \text{ et } (e|x)R) \\ \Gamma \vdash \forall x \forall y \forall z [((R \text{ et } (y|x)R) \text{ et } Gxyz) \Rightarrow (z|x)R] \end{array}}{\Gamma \vdash \forall x R.} \quad (3)$$

Une telle déduction se fait normalement avec une proposition **R** dans laquelle la variable **x** figure librement, et la proposition exprime une propriété de l'être désigné par **x**. La proposition $((a|x)R \text{ et } (e|x)R)$ signifie que le premier homme et la première femme ont cette propriété. L'assertion de la deuxième prémisses signifie que cette propriété est *héréditaire*, en ce sens que si les deux parents d'un être humain ont cette propriété, alors cet être humain l'a aussi. On en déduit que tout être humain a cette propriété, parce que, dans cette vision de l'humanité, tout être humain descend des deux premiers (**a**, **e**). C'est le schéma de déduction (3) qui traduit dans la théorie cette vision de l'humanité descendant de deux êtres initiaux. Ce ne sont pas les axiomes de la théorie.

Prenons par exemple pour $\mathcal{R}$ la proposition $(Fx \text{ ou } Hx)$. Aucun axiome de la théorie n'affirme que tout être humain est un homme ou une femme. Mais cela peut se démontrer par une déduction suivant le schéma (3), car $\mathbf{a}$ et $\mathbf{e}$ ont cette propriété (premier axiome), et cette propriété est héréditaire (septième axiome).

5.4 Démonstrations et déductions

Les déductions élémentaires d'une théorie $\mathcal{T}$ font partie de la donnée de cette théorie, ainsi que son langage et ses axiomes. Mais elles ne constituent que le répertoire *initial* de déductions de $\mathcal{T}$. À ce répertoire viennent s'ajouter d'autres déductions, dites *dérivées*, qui s'obtiennent en combinant des déductions élémentaires ou des déductions dérivées déjà démontrées. Il y a donc un mécanisme d'extension du répertoire des déductions d'une théorie, et c'est ce mécanisme que nous décrivons dans ce paragraphe.

Dans cette description, nous utiliserons la notation horizontale $J_1, \dots, J_n \mid J$ pour désigner une déduction dont les prémisses sont n jugements $J_1, \dots, J_n$ et la conclusion est un jugement J . La barre verticale « $\mid$ » est la barre de déduction. Sauf indication contraire, il est admis dans cette notation que le nombre de prémisses peut être égal à zéro, auquel cas la déduction peut être notée aussi $\mid J$.

Lorsque nous parlons d'une déduction

$$J_1, \dots, J_n \mid J \tag{1}$$

d'une théorie $\mathcal{T}$, il s'agit donc soit d'une déduction élémentaire de $\mathcal{T}$, soit d'une déduction appartenant à un répertoire de déductions plus étendu, mais construit suivant le mécanisme d'extension décrit dans ce paragraphe. Au lieu de dire que la séquence de jugements (1) est une déduction de $\mathcal{T}$, on s'exprime aussi en disant que le jugement J *se déduit de* $J_1, \dots, J_n$ *dans* $\mathcal{T}$, ou encore que l'on peut déduire J de $J_1, \dots, J_n$ dans $\mathcal{T}$.

Lorsque le nombre de prémisses d'une déduction (1) de $\mathcal{T}$ est nul, on peut dire aussi que J se déduit de «rien» dans $\mathcal{T}$. Mais on dira plutôt que J est un *théorème* de $\mathcal{T}$, car c'est précisément la définition des théorèmes d'une théorie qui sera donnée plus loin (§5.5). Un théorème d'une théorie $\mathcal{T}$ est par définition un jugement J dans le langage de $\mathcal{T}$ tel que $\mid J$ soit une déduction sans prémisses de $\mathcal{T}$. Si l'on assimile une déduction sans prémisses $\mid J$ au jugement J , on peut dire que la notion de théorème d'une théorie $\mathcal{T}$ est un cas particulier de celle de déduction de $\mathcal{T}$.

Démonstrations. Ce que nous avons appelé ci-dessus de manière assez vague le «mécanisme d'extension» du répertoire des déductions d'une théorie $\mathcal{T}$ repose sur la notion de démonstration. Nous en donnons plus bas une définition précise. L'idée générale est qu'une démonstration est une combinaison de déductions, prises dans le répertoire de déductions dont on dispose pour une

théorie $\mathcal{T}$, et que cet assemblage de déductions fournit une déduction résultante dont il est la démonstration. Cette déduction résultante est ainsi démontrée et peut être ajoutée au répertoire.

Supposons par exemple que J_1, J_2, J_3, J_4 soient quatre jugements dans le langage d'une théorie $\mathcal{T}$ et que les deux séquences de jugements

$$J_1, J_2 \mid J_3, \quad J_1, J_3 \mid J_4 \quad (2)$$

soient des déductions de $\mathcal{T}$, élémentaires ou non. En utilisant la terminologie introduite ci-dessus on peut dire que, dans cette théorie, J_3 se déduit de J_1 et de J_2 , et que J_4 se déduit de J_1 et de J_3 . Le sens intuitif du mot « déduction » suggère alors de dire que J_4 se déduit indirectement de J_1 et de J_2 dans $\mathcal{T}$. La séquence de jugements

$$J_1, J_2 \mid J_4 \quad (3)$$

peut être admise dans $\mathcal{T}$, car elle résulte de la combinaison des déductions (2). La combinaison des déductions (2) est la démonstration de la déduction (3). Il reste à préciser exactement comment se combinent des déductions, dans le cas général. C'est le but de la définition 1 ci-dessous. Une démonstration comporte deux parties, appelées la base et le corps de la démonstration. Dans le présent exemple, ce sont

$$\begin{aligned} \text{Base: } & J_1, J_2 \\ \text{Corps: } & J_1, J_2, J_3, J_4. \end{aligned} \quad (4)$$

La base et le corps d'une démonstration sont composés de jugements. Le corps est une séquence de jugements qui possède la propriété suivante: chaque jugement de la séquence est un jugement de la base ou se déduit de jugements qui figurent avant lui dans le corps. J_3 se déduit de J_1, J_2 et J_4 se déduit de J_1, J_3 . La déduction résultante (3) est la séquence de jugements qui se compose des jugements de la base de la démonstration comme prémisses et du dernier jugement du corps de la démonstration comme conclusion.

Définition 1. Soit $\mathcal{L}$ le langage d'une théorie $\mathcal{T}$. Une *démonstration* $\mathcal{D}$ de $\mathcal{T}$, ou dans $\mathcal{T}$, est un document qui se compose de deux parties, appelées respectivement la *base* et le *corps* de $\mathcal{D}$.

La *base* de $\mathcal{D}$ est une liste, éventuellement vide, de jugements de $\mathcal{L}$.

Le *corps* de $\mathcal{D}$ est une séquence $K_1, \dots, K_p$ de jugements de $\mathcal{L}$ dans laquelle, pour chaque jugement K_i ($i = 1, \dots, p$), l'une au moins des conditions suivantes est vérifiée:

- (a) K_i est un jugement qui figure dans la base de la démonstration.
- (b) $\mid K_i$ est une déduction sans prémisses de $\mathcal{T}$.

(c) K_i se déduit de jugements $K_{i_1}, \dots, K_{i_m}$ qui le *précèdent* dans le corps de la démonstration, suivant une déduction de $\mathcal{T}$. En d'autres termes, il y a parmi les jugements $K_1, K_2, \dots, K_{i-1}$ des jugements $K_{i_1}, \dots, K_{i_m}$ tels que la séquence de jugements $K_{i_1}, \dots, K_{i_m} \mid K_i$ soit une déduction de $\mathcal{T}$.

Lorsque la base de $\mathcal{D}$ est vide, chacun des jugements K_i du corps de $\mathcal{D}$ vérifie l'une au moins des conditions (b), (c).

Le plan général du contenu d'un tel document est représenté dans la figure 5.1. La manière dont sont disposés dans une démonstration réelle la liste des jugements de la base et la séquence des jugements du corps de la démonstration n'a pas d'importance. Nous utiliserons d'ailleurs, dans les démonstrations des chapitres ultérieurs, une disposition qui n'est pas celle de la figure 5.1. Mais l'important est que la base et le corps de la démonstration soient clairement donnés, et pour bien comprendre le présent chapitre, il est commode de se représenter ce document sous la forme de deux tableaux intitulés *base* et *corps*, divisés en lignes contenant chacune un jugement.

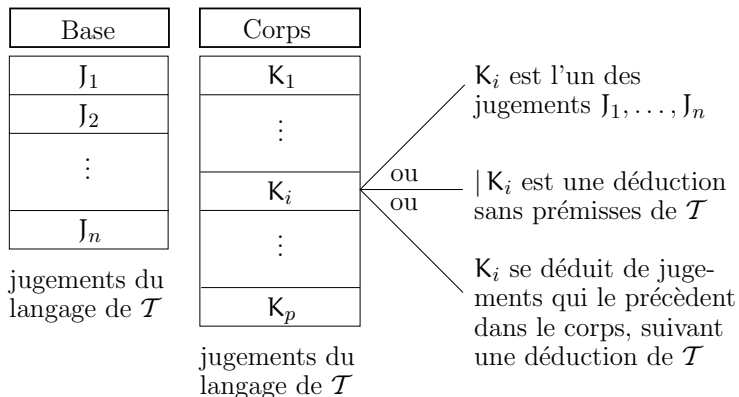


Figure 5.1. Plan général du contenu d'une démonstration $\mathcal{D}$ d'une théorie $\mathcal{T}$. La liste de jugements $J_1, \dots, J_n$ de la base de $\mathcal{D}$ peut être vide. Chaque jugement K_i du corps de $\mathcal{T}$ vérifie l'une au moins des trois conditions indiquées.

Déductions. Nous avons introduit plus haut la notion de démonstration comme étant une combinaison de déductions, *dont on tire* une déduction qui est le résultat de cette combinaison. La définition exacte de cette déduction résultante est la suivante: si $J_1, \dots, J_n$ sont les jugements de la base de la démonstration et si J est le dernier jugement du corps de celle-ci, la séquence de jugements $J_1, \dots, J_n \mid J$ est la déduction résultante. En fait il n'y a pas de raison de se limiter au dernier jugement du corps de la démonstration. On peut dire que pour chaque jugement J du corps de la démonstration la séquence $J_1, \dots, J_n \mid J$ est une déduction qui résulte de cette démonstration.

Comme nous l'avons expliqué plus haut, ceci est un « mécanisme d'extension » du répertoire des déductions d'une théorie $\mathcal{T}$. (a) On part d'un répertoire initial qui se compose des déductions élémentaires de $\mathcal{T}$. (b) On élargit une première fois ce répertoire en admettant comme déduction sans prémisses de $\mathcal{T}$ toute séquence $\mid J$ où J est un jugement du langage de $\mathcal{T}$ dont l'assertion est un axiome de $\mathcal{T}$, autrement dit un jugement de la forme $\Gamma \vdash A$ où A est un axiome de $\mathcal{T}$. L'idée est qu'un jugement de cette forme se déduit de « rien »

puisque'un axiome est une proposition dont la vérité est admise. (c) Ensuite, en combinant des déductions de ce répertoire, on peut écrire des démonstrations, dont les déductions résultantes s'ajoutent au répertoire et peuvent être utilisées dans de nouvelles démonstrations dont les déductions résultantes s'ajoutent au répertoire, et ainsi de suite. La définition des déductions d'une théorie $\mathcal{T}$ est donc donnée sous la forme de trois règles qui définissent le répertoire initial, son premier élargissement utilisant les axiomes de $\mathcal{T}$ et le mécanisme d'extension au moyen des démonstrations.

Définition 2. Soit $\mathcal{L}$ le langage d'une théorie $\mathcal{T}$. Les déductions de $\mathcal{T}$ sont les séquences de jugement de $\mathcal{L}$ définies inductivement par les règles suivantes :

- (a) Toute déduction élémentaire de $\mathcal{T}$ est une déduction de $\mathcal{T}$.
- (b) Si l'assertion d'un jugement J de $\mathcal{L}$ est un axiome de $\mathcal{T}$, autrement dit si J est de la forme $\Gamma \vdash A$ où A est un axiome de $\mathcal{T}$, la séquence de jugements $| J$ est une déduction sans prémisses de $\mathcal{T}$.
- (c) Si un jugement J de $\mathcal{L}$ figure dans le corps d'une démonstration de $\mathcal{T}$ de base $J_1, \dots, J_n$, la séquence de jugements $J_1, \dots, J_n | J$ est une déduction de $\mathcal{T}$. En particulier, si J figure dans le corps d'une démonstration de $\mathcal{T}$ de base vide, la séquence de jugements $| J$ est une déduction sans prémisses de $\mathcal{T}$.

Interprétation des définitions 1 et 2. Il y a dans les définitions 1 et 2 une « circularité » apparente qui n'aura pas échappé au lecteur. La définition des démonstrations d'une théorie $\mathcal{T}$ se réfère aux déductions de $\mathcal{T}$ et la définition des déductions de $\mathcal{T}$ se réfère aux démonstrations de $\mathcal{T}$. Nous devons donc mettre les choses au point, c'est-à-dire préciser l'interprétation correcte de ces définitions évitant toute circularité.

Les définitions 1 et 2 forment ensemble une définition simultanée de deux classes d'objets de nature textuelle: celle des démonstrations et celle des déductions d'une théorie $\mathcal{T}$. En termes techniques, on dit que ces deux classes sont définies par induction mutuelle. Pratiquement cela veut dire que les deux définitions décrivent de quelle manière on peut étendre le répertoire initial de déductions de $\mathcal{T}$ défini par les clauses (a) et (b) de la définition 2.

La définition 1 doit être interprétée *par rapport* à un répertoire déterminé de déductions de $\mathcal{T}$ que nous appelons le répertoire « actuel ». Celui-ci peut être le répertoire initial ou un répertoire qui est déjà obtenu par des extensions du répertoire initial. Dans les conditions (b) et (c) de la définition 1, le mot « déduction » doit être interprété dans le sens d'une déduction appartenant au répertoire actuel. Dans la clause (c) de la définition 2, le mot « démonstration » doit être interprété dans le sens d'une démonstration par rapport au répertoire de déductions actuel, et lorsqu'on dit que $J_1, \dots, J_n | J$ est une déduction de $\mathcal{T}$, on veut dire que cette séquence de jugements peut être *ajoutée* au répertoire actuel de déductions de $\mathcal{T}$, lequel subit ainsi une extension.

Déductions triviales. Soit $\mathcal{L}$ le langage d'une théorie $\mathcal{T}$. Si $J_1, \dots, J_n$ est une liste quelconque de jugements de $\mathcal{L}$, alors pour chaque jugement J_k de cette liste la séquence $J_1, \dots, J_n | J_k$ est une déduction de $\mathcal{T}$. Car le document formé

comme suit :

Base: $J_1, \dots, J_n$
Corps: J_k

est une démonstration de $\mathcal{T}$ puisque le seul jugement du corps figure dans la base. Nous dirons qu'une telle déduction est *triviale*. Un cas particulier est celui d'une *déduction identique* $J \mid J$.

Règle d'inclusion des prémisses. Si $J_1, \dots, J_m \mid J$ est une déduction d'une théorie $\mathcal{T}$, et si $J'_1, \dots, J'_n$ est une liste de jugements qui contient la liste $J_1, \dots, J_m$ en ce sens que chacun des jugements $J_1, \dots, J_m$ figure dans la liste $J'_1, \dots, J'_n$, alors la séquence $J'_1, \dots, J'_n \mid J$ est une déduction de $\mathcal{T}$. Car le document formé comme suit :

Base: $J'_1, \dots, J'_n$
Corps: $J_1, \dots, J_m, J$

est une démonstration de $\mathcal{T}$. Chacun des jugements $J_1, \dots, J_m$ du corps figure dans la base et $J_1, \dots, J_m \mid J$ est une déduction de $\mathcal{T}$.

Comme cas particulier, la liste $J'_1, \dots, J'_n$ peut être une liste construite en effectuant des permutations de jugements dans la liste $J_1, \dots, J_n$. On peut donc dire que l'ordre des prémisses d'une déduction n'a pas d'importance.

Jugements équivalents. Nous dirons que deux jugements J, J' du langage d'une théorie $\mathcal{T}$ sont *équivalents dans $\mathcal{T}$* si chacun se déduit de l'autre dans $\mathcal{T}$, autrement dit si les deux séquences de jugements $J \mid J'$ et $J' \mid J$ sont des déductions de $\mathcal{T}$.

5.5 Théorèmes

Définition. Un jugement J dans le langage $\mathcal{L}$ d'une théorie $\mathcal{T}$ est appelé un *théorème* de $\mathcal{T}$ si la séquence de jugements $\mid J$ est une déduction sans prémisses de $\mathcal{T}$.

Règle 1. Soit $\mathcal{L}$ le langage d'une théorie $\mathcal{T}$. Si A est un axiome de $\mathcal{T}$ et si Γ est une liste quelconque de propositions de $\mathcal{L}$, le jugement $\Gamma \vdash A$ est un théorème de $\mathcal{T}$.

Cette assertion est évidente d'après la définition ci-dessus et d'après la lettre (b) de la définition des déductions de $\mathcal{T}$ (§5.4, Déf. 2). La règle suivante découle semblablement de la lettre (c) de cette définition.

Règle 2. Si J est un jugement de $\mathcal{L}$ qui figure dans le corps d'une démonstration de $\mathcal{T}$ de base *vide*, J est un théorème de $\mathcal{T}$.

Règle 3. Si $J_1, \dots, J_n \mid J$ est une déduction de $\mathcal{T}$ et si chacune des prémisses $J_1, \dots, J_n$ est un théorème de $\mathcal{T}$, le jugement J est un théorème de $\mathcal{T}$.

Supposons en effet que $J_1, \dots, J_n \mid J$ soit une déduction de $\mathcal{T}$ et que chacune des prémisses $J_1, \dots, J_n$ soit un théorème de $\mathcal{T}$. Alors la liste de jugements *vide*, prise comme base, et la séquence de jugements $J_1, \dots, J_n, J$, prise comme corps,

forment une démonstration de $\mathcal{T}$. Pour bien le voir, il suffit de désigner la même séquence de jugements par $K_1, \dots, K_n, K_{n+1}$ et de se référer à la définition d'une démonstration (§5.4, Déf. 1). Chacun des jugements K_1 à K_n (J_1 à J_n) vérifie la condition (b) de cette définition puisqu'il est un théorème de $\mathcal{T}$. Le jugement K_{n+1} (J) vérifie la condition (c), puisque la séquence $J_1, \dots, J_n \mid J$ est une déduction de $\mathcal{T}$. Le jugement J figure ainsi dans le corps d'une démonstration de $\mathcal{T}$ de base vide. Il est donc un théorème de $\mathcal{T}$ d'après la règle 2.

Théorèmes sans hypothèses. Un théorème sans hypothèses d'une théorie $\mathcal{T}$, c'est-à-dire un théorème de $\mathcal{T}$ de la forme $\vdash A$, est parfois assimilé à la proposition A qui est l'assertion de ce théorème. On parle ainsi du « théorème » A , par abus de langage.

Axiomes en tant que théorèmes. Un axiome A d'une théorie $\mathcal{T}$ est parfois assimilé au théorème $\vdash A$ de $\mathcal{T}$ (Règle 1). On parle ainsi de « l'axiome » $\vdash A$, par abus de langage. Un axiome est vu ainsi comme un cas particulier de théorème sans hypothèses.

Démonstration d'une déduction, d'un théorème, d'une proposition. Nous précisons ici le sens du verbe « démontrer » que l'on utilise aussi bien à propos d'une déduction que d'un théorème, ou encore d'une proposition.

(a) Démontrer une déduction $J_1, \dots, J_n \mid J$ d'une théorie $\mathcal{T}$ par rapport à un répertoire actuel de déductions de $\mathcal{T}$ signifie construire une démonstration de $\mathcal{T}$ de base $J_1, \dots, J_n$, n'utilisant que des déductions de ce répertoire, dans le corps de laquelle figure le jugement J . En particulier, démontrer une déduction sans prémisses $\mid J$ de $\mathcal{T}$ signifie construire une démonstration de $\mathcal{T}$ de base vide dans le corps de laquelle figure le jugement J .

(b) Démontrer un théorème J d'une théorie $\mathcal{T}$, par rapport à un répertoire actuel de déductions de $\mathcal{T}$ signifie démontrer la déduction sans prémisses $\mid J$ de $\mathcal{T}$ au sens de (a).

(c) Démontrer une proposition A sous une liste d'hypothèses Γ dans une théorie $\mathcal{T}$ signifie démontrer le théorème $\Gamma \vdash A$ de $\mathcal{T}$ au sens de (b).

5.6 Propositions vraies, propositions fausses

Le but du présent chapitre, annoncé au début (§5.1), était de donner un sens précis aux notions de vérité et de contexte, sachant que la valeur de vérité d'une proposition dépend toujours du contexte. Nous sommes maintenant en mesure de le faire.

Propositions vraies. Un *contexte*, au sens de la logique, est défini par la donnée d'une *théorie* $\mathcal{T}$ et d'une liste Γ , éventuellement vide, de propositions du langage de $\mathcal{T}$ prises comme *hypothèses*. Nous rappelons que la donnée d'une théorie $\mathcal{T}$ comprend celle de son langage — un langage du premier ordre $\mathcal{L}$; d'une liste de propositions de $\mathcal{L}$ appelées axiomes de $\mathcal{T}$; et d'une classe de séquences de jugements de $\mathcal{L}$ appelées déductions élémentaires de $\mathcal{T}$.

Dire qu'une proposition A du langage $\mathcal{L}$ est *vraie* dans le contexte de la théorie $\mathcal{T}$ et de la liste d'hypothèses Γ signifie, par définition, que le jugement $\Gamma \vdash A$ est un théorème de $\mathcal{T}$. La preuve de cette vérité réside dans une démonstration de ce théorème. Le sens du mot « vrai » est ramené ainsi à celui de « théorème d'une théorie », défini au paragraphe 5.5.

Dans les ouvrages de mathématique, on n'utilise le nom de théorème que pour les théorèmes « importants » d'une théorie. Un jugement banal tel que $\vdash 1 + 1 = 2$ est rarement présenté comme un « théorème » d'analyse ou d'arithmétique. Mais au sens de la logique c'en est bien un. Chaque fois que le mot « vrai » peut être appliqué, le mot « théorème » peut l'être aussi, et vice-versa. La logique ne sait pas définir ce qu'est l'importance d'un théorème. Elle laisse cela à l'appréciation de l'utilisateur.

Propositions fausses. Une proposition A du langage $\mathcal{L}$ d'une théorie $\mathcal{T}$ est dite *fausse* dans le contexte de $\mathcal{T}$ et d'une liste d'hypothèses Γ si sa négation, la proposition $\text{non}A$, est vraie dans ce contexte, autrement dit si $\Gamma \vdash \text{non}A$ est un théorème de $\mathcal{T}$.

Contextes et théories contradictoires. Il ne faut pas confondre le sens *opérationnel* donné ici aux adjectifs « vrai » et « faux » avec le sens intuitif que l'on peut avoir de ces mots. Suivant la définition opérationnelle qui précède, une proposition A du langage $\mathcal{L}$ de $\mathcal{T}$ est dite vraie dans le contexte de $\mathcal{T}$ et de Γ lorsqu'on *peut démontrer* le théorème $\Gamma \vdash A$ dans $\mathcal{T}$. La proposition A est dite fausse dans ce contexte lorsqu'on *peut démontrer* le théorème $\Gamma \vdash \text{non}A$ dans $\mathcal{T}$. Or il y a bien des cas où l'on peut démontrer les deux. Dans de tels cas, la proposition A est *à la fois* vraie et fausse dans le contexte de $\mathcal{T}$ et de Γ . On dit alors que ce contexte est *contradictoire*. Nous verrons au chapitre 6 que suivant les schémas de déduction de la logique classique *toute* proposition de $\mathcal{L}$ est vraie et fausse dans un tel contexte. Ce n'est pas une situation rare. C'est même une situation que l'on recherche volontairement dans les démonstrations dites par réduction à l'absurde (§6.1).

Si deux jugements *sans hypothèses* de la forme $\vdash A$, $\vdash \text{non}A$ sont tous les deux des théorèmes d'une théorie $\mathcal{T}$, alors cette théorie elle-même est dite contradictoire. Une telle théorie est sans intérêt, car n'importe quel jugement du langage de $\mathcal{T}$ est un théorème de $\mathcal{T}$. On peut dire que celle-ci ne discerne pas le faux du vrai. La cause réside dans un mauvais choix d'axiomes ou de schémas de déduction élémentaires, qui demande à être corrigé.

En ce qui concerne les théories mathématiques usuelles, on ne sait pas si elles sont contradictoires ou non. Mais plus on accumule les démonstrations de théorèmes, plus la découverte d'une contradiction semble improbable. Les premières versions de la théorie des ensembles, à la fin du dix-neuvième siècle, étaient contradictoires. Ces « défauts de jeunesse » furent rapidement corrigés et de nos jours personne ne s'attend à l'apparition de nouvelles contradictions.

Propositions indécidables. Il y a bien des cas aussi où dans le contexte d'une théorie $\mathcal{T}$ et d'une liste d'hypothèses Γ une proposition A n'est ni vraie

ni fausse au sens opérationnel défini ci-dessus, c'est-à-dire où ni l'un ni l'autre des jugements $\Gamma \vdash A$, $\Gamma \vdash \text{non}A$ n'est un théorème de $\mathcal{T}$. On dit que la proposition A est *indécidable* dans ce contexte. Des exemples évidents sont fournis notamment par la plupart des propositions qui possèdent une variable libre sur laquelle on n'a fait aucune hypothèse dans le contexte. Par exemple, dans la théorie des nombres réels, on ne peut démontrer ni le jugement $\vdash x \geq 0$, ni le jugement $\vdash \text{non}(x \geq 0)$, puisque ces jugements ne font aucune hypothèse sur x . Plus exactement, il y a tout lieu de penser que ces jugements ne sont pas démontrables, car si l'on parvenait à démontrer l'un ou l'autre, la théorie en question se révélerait contradictoire, et c'est une possibilité considérée actuellement comme très improbable.

Par contre, que ce soit dans la théorie des nombres réels ou dans une autre théorie mathématique courante $\mathcal{T}$, il est très difficile de donner un exemple d'une proposition A sans variables libres pour laquelle on puisse en dire autant, c'est-à-dire pour laquelle on ait de bonnes raisons de penser que ni l'un ni l'autre des jugements sans hypothèses $\vdash A$, $\vdash \text{non}A$ n'est un théorème de $\mathcal{T}$. Pour une proposition A sans variables libres, notre intuition veut au contraire que l'une ou l'autre des propositions A , $\text{non}A$ soit vraie, et lorsque nous ne sommes pas parvenus à démontrer l'une ou l'autre, nous attribuons cet insuccès à un manque d'habileté. C'est la raison pour laquelle on persiste à chercher depuis des siècles la démonstration de conjectures célèbres comme celle de Goldbach sur les nombres entiers: tout nombre pair supérieur à 3 est la somme de deux nombres premiers. Une assertion sans variables libres comme celle-ci nous semble devoir être vraie ou fausse, donc que l'on doit pouvoir démontrer $\vdash A$ ou démontrer $\vdash \text{non}A$.

La découverte que cette intuition est erronée est considérée comme étant la grande découverte du vingtième siècle en logique. Elle est due au logicien autrichien Kurt Gödel, qui a montré dans un célèbre article en 1931 que si une théorie mathématique $\mathcal{T}$ n'est pas trop rudimentaire et si elle n'est pas contradictoire, il y a des propositions sans variables libres du langage de $\mathcal{T}$ qui sont indécidables, en ce sens que pour une telle proposition A ni l'un ni l'autre des jugements $\vdash A$, $\vdash \text{non}A$ n'est un théorème de $\mathcal{T}$.

5.7 Extensions d'une théorie

Définition. Nous dirons qu'une théorie $\mathcal{T}'$ est une *extension* d'une théorie $\mathcal{T}$ lorsque $\mathcal{T}$ et $\mathcal{T}'$ ont le même langage et lorsque les conditions suivantes sont satisfaites:

- (a) Tout axiome de $\mathcal{T}$ est un axiome de $\mathcal{T}'$.
- (b) Toute déduction élémentaire de $\mathcal{T}$ est une déduction élémentaire de $\mathcal{T}'$.

Le sens du mot « extension » est clair: on passe de $\mathcal{T}$ à $\mathcal{T}'$ en ajoutant des axiomes et/ou des déductions élémentaires. Comme cas limite, on peut dire que $\mathcal{T}$ est une extension d'elle-même, puisque les conditions (a) et (b) sont

satisfaites si l'on prend pour $\mathcal{T}'$ la théorie $\mathcal{T}$ elle-même. Il est clair d'autre part que la relation « est une extension » est transitive: si $\mathcal{T}'$ est une extension de $\mathcal{T}$ et si $\mathcal{T}''$ est une extension de $\mathcal{T}'$, alors $\mathcal{T}''$ est une extension de $\mathcal{T}$.

Propriété. Soient $\mathcal{T}$ et $\mathcal{T}'$ deux théories de même langage $\mathcal{L}$. Si $\mathcal{T}'$ est une extension de $\mathcal{T}$, toute déduction de $\mathcal{T}$ est une déduction de $\mathcal{T}'$, et par suite, tout théorème de $\mathcal{T}$ est un théorème de $\mathcal{T}'$.

Pour montrer cela, nous nous référons à la définition des déductions de $\mathcal{T}$ (§5.4, Déf. 2). Si $\mathcal{T}'$ est une extension de $\mathcal{T}$, les déductions de $\mathcal{T}$ mentionnées dans les lettres (a) et (b) de la définition en question sont des déductions de $\mathcal{T}'$, par définition d'une extension de $\mathcal{T}$. Donc toutes les déductions du répertoire initial de déductions de $\mathcal{T}$ sont des déductions de $\mathcal{T}'$. Pour les autres, il suffit de voir que si toutes les déductions d'un répertoire actuel de déductions de $\mathcal{T}$ sont des déductions de $\mathcal{T}'$, alors une démonstration de $\mathcal{T}$ par rapport à ce répertoire actuel est aussi une démonstration de $\mathcal{T}'$, et la déduction qui en résulte est donc aussi une déduction de $\mathcal{T}'$. Donc tout répertoire de déductions de $\mathcal{T}$ construit à partir du répertoire initial de la manière décrite au paragraphe 5.4 se compose uniquement de déductions de $\mathcal{T}'$.

6

Logique propositionnelle

6.1 Dédutions élémentaires

Pour tout langage du premier ordre $\mathcal{L}$, nous allons définir une théorie appelée *logique propositionnelle de langage $\mathcal{L}$* , que nous désignerons brièvement par $\mathbf{Lprop}(\mathcal{L})$. Nous rappelons qu'une théorie (§5.3) est caractérisée par son langage, ses axiomes, et ses déductions élémentaires. La théorie $\mathbf{Lprop}(\mathcal{L})$ a pour langage le langage $\mathcal{L}$. Elle n'a *pas d'axiomes*, ce qui est le cas limite d'une liste d'axiomes: la liste vide. Ses déductions élémentaires sont décrites par les 14 schémas de déduction S1–S14 de la figure 6.1.

Ces schémas de déduction, de même que ceux de la logique des prédicats (Chap. 9) et ceux de la logique des prédicats avec égalité (Chap. 11), font partie des schémas de déduction élémentaires de toutes les théories dites *logiques* (§5.3). Une théorie logique particulière, par exemple celle des ensembles, admet encore d'autres déductions élémentaires spécifiques, décrites par des schémas de déduction spécifiques.

Les schémas de déduction de la théorie $\mathbf{Lprop}(\mathcal{L})$ concernent l'emploi des connecteurs logiques (non, ou, et, $\Rightarrow$, $\Leftrightarrow$). Ces connecteurs sont communs à tous les langages du premier ordre, de sorte que l'on parle souvent de la «logique propositionnelle» sans en indiquer le langage, qui peut être n'importe quel langage du premier ordre $\mathcal{L}$. Le lecteur peut admettre que $\mathcal{L}$ désigne dans tout ce chapitre un langage du premier ordre indéterminé mais fixe et que «logique propositionnelle» signifie logique propositionnelle de langage $\mathcal{L}$.

Les schémas de déduction de la logique en général, non seulement la logique propositionnelle mais encore la logique des prédicats et la logique des prédicats avec égalité, concernent l'usage des symboles communs à tous les langages du premier ordre — connecteurs logiques, variables, quantificateurs — et du symbole relationnel d'égalité qui appartient à tous les langages usuels. Par contre les schémas de déduction spécifiques d'une théorie logique particulière concernent l'emploi des symboles qui lui sont spécifiques, par exemple le symbole d'appartenance $\in$ pour la théorie des ensembles, ou le symbole σ (successeur)

S1. *Vérité d'une hypothèse*

$$\frac{}{\Gamma; A \vdash A}$$

S2. *Vérité a fortiori*

$$\frac{\Gamma \vdash B}{\Gamma; A \vdash B}$$

S3. *Introduction de et*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \text{ et } B}$$

S4. *Élimination de et*

$$\frac{\Gamma \vdash A \text{ et } B}{\Gamma \vdash A} \quad \frac{\Gamma \vdash A \text{ et } B}{\Gamma \vdash B}$$

S5. *Introduction de ou*

$$\frac{\Gamma \vdash A}{\Gamma \vdash A \text{ ou } B} \quad \frac{\Gamma \vdash B}{\Gamma \vdash A \text{ ou } B}$$

S6. *Disjonction des cas*

$$\frac{\Gamma \vdash A \text{ ou } B \quad \Gamma; A \vdash C \quad \Gamma; B \vdash C}{\Gamma \vdash C}$$

S7. *Introduction de $\Rightarrow$*

$$\frac{\Gamma; A \vdash B}{\Gamma \vdash A \Rightarrow B}$$

S8. *Modus ponens*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash A \Rightarrow B}{\Gamma \vdash B}$$

S9. *Introduction de $\perp$*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash \text{non}A}{\Gamma \vdash \perp}$$

S10. *Ex falso quodlibet*

$$\frac{\Gamma \vdash \perp}{\Gamma \vdash B}$$

S11. *Réduction à l'absurde J*

$$\frac{\Gamma; A \vdash \perp}{\Gamma \vdash \text{non}A}$$

S12. *Réduction à l'absurde K*

$$\frac{\Gamma; \text{non}A \vdash \perp}{\Gamma \vdash A}$$

S13. *Introduction de $\Leftrightarrow$*

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash B \Rightarrow A}{\Gamma \vdash A \Leftrightarrow B}$$

S14. *Élimination de $\Leftrightarrow$*

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash A \Rightarrow B} \quad \frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash B \Rightarrow A}$$

Figure 6.1 Schémas de déduction élémentaires de la logique propositionnelle.

pour l'arithmétique du premier ordre (§5.3).

Nous rappelons que les déductions d'une théorie sont des séquences de jugements dans le langage de cette théorie. Un schéma de déduction définit une classe de déductions qui sont de la même forme. Par exemple, si A et B sont deux propositions du langage $\mathcal{L}$ et si Γ est une liste quelconque de propositions de $\mathcal{L}$, la séquence de jugements

$$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \text{ et } B.}$$

est une déduction élémentaire de la théorie $\mathbf{Lprop}(\mathcal{L})$, car elle est conforme au schéma S3. Si le langage $\mathcal{L}$ contient le symbole relationnel binaire $\in$, un exemple de déduction de cette forme est le suivant :

$$\frac{\begin{array}{l} \vdash x \in a \\ \vdash x \in b \end{array}}{\vdash x \in a \text{ et } x \in b.}$$

Tous les schémas de déduction logiques présentés dans ce chapitre et dans les suivants seront numérotés. Le schéma de déduction numéro n ($n = 1, 2, \dots$) sera désigné par S_n . Certains schémas, notamment tous ceux de la figure 6.1, portent un nom, écrit en italique à côté du numéro, plus facile à mémoriser peut-être que leur numéro. On pourra se référer à un schéma par son nom ou par son numéro.

Deux schémas analogues sont groupés parfois sous le même numéro et le même nom. Il y a ainsi deux schémas appelés *élimination de et* (S4), deux schémas *introduction de ou* (S5), deux schémas *élimination de $\Leftrightarrow$* (S14).

Lecture des schémas de déduction. Si nous représentons par

$$\frac{\Gamma_1 \vdash A_1 \quad \dots \quad \Gamma_p \vdash A_p}{\Gamma \vdash B} \quad (1)$$

la forme générale d'un schéma de déduction S_n , celui-ci peut se lire de deux manières :

(a) La première manière n'est qu'une lecture de (1) sans interprétation, en disant par exemple: $\Gamma \vdash B$ se déduit de $\Gamma_1 \vdash A_1, \dots, \Gamma_p \vdash A_p$. Dans le cas d'un schéma sans prémisses ($n = 0$), on pourra dire: $\Gamma \vdash B$ se déduit de rien.

(b) La seconde manière de lire un tel schéma se réfère à la définition des théorèmes d'une théorie et à la règle 3 du paragraphe 5.5. Cette règle permet de lire un schéma de déduction (1) de $\mathbf{Lprop}(\mathcal{L})$ en disant: si $\Gamma_1 \vdash A_1, \dots, \Gamma_p \vdash A_p$ sont des théorèmes de $\mathbf{Lprop}(\mathcal{L})$, $\Gamma \vdash B$ est un théorème de cette théorie. Plus généralement, puisque les déductions de $\mathbf{Lprop}(\mathcal{L})$ sont aussi des déductions de toute théorie logique de langage $\mathcal{L}$, on peut dire: si $\Gamma_1 \vdash A_1, \dots, \Gamma_p \vdash A_p$ sont des théorèmes d'une théorie logique $\mathcal{T}$, $\Gamma \vdash B$ est un théorème de $\mathcal{T}$.

En parlant de propositions vraies au lieu de théorèmes (§5.6), ceci peut s'énoncer: dans le contexte d'une théorie logique $\mathcal{T}$, si la proposition A_1 est vraie sous les hypothèses Γ_1 et si A_2 est vraie sous les hypothèses Γ_2 et ... et si A_p est vraie sous les hypothèses Γ_p , alors B est vraie sous les hypothèses Γ .

Lus de cette manière, la plupart des schémas de déduction de la logique propositionnelle présenteront pour certains lecteurs un caractère d'évidence, en raison du sens intuitif des connecteurs logiques qu'ils ont acquis dans leur formation. Mais dans ce domaine la formation des individus est très inégale et nous *n'attendons pas* du lecteur qu'il trouve d'emblée «évidents» tous les schémas de déduction que nous présentons. Dans ce domaine, comme dans bien d'autres, c'est l'étude et la pratique qui finissent par rendre les choses évidentes.

Comme exemple de cette lecture des schémas de déduction, le schéma S3 peut se lire comme suit: si des propositions A , B sont toutes deux vraies sous des hypothèses Γ , alors la proposition (A et B) est vraie sous ces hypothèses. Pour un schéma comme celui-ci dans lequel tous les jugements, prémisses et

conclusion, font les mêmes hypothèses Γ , on peut même s'abstenir de mentionner celles-ci et dire simplement : si deux propositions A , B sont vraies dans un même contexte, la proposition $(A \text{ et } B)$ est vraie dans ce contexte.

Introduction et élimination de symboles logiques. Dans le nom de certains schémas de déduction, il est question de *l'introduction* ou de *l'élimination* d'un certain symbole logique. Cette terminologie obéit à la règle suivante : il y a introduction de ce symbole lorsqu'il figure seulement dans la conclusion, et élimination lorsqu'il figure seulement dans une prémisse. Dans le premier cas il apparaît lorsqu'on passe des prémisses à la conclusion. Dans le second, il disparaît.

Ce genre de dénomination est utilisé plus systématiquement par certains auteurs. Ainsi le schéma de *disjonction des cas* est appelé aussi *élimination de ou*, le schéma *modus ponens* est nommé *élimination de $\Rightarrow$* , le schéma de *réduction à l'absurde* J est nommé *introduction de non*, et au lieu de *ex falso quodlibet* on dit *élimination de $\perp$* .

Nous passons maintenant en revue les 14 schémas de déduction élémentaires de la logique propositionnelle.

Schémas structurels S1–S2. Contrairement aux schémas suivants (S3 à S14) qui concernent chacun l'un ou l'autre des connecteurs logiques, les schémas S1 et S2 font partie du mécanisme général de la déduction, indépendamment de symboles logiques particuliers. Le premier (S1) traduit la notion même d'hypothèse : les hypothèses sont des propositions *admises comme vraies*. Donc si une proposition A figure dans une liste d'hypothèses, en particulier comme dernière proposition de cette liste, elle est vraie dans un contexte caractérisé par cette liste d'hypothèses. Le schéma S2 correspond à l'idée que la vérité d'une proposition B dans un contexte se conserve lorsqu'on ajoute une hypothèse supplémentaire A . Intuitivement, plus on fait d'hypothèses, plus il y a de propositions vraies. D'où le nom de la règle : si B est vraie sous une liste d'hypothèses Γ , alors B est vraie à *plus forte raison* (en latin : *a fortiori*) sous les hypothèses $\Gamma; A$.

Schémas sur la conjonction S3–S4. Ces schémas correspondent au sens usuel du mot « et ». Si deux propositions A , B sont vraies dans un même contexte, la proposition $(A \text{ et } B)$ est vraie dans celui-ci (S3). Inversement, si $(A \text{ et } B)$ est vraie dans un certain contexte, chacune des propositions A , B est vraie dans celui-ci (S4).

Substitutions dans les schémas de déduction. Dans tous nos schémas de déduction, les métasymboles A , B , C désignent des propositions quelconques d'un langage $\mathcal{L}$ et Γ désigne une liste quelconque de propositions de $\mathcal{L}$. Dans chaque déduction concrète suivant l'un ou l'autre de ces schémas, au cours d'une démonstration particulière, on aura à la place des métasymboles A , B , C des propositions déterminées d'un langage particulier.

Mais tout en utilisant les métasymboles A , B , C ou d'autres, désignant

des propositions *indéterminées*, on peut décrire des déductions *particulières*, conformes aux schémas généraux. Par exemple, les déductions décrites par les schémas S3 et S4 peuvent se faire avec des propositions A, B quelconques. Elles peuvent se faire, en particulier, avec des propositions A, B qui sont la même proposition. Autrement dit, les déductions de la forme suivante sont *conformes* aux schémas S3, S4:

$$\frac{\Gamma \vdash A}{\Gamma \vdash A \text{ et } A} \qquad \frac{\Gamma \vdash A \text{ et } A}{\Gamma \vdash A} \quad (2)$$

On obtient ainsi de « nouveaux » schémas de déductions permises par la logique propositionnelle, mais qui ne sont pas vraiment nouveaux, puisqu'ils ne décrivent qu'une sous-classe de la classe de déductions décrite par les schémas S3 et S4. En lisant ces schémas la manière indiquée plus haut dans *Lecture des schémas de déduction*, lettre (b), on peut dire ceci: lorsqu'une proposition A est vraie dans un contexte, la proposition (A et A) est vraie dans ce contexte, et vice-versa.

Les schémas de déduction (2) s'obtiennent simplement en remplaçant toutes les occurrences de la métavariable B par A dans les schémas S3 et S4. On peut faire d'autres substitutions dans ces schémas, par exemple remplacer A en chacune de ses occurrences par (B ou A) et remplacer simultanément chacune des occurrences de B par (B ou C). En faisant cela dans le schéma S3, on obtient le schéma suivant qui décrit aussi une sous-classe de la classe de déductions décrite par S3:

$$\frac{\Gamma \vdash B \text{ ou } A \quad \Gamma \vdash B \text{ ou } C}{\Gamma \vdash (B \text{ ou } A) \text{ et } (B \text{ ou } C)} \quad (3)$$

Car si A, B, C sont des propositions quelconques d'un langage $\mathcal{L}$ et si Γ est une liste quelconque de propositions de ce langage, la séquence de jugements

$$\frac{\Gamma \vdash B \text{ ou } A \quad \Gamma \vdash B \text{ ou } C}{\Gamma \vdash (B \text{ ou } A) \text{ et } (B \text{ ou } C)}$$

appartient à la classe de séquences de jugements décrite par le schéma S3.

En résumé, les métasymboles A, B, C, etc. des schémas de déduction sont des *métavariables* représentant des propositions indéterminées, auxquelles on peut substituer n'importe quelles expressions représentant aussi des propositions. Mais bien entendu, dans un même schéma, toutes les occurrences d'une même métavariable A, B, C, etc. représentent la même proposition, et l'on doit remplacer toutes ces occurrences par la même expression.

Schémas sur la disjonction S5–S6. Le schéma *introduction de ou* (S5) correspond de manière assez évidente au sens du mot « ou ». Si l'une des propositions A, B est vraie dans un contexte, la proposition A ou B est vraie dans ce contexte. Le schéma *disjonction des cas* (S6) correspond à une méthode de démonstration dont le lecteur aura sans doute rencontré de nombreux exemples

en mathématique. Lorsque, sous certaines hypothèses, on doit démontrer une proposition C , on se sert souvent d'une proposition de la forme A ou B , dont on sait qu'elle est vraie sous ces hypothèses. On démontre alors C dans le « cas » A , c'est-à-dire en faisant l'hypothèse supplémentaire A , et séparément dans le « cas » B , c'est-à-dire en prenant B comme hypothèse supplémentaire. On en conclut que C est vraie indépendamment de ces hypothèses supplémentaires, puisque A ou B est vraie. On dit alors que l'on a démontré C *par disjonction des cas* A , B .

Exemple. Supposons que dans le cadre de la théorie des nombres entiers on ait à démontrer, pour un entier x quelconque, la proposition C suivante :

$$\underbrace{\text{le produit } x(x+1) \text{ est un nombre pair}}_C.$$

On sait que la disjonction A ou B suivante est vraie :

$$\underbrace{(x \text{ est pair})}_A \text{ ou } \underbrace{(x+1 \text{ est pair})}_B.$$

Supposons que x soit pair (hypothèse A). Le produit d'un nombre pair par un nombre entier quelconque est pair. Donc $x(x+1)$ est pair. On a démontré C sous l'hypothèse A et celle-ci est abrogée (n'est plus en vigueur) à partir d'ici. Supposons que $x+1$ soit pair (hypothèse B). Alors $x(x+1)$ est pair pour la même raison. On a démontré C sous l'hypothèse B et celle-ci est abrogée à partir d'ici. Conclusion : $x(x+1)$ est vraie dans tous les cas, c'est-à-dire sans hypothèse sur x . On a démontré C par disjonction des cas : x est pair, $x+1$ est pair.

La liste d'hypothèses Γ sous laquelle se déroule cette démonstration peut être la liste vide, pour fixer les idées. On admet que $\Gamma \vdash (A \text{ ou } B)$ est un théorème démontré antérieurement. On a démontré ici le théorème $\Gamma; A \vdash C$, puis le théorème $\Gamma; B \vdash C$, et de ces *trois* théorèmes on a déduit $\Gamma \vdash C$, d'après le schéma S6.

Dans cet exemple les propositions A , B sont mutuellement exclusives, c'est-à-dire qu'on a le théorème $\Gamma \vdash \text{non}(A \text{ et } B)$. Les deux « cas » s'excluent l'un l'autre. Cette situation est très fréquente, cependant le schéma *disjonction des cas* (S6) ne l'exige pas.

Schémas sur l'implication S7–S8. Le premier de ces schémas, *introduction de $\Rightarrow$* , a été décrit et illustré précédemment (§5.2, Exemple 3). Il justifie la méthode la plus fréquemment utilisée pour démontrer une implication $A \Rightarrow B$, qui consiste à démontrer B sous l'hypothèse supplémentaire A . Le schéma *modus ponens* correspond au sens usuel des mots « si ... alors » représentés par le symbole $\Rightarrow$. Lorsqu'une proposition $A \Rightarrow B$ et la proposition A sont vraies dans un contexte, la proposition B est vraie dans ce contexte.

Schémas sur la proposition $\perp$. Nous rappelons que le symbole $\perp$ (§3.6) est un connecteur logique nulnaire formant à lui seul une proposition particulière :

la proposition *fausse* par excellence. Nous avons défini (§5.6) une proposition fausse dans un contexte comme étant une proposition A dont la négation $\text{non}A$ est vraie dans ce contexte. En ce qui concerne la proposition $\perp$, nous verrons plus loin que si Γ est une liste d'hypothèses quelconques dans le langage d'une théorie logique $\mathcal{T}$, le jugement $\Gamma \vdash \text{non}\perp$ est un théorème de $\mathcal{T}$. Autrement dit, la proposition $\perp$ est en effet fausse, au sens de notre définition de ce mot, dans n'importe quel contexte logique. C'est ce qui lui vaut le nom de proposition fausse *par excellence*.

Mais cela n'empêche pas la proposition $\perp$ d'être aussi vraie dans certains contextes logiques. Dans ceux-ci, elle est à la fois vraie et fausse. Suivant le schéma S9, la proposition $\perp$ peut se déduire de deux propositions contradictoires, c'est-à-dire deux propositions dont l'une est la négation de l'autre. À part cela, ces deux propositions peuvent être quelconques, seule leur contradiction importe. Sitôt que l'on a démontré une proposition A et sa négation $\text{non}A$ dans un même contexte, la proposition $\perp$ est vraie dans ce contexte.

Suivant le schéma S10, si la proposition $\perp$ est vraie dans un contexte logique, *n'importe quelle proposition* B est vraie dans ce contexte. Le nom latin de ce schéma, *ex falso quodlibet*, signifie: du faux (on déduit) ce que l'on veut. Un lecteur qui n'a pas déjà étudié un peu la logique peut trouver ce schéma surprenant et sans évidence intuitive. Cela provient du fait que le symbole $\perp$, dans le sens du « faux », n'est pas couramment employé en mathématique, contrairement aux connecteurs logiques et, ou, non, $\Rightarrow$, $\Leftrightarrow$. La mathématique usuelle n'a pas de symbole du faux, elle ne manipule pas formellement ce concept.

Schémas de réduction à l'absurde. Les deux schémas S11 et S12, appelés *réduction à l'absurde J* et *réduction à l'absurde K*, sont symétriques — par permutation de A et de $\text{non}A$. Les lettres J et K sont des abréviations germaniques des adjectifs « intuitionniste » et « classique »¹. La classe de déductions élémentaires que nous présentons dans ce chapitre caractérise la logique propositionnelle dite *classique*. La logique propositionnelle *intuitionniste* n'en admet qu'une sous-classe, celle que l'on obtient en supprimant le schéma de réduction à l'absurde classique (S12). Nous expliquerons plus loin (§6.4) les raisons que l'on peut avoir de rejeter cette forme de déduction, et les conséquences de ce rejet sur les mathématiques.

Les schémas S11 et S12 peuvent se lire comme suit: si, en ajoutant aux hypothèses Γ d'un contexte logique une proposition A (resp. la proposition $\text{non}A$), on parvient à démontrer la proposition notoirement fausse $\perp$, alors la proposition $\text{non}A$ (resp. A), c'est-à-dire le contraire de l'hypothèse faite, est vraie dans ce contexte.

Schémas sur l'équivalence S13–S14. Si une implication $A \Rightarrow B$ et l'implication réciproque $B \Rightarrow A$ sont vraies dans un même contexte logique, alors la

¹ Ce sont les symboles utilisés dans un travail influent du logicien allemand Gerhard Gentzen en 1934 pour se référer à la logique intuitionniste et à la logique classique.

proposition $A \Leftrightarrow B$ (A si et seulement si B , A équivaut à B) est vraie dans ce contexte, et vice-versa.

6.2 Quelques démonstrations

Exemple 1. Si A , B sont des propositions quelconques du langage $\mathcal{L}$, et si Γ est une liste quelconque de propositions de $\mathcal{L}$, le document

Base	Corps	(1)
$\Gamma \vdash A$	$\Gamma \vdash A$	
$\Gamma \vdash \text{non}A$	$\Gamma \vdash \text{non}A$	
	$\Gamma \vdash \perp$	
	$\Gamma \vdash B$	

est une démonstration de la théorie $\mathbf{Lprop}(\mathcal{L})$, au sens de la définition de ce terme (§5.4, Déf. 1). Chacun des quatre jugements du corps de ce document vérifie en effet l'une des conditions (a), (b), (c) de la définition mentionnée. Les deux premiers vérifient la condition (a): ils figurent dans la base. Le troisième, $\Gamma \vdash \perp$, vérifie la condition (c): il se déduit des deux premiers suivant une déduction de la théorie en question, à savoir la déduction

$$\frac{\Gamma \vdash A \quad \Gamma \vdash \text{non}A}{\Gamma \vdash \perp}$$

de schéma S9. Le quatrième, $\Gamma \vdash B$, vérifie aussi la condition (c): il se déduit du troisième suivant une déduction de la théorie, la déduction

$$\frac{\Gamma \vdash \perp}{\Gamma \vdash B}$$

de schéma S10. D'après la définition des déductions d'une théorie (§5.4, Déf. 2), la séquence de jugements

$$\frac{\Gamma \vdash \text{non}A \quad \Gamma \vdash A}{\Gamma \vdash B} \quad (2)$$

est une déduction de la théorie $\mathbf{Lprop}(\mathcal{L})$, puisqu'elle se compose des jugements de la base de la démonstration (1), suivis d'un jugement qui figure dans le corps de celle-ci. La séquence de jugements (2) est une *déduction dérivée* (§5.4) de la théorie $\mathbf{Lprop}(\mathcal{L})$.

Une séquence de jugements de la forme (2) est une déduction de la théorie $\mathbf{Lprop}(\mathcal{L})$ *quelles que soient* les propositions A , B et la liste de propositions Γ de $\mathcal{L}$. C'est pourquoi nous pouvons ajouter au répertoire initial de schémas de déductions de cette théorie, en tant que *schéma de déduction dérivé*, le

nouveau schéma

$$\frac{\Gamma \vdash A \quad \Gamma \vdash \text{non}A}{\Gamma \vdash B.} \quad (3)$$

Au chapitre 7, nous établirons un catalogue complet de schémas de déduction dérivés de logique propositionnelle. Rigoureusement parlant, un tel catalogue n'est jamais complet. On peut toujours démontrer de nouvelles déductions qui n'entrent pas dans un schéma du catalogue. Ce que nous entendons ici par « complet », c'est un recueil couvrant la grande majorité des déductions usuelles et suffisamment riche pour que l'on n'ait jamais à faire par la suite de longues chaînes de déductions de logique propositionnelle.

Le schéma (3) figurera notamment dans notre répertoire futur de schémas de déduction de logique propositionnelle, sous le nom: *tout se déduit d'une contradiction*. Car il dit ceci: si l'on a une contradiction dans un contexte logique, c'est-à-dire qu'une certaine proposition A y est à la fois vraie et fausse ($\text{non}A$ est vraie), *n'importe quelle proposition* B est vraie dans ce contexte.

Notation des démonstrations. Une idée qui vient immédiatement à l'esprit lorsqu'on regarde une démonstration telle que (1) est d'en simplifier la notation, vu que les jugements de la base figurent aussi dans le corps. Il suffit de marquer ces jugements d'un signe spécial dans le corps de la démonstration. C'est ce que nous ferons, en supprimant la colonne intitulée « base ». Mais la distinction entre base et corps restera présente conceptuellement. Par ailleurs, le corps des démonstrations sera muni dorénavant d'une *numérotation* des jugements. Enfin chaque jugement sera muni d'une *justification*, écrite à sa droite, indiquant laquelle des conditions (a), (b), (c) de la définition d'une démonstration (§5.4, Déf. 1) il vérifie. Cette indication sera donnée *explicitement*, non par les lettres (a), (b), (c), mais comme dans l'exemple suivant, qui sera la notation adoptée pour la démonstration (1):

$$\begin{array}{lll} 1^\circ & \Gamma \vdash A & \textit{base} \\ 2^\circ & \Gamma \vdash \text{non}A & \textit{base} \\ 3^\circ & \Gamma \vdash \perp & 1^\circ, 2^\circ, \perp\text{-intro} \\ 4^\circ & \Gamma \vdash B & 3^\circ, \textit{ex falso}. \end{array} \quad (4)$$

Le mot « base » indique évidemment qu'un jugement du corps de la démonstration figure dans la base (omise) de celle-ci. La justification du jugement 3° indique que celui-ci se déduit des jugements 1° et 2° suivant le schéma de déduction *introduction de $\perp$* (S9). Les noms des schémas de déduction seront généralement abrégés, comme dans cet exemple. Le jugement 4° se déduit de 3° suivant le schéma *ex falso quodlibet*.

La définition d'une démonstration (§5.4, Déf. 1) ne demande pas que tous les jugements de la base d'une démonstration figurent dans le corps de celle-ci. Aussi nous convenons que sauf indication contraire la base des démonstrations données dans la suite est composée seulement des jugements portant la mention « base » dans le corps.

Lorsque nous parlerons par la suite d'un jugement d'une démonstration, sans préciser s'il s'agit d'un jugement de la base ou du corps de celle-ci, il s'agira toujours d'un jugement du corps. Dans notre présentation des démonstrations, chaque jugement occupera toujours une ligne. C'est pourquoi, au lieu de parler des jugements d'une démonstration, nous parlerons souvent des *lignes* de celle-ci. À propos de la démonstration (4) par exemple, nous dirons que la ligne 3^o se déduit des lignes 1^o et 2^o suivant le schéma S9.

Permutation de lignes. Dans le corps d'une démonstration, on peut souvent permuter certains jugements sans que le document perde les propriétés d'une démonstration. Il suffit que chaque ligne qui se déduit d'autres lignes vienne après celles-ci. Dans l'exemple de la démonstration (4), la seule permutation possible est celle des deux premiers jugements, ce qui donne la démonstration

$$\begin{array}{llll}
 1^o & \Gamma \vdash \text{non}A & \text{base} & \\
 2^o & \Gamma \vdash A & \text{base} & \\
 3^o & \Gamma \vdash \perp & 2^o, 1^o, \perp\text{-intro} & (5) \\
 4^o & \Gamma \vdash B & 3^o, \text{ex falso.} &
 \end{array}$$

Nous avons permuté du même coup les numéros 1^o et 2^o dans la justification de la troisième ligne, vu que les prémisses du schéma de déduction S9 ($\perp$ -intro) sont A , $\text{non}A$ dans cet ordre. Mais ce respect de l'ordre des prémisses d'un schéma de déduction dans la justification d'une ligne de démonstration est facultatif. Ces justifications ne servent finalement qu'à faciliter le travail d'un contrôleur qui, disposant d'un catalogue de schémas de déductions, est chargé de vérifier si une séquence de jugements donnée, dans laquelle certains jugements portent la mention « base », est ou n'est pas une démonstration d'une théorie $\mathcal{T}$. Même si l'on s'abstient de toute justification, en ne laissant subsister que la mention « base », un contrôleur patient et méthodique pourra toujours discerner les séquences de jugements qui sont des démonstrations de celles qui n'en sont pas. Il devra seulement chercher lui-même, pour chaque ligne de la séquence qui lui est soumise, de quelles lignes elle est déduite et suivant quel schéma de déduction de son catalogue. Le lecteur informaticien se rendra bientôt compte que les schémas de déduction de la logique sont tels que cette vérification peut être effectuée par un programme.

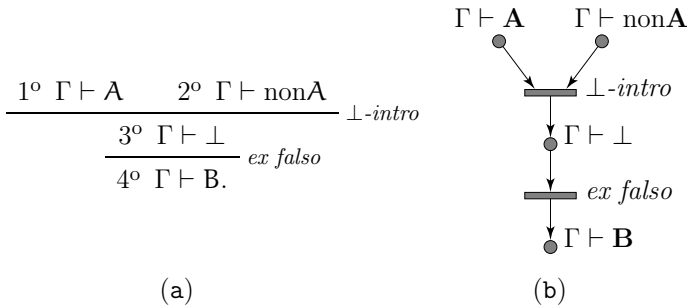


Figure 6.2 Arbre de la démonstration (4).

Arbres de démonstrations. Une présentation graphique des démonstrations consiste à noter chacune de leurs déductions sous la forme d'une « fraction » dont la barre est une barre de déduction et à combiner ces fractions entre elles. La démonstration (4) est présentée sous cette forme dans la figure 6.2 (a). Cette expression est appelée un *arbre de démonstration*. Nous en verrons plus loin des exemples plus compliqués. La structure d'arbre de cette expression est explicitée dans la figure 6.2 (b), c'est-à-dire en distinguant les nœuds de l'arbre des expressions qui leur sont associées. Il y a dans un arbre de démonstration deux espèces de nœuds, dessinés dans cette figure sous la forme de petits cercles et de barres épaisses. Les jugements de la démonstration sont associés aux cercles et les déductions (ou plus exactement les noms des schémas de déduction auxquelles elles appartiennent) sont associés aux barres. La racine de l'arbre est en bas. Le sens des arcs est inversé par rapport au sens normal d'un arbre: ils vont des feuilles de l'arbre vers la racine, ce qui correspond au sens de la déduction. Les deux types de nœuds alternent: les arcs vont d'un nœud de jugement à un nœud de déduction ou inversement. La racine de l'arbre est toujours un nœud de jugement. Il porte le dernier jugement de la démonstration. Les nœuds de jugement qui sont des feuilles de l'arbre (extrémités des branches) portent les jugements de la base de la démonstration.

Exemple 2. Si A, B sont deux propositions quelconques d'un langage $\mathcal{L}$ et si Γ est une liste quelconque de propositions de $\mathcal{L}$, les deux séquences de jugements qui suivent sont des démonstrations de la théorie $\mathbf{Lprop}(\mathcal{L})$, ou plus exactement les corps de deux démonstrations, dont les bases se composent chacune d'un seul jugement.

$$\begin{array}{ll}
 1^\circ & \Gamma \vdash A \Leftrightarrow B & \text{base} \\
 2^\circ & \Gamma \vdash A \Rightarrow B & 1^\circ, \Leftrightarrow\text{-élim (S14)} \\
 3^\circ & \Gamma \vdash B \Rightarrow A & 1^\circ, \Leftrightarrow\text{-élim} \\
 4^\circ & \Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A) & 2^\circ, 3^\circ, \text{et-intro (S3)}
 \end{array} \tag{6}$$

$$\begin{array}{ll}
 1^\circ & \Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A) & \text{base} \\
 2^\circ & \Gamma \vdash A \Rightarrow B & 1^\circ, \text{et-élim (S4)} \\
 3^\circ & \Gamma \vdash B \Rightarrow A & 1^\circ, \text{et-élim} \\
 4^\circ & \Gamma \vdash A \Leftrightarrow B & 2^\circ, 3^\circ, \Leftrightarrow\text{-intro (S13)}
 \end{array} \tag{7}$$

Par conséquent les deux séquences de jugements

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A)} \quad \frac{\Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A)}{\Gamma \vdash A \Leftrightarrow B} \tag{8}$$

sont des déductions de cette théorie, puisque chacune se compose du jugement de la base de la démonstration correspondante, suivi d'un jugement du corps de celle-ci (§5.4, Déf. 2, lettre (c)). Les fractions (8) peuvent être notées comme schémas de déduction dérivés de logique propositionnelle.

Les arbres des démonstrations (6) et (7) sont donnés dans la figure 6.3. Nous pouvons indiquer ici la manière générale de construire ces arbres. La

construction se fait de bas en haut, en remontant le corps de la démonstration. On commence par la racine de l'arbre, à laquelle on associe le dernier jugement de la démonstration. Ce jugement est surmonté d'une barre de déduction au-dessus de laquelle on écrit les jugements dont il est déduit. On procède de même pour chacun de ceux-ci.

$$\begin{array}{c}
 \frac{1^\circ \Gamma \vdash A \Leftrightarrow B}{2^\circ \Gamma \vdash A \Rightarrow B} \Leftrightarrow\text{-}\acute{e}lim \quad \frac{1^\circ \Gamma \vdash A \Leftrightarrow B}{3^\circ \Gamma \vdash B \Rightarrow A} \Leftrightarrow\text{-}\acute{e}lim \\
 \hline
 4^\circ \Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A) \quad \text{et-intro} \\
 \\
 \frac{1^\circ \Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A)}{2^\circ \Gamma \vdash A \Rightarrow B} \text{et-}\acute{e}lim \quad \frac{1^\circ \Gamma \vdash (A \Rightarrow B) \text{ et } (B \Rightarrow A)}{3^\circ \Gamma \vdash B \Rightarrow A} \text{et-}\acute{e}lim \\
 \hline
 4^\circ \Gamma \vdash A \Leftrightarrow B \quad \Leftrightarrow\text{-intro}
 \end{array}$$

Figure 6.3. Arbres des démonstrations (6) et (7).

On observe que dans l'arbre d'une démonstration il peut y avoir plusieurs nœuds qui correspondent à la même ligne de la démonstration. Dans les exemples présents, il s'agit de la ligne 1°. On est conduit au dédoublement de ce jugement, en construisant l'arbre de bas en haut comme nous venons de l'indiquer. Cela provient de ce que la ligne 1° est utilisée deux fois comme prémisses de déduction dans la démonstration. Si la ligne 1° était elle-même déduite de lignes précédentes, c'est tout un sous-arbre de la démonstration qu'il faudrait répéter. Pour cette raison, les arbres de démonstrations un peu complexes prennent rapidement des dimensions qui les rendent inutilisables en tant que notation pratique.

Exemple 3. Étant donné une proposition A quelconque du langage $\mathcal{L}$ et une liste de propositions Γ quelconque de ce langage, on peut écrire dans la théorie $\mathbf{Lprop}(\mathcal{L})$ la démonstration suivante, de base vide :

$$\begin{array}{llll}
 1^\circ & (\Gamma_1 \vdash A_1) & \Gamma; A \vdash A & \text{hyp} \\
 2^\circ & (\Gamma_2 \vdash A_2) & \Gamma; A; \text{non}A \vdash \text{non}A & \text{hyp} \\
 3^\circ & (\Gamma_3 \vdash A_3) & \Gamma; A; \text{non}A \vdash A & 1^\circ, \text{a fortiori} \\
 4^\circ & (\Gamma_4 \vdash A_4) & \Gamma; A; \text{non}A \vdash \perp & 2^\circ, 3^\circ, \perp\text{-intro} \\
 5^\circ & (\Gamma_5 \vdash A_5) & \Gamma; A \vdash \text{nonnon}A & 4^\circ, \text{red abs J} \\
 6^\circ & (\Gamma_6 \vdash A_6) & \Gamma \vdash A \Rightarrow \text{nonnon}A & 5^\circ, \Rightarrow\text{-intro}
 \end{array} \tag{9}$$

C'est un exemple de démonstration dans laquelle, contrairement aux précédentes, les jugements n'ont pas tous la même liste d'hypothèses. Pour commenter cette démonstration, nous avons désigné ses jugements par $\Gamma_i \vdash A_i$ ($i = 1, \dots, 6$). Ainsi Γ_4 désigne la liste d'hypothèses $\Gamma; A; \text{non}A$ (§5.2, Notations) et A_4 désigne la proposition $\perp$.

Les déductions faites dans cette démonstration devraient être claires. Le schéma **S1** (*vérité d'une hypothèse*), dont nous abrégeons le nom par « hyp », affirme que la dernière proposition d'une liste d'hypothèses est vraie sous ces

hypothèses. En particulier, la proposition A est vraie sous la liste d'hypothèses $\Gamma; A$ (ligne 1°) et la proposition $\text{non}A$ est vraie sous les hypothèses $\Gamma; A; \text{non}A$ (ligne 2°). Le schéma de déduction **S1** est un schéma sans prémisses. Les lignes 1°, 2° se déduisent de *rien* suivant ce schéma. Donc la seule mention de celui-ci constitue leur justification.

Le schéma de déduction **S2** (*vérité a fortiori*), affirme que si une proposition est vraie sous une liste d'hypothèses, cette proposition est vraie sous la même liste d'hypothèses augmentée d'une hypothèse quelconque. Dans la présente démonstration, la ligne 3° se déduit de la ligne 1° suivant ce schéma, car les propositions A_1 et A_3 sont identiques et la liste Γ_3 est la liste $\Gamma_1; \text{non}A$.

Le schéma **S9** (*introduction de $\perp$*) affirme que si deux propositions dont l'une est la négation de l'autre sont vraies sous une même liste d'hypothèses, la proposition $\perp$ est vraie sous celle-ci. La ligne 4° se déduit donc des lignes 2° et 3° suivant ce schéma, puisque ces trois lignes ont la même liste d'hypothèses et la proposition A_2 est la négation de A_3 .

Le schéma de *réduction à l'absurde* **J** (**S11**) affirme que si la proposition $\perp$ est vraie sous une liste d'hypothèses (non vide), alors la négation de la dernière proposition de cette liste est vraie sous les autres hypothèses de la liste. Donc, de la ligne 4°, on déduit que la négation de la dernière proposition de la liste Γ_4 , autrement dit la proposition $\text{nonnon}A$, est vraie sous les autres hypothèses de Γ_4 , c'est-à-dire sous la liste d'hypothèses $\Gamma; A$. C'est ce que contient la ligne 5°. Le passage de la ligne 5° à la ligne 6° est une déduction conforme au schéma *introduction de $\Rightarrow$* (**S7**).

$$\begin{array}{c}
 \frac{}{1^\circ \Gamma; A \vdash A} \text{hyp} \\
 \frac{}{3^\circ \Gamma; A; \text{non}A \vdash A} \text{a fortiori} \quad \frac{}{2^\circ \Gamma; A; \text{non}A \vdash \text{non}A} \text{hyp} \\
 \hline
 \frac{}{4^\circ \Gamma; A; \text{non}A \vdash \perp} \perp\text{-intro} \\
 \frac{}{5^\circ \Gamma; A \vdash \text{nonnon}A} \text{red abs J} \\
 \hline
 6^\circ \Gamma \vdash A \Rightarrow \text{nonnon}A \quad \Rightarrow\text{-intro}
 \end{array}$$

Figure 6.4. Arbre de la démonstration (9).

L'arbre de cette démonstration est donné dans la figure 6.4. La base de la démonstration étant vide, toutes les feuilles de cet arbre sont des barres de déduction sans prémisses.

La séquence de jugements (9) étant une démonstration de base vide de la théorie **Lprop**($\mathcal{L}$), le jugement $\Gamma \vdash A \Rightarrow \text{nonnon}A$, qui figure dans le corps de cette démonstration est un théorème de cette théorie (§5.5, Règle 2). D'après la définition d'un théorème (§5.5), il revient au même de dire que la séquence de jugements qui se compose de ce seul jugement est une déduction sans prémisses de **Lprop**($\mathcal{L}$). Nous pouvons noter comme schéma de déduction (dérivé) de

logique propositionnelle:

$$\frac{}{\Gamma \vdash A \Rightarrow \text{nonnon}A.} \quad (10)$$

6.3 Démonstrations naturelles

Exemple 1. La démonstration précédente (§6.2, Exemple 3) est recopiée dans la figure 6.5, en séparant visiblement chacune des listes d'hypothèses Γ_i de l'assertion A_i qui lui est associée.

	(Γ_0)	Γ		
1°	$(\Gamma_1 \vdash A_1)$	$\Gamma; A \vdash$	A	hyp
2°	$(\Gamma_2 \vdash A_2)$	$\Gamma; A; \text{non}A \vdash$	$\text{non}A$	hyp
3°	$(\Gamma_3 \vdash A_3)$	$\Gamma; A; \text{non}A \vdash$	A	1°, a fortiori
4°	$(\Gamma_4 \vdash A_4)$	$\Gamma; A; \text{non}A \vdash$	$\perp$	2°, 3°, $\perp$ -intro
5°	$(\Gamma_5 \vdash A_5)$	$\Gamma; A \vdash$	$\text{nonnon}A$	4°, red abs J
6°	$(\Gamma_6 \vdash A_6)$	$\Gamma \vdash$	$A \Rightarrow \text{nonnon}A$	5°, $\Rightarrow$ -intro .

Figure 6.5. Démonstration naturelle.

Cette démonstration est d'une forme que nous qualifions de *naturelle*, et nous allons mettre en évidence les traits qui lui valent ce qualificatif. Nous définirons ensuite de façon générale ce que nous appelons une démonstration naturelle.

Nous supposons toujours que A est une proposition quelconque d'un langage $\mathcal{L}$, et que Γ est une liste quelconque (éventuellement vide) de propositions de $\mathcal{L}$. Nous désignons aussi par Γ_0 la liste d'hypothèses Γ , que nous avons notée au-dessus de la démonstration pour les besoins de la discussion.

Premièrement, on peut observer que chacune des listes d'hypothèses Γ_i (pour $i \geq 1$) est une *extension* de la liste Γ_0 , et qu'il y a au moins une ligne dont la liste d'hypothèses est exactement Γ_0 (ligne 6°). Nous rappelons que selon notre terminologie (§5.2, Notations) une liste d'hypothèses est une extension d'elle-même, c'est pourquoi nous pouvons dire que *toutes* les lignes ont une liste d'hypothèse qui est une extension de Γ_0 , y compris la ligne 6°.

Deuxièmement, lorsque la liste Γ_{i+1} est différente de la liste Γ_i , on passe de Γ_i à Γ_{i+1} soit par adjonction d'une proposition à Γ_i (augmentation), soit par suppression de la dernière proposition de Γ_i (diminution):

Γ_1 est identique à:	$\Gamma_0; A$	augmentation
Γ_2 est identique à:	$\Gamma_1; \text{non}A$	augmentation
Γ_4 est identique à:	$\Gamma_5; \text{non}A$	diminution
Γ_5 est identique à:	$\Gamma_6; A$	diminution.

En termes d'informaticien, on peut dire que lorsque la liste d'hypothèses varie d'une ligne à la suivante, dans une démonstration naturelle, c'est toujours

comme une « pile » de propositions (en anglais *stack*), c'est-à-dire par adjonction ou suppression d'une proposition au sommet de la pile, si nous appelons ainsi l'extrémité droite de la liste.

Troisièmement, lorsqu'on passe de la liste d'hypothèses Γ_i à la liste Γ_{i+1} par augmentation (adjonction d'une proposition), la ligne $i + 1$ est toujours une *ligne d'hypothèse*, c'est-à-dire un jugement qui est déduit de rien suivant le schéma **S1** (*vérité d'une hypothèse*). La proposition qui est ajoutée à Γ_i dans la liste Γ_{i+1} est donc aussi l'assertion qui se trouve à droite du signe $\vdash$ dans la ligne $i + 1$.

Telles sont les trois propriétés caractéristiques d'une démonstration dite *naturelle*. L'intérêt de ces propriétés est qu'elles permettent une notation abrégée commode des démonstrations, dans laquelle on n'écrit pas explicitement la liste d'hypothèses de chaque ligne, mais on se contente de noter la liste Γ_0 et ensuite seulement les *variations* de la liste, d'une ligne à la suivante. Cette notation est décrite plus bas sous le titre « *Blocs d'une démonstration naturelle* ». Nous formulons d'abord la définition générale d'une démonstration naturelle.

Définition. Une démonstration d'une théorie logique $\mathcal{T}$ dont le corps est la séquence de jugements

$$\begin{array}{c} \Gamma_1 \vdash A_1 \\ \dots \\ \Gamma_p \vdash A_p \end{array}$$

est dite *naturelle* si toutes les listes d'hypothèses $\Gamma_1, \dots, \Gamma_p$ de ces jugements sont des extensions de l'une d'entre elles, et si, lorsqu'on désigne cette liste minimale par Γ_0 , chacune des listes Γ_i ($i = 1, \dots, n$) vérifie l'une des conditions suivantes :

- (a) Γ_i est identique à Γ_{i-1} ;
- (b) Γ_i est identique à $\Gamma_{i-1}; A_i$;
- (c) Γ_i s'obtient en supprimant la dernière hypothèse de la liste Γ_{i-1} .

Dans le cas (b), le jugement $\Gamma_i \vdash A_i$ est identique à $\Gamma_{i-1}; A_i \vdash A_i$ et se déduit donc de rien suivant le schéma **S1**.

Blocs d'une démonstration naturelle. Une notation abrégée des démonstrations naturelles est facile à imaginer lorsqu'on se rend compte que celle de la figure 6.5 est entièrement déterminée par la partie qui se trouve à droite de la ligne verticale pointillée, et par la donnée de la liste d'hypothèses Γ . Il suffit d'adopter les conventions typographiques suivantes : les assertions de deux lignes consécutives ayant la même liste d'hypothèses sont alignées verticalement; l'assertion d'une ligne d'hypothèse est décalée à droite par rapport à celle de la ligne précédente ou par rapport à la marge gauche; l'assertion d'une ligne dont la liste d'hypothèses est diminuée par rapport à la précédente est décalée à gauche par rapport à la précédente. C'est ainsi qu'est disposée la partie droite de la figure 6.5. En vertu de ces conventions, on peut effacer toute la partie gauche, en ne laissant subsister que la numérotation des lignes,

et la liste d'hypothèses Γ . La partie gauche peut être reconstituée au besoin d'après la partie droite et la donnée de Γ .

		Γ																
$\Gamma; A$	1°	<table> <tr> <td colspan="2">A</td> <td>hyp</td> </tr> <tr> <td colspan="2">nonA</td> <td>hyp</td> </tr> <tr> <td>A</td> <td></td> <td>1°, a fortiori</td> </tr> <tr> <td>$\perp$</td> <td></td> <td>2°, 3°, $\perp$-intro</td> </tr> <tr> <td colspan="2">nonnonA</td> <td>4°, red abs J</td> </tr> </table>	A		hyp	non A		hyp	A		1°, a fortiori	$\perp$		2°, 3°, $\perp$ -intro	nonnon A		4°, red abs J	
A		hyp																
non A		hyp																
A		1°, a fortiori																
$\perp$		2°, 3°, $\perp$ -intro																
nonnon A		4°, red abs J																
$\Gamma; A; \text{non}A$	2°																	
$\Gamma; A; \text{non}A$	3°																	
$\Gamma; A; \text{non}A$	4°																	
$\Gamma; A$	5°																	
Γ	6°	$A \Rightarrow \text{nonnon}A$	5°, $\Rightarrow$ -intro.															

Figure 6.6 Notation emboîtée d'une démonstration naturelle.

La disposition typographique que nous utiliserons finalement est donnée dans la figure 6.6. Nous avons laissé subsister encore les listes d'hypothèses de toutes les lignes pour les besoins de la discussion, mais elles disparaîtront par la suite. La notation utilise des *cadres*, qui entourent ce que nous appelons les *blocs de lignes* de la démonstration. Un bloc de lignes est un groupe de lignes *consécutives*, dont les listes d'hypothèses sont toutes des extensions d'une même liste, associée au groupe.

Le bloc *extérieur* se compose notamment de toutes les lignes de la démonstration, et la liste d'hypothèses Γ_0 (Γ) lui est associée. Il est entouré par le cadre extérieur dans la figure 6.6. Les lignes 1° à 5° forment un bloc auquel est associée la liste d'hypothèses $\Gamma; A$. Les lignes 2° à 4° forment un troisième bloc, auquel est associée la liste $\Gamma; A; \text{non}A$. Ce sont les seuls blocs de cette démonstration naturelle, car nous devons ajouter dans la définition d'un bloc que celui-ci est un groupe *maximal* de lignes consécutives dont les listes d'hypothèses sont des extensions d'une même liste. Par exemple les lignes 1° à 4° ont des listes d'hypothèses qui sont toutes des extensions de la liste $\Gamma; A$, mais elles ne forment pas un groupe maximal avec cette propriété, puisque le groupe de lignes 1° à 5° a la même propriété.

Une autre manière de définir un bloc de lignes d'une démonstration naturelle est de dire qu'il s'agit du bloc extérieur (toutes les lignes) ou alors d'un groupe de lignes consécutives satisfaisant aux trois conditions suivantes :

- (a) La première ligne du bloc est une ligne d'hypothèse (lignes 1° et 2°).
- (b) La liste d'hypothèses de chacune des lignes du bloc est une extension de celle de la première ligne du bloc. Cette liste est par définition la liste d'hypothèses associée au bloc. Les hypothèses de cette liste sont donc présentes dans tout le bloc et l'on dira qu'elles sont *en vigueur* dans tout le bloc.
- (c) La dernière ligne du bloc est la dernière ligne de la démonstration, ou alors elle est suivie d'une ligne dont la liste d'hypothèses n'est pas une extension de celle du bloc.

Chaque fois que l'on fait une hypothèse, dans une démonstration naturelle, on entre dans un bloc, et lorsqu'on sort de ce bloc, c'est-à-dire lorsqu'on passe

de la dernière ligne du bloc à la ligne suivante de la démonstration, cette hypothèse est retirée de la liste.

La notation abrégée d'une démonstration naturelle consiste à mettre en évidence les blocs de la démonstration, ce que nous faisons au moyen de cadres, et à omettre les listes d'hypothèses de toutes les lignes, en ne mentionnant que la liste minimale Γ associée au bloc extérieur. La liste d'hypothèses d'une ligne qui appartient à un bloc intérieur peut être reconstituée en ajoutant à Γ la première assertion de chacun des blocs intérieurs auxquels appartient la ligne, par ordre décroissant de ceux-ci. Nous appellerons cette notation la *notation emboîtée* d'une démonstration naturelle.

La structure des démonstrations naturelles ainsi représentées correspond à celle des démonstrations mathématiques usuelles, et c'est ce qui leur vaut le qualificatif « naturelles ». En effet, les hypothèses faites au cours d'une démonstration mathématique usuelle, signalées par des déclarations telles que « supposons que ... » ou « soit ... », sont *mémorisées* par le lecteur. Celui-ci possède donc à chaque instant (mentalement) une liste d'hypothèses en vigueur, laquelle s'allonge chaque fois qu'une nouvelle hypothèse est déclarée et se raccourcit chaque fois que le lecteur réalise qu'une hypothèse cesse d'être en vigueur. Il se trouve seulement que cette fin de validité n'est pas signalée explicitement, contrairement à notre notation. Chaque assertion d'une démonstration mathématique usuelle est ainsi appréciée par rapport à une liste d'hypothèses bien déterminée, et si l'on prenait la peine d'écrire cette liste à côté de chacune des assertions, on obtiendrait une séquence de jugements, autrement dit un corps de démonstration formelle comme nous l'avons défini.

Il nous reste à mentionner une propriété importante des théories logiques, sans laquelle la notion de démonstration naturelle perdrait de son intérêt. C'est que tout théorème d'une théorie logique, de même que toute déduction dérivée d'une telle théorie, admet une démonstration de forme naturelle. Cette propriété deviendra rapidement évidente. Sa preuve est le sujet d'un exercice du paragraphe 6.5.

En raison de cette propriété des démonstrations naturelles et de leur relation avec les démonstrations usuelles des mathématiques, nous ne donnerons dans les chapitres suivants que des démonstrations de cette forme.

Méthodes de démonstration. Il est utile de se représenter la manière dont les schémas de déduction de la logique propositionnelle sont utilisés dans les démonstrations naturelles au moyen de plans ou schémas de démonstration que l'on peut appeler des méthodes de démonstration. Cela est fait dans la figure 6.7, pour les schémas de déduction S7 (introduction de $\Rightarrow$), S11 (réduction à l'absurde J), S2 (a fortiori), et S6 (disjonction des cas).

Le premier schéma de démonstration de la figure indique comment on démontre une proposition $A \Rightarrow B$ sous une liste d'hypothèses Γ , par *introduction de $\Rightarrow$* . Dans un bloc auquel est associée la liste d'hypothèses Γ , on

fait l'hypothèse supplémentaire A , ce qui constitue la première ligne d'un bloc intérieur au premier. Dans ce bloc intérieur, auquel est associée la liste d'hypothèses $\Gamma; A$, on démontre la proposition B . Lorsqu'on y est parvenu (ligne n^o), on peut marquer la fin de ce bloc intérieur par la fermeture du cadre correspondant et écrire $A \Rightarrow B$ à la ligne suivante. Cette ligne se justifie en donnant le numéro de celle dont elle est déduite (n^o) et le nom du schéma de déduction. Dans ce schéma de démonstration, le bloc auquel est associée la liste d'hypothèses Γ n'est pas nécessairement le bloc extérieur d'une démonstration. Ce peut être un bloc intérieur. Par ailleurs, la démonstration de B dans le bloc correspondant à la liste d'hypothèses $\Gamma; A$ peut nécessiter l'ouverture et la fermeture de blocs intérieurs à celui-ci. Ces remarques s'appliquent, *mutatis mutandis*, aux trois autres schémas de démonstration de la figure.

Le deuxième schéma de démonstration indique comment on démontre une proposition $\text{non}A$ sous une liste d'hypothèses Γ , par *réduction à l'absurde* J . Dans un bloc auquel est associée la liste d'hypothèses Γ , on fait l'hypothèse supplémentaire A , ce qui constitue la première ligne d'un bloc intérieur au premier. Dans ce bloc intérieur, c'est-à-dire sous les hypothèses $\Gamma; A$, on démontre la proposition $\perp$. Lorsqu'on y est parvenu (ligne n^o) on peut marquer la fin de ce bloc intérieur et écrire $\text{non}A$ à la ligne suivante. Cette ligne se justifie en donnant le numéro de celle dont elle est déduite (n^o) et le nom du schéma de déduction.

Le schéma de déduction $S2$ (*vérité a fortiori*) permet de copier une assertion B d'un bloc d'une démonstration naturelle dans un bloc intérieur au premier. Cela est représenté dans le troisième schéma de démonstration de la figure 6.7. Une proposition B ayant été démontrée sous une liste d'hypothèses Γ , cette proposition est encore vraie sous une hypothèse supplémentaire A . La ligne du bloc intérieur contenant B doit venir bien entendu après celle du bloc extérieur contenant B (ligne n^o), dont elle se déduit par *vérité a fortiori*.

Enfin, la quatrième partie de la figure 6.7 présente la démonstration d'une proposition C sous une liste d'hypothèses Γ , par *disjonction des cas* à partir d'une proposition A ou B qui est vraie sous ces hypothèses. Dans un bloc auquel est associée la liste Γ , et après une ligne contenant la proposition A ou B (ligne k^o), on fait d'abord l'hypothèse A , et dans le bloc intérieur qui commence ainsi, on démontre C (ligne m^o). On peut alors marquer la fin de ce bloc, puis le début d'un autre bloc intérieur commençant par l'hypothèse B . Dans ce deuxième bloc intérieur, on démontre aussi C (ligne n^o), puis on marque la fin de ce bloc. À ce stade, on a démontré les trois jugements: $\Gamma \vdash A$ ou B , $\Gamma; A \vdash C$, $\Gamma; B \vdash C$. On peut alors écrire C à la ligne suivante, dans le bloc correspondant à Γ , en justifiant cette ligne par les numéros k^o , m^o , n^o des lignes dont elle est déduite et par le nom du schéma de déduction.

À vrai dire, ce dernier schéma de démonstration n'est pas rigoureusement de forme naturelle au sens de la définition qui précède. Car lorsqu'on passe de la ligne m^o à la suivante, la liste d'hypothèses passe de $\Gamma; A$ à $\Gamma; B$. Ce

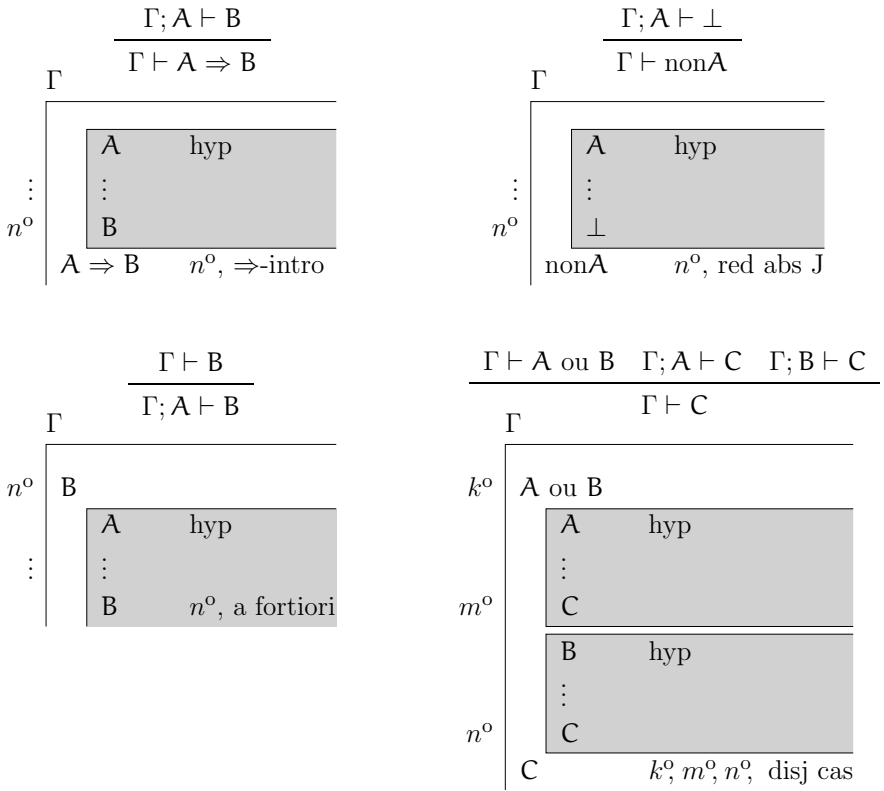


Figure 6.7. Méthodes de démonstration naturelle.

changement n'est pas de l'une des trois formes admises dans la définition d'une démonstration naturelle, selon laquelle la liste d'hypothèses ne peut varier d'une ligne à la suivante que par adjonction ou suppression d'une hypothèse. Il y a ici d'un seul coup suppression de A et adjonction de B , autrement dit *changement* de la dernière hypothèse de la liste. Cependant, on peut facilement mettre cette démonstration sous forme naturelle en insérant entre les deux cadres intérieurs une ligne intermédiaire avec la proposition $A \Rightarrow C$, déduite de la ligne m^o suivant $\Rightarrow\text{-introduction}$.

Nous admettons les démonstrations comme celle-ci, avec changement de la dernière hypothèse entre deux lignes consécutives, ce qui se produit toujours dans une démonstration par disjonction des cas. Mais on peut remarquer à ce sujet que lorsqu'on s'en tient strictement à la forme naturelle, le dessin des cadres est en principe superflu, car alors l'indentation (décalage et alignement) des propositions suffit à marquer les variations de la liste d'hypothèses d'une ligne à l'autre, donc le début et la fin des blocs.

Exemple 2. Étant donné des propositions A, B, C quelconques d'un langage $\mathcal{L}$ et une liste quelconque Γ de propositions de $\mathcal{L}$, la séquence de jugements de

1°	$(\Gamma_1 \vdash A_1)$	$\Gamma \vdash A \Rightarrow B$	<i>base</i>	
2°	$(\Gamma_2 \vdash A_2)$	$\Gamma \vdash B \Rightarrow C$	<i>base</i>	
3°	$(\Gamma_3 \vdash A_3)$	$\Gamma; A \vdash A$	<i>hyp</i>	
4°	$(\Gamma_4 \vdash A_4)$	$\Gamma; A \vdash A \Rightarrow B$	1°, a fortiori	
5°	$(\Gamma_5 \vdash A_5)$	$\Gamma; A \vdash B$	3°, 4°, modus ponens	(a)
6°	$(\Gamma_6 \vdash A_6)$	$\Gamma; A \vdash B \Rightarrow C$	2°, a fortiori	
7°	$(\Gamma_7 \vdash A_7)$	$\Gamma; A \vdash C$	5°, 6°, modus ponens	
8°	$(\Gamma_8 \vdash A_8)$	$\Gamma \vdash A \Rightarrow C$	7°, $\Rightarrow$ -intro.	

$\frac{}{3^\circ \Gamma; A \vdash A} \text{hyp}$		$\frac{1^\circ \Gamma \vdash A \Rightarrow B}{4^\circ \Gamma; A \vdash A \Rightarrow B} \text{a fortiori}$		
		$\frac{}{5^\circ \Gamma; A \vdash B} \text{modus ponens}$	$\frac{2^\circ \Gamma \vdash B \Rightarrow C}{6^\circ \Gamma; A \vdash B \Rightarrow C} \text{a fortiori}$	
				$\frac{}{7^\circ \Gamma; A \vdash C} \text{modus ponens}$
				$\frac{}{8^\circ \Gamma \vdash A \Rightarrow C} \Rightarrow\text{-intro}$

Γ			
1°	$A \Rightarrow B$	<i>base</i>	
2°	$B \Rightarrow C$	<i>base</i>	
3°	A	<i>hyp</i>	
4°	$A \Rightarrow B$	1°, a fortiori	
5°	B	3°, 4°, modus ponens	
6°	$B \Rightarrow C$	2°, a fortiori	
7°	C	5°, 6°, modus ponens	
8°	$A \Rightarrow C$	7°, $\Rightarrow$ -intro	(c)

Figure 6.8. Démonstration naturelle de la transitivité de l'implication: (a) corps, (b) arbre, (c) notation emboîtée.

la figure 6.8 (a) est une démonstration de base $\Gamma \vdash A \Rightarrow B$, $\Gamma \vdash B \Rightarrow C$ de la théorie $\mathbf{Lprop}(\mathcal{L})$, et plus généralement de toute théorie logique $\mathcal{T}$ de langage $\mathcal{L}$. L'arbre de cette démonstration est donné dans la figure 6.8 (b). Les feuilles de cet arbre sont la barre de déduction sans prémisse (hyp) et les deux jugements de la base. On s'intéresse particulièrement au jugement $\Gamma \vdash A \Rightarrow C$ (8°) du corps de cette démonstration. La séquence de jugements constituée par les deux jugements de la base et ce dernier jugement est une déduction (dérivée) de la théorie considérée (§5.4, Déf. 2, lettre (c)). Nous pouvons donc noter, comme schéma de déduction dérivé de logique propositionnelle, le schéma suivant dit *transitivité de l'implication*:

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash B \Rightarrow C}{\Gamma \vdash A \Rightarrow C}.$$

La démonstration de la figure 6.8 (a) est de forme naturelle. Toutes les listes d'hypothèses Γ_i ($i = 1, \dots, 8$) sont des extensions de Γ , et il y a au moins une

	Γ	
1°	$A \Rightarrow B$	hyp
2°	nonB	hyp
3°	$A \Rightarrow B$	1°, a fortiori
4°	A	hyp
5°	$A \Rightarrow B$	3°, a fortiori
6°	B	4°, 5°, mod ponens
7°	nonB	2°, a fortiori
8°	$\perp$	6°, 7°, $\perp$ -intro
9°	nonA	8°, red abs J
10°	$\text{nonB} \Rightarrow \text{nonA}$	9°, $\Rightarrow$ -intro
11°	$(A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA})$	10°, $\Rightarrow$ -intro.

Figure 6.9. Démonstration naturelle d'un théorème de contraposition $\Gamma \vdash (A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA})$.

liste Γ_i identique à Γ (voir $\Gamma_1, \Gamma_2, \Gamma_8$). Chacune des listes Γ_i ($i = 1, \dots, 8$) est soit identique à Γ_{i-1} (on définit Γ_0 comme étant la liste Γ), soit identique à $\Gamma_{i-1}; A_i$, soit identique à Γ_{i-1} privée de sa dernière proposition. La notation emboîtée de cette démonstration est donnée dans la figure 6.8 (c).

Exemple 3. On appelle *contraposée* d'une proposition de la forme $A \Rightarrow B$ la proposition $\text{nonB} \Rightarrow \text{nonA}$. Si A, B sont des propositions quelconques d'un langage $\mathcal{L}$, et si Γ est une liste quelconque de propositions de $\mathcal{L}$, le jugement

$$\Gamma \vdash (A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA}) \quad (1)$$

est un théorème de la théorie $\mathbf{Lprop}(\mathcal{L})$, et plus généralement de toute théorie logique $\mathcal{T}$ de langage $\mathcal{L}$. La preuve en est fournie par la démonstration naturelle de la figure 6.9. La base de cette démonstration est vide, et le jugement (1) figure dans le corps de celle-ci (§5.5, Règle 2).

Dire que le jugement (1) est un théorème de $\mathbf{Lprop}(\mathcal{L})$ signifie par définition (§5.5) que la séquence de jugements composée de ce seul jugement est une déduction sans prémisses de cette théorie. Comme les propositions A, B et la liste de propositions Γ peuvent être quelconques, nous pouvons noter comme schéma de déduction (dérivé) de logique propositionnelle :

$$\frac{}{\Gamma \vdash (A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA})}.$$

Exercice. Écrire la séquence de jugements complets $\Gamma_1 \vdash A_1, \dots, \Gamma_{11} \vdash A_{11}$ qui constitue le corps de la démonstration de la figure 6.9 (cf. figure 6.8 (a)), et construire l'arbre de cette démonstration (cf. figure 6.8 (b)).

Omission des déductions a fortiori. Nous voulons introduire, avec l'exemple qui précède, une simplification de la notation des démonstrations naturelles que

Γ		
1°	$A \Rightarrow B$	hyp
2°	nonB	hyp
3°	$A \Rightarrow B$	1°, a fortiori
4°	A	hyp
5°	$A \Rightarrow B$	3°, a fortiori
6°	B	4°, 5°, mod ponens
7°	nonB	2°, a fortiori
8°	$\perp$	6°, 7°, $\perp$ -intro
9°	nonA	8°, red abs J
10°	$\text{nonB} \Rightarrow \text{nonA}$	9°, $\Rightarrow$ -intro
11°	$(A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA})$	10°, $\Rightarrow$ -intro.

Γ		
1°	$A \Rightarrow B$	hyp
2°	nonB	hyp
3°	A	hyp
4°	B	3°, 1°, mod ponens
5°	$\perp$	4°, 2°, $\perp$ -intro
6°	nonA	5°, red abs J
7°	$\text{nonB} \Rightarrow \text{nonA}$	6°, $\Rightarrow$ -intro
8°	$(A \Rightarrow B) \Rightarrow (\text{nonB} \Rightarrow \text{nonA})$	7°, $\Rightarrow$ -intro.

Figure 6.10. Simplification de la notation d'une démonstration naturelle: omission des déductions *a fortiori*.

nous pratiquerons régulièrement dans la suite. Cette simplification consiste à omettre les lignes d'une démonstration déduites suivant *a fortiori* (S2). La démonstration de la figure 6.9 est recopiée dans la figure 6.10, et les trois lignes déduites suivant ce schéma y sont mises en évidence. Ce schéma permet, comme nous l'avons vu (Figure 6.7), de recopier dans un cadre une proposition figurant plus haut dans le cadre environnant. Dans cet exemple, la proposition $A \Rightarrow B$ de la ligne 1° est recopiée successivement du deuxième cadre dans le troisième (ligne 3°) et du troisième dans le quatrième (ligne 5°). Ces copies ont été faites uniquement pour pouvoir appliquer le schéma *modus ponens* aux propositions 4° et 5°, et en tirer B (6°). La simplification consistera à omettre ces copies, qui resteront *sous-entendues*, et au lieu de se référer à une copie, de se référer à l'original: la référence à 5° dans la justification de la ligne 6° sera remplacée par une référence à 1°. On omettra de même de recopier la proposition nonB de la ligne 2° dans le quatrième cadre (7°), et à la ligne 8°, on se référera à 2° au lieu de 7°. La démonstration simplifiée, renumérotée, est

Γ

$\vdots$
$(A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)$

(a) On écrit la proposition à démontrer au bas du cadre extérieur (Γ).

 Γ

$A \Rightarrow B$	hyp
$\vdots$	
$\text{non}B \Rightarrow \text{non}A$	
$(A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)$	$\Rightarrow$ -intro

(b) Ce but est une implication. Pour la démontrer, on choisit la méthode *introduction de $\Rightarrow$* . On fait l'hypothèse $A \Rightarrow B$ et l'on se fixe comme but à atteindre, dans le nouveau cadre, la proposition $\text{non}B \Rightarrow \text{non}A$, que l'on note au bas de ce cadre.

 Γ

$A \Rightarrow B$	hyp
$\text{non}B$	hyp
$\vdots$	
$\text{non}A$	
$\text{non}B \Rightarrow \text{non}A$	$\Rightarrow$ -intro
$(A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)$	$\Rightarrow$ -intro

(c) Ce nouveau but est aussi une implication, et l'on choisit la même méthode pour la démontrer : on fait l'hypothèse $\text{non}B$ et l'on se fixe comme nouveau but, dans le cadre correspondant, la proposition $\text{non}A$.

 Γ

$A \Rightarrow B$	hyp
$\text{non}B$	hyp
A	hyp
$\vdots$	
$\perp$	
$\text{non}A$	red abs J
$\text{non}B \Rightarrow \text{non}A$	$\Rightarrow$ -intro
$(A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)$	$\Rightarrow$ -intro

(d) On choisit pour démontrer $\text{non}A$ la méthode de *réduction à l'absurde J*. Parmi les schémas de déduction dont nous disposons actuellement (Fig. 6.1), celui de *réduction à l'absurde J* se signale de lui-même lorsqu'il s'agit de démontrer une négation $\text{non}A$. On fait l'hypothèse A et l'on se fixe comme but à atteindre sous cette hypothèse la proposition $\perp$. La manière d'y parvenir saute aux yeux (Fig. 6.10), et l'on peut achever la démonstration.

Figure 6.11 Construction remontante d'une démonstration.

présentée aussi dans la figure 6.10.

Construction remontante de démonstrations. Une démonstration achevée se lit de haut en bas, dans l'ordre de la numérotation de ses lignes. Mais la *construction* de la démonstration s'effectue très souvent partiellement en sens inverse, c'est-à-dire en partant de la dernière ligne. Cette ligne est connue au départ — c'est le théorème ou la conclusion de la déduction à démontrer — et constitue le *but* à atteindre. La méthode que nous appelons *remontante* consiste à partir de ce but et à se fixer des buts intermédiaires à partir desquels, si on les atteint, on est sûr de pouvoir atteindre le but final. Cette démarche est illustrée par la figure 6.11, montrant les quatre premières étapes de la construction de

la démonstration (simplifiée) de la figure 6.10.

6.4 La loi du tiers exclu et la logique intuitionniste

Il est communément admis qu'une proposition de la forme $(A \text{ ou } \text{non}A)$ est vraie dans n'importe quel contexte. Par exemple :

$$\begin{aligned} & x = y \text{ ou } x \neq y \\ & \text{Socrate est mortel ou Socrate n'est pas mortel} \\ & \exists x(x^2 = a) \text{ ou } \text{non}\exists x(x^2 = a). \end{aligned}$$

C'est un fait que si A est une proposition quelconque d'un langage $\mathcal{L}$, et si Γ est une liste quelconque de propositions de $\mathcal{L}$, le jugement $\Gamma \vdash A \text{ ou } \text{non}A$ est un théorème de la théorie $\mathbf{Lprop}(\mathcal{L})$, comme on le voit d'après la démonstration (de base vide) donnée dans la figure 6.12.

Γ		
1°	$\text{non}(A \text{ ou } \text{non}A)$	hyp
2°	A	hyp
3°	$A \text{ ou } \text{non}A$	2°, ou-intro
4°	$\perp$	1°, 3°, $\perp$ -intro
5°	$\text{non}A$	4°, red abs J
6°	$\text{non}A$	hyp
7°	$A \text{ ou } \text{non}A$	6°, ou-intro
8°	$\perp$	1°, 7°, $\perp$ -intro
9°	$\text{nonnon}A$	8°, red abs J
10°	$\perp$	5°, 9°, $\perp$ -intro
11°	$A \text{ ou } \text{non}A$	10°, red abs K

Figure 6.12. Démonstration de la loi du tiers exclu.

Les déductions de la forme

$$\frac{}{\Gamma \vdash A \text{ ou } \text{non}A} \quad (1)$$

appartiennent donc à la logique propositionnelle classique. Le schéma de déduction (1) est appelé *principe* ou *loi* du *tiers exclu*. Le mot « tiers » est le substantif signifiant : une troisième chose, une troisième personne, une troisième possibilité. Par exemple, pour un nombre x , il y a deux possibilités : $x = 0$ ou $x \neq 0$; une troisième possibilité est *exclue*. En anglais, le schéma est appelé *law of excluded middle*, ce qui peut se traduire par « il n'y a pas de milieu ». En allemand, on utilise généralement le nom latin *Tertium non datur* — littéralement « un tiers n'est pas donné ».

Nous verrons plus loin que ce schéma est rejeté par une école de logiciens et de mathématiciens, qui n'utilise qu'une partie de la logique classique appelée la *logique intuitionniste*.

Tiers exclu et réduction à l'absurde K. La démonstration du schéma du tiers exclu qui est donnée dans la figure 6.12 utilise cinq de nos schémas de déductions élémentaires de la logique propositionnelle classique: vérité d'une hypothèse (S1), introduction de ou (S5), introduction de $\perp$ (S9), réduction à l'absurde J (S11), et réduction à l'absurde K (S12).

L'un au moins de ces cinq schémas doit être éliminé du répertoire de schémas élémentaires si l'on rejette le schéma du tiers exclu, afin que la démonstration de la figure 6.12 ne soit plus possible. Le choix qui est fait par la logique dite *intuitionniste* (J) est la suppression du schéma de réduction à l'absurde K (classique).

Le schéma du tiers exclu et celui de réduction à l'absurde K sont logiquement liés l'un à l'autre, en ce sens que l'on peut choisir indifféremment l'un ou l'autre comme schéma élémentaire de la logique propositionnelle classique. Nous avons choisi comme schéma élémentaire (S12) celui de la réduction à l'absurde K. D'autres auteurs prennent à sa place le schéma du tiers exclu. Dans notre cas, le schéma du tiers exclu s'obtient comme schéma *dérivé*, au moyen de la démonstration de la figure 6.12. Ceux qui prennent le schéma du tiers exclu comme schéma élémentaire, obtiennent le schéma de réduction à l'absurde K comme schéma dérivé, d'après l'une ou l'autre des démonstrations de la figure 6.13.

		Γ
1° $\Gamma; \text{non}A \vdash \perp$	<i>base</i>	1° $A \text{ ou non}A$ tiers exclu
2° $\Gamma \vdash A \text{ ou non}A$	tiers exclu	2° A hyp
3° $\Gamma; A \vdash A$	hyp	3° $A \Rightarrow A$ 2°, $\Rightarrow$ -intro
4° $\Gamma; \text{non}A \vdash A$	1°, ex falso	4° $\perp$ hyp
5° $\Gamma \vdash A$	2°, 3°, 4°, disj cas.	5° A base
		6° A 4°, ex falso
		6° A 1°, 2°, 5° disj cas

(a)

(b)

Figure 6.13 Démonstrations du schéma de réduction à l'absurde K en admettant celui du tiers exclu: (a) non naturelle; (b) naturelle.

On omet généralement dans (b) la ligne non numérotée.

Si A , B sont des propositions quelconques d'un langage $\mathcal{L}$ et si Γ est une liste quelconque de propositions de $\mathcal{L}$, les deux séquences de jugements de la figure 6.13 sont des démonstrations de la théorie $\mathbf{Lprop}(\mathcal{L})$, lorsqu'on remplace son schéma élémentaire S12 par celui du tiers exclu. La base de ces deux démonstrations se compose du seul jugement $\Gamma; \text{non}A \vdash \perp$ et dans le corps de

chacune d'elles se trouve le jugement $\Gamma \vdash A$. La séquence de jugements

$$\frac{\Gamma; \text{non}A \vdash \perp}{\Gamma \vdash A}$$

est donc une déduction de $\text{Lprop}(\mathcal{L})$ (§5.4, Déf. 2, lettre (c)).

Exercices. (a) Expliquer pourquoi la première démonstration de la figure 6.13 n'est pas de forme naturelle. (b) Dessiner l'arbre de cette démonstration. (c) Écrire la séquence des jugements complets de la deuxième démonstration. Observer que certains jugements ne sont présents dans cette séquence que pour satisfaire aux conditions d'une démonstration naturelle, et ne sont utilisés dans aucune déduction. (d) Montrer que cette démonstration ne peut pas être représentée complètement par un seul arbre de démonstration, mais qu'il faut une « forêt » de trois arbres.

Réduction à l'absurde K et nonnon-élimination. Un autre schéma de déduction qui peut être pris à la place de celui de la réduction à l'absurde K dans le répertoire des schémas de déduction élémentaires de la logique propositionnelle classique est le schéma *d'élimination d'une double négation*:

$$\frac{\Gamma \vdash \text{nonnon}A}{\Gamma \vdash A.} \quad (2)$$

Nous le désignerons de façon abrégée par *nonnon-élimination*. Les démonstrations de la figure 6.14 montrent: (a) que le schéma d'élimination (2) s'obtient comme schéma dérivé lorsque celui de la réduction à l'absurde K est pris comme schéma élémentaire, et (b) que le schéma de réduction à l'absurde K s'obtient inversement comme schéma dérivé lorsqu'on prend (2) comme schéma élémentaire.

Γ		Γ	
1°	$\text{nonnon}A \quad \text{base}$	1°	$\text{non}A \quad \text{hyp}$
2°	$\text{non}A \quad \text{hyp}$	2°	$\perp \quad \text{base}$
3°	$\perp \quad 1^\circ, 2^\circ, \perp\text{-intro}$	3°	$\text{nonnon}A \quad 2^\circ, \text{red abs J}$
4°	$A \quad 3^\circ, \text{red abs K}$	4°	$A \quad \text{nonnon-élim}$
(a)		(b)	

Figure 6.14 (a) De $\Gamma \vdash \text{nonnon}A$ on déduit $\Gamma \vdash A$ par réduction à l'absurde K. (b) De $\Gamma; \text{non}A \vdash \perp$ on déduit $\Gamma \vdash A$ par réduction à l'absurde J et nonnon-élimination.

La logique intuitionniste. La logique appelée *intuitionniste* rejette le schéma du tiers exclu, et rejette donc le schéma élémentaire *réduction à l'absurde K* permettant de le démontrer. Tous les autres schémas de notre répertoire (figure 6.1) sont conservés. Ce que l'école intuitionniste critique, dans le schéma

du tiers exclu, est qu'il permet de démontrer des théorèmes *d'existence*, c'est-à-dire des théorèmes dont l'assertion est de la forme

il existe un x tel que $\dots$,

affirmant l'existence d'un objet possédant une certaine propriété, sans avoir besoin pour cela de *trouver* un objet particulier possédant la propriété en question. L'exemple qui suit est l'un des plus simples, et il est traditionnellement utilisé pour faire comprendre ce qui gêne les intuitionnistes. Mais la démonstration que nous allons présenter ne sera pas entièrement formelle, car nous ne disposons pas encore des schémas de déduction de la logique des prédicats, relatifs aux quantificateurs existentiels.

Exemple. On rappelle qu'un nombre réel x est dit rationnel s'il existe deux nombres entiers p et q tels que $x = p/q$. Sinon, x est dit irrationnel. Autrement dit, la proposition « x est irrationnel» est définie comme une manière d'exprimer la négation: $\text{non}(x \text{ est rationnel})$. Donc, suivant la logique classique (tiers exclu), la proposition

x est rationnel ou x est irrationnel

est vraie. En nous basant sur cela, nous allons démontrer le théorème suivant de la théorie des nombres réels:

$$\vdash \text{il existe deux nombres irrationnels } a \text{ et } b \quad (3)$$

tels que a^b soit *rationnel*.

Nous donnons d'abord cette démonstration de manière informelle, en quelques phrases. La démonstration formelle correspondante se trouve dans la figure 6.15.

On sait que le nombre $\sqrt{2}$ est irrationnel. Considérons le nombre $\sqrt{2}^{\sqrt{2}}$. Ce nombre est rationnel ou irrationnel (tiers exclu). S'il est rationnel, l'assertion (3) est vérifiée par $a = b = \sqrt{2}$. Sinon, elle est vérifiée par $a = \sqrt{2}^{\sqrt{2}}$ et $b = \sqrt{2}$, car ces égalités entraînent $a^b = 2$ (rationnel).

Pour faire cette démonstration, il n'y a pas besoin de *savoir* si le nombre $\sqrt{2}^{\sqrt{2}}$ est rationnel ou non. Il suffit de savoir qu'il est l'un ou l'autre. Personne ne peut s'empêcher d'éprouver une certaine «frustration» à la lecture de cette démonstration. Car on s'attend normalement à *connaître*, après une démonstration de (3), un couple de nombres irrationnels a, b tel que a^b soit rationnel. Or après *cette* démonstration, nous n'en connaissons toujours pas. Tout ce que nous pouvons dire est que l'un des deux couples de nombres

$$\begin{aligned} a &= \sqrt{2}, & b &= \sqrt{2} \\ a &= \sqrt{2}^{\sqrt{2}}, & b &= \sqrt{2} \end{aligned}$$

a cette propriété, mais nous ne savons pas lequel.

La démonstration de la figure 6.15 n'est évidemment qu'à moitié formelle. Les lignes 3^o et 7^o notamment nécessitent des déductions qui relèvent de la

Γ		
1°	$\sqrt{2}\sqrt{2}$ est rationnel ou $\sqrt{2}\sqrt{2}$ est irrationnel	tiers exclu
2°	$\sqrt{2}\sqrt{2}$ est rationnel	hyp
3°	il existe deux nombres irrationnels a, b tels que a^b soit rationnel	2°, $a=\sqrt{2}, b=\sqrt{2}$
4°	$\sqrt{2}\sqrt{2}$ est irrationnel	hyp
5°	$(\sqrt{2}\sqrt{2})^{\sqrt{2}} = 2$	calcul
6°	$(\sqrt{2}\sqrt{2})^{\sqrt{2}}$ est rationnel	5°
7°	il existe deux nombres irrationnels a, b tels que a^b soit rationnel	6°, $a=(\sqrt{2}\sqrt{2}), b=\sqrt{2}$
8°	il existe deux nombres irrationnels a, b tels que a^b soit rationnel	1°, 3°, 7°, disj cas.

Figure 6.15 Démonstration classique non agréée par la logique intuitionniste. Le contexte est celui de la théorie des nombres réels. La liste d'hypothèses Γ est vide.

logique des prédicats (chapitre 9). Elles ne sont justifiées ici qu'intuitivement. Mais la démonstration met bien en évidence les déductions de logique propositionnelle qui sont faites dans ce raisonnement : tiers exclu et disjonction des cas.

Les mathématiques dites *constructives* exigent que toute démonstration d'existence fournisse un exemple « concret » de ce dont on affirme l'existence. C'est ce qu'on appelle une preuve d'existence *constructive*. Elles se limitent donc à l'emploi de la logique intuitionniste qui ne permet pas les raisonnements tels que celui que nous venons de faire. Cette preuve est non constructive. Le schéma de *réduction à l'absurde* K lui-même permet aussi des preuves d'existence non constructives, de la forme

Γ		
1°	non $\exists xA$	hyp
$\vdots$	$\vdots$	
n°	$\perp$	...
	$\exists xA$	n° , red abs K

où l'on montre seulement que l'hypothèse de la non-existence entraîne une contradiction.

Le théorème (3) n'est pas des plus importants. Mais les mathématiques classiques ont accumulé depuis un siècle de nombreux théorèmes d'existence importants en se contentant de démonstrations non constructives. Nous y avons fait allusion au paragraphe 3.7 à propos de l'expression « on peut trouver un x tel que A » que certains utilisent souvent à la place de « il existe un x tel

que A ». Les mathématiques classiques ont démontré l'existence de choses que personne n'a jamais pu trouver jusqu'ici.

La critique des démonstrations non constructives de la logique classique date d'une centaine d'années environ. En 1890, la démonstration par réduction à l'absurde K d'un célèbre théorème d'existence, par le mathématicien David Hilbert, suscita l'exclamation d'un confrère: – ce n'est pas de la mathématique, c'est de la théologie!² Mais ce n'est que depuis quelques dizaines d'années qu'une école constructiviste s'est donné pour tâche d'explorer systématiquement ce qui subsiste des mathématiques classiques lorsqu'on se limite aux méthodes constructives. Il s'est avéré que ce « sous-ensemble » des mathématiques classiques est loin d'être négligeable³.

La logique classique reste cependant l'instrument principal de la déduction mathématique. Son étude est prioritaire, et il ne sera plus question de logique intuitionniste dans le reste de cet ouvrage.

6.5 Dédutions structurelles

Certaines séquences de jugements d'un langage $\mathcal{L}$ sont appelées des *dédutions structurelles* de la théorie $\mathbf{Lprop}(\mathcal{L})$. Ces déductions n'ont pas de rapport avec un connecteur logique particulier, mais consistent uniquement en des manipulations de listes d'hypothèses: adjonction, permutation d'hypothèses, etc.

On classe parmi les déductions structurelles de $\mathbf{Lprop}(\mathcal{L})$ les déductions élémentaires de schémas $S1$ et $S2$:

$$\frac{}{\Gamma; A \vdash A} \text{hyp} \qquad \frac{\Gamma \vdash B}{\Gamma; A \vdash B.} \text{a fortiori} \quad (1)$$

On dira que $S1$ et $S2$ sont des schémas de déduction *structurels*. Le but de ce paragraphe est d'établir quelques autres schémas structurels, dérivés. Mais les démonstrations que nous en donnerons contiendront aussi d'autres déductions élémentaires que les déductions structurelles de la forme (1).

Le schéma d'introduction de $\Rightarrow$ généralisé et son inverse. Nous rappelons d'abord le schéma d'introduction de $\Rightarrow$ ($S7$) et nous notons à côté de lui son inverse, que nous appelons *élimination de $\Rightarrow$* :

$$\frac{\Gamma; A \vdash B}{\Gamma \vdash A \Rightarrow B} \Rightarrow\text{-intro} \qquad \frac{\Gamma \vdash A \Rightarrow B}{\Gamma; A \vdash B.} \Rightarrow\text{-élim} \quad (2)$$

Certains auteurs appellent « élimination de $\Rightarrow$ » le schéma *modus ponens* ($S8$). Les schémas (2) sont repris dans la figure 6.16, où ils sont munis de l'indice 1, car chacun d'eux est le premier schéma d'une famille de schémas de déduction,

² Hilbert répliqua plus tard aux intuitionnistes que priver un mathématicien de la loi du tiers exclu était comme interdire à un boxeur l'usage de ses poings!

³ Voir par exemple le livre d'Errett Bishop et Douglas Bridges, *Constructive Analysis*, Springer-Verlag, 1985.

$\frac{\Gamma; A_1 \vdash B}{\Gamma \vdash A_1 \Rightarrow B} \Rightarrow\text{-intro}_1$		$\frac{\Gamma \vdash A_1 \Rightarrow B}{\Gamma; A_1 \vdash B} \Rightarrow\text{-élim}_1$	
$\frac{\Gamma; A_1; A_2 \vdash B}{\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow B)} \Rightarrow\text{-intro}_2$		$\frac{\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow B)}{\Gamma; A_1; A_2 \vdash B} \Rightarrow\text{-élim}_2$	
$\frac{\Gamma; A_1; A_2; A_3 \vdash B}{\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B))} \Rightarrow\text{-intro}_3$		$\frac{\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B))}{\Gamma; A_1; A_2; A_3 \vdash B} \Rightarrow\text{-élim}_3$	
$\Rightarrow\text{-intro}_2$		$base$	
1° $\Gamma; A_1; A_2 \vdash B$		1°, $\Rightarrow\text{-intro}$	
2° $\Gamma; A_1 \vdash A_2 \Rightarrow B$		2°, $\Rightarrow\text{-intro}$	
3° $\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow B)$			
$\Rightarrow\text{-intro}_3$		$base$	
1° $\Gamma; A_1; A_2; A_3 \vdash B$		1°, $\Rightarrow\text{-intro}_2$	
2° $\Gamma; A_1 \vdash A_2 \Rightarrow (A_3 \Rightarrow B)$		2°, $\Rightarrow\text{-intro}$	
3° $\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B))$			
$\Rightarrow\text{-élim}_1$		$base$	
1° $\Gamma \vdash A_1 \Rightarrow B$		hyp	
2° $\Gamma; A_1 \vdash A_1$		1°, a fortiori	
3° $\Gamma; A_1 \vdash A_1 \Rightarrow B$		2°, 3°, mod ponens.	
4° $\Gamma; A_1 \vdash B$			
$\Rightarrow\text{-élim}_2$		$base$	
1° $\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow B)$		1°, $\Rightarrow\text{-élim}$	
2° $\Gamma; A_1 \vdash A_2 \Rightarrow B$		2°, $\Rightarrow\text{-élim}$	
3° $\Gamma; A_1; A_2 \vdash B$			
$\Rightarrow\text{-élim}_3$		$base$	
1° $\Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B))$		2°, $\Rightarrow\text{-élim}_2$	
2° $\Gamma; A_1; A_2 \vdash A_3 \Rightarrow B$		3°, $\Rightarrow\text{-élim}$	
3° $\Gamma; A_1; A_2; A_3 \vdash B$			

Figure 6.16 Schémas de déduction $\Rightarrow\text{-intro}_{1-3}$, $\Rightarrow\text{-élim}_{1-3}$.

dont la figure contient les trois premiers. Tous, sauf $\Rightarrow\text{-intro}_1$, sont des schémas dérivés de $\mathbf{Lprop}(\mathcal{L})$ et la figure contient les démonstrations qui le prouvent.

Pour rassembler tous ces schémas en un seul, on convient que la notation sans parenthèses $A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B$ représente

pour $n = 1$: $A_1 \Rightarrow B$
pour $n = 2$: $A_1 \Rightarrow (A_2 \Rightarrow B)$
pour $n = 3$: $A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B))$
etc.

Cette convention est appelée la convention *d'association à droite* pour l'implication. Avec celle-ci on peut représenter la classe des déductions de schémas $\Rightarrow\text{-intro}_1$, $\Rightarrow\text{-intro}_2$, $\Rightarrow\text{-intro}_3$, etc. et celle des déductions de schémas $\Rightarrow\text{-élim}_1$, $\Rightarrow\text{-élim}_2$, $\Rightarrow\text{-élim}_3$, etc. par les deux schémas suivants, que nous appellerons

les schémas d'introduction et d'élimination de $\Rightarrow$ *généralisés*:

$$\frac{\Gamma; A_1; \dots; A_n \vdash B}{\Gamma \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B} \quad \frac{\Gamma \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B}{\Gamma; A_1; \dots; A_n \vdash B.} \quad (3)$$

Nous pouvons même étendre la classe de déductions représentée par ces schémas, en y incluant les *dédutions identiques* (§5.4):

$$\frac{\Gamma \vdash B}{\Gamma \vdash B.} \quad (4)$$

Ces déductions appartiennent à la classe décrite par les schémas (3), si nous convenons que pour $n = 0$, la notation $A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B$ représente B .

Les schémas de déduction (3) de $\mathbf{Lprop}(\mathcal{L})$ permettent de dire que deux jugements de la forme

$$\Gamma; A_1; \dots; A_n \vdash B, \quad \Gamma \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B$$

sont *équivalents* (§5.4) dans $\mathbf{Lprop}(\mathcal{L})$. Il en est ainsi en particulier de jugements de la forme

$$A_1; \dots; A_n \vdash B, \quad \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B.$$

On voit donc que dans une théorie logique tout jugement est équivalent à un jugement *sans hypothèses*.

Généralisations du schéma de déduction a fortiori. Nous allons décrire maintenant diverses classes de déductions structurelles, contenant la classe décrite par le schéma S2. La première de ces classes est celle des déductions de la forme

$$\frac{\Gamma \vdash B}{\Gamma; A_1; \dots; A_n \vdash B.} \quad (5)$$

La preuve qu'une telle séquence de jugements, pour un n donné, est bien une déduction de $\mathbf{Lprop}(\mathcal{L})$ est fournie par la démonstration

$$\begin{array}{ll} 1^\circ & \Gamma \vdash B \quad \text{base} \\ 2^\circ & \Gamma; A_1 \vdash B \quad 1^\circ, \text{ a fortiori} \\ 3^\circ & \Gamma; A_1; A_2 \vdash B \quad 2^\circ, \text{ a fortiori} \\ & \dots \\ (n+1)^\circ & \Gamma; A_1; A_2; \dots; A_n \vdash B \quad n^\circ, \text{ a fortiori.} \end{array}$$

Cette classe de déductions inclut les déductions identiques (4), si nous convenons que pour $n = 0$ la notation $\Gamma; A_1; \dots; A_n \vdash B$ représente $\Gamma \vdash B$.

Une sous-classe de la classe des déductions de la forme (5) est celle des déductions de la forme

$$\frac{\vdash B}{A_1; \dots; A_n \vdash B.} \quad (6)$$

Ce sont les déductions de la forme (5) dans lesquelles la liste de propositions Γ est vide. Comme la liste $A_1; \dots; A_n$ peut être quelconque, y compris vide

($n = 0$), nous pouvons décrire plus simplement la classe des déductions de la forme (6) comme étant celle des déductions de la forme

$$\frac{\vdash B}{\Gamma \vdash B.} \quad (7)$$

Il suffit de redésigner par Γ la liste arbitraire $A_1; \dots; A_n$ de (6).

Le schéma de déduction (5) peut être interprété en disant que l'on peut allonger arbitrairement à droite la liste d'hypothèses d'un jugement. C'est évidemment une manière abrégée de dire que l'on peut déduire d'un jugement J le jugement obtenu en allongeant la liste d'hypothèses de J . Il est aussi possible d'allonger cette liste à gauche. On peut voir en effet que les séquences de jugements de la forme

$$\frac{A_1; \dots; A_n \vdash B}{\Gamma; A_1; \dots; A_n \vdash B} \quad (8)$$

sont des déductions de $\mathbf{Lprop}(\mathcal{L})$. La preuve de ce fait, pour un n donné, est fournie par la démonstration suivante:

$$\begin{array}{ll} 1^\circ & A_1; \dots; A_n \vdash B \quad \text{base} \\ 2^\circ & \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B \quad 1^\circ, \Rightarrow\text{-intro}_n \\ 3^\circ & \Gamma \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow B \quad 2^\circ, \text{a fortiori gén (7)} \\ 4^\circ & \Gamma; A_1; \dots; A_n \vdash B \quad \Rightarrow\text{-élim}_n. \end{array}$$

Dans (8), la liste d'hypothèses de la conclusion est celle de la prémisse allongée arbitrairement à gauche. En changeant de notation, c'est-à-dire en redésignant par Γ la liste arbitraire $A_1; \dots; A_n$ de (8) et par $A_1; \dots; A_m$ la liste arbitraire Γ de (8), nous pouvons décrire la même classe de déductions par le schéma

$$\frac{\Gamma \vdash B}{A_1; \dots; A_m; \Gamma \vdash B.} \quad (9)$$

Finalement, la liste d'hypothèses d'un jugement peut être allongée à gauche et à droite, c'est-à-dire que les séquences de jugements de la forme

$$\frac{\Gamma \vdash C}{A_1; \dots; A_m; \Gamma; B_1; \dots; B_n \vdash C} \quad (10)$$

sont encore des déductions de $\mathbf{Lprop}(\mathcal{L})$. Une telle déduction admet en effet la démonstration

$$\begin{array}{ll} 1^\circ & \Gamma \vdash C \quad \text{base} \\ 2^\circ & A_1; \dots; A_m; \Gamma \vdash C \quad 1^\circ, \text{a fortiori gén (9)} \\ 3^\circ & A_1; \dots; A_m; \Gamma; B_1; \dots; B_n \vdash C \quad 2^\circ, \text{a fortiori gén (5)}. \end{array}$$

Tautologies. On appelle *tautologie* d'un langage $\mathcal{L}$ toute proposition B de $\mathcal{L}$ telle que le jugement $\vdash B$ soit un théorème de la théorie $\mathbf{Lprop}(\mathcal{L})$. Pour toute liste Γ de propositions de $\mathcal{L}$, le jugement $\Gamma \vdash B$ est alors aussi un théorème de $\mathbf{Lprop}(\mathcal{L})$, qui se déduit du premier suivant (7).

Les schémas de déduction dérivés sans prémisses de la logique propositionnelle, dont nous verrons de nombreux exemples au chapitre 7, décrivent la forme de tautologies. Nous avons déjà vu par exemple (§6.2, §6.3 Exemple 3, §6.4) les schémas sans prémisses

$$\frac{}{\Gamma \vdash A \Rightarrow \text{nonnon}A}$$

$$\frac{}{\Gamma \vdash (A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)} \quad \frac{}{\Gamma \vdash A \text{ ou } \text{non}A}.$$

On peut exprimer ces schémas de déduction en disant que les propositions de la forme $A \Rightarrow \text{nonnon}A$, $(A \Rightarrow B) \Rightarrow (\text{non}B \Rightarrow \text{non}A)$, A ou $\text{non}A$ d'un langage $\mathcal{L}$ sont des tautologies de $\mathcal{L}$.

Schéma de vérité d'une hypothèse généralisé. En appliquant le schéma de déduction (10), on peut montrer que les déductions sans prémisses de la forme suivante sont des déductions de $\text{Lprop}(\mathcal{L})$:

$$\frac{}{A_1; \dots; A_i; \dots; A_n \vdash A_i.} \quad (11)$$

Nous nous référerons à ce schéma par l'abréviation *hyp gén*, car il s'agit d'une généralisation du schéma *vérité d'une hypothèse* (S1). Une déduction de cette forme admet la démonstration

$$\begin{array}{ll} 1^\circ & A_i \vdash A_i \quad \text{hyp (S1 avec } \Gamma \text{ vide)} \\ 2^\circ & A_1; \dots; A_{i-1}; A_i; A_{i+1}; \dots; A_n \vdash A_i \quad 1^\circ, \text{ a fortiori gén (10).} \end{array}$$

Permutations et réductions dans une liste d'hypothèses. Les séquences de jugements de la forme

$$\frac{\Gamma; A; B \vdash C}{\Gamma; B; A \vdash C} \quad (12)$$

sont des déductions de $\text{Lprop}(\mathcal{L})$. Elles admettent en effet la démonstration

$$\begin{array}{ll} 1^\circ & \Gamma; A; B \vdash C \quad \text{base} \\ 2^\circ & \Gamma \vdash A \Rightarrow (B \Rightarrow C) \quad 1^\circ, \Rightarrow\text{-intro}_2 \\ 3^\circ & \Gamma; B; A \vdash A \Rightarrow (B \Rightarrow C) \quad 2^\circ, \text{ a fortiori gén (5)} \\ 4^\circ & \Gamma; B; A \vdash A \quad \text{hyp} \\ 5^\circ & \Gamma; B; A \vdash B \Rightarrow C \quad 4^\circ, 3^\circ, \text{ mod ponens} \\ 6^\circ & \Gamma; B; A \vdash B \quad \text{hyp gén (11)} \\ 7^\circ & \Gamma; B; A \vdash C \quad 6^\circ, 5^\circ, \text{ mod ponens.} \end{array}$$

On peut donc permuter les deux dernières hypothèses d'un jugement. Mais on peut aussi permuter deux hypothèses consécutives quelconques de la liste. Les séquences de jugements de la forme

$$\frac{A_1; \dots; A_i; A_{i+1}; \dots; A_n \vdash B}{A_1; \dots; A_{i+1}; A_i; \dots; A_n \vdash B} \quad (13)$$

sont en effet des déductions de $\mathbf{Lprop}(\mathcal{L})$. Leur démonstration est laissée au lecteur.

La liste d'hypothèses d'un jugement peut contenir des répétitions, c'est-à-dire que la même proposition peut figurer plusieurs fois dans cette liste. Les déductions de la forme

$$\frac{\Gamma; A; A \vdash B}{\Gamma; A \vdash B} \quad (14)$$

permettent d'éliminer de telles répétitions — c'est ce que nous appelons une *réduction* de la liste. Les séquences de jugements de la forme (14) sont des déductions de $\mathbf{Lprop}(\mathcal{L})$ car elles admettent la démonstration

$$\begin{array}{ll} 1^\circ & \Gamma; A; A \vdash B \quad \text{base} \\ 2^\circ & \Gamma; A \vdash A \quad \text{hyp} \\ 3^\circ & \Gamma; A \vdash A \Rightarrow B \quad 1^\circ, \Rightarrow\text{-intro} \\ 4^\circ & \Gamma; A \vdash B \quad 2^\circ, 3^\circ \text{ mod ponens.} \end{array}$$

Le schéma (14) permet d'éliminer une répétition d'hypothèse en *fin de liste*. Toute autre forme de répétition dans la liste peut être éliminée ainsi, en effectuant d'abord des permutations qui amènent les propositions répétées en fin de liste, autrement dit en combinant des déductions de la forme (13) et de la forme (14).

Transitivité des jugements. Les séquences de jugements de la forme

$$\frac{\Gamma \vdash A_1 \quad \dots \quad \Gamma \vdash A_n \quad A_1; \dots; A_n \vdash B}{\Gamma \vdash B} \quad (15)$$

sont des déductions de $\mathbf{Lprop}(\mathcal{L})$. Pour $n = 3$, par exemple, une telle déduction admet la démonstration suivante :

$$\begin{array}{ll} 1^\circ & \Gamma \vdash A_1 \quad \text{base} \\ 2^\circ & \Gamma \vdash A_2 \quad \text{base} \\ 3^\circ & \Gamma \vdash A_3 \quad \text{base} \\ 4^\circ & A_1; A_2; A_3 \vdash B \quad \text{base} \\ 5^\circ & \Gamma; A_1; A_2; A_3 \vdash B \quad 4^\circ, \text{ a fortiori gén (8)} \\ 6^\circ & \Gamma \vdash A_1 \Rightarrow (A_2 \Rightarrow (A_3 \Rightarrow B)) \quad 5^\circ, \Rightarrow\text{-intro}_3 \\ 7^\circ & \Gamma \vdash A_2 \Rightarrow (A_3 \Rightarrow B) \quad 1^\circ, 6^\circ, \text{ mod ponens} \\ 8^\circ & \Gamma \vdash A_3 \Rightarrow B \quad 2^\circ, 7^\circ, \text{ mod ponens} \\ 9^\circ & \Gamma \vdash B \quad 3^\circ, 8^\circ, \text{ mod ponens.} \end{array}$$

Applications. Pour démontrer une déduction

$$\frac{A_1; \dots; A_n \vdash B}{\Gamma \vdash B} \quad (16)$$

par transitivité des jugements, il suffit de démontrer chacun des jugements $\Gamma \vdash A_i$ ($i = 1, \dots, n$), autrement dit de démontrer que *chacune des hypothèses*

de la prémisse est vraie sous les hypothèses de la conclusion. Car des n jugements $\Gamma \vdash A_i$ et de la prémisse de (16) on déduit $\Gamma \vdash B$ par transitivité des jugements (15).

Par exemple, d'après le schéma de transitivité des jugements (15), il est évident que les séquences de jugements de la forme

$$\frac{\Gamma; A; B \vdash C}{\Gamma; A \text{ et } B \vdash C} \qquad \frac{\Gamma; A \text{ et } B \vdash C}{\Gamma; A; B \vdash C.} \quad (17)$$

sont des déductions de $\mathbf{Lprop}(\mathcal{L})$, car chaque proposition de la liste $\Gamma; A; B$ est vraie sous la liste d'hypothèses $\Gamma; A$ et B d'après le schéma (11) et le schéma d'élimination de et ; inversement, chaque proposition de la liste $\Gamma; A$ et B est vraie sous la liste d'hypothèses $\Gamma; A; B$ d'après (11) et introduction de et .

Les schémas (17) permettent de remplacer les deux dernières hypothèses d'un jugement par leur conjonction, et inversement, lorsque la dernière hypothèse d'un jugement est la conjonction de deux propositions, de remplacer cette hypothèse par ces deux propositions. Ces opérations peuvent être répétées, de sorte qu'une liste de n hypothèses $A_1; \dots; A_n$ peut être remplacée par la conjonction $A_1 \text{ et } \dots \text{ et } A_n$, si nous faisons pour la conjonction la même convention d'association à droite que pour l'implication, à savoir :

$A_1 \text{ et } \dots \text{ et } A_n \text{ et } B$ représente
pour $n = 0$: B
pour $n = 1$: $A_1 \text{ et } B$
pour $n = 2$: $A_1 \text{ et } (A_2 \text{ et } B)$
pour $n = 3$: $A_1 \text{ et } (A_2 \text{ et } (A_3 \text{ et } B))$
etc.

On peut donc faire dans $\mathbf{Lprop}(\mathcal{L})$ les déductions

$$\frac{A_1; \dots; A_n \vdash C}{A_1 \text{ et } \dots \text{ et } A_n \vdash C} \qquad \frac{A_1 \text{ et } \dots \text{ et } A_n \vdash C}{A_1; \dots; A_n \vdash C.}$$

En d'autres termes, deux jugements de la forme

$$A_1; \dots; A_n \vdash C, \qquad A_1 \text{ et } \dots \text{ et } A_n \vdash C$$

sont équivalents (§5.4) dans $\mathbf{Lprop}(\mathcal{L})$.

Exercice 1. Donner une démonstration d'une déduction de la forme (13), combinant des déductions de la forme $\Rightarrow\text{-intro}_k$, (12), $\Rightarrow\text{-élim}_k$.

Exercice 2. Montrer que deux jugements de la forme $\Gamma \vdash A_1 \Rightarrow \dots \Rightarrow A_n \Rightarrow C$, $\Gamma \vdash (A_1 \text{ et } \dots \text{ et } A_n) \Rightarrow C$ sont équivalents dans $\mathbf{Lprop}(\mathcal{L})$.

Exercice 3. (a) La plupart des démonstrations de ce paragraphe ne sont pas de forme naturelle. Transformer chacune d'elles en une démonstration naturelle en y adjoignant ou en y insérant seulement des lignes appropriées, qui sont soit des lignes d'hypothèse, soit des lignes déduites de précédentes par introduction ou élimination de $\Rightarrow$.

(b) Généraliser le traitement des exemples de l'exercice (a) en donnant une méthode systématique (algorithme) permettant de mettre toute démonstration d'une théorie logique $\mathcal{T}$ sous forme naturelle par des adjonctions et insertions de lignes supplémentaires.

Déductions dérivées de logique propositionnelle

7.1 Proposition fausse et contradictions

Les démonstrations du chapitre 6 ont permis au lecteur de se familiariser avec certains des schémas de déduction élémentaires S1–S14 de la logique propositionnelle (figure 6.1) et avec leur usage dans les démonstrations. Dans la pratique du raisonnement, on utilise couramment un assez grand nombre de schémas dérivés. Le but de ce chapitre est d'en présenter un répertoire suffisant, ainsi que certaines méthodes générales de démonstration.

Ce répertoire se compose des schémas S15 à S60 rassemblés dans le *Formulaire* annexé à cet ouvrage. C'est un *exercice indispensable* à la maîtrise de la logique en vue de son utilisation que de démontrer soi-même rigoureusement chacun de ces schémas. Le lecteur doit prendre cela comme *l'exercice principal* de ce chapitre. Les démonstrations données dans ce paragraphe et dans les suivants constituent la *solution* de cet exercice et le lecteur est invité à ne les lire que pour vérifier son propre travail.

La présentation des schémas est divisée en paragraphes correspondant à des groupes de schémas. Mais cette classification est en partie arbitraire et il ne faut pas lui attacher trop d'importance. Le présent paragraphe est consacré à quelques formes de déduction qui dérivent de manière très simple des schémas élémentaires relatifs à la proposition fausse $\perp$.

La proposition fausse $\perp$. Suivant notre définition du paragraphe 5.6, une proposition A est dite *fausse* dans un contexte logique lorsque sa négation, $\text{non}A$, est vraie dans ce contexte. Le schéma dérivé suivant justifie donc l'appellation de proposition fausse *par excellence* donnée à la proposition $\perp$, puisque sa négation, $\text{non}\perp$, est vraie dans tout contexte :

S15. $\perp$ est fausse

$$\frac{}{\Gamma \vdash \text{non}\perp}.$$

La démonstration naturelle de cette déduction sans prémisse est d'une simplicité qui peut être déroutante. Nous la présentons donc sous les deux formes : séquence de jugements complets et notation emboîtée. C'est cette dernière qui peut être déroutante. Le lecteur peut se reporter aux plans de démonstration de la figure 6.7 (§6.3), et voir ce que devient celui qui concerne le schéma de réduction à l'absurde J, lorsque la proposition A est la proposition $\perp$ elle-même. Il arrive ainsi qu'un bloc intérieur d'une démonstration naturelle se compose d'une seule ligne, contenant l'hypothèse de ce bloc.

			Γ				
1°	$\Gamma; \perp \vdash \perp$	hyp	1°				
2°	$\Gamma \vdash \text{non}\perp$	1°, red abs J.	2°				
			<table> <tr> <td>$\perp$</td> <td>hyp</td> </tr> <tr> <td>$\text{non}\perp$</td> <td>1°, red abs J.</td> </tr> </table>	$\perp$	hyp	$\text{non}\perp$	1°, red abs J.
$\perp$	hyp						
$\text{non}\perp$	1°, red abs J.						

Contradictions. Deux propositions sont dites *contradictoires* lorsque l'une d'elles est la négation de l'autre. Autrement dit ce sont deux propositions de la forme A , $\text{non}A$. Si deux propositions contradictoires sont vraies dans un contexte, on dit qu'il y a une contradiction dans celui-ci, ou encore que ce contexte est contradictoire (§5.6).

S16. *Tout se déduit d'une contradiction*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash \text{non}A}{\Gamma \vdash B.}$$

Ce schéma confirme ce que nous avons annoncé précédemment (§5.6) : dans un contexte contradictoire, c'est-à-dire dans lequel on peut démontrer une proposition A et sa négation, toute proposition est vraie. La démonstration de ces déductions dérivées a déjà été présentée (§6.2, Exemple 1). C'est une démonstration naturelle sans bloc intérieur, que nous reproduisons ici dans la notation emboîtée.

	Γ
1°	A base
2°	$\text{non}A$ base
3°	$\perp$ 1°, 2°, $\perp$ -intro
4°	B 3°, ex falso.

S17.

$$\frac{}{\Gamma \vdash \text{non}(A \text{ et } \text{non}A).}$$

Une relation de la forme $(A \text{ et } \text{non}A)$ est fausse dans tout contexte logique.

	Γ
1°	$A \text{ et } \text{non}A$ hyp
2°	A 1°, et-élim
3°	$\text{non}A$ 1°, et-élim
4°	$\perp$ 2°, 3°, $\perp$ -intro
5°	$\text{non}(A \text{ et } \text{non}A)$ 4°, red abs J.

La démonstration par réduction à l'absurde J est donnée ci-contre. Dorénavant, comme ici, le cadre extérieur de la notation emboîtée d'une démonstration ne sera plus marqué que par son bord gauche.

S18.

$$\frac{\Gamma \vdash A \text{ ou } B \quad \Gamma \vdash \text{non}A}{\Gamma \vdash B} \qquad \frac{\Gamma \vdash \text{non}A \text{ ou } B \quad \Gamma \vdash A}{\Gamma \vdash B.}$$

Une déduction fréquente, dans le raisonnement mathématique comme dans la vie courante. Sachant par exemple qu'il a plu hier ou avant-hier (ou les deux jours), et apprenant qu'il n'a pas plu avant-hier, on en déduit qu'il a plu hier. Cela fait partie du sens du mot « ou ». Cette déduction se démontre par disjonction des cas. C'est ici une application particulière du schéma **S6**, dans laquelle on prend pour C la proposition B elle-même. Le deuxième bloc intérieur de la démonstration n'a donc qu'une ligne. Le lecteur comparera cet exemple avec le plan général de la figure 6.7. Nous avons appliqué

ici, comme nous le ferons presque toujours, la convention d'omission des déductions par *a fortiori*. Entre les lignes 3^o et 4^o du premier bloc intérieur est omise ici une ligne contenant la proposition $\text{non}A$, déduite de la ligne 2^o par *a fortiori*, et c'est de cette ligne qu'est déduite la quatrième suivant **S16**.

Le schéma de déduction symétrique

$$\frac{\Gamma \vdash \text{non}A \text{ ou } B \quad \Gamma \vdash A}{\Gamma \vdash B}$$

admet une démonstration analogue. Il suffit de permuter $\text{non}A$ et A dans la précédente.

7.2 Schémas d'implication

Propriétés de l'implication. Les schémas de déduction rassemblés sous ce titre ne contiennent que le connecteur logique $\Rightarrow$. Seul le schéma **S22** contient encore une négation.

S19. *Réflexivité de l'implication*

$$\frac{}{\Gamma \vdash A \Rightarrow A.}$$

La démonstration en deux lignes qui justifie ce schéma est donnée ci-dessous de deux manières, comme celle de **S15** qui précède. Le lecteur peut se reporter au plan d'application du schéma $\Rightarrow$ -*introduction* de la figure 6.7 et considérer le cas particulier où B est la proposition A elle-même.

$$\begin{array}{ll} \begin{array}{l} 1^{\circ} \quad \Gamma; A \vdash A \quad \text{hyp} \\ 2^{\circ} \quad \Gamma \vdash A \Rightarrow A \quad 1^{\circ}, \Rightarrow\text{-intro.} \end{array} & \begin{array}{l} \Gamma \\ \hline 1^{\circ} \quad \boxed{\begin{array}{l} A \quad \text{hyp} \end{array}} \\ \hline 2^{\circ} \quad A \Rightarrow A \quad 1^{\circ}, \Rightarrow\text{-intro.} \end{array} \end{array}$$

S20. *Transitivité de l'implication*

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash B \Rightarrow C}{\Gamma \vdash A \Rightarrow C.}$$

Ce schéma a déjà été présenté (§6.3, Exemple 2). La démonstration complète, donnée dans la figure 6.8, se simplifie comme suit par omission des lignes de déduction a fortiori:

Γ		
1°	$A \Rightarrow B$	<i>base</i>
2°	$B \Rightarrow C$	<i>base</i>
3°	A	<i>hyp</i>
4°	B	3°, 1°, mod ponens
5°	C	4°, 2°, mod ponens
6°	$A \Rightarrow C$	5°, $\Rightarrow$ -intro.

S21. *Tout implique le vrai*

$$\frac{\Gamma \vdash B}{\Gamma \vdash A \Rightarrow B.}$$

Si la proposition B est vraie dans un contexte logique, alors pour n'importe quelle proposition A, la proposition $A \Rightarrow B$ est vraie dans ce contexte. D'où le nom du schéma. La démonstration ci-contre contient une ligne déduite par a fortiori qu'il n'est évidemment pas judicieux d'omettre. La ligne 2° n'est utilisée dans aucune déduction. Elle n'est là que pour la forme naturelle de la démonstration. Trois lignes suffisent pour une démonstration non naturelle:

Γ		
1°	B	<i>base</i>
2°	A	<i>hyp</i>
3°	B	1°, a fortiori
4°	$A \Rightarrow B$	3°, $\Rightarrow$ -intro.

1°	$\Gamma \vdash B$	<i>base</i>
2°	$\Gamma; A \vdash B$	1°, a fortiori
3°	$\Gamma \vdash A \Rightarrow B$	2°, $\Rightarrow$ -intro.

S22. *Le faux implique tout*

$$\frac{\Gamma \vdash \text{non}A}{\Gamma \vdash A \Rightarrow B} \quad \frac{\Gamma \vdash A}{\Gamma \vdash \text{non}A \Rightarrow B.}$$

Si une proposition A est fausse (nonA est vraie) dans un contexte logique, alors quelle que soit la proposition B, l'implication $A \Rightarrow B$ est vraie dans ce contexte.

Γ		
1°	nonA	<i>base</i>
2°	A	<i>hyp</i>
3°	B	1°, 2°, S16
4°	$A \Rightarrow B$	3°, $\Rightarrow$ -intro.

Schémas d'implication et disjonction. Le schéma suivant est la variante sans prémisse du schéma d'introduction de ou (S5).

S23.

$$\frac{}{\Gamma \vdash A \Rightarrow (A \text{ ou } B)} \quad \frac{}{\Gamma \vdash B \Rightarrow (A \text{ ou } B)}.$$

Les propositions de la forme $A \Rightarrow (A \text{ ou } B)$, resp. $B \Rightarrow (A \text{ ou } B)$, sont vraies dans tout contexte logique d'après les démonstrations suivantes :

$$\begin{array}{l} \Gamma \\ \begin{array}{l} 1^\circ \left| \begin{array}{cc} A & \text{hyp} \\ \hline A \text{ ou } B & 1^\circ, \text{ou-intro} \end{array} \right. \\ 2^\circ \left| \begin{array}{cc} A \text{ ou } B & 1^\circ, \text{ou-intro} \\ \hline A \Rightarrow (A \text{ ou } B) & 2^\circ, \Rightarrow\text{-intro} \end{array} \right. \end{array} \end{array} \quad \begin{array}{l} \Gamma \\ \begin{array}{l} 1^\circ \left| \begin{array}{cc} B & \text{hyp} \\ \hline A \text{ ou } B & 1^\circ, \text{ou-intro} \end{array} \right. \\ 2^\circ \left| \begin{array}{cc} A \text{ ou } B & 1^\circ, \text{ou-intro} \\ \hline A \Rightarrow (A \text{ ou } B) & 2^\circ, \Rightarrow\text{-intro} \end{array} \right. \end{array} \end{array}$$

Un bon nombre de schémas dérivés de la logique propositionnelle sont des schémas sans prémisses. Ils décrivent la forme de propositions qui sont vraies dans tout contexte logique. Nous rappelons (§6.5) qu'on appelle *tautologie* d'un langage $\mathcal{L}$ toute proposition A de $\mathcal{L}$ telle que le jugement $\vdash A$ soit un théorème de la théorie $\mathbf{Lprop}(\mathcal{L})$. Nous pouvons donc dire que toutes les propositions de la forme $A \Rightarrow (A \text{ ou } B)$ sont des tautologies. En effet, d'après S23, tous les jugements de la forme $\Gamma \vdash A \Rightarrow (A \text{ ou } B)$ sont des théorèmes de $\mathbf{Lprop}(\mathcal{L})$, y compris ceux dans lesquels Γ est vide.

S24.

$$\frac{\Gamma \vdash A \Rightarrow C \quad \Gamma \vdash B \Rightarrow C}{\Gamma \vdash (A \text{ ou } B) \Rightarrow C} \quad \frac{\Gamma \vdash A \Rightarrow C \quad \Gamma \vdash B \Rightarrow D}{\Gamma \vdash (A \text{ ou } B) \Rightarrow (C \text{ ou } D)}.$$

La démonstration de ces schémas est donnée dans la figure 7.1. Le premier est utilisé dans la démonstration du deuxième.

Nous verrons au chapitre 10 que les schémas de déduction de la logique propositionnelle sont utilisés pour démontrer les propriétés de la réunion et de l'intersection d'ensembles. Par exemple, en prenant pour A, B, C les propositions $x \in a, x \in b, x \in c$, dans lesquelles a, b, c désignent des ensembles, on pourra faire la déduction suivante, de la forme S24 avec Γ vide :

$$\frac{\vdash x \in a \Rightarrow x \in c \quad \vdash x \in b \Rightarrow x \in c}{\vdash (x \in a \text{ ou } x \in b) \Rightarrow x \in c}.$$

Dans le contexte de la théorie des ensembles, il en découlera la propriété suivante de la réunion d'ensembles : si deux ensembles a, b sont contenus dans un ensemble c , l'ensemble $a \cup b$ est contenu dans c .

S25.

$$\frac{\Gamma \vdash (A \text{ ou } B) \quad \Gamma \vdash A \Rightarrow C \quad \Gamma \vdash B \Rightarrow C}{\Gamma \vdash C}$$

Ce schéma dérivé peut être considérée comme une variante de *disjonction des cas*. On peut le substituer à S6 comme schéma élémentaire et obtenir inversement celui-ci comme schéma dérivé, ainsi que le lecteur pourra le vérifier.

Γ		
1°	$A \Rightarrow C$	<i>base</i>
2°	$B \Rightarrow C$	<i>base</i>
3°	$A \text{ ou } B$	hyp
4°	A	hyp
5°	C	4°, 1°, mod ponens
6°	B	hyp
7°	C	6°, 2°, mod ponens
8°	C	3°, 5°, 7°, disj cas
9°	$(A \text{ ou } B) \Rightarrow C$	8°, $\Rightarrow$ -intro.

Γ		
1°	$A \Rightarrow C$	<i>base</i>
2°	$C \Rightarrow (C \text{ ou } D)$	S23
3°	$A \Rightarrow (C \text{ ou } D)$	1°, 2°, S20
4°	$B \Rightarrow D$	<i>base</i>
5°	$D \Rightarrow (C \text{ ou } D)$	S23
6°	$B \Rightarrow (C \text{ ou } D)$	1°, 2°, S20
7°	$(A \text{ ou } B) \Rightarrow (C \text{ ou } D)$	3°, 6°, S24 (premier).

Figure 7.1 Démonstration de S24.

Le schéma S25 s'obtient immédiatement par combinaison de S24 et modus ponens:

Γ		
1°	$A \text{ ou } B$	<i>base</i>
2°	$A \Rightarrow C$	<i>base</i>
3°	$B \Rightarrow C$	<i>base</i>
4°	$(A \text{ ou } B) \Rightarrow C$	2°, 3°, S24
5°	C	1°, 4°, mod ponens.

S26.

$$\frac{\Gamma \vdash A \text{ ou } B \quad \Gamma \vdash A \Rightarrow B}{\Gamma \vdash B}.$$

Si une proposition A ou B est vraie dans un contexte, et si $A \Rightarrow B$ est vraie dans ce contexte, alors B est vraie dans ce contexte. Il suffit d'appliquer S25 en prenant pour C la proposition B elle-même.

Γ		
1°	$A \text{ ou } B$	<i>base</i>
2°	$A \Rightarrow B$	<i>base</i>
3°	$B \Rightarrow B$	S19
4°	B	1°, 2°, 3°, S25.

S27.

$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (A \text{ ou } C) \Rightarrow (B \text{ ou } C)} \quad \frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (C \text{ ou } A) \Rightarrow (C \text{ ou } B)}.$$

Ces schémas dérivent immédiatement de S24. La démonstration est donnée ici pour le premier :

Γ		
1 ^o	$A \Rightarrow B$	<i>base</i>
2 ^o	$C \Rightarrow C$	S19
3 ^o	$(A \text{ ou } C) \Rightarrow (B \text{ ou } C)$	1 ^o , 2 ^o , S24 (2ème).

Schémas d'implication et conjonction.

S28.

$$\frac{}{\Gamma \vdash (A \text{ et } B) \Rightarrow A} \quad \frac{}{\Gamma \vdash (A \text{ et } B) \Rightarrow B}.$$

On observe dans le système des schémas élémentaires (fig. 6.1) une symétrie entre les schémas d'élimination de *et* (S4) et ceux d'introduction de *ou* (S5). Il y a une symétrie correspondante entre les schémas S23 et les schémas S28, qui dérivent immédiatement du schéma élémentaire d'élimination de *et* :

Γ			Γ		
1 ^o	$\frac{A \text{ et } B \quad \text{hyp}}{A}$	hyp	1 ^o	$\frac{A \text{ et } B \quad \text{hyp}}{B}$	hyp
2 ^o		1 ^o , et-élim	2 ^o		1 ^o , et-élim
3 ^o	$(A \text{ et } B) \Rightarrow A$	2 ^o , $\Rightarrow$ -intro.	3 ^o	$(A \text{ et } B) \Rightarrow B$	2 ^o , $\Rightarrow$ -intro.

S29.

$$\frac{\Gamma \vdash C \Rightarrow A \quad \Gamma \vdash C \Rightarrow B}{\Gamma \vdash C \Rightarrow (A \text{ et } B)} \quad \frac{\Gamma \vdash C \Rightarrow A \quad \Gamma \vdash D \Rightarrow B}{\Gamma \vdash (C \text{ et } D) \Rightarrow (A \text{ et } B)}.$$

Le premier de ces schémas dérive essentiellement de *et-introduction*. Le deuxième dérive du premier :

Γ		
1 ^o	$C \Rightarrow A$	<i>base</i>
2 ^o	$C \Rightarrow B$	<i>base</i>
3 ^o	C	hyp
4 ^o	A	1 ^o , 3 ^o , mod ponens
5 ^o	B	2 ^o , 3 ^o , mod ponens
6 ^o	$A \text{ et } B$	4 ^o , 5 ^o , et-intro
7 ^o	$C \Rightarrow (A \text{ et } B)$	6 ^o , $\Rightarrow$ -intro.

Γ		
1 ^o	$C \Rightarrow A$	<i>base</i>
2 ^o	$(C \text{ et } D) \Rightarrow C$	S28
3 ^o	$(C \text{ et } D) \Rightarrow A$	2 ^o , 1 ^o , S20
4 ^o	$D \Rightarrow B$	<i>base</i>
5 ^o	$(C \text{ et } D) \Rightarrow D$	S28
6 ^o	$(C \text{ et } D) \Rightarrow B$	5 ^o , 4 ^o , S20
7 ^o	$(C \text{ et } D) \Rightarrow (A \text{ et } B)$	3 ^o , 6 ^o , S29 (premier).

S'il y a une symétrie entre les schémas élémentaires *et-élimination* et *ou-introduction*, il n'y en a pas entre le schéma élémentaire *et-introduction* et celui de *disjonction des cas*, que certains appellent aussi *ou-élimination*. Si l'on voulait un système de schémas élémentaires symétriques pour la disjonction et la conjonction, on pourrait prendre le premier des schémas S24 et le premier des schémas S29 comme schémas élémentaires, à la place des schémas S6 et S3. Ces derniers pourraient s'obtenir alors comme schémas dérivés, ainsi que le lecteur peut le vérifier.

S30.

$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (A \text{ et } C) \Rightarrow (B \text{ et } C)} \quad \frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (C \text{ et } A) \Rightarrow (C \text{ et } B)}.$$

Les justifications des lignes sont les mêmes dans les deux démonstrations:

Γ	Γ
1° $A \Rightarrow B$	1° $A \Rightarrow B$ <i>base</i>
2° $A \text{ et } C$	2° $C \text{ et } A$ <i>hyp</i>
3° A	3° A 2°, <i>et-élim</i>
4° B	4° B 1°, 3°, <i>mod ponens</i>
5° C	5° C 2°, <i>et-élim</i>
6° $B \text{ et } C$	6° $C \text{ et } B$ 4°, 5°, <i>et-intro</i>
7° $(A \text{ et } C) \Rightarrow (B \text{ et } C)$	7° $(A \text{ et } C) \Rightarrow (C \text{ et } B)$ 6°, $\Rightarrow$ -intro.

Schémas d'implication et négation.

S31. *Contraposition*

$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash \text{non}B \Rightarrow \text{non}A} \quad \frac{\Gamma \vdash \text{non}B \Rightarrow \text{non}A}{\Gamma \vdash A \Rightarrow B}.$$

Les déductions de la première forme S31 se démontrent comme suit et celles de la deuxième forme se démontrent de manière semblable par réduction à l'absurde K.

Γ	
1° $A \Rightarrow B$	<i>base</i>
2° $\text{non}B$	<i>hyp</i>
3° A	<i>hyp</i>
4° B	1°, 3°, <i>mod ponens</i>
5° $\perp$	2°, 4°, $\perp$ -intro
6° $\text{non}A$	5°, <i>red abs J</i>
7° $\text{non}B \Rightarrow \text{non}A$	6°, $\Rightarrow$ -intro.

Selon la terminologie du paragraphe 5.4, nous pouvons dire qu'un jugement de la forme $\Gamma \vdash A \Rightarrow B$ est équivalent, dans une théorie logique, au jugement $\Gamma \vdash \text{non}B \Rightarrow \text{non}A$, puisque chacun d'eux se déduit de l'autre. Il faut prendre garde à la permutation des positions de A et B entre ces deux jugements.

C'est une erreur assez fréquente que de conclure à la vérité d'une proposition $\text{non}A \Rightarrow \text{non}B$ lorsque la proposition $A \Rightarrow B$ est vraie. C'est la proposition $\text{non}B \Rightarrow \text{non}A$ qui est vraie, et elle est appelée la *contraposée* de $A \Rightarrow B$.

S32.

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash \text{non}B}{\Gamma \vdash \text{non}A.}$$

Ce schéma dérive immédiatement du précédent par modus ponens:

$$\begin{array}{l|ll} \Gamma & & \\ 1^\circ & A \Rightarrow B & \text{base} \\ 2^\circ & \text{non}B & \text{base} \\ 3^\circ & \text{non}B \Rightarrow \text{non}A & 1^\circ, \text{S31} \\ 4^\circ & \text{non}A & 2^\circ, 3^\circ, \text{mod ponens.} \end{array}$$

S33.

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash \text{non}A \Rightarrow \text{non}B}{\Gamma \vdash A \Leftrightarrow B.}$$

Ce schéma fournit une méthode assez fréquemment employée pour démontrer une équivalence $A \Leftrightarrow B$ dans un contexte donné. Elle consiste à démontrer dans ce contexte les deux implications $A \Rightarrow B$, $\text{non}A \Rightarrow \text{non}B$. Car alors $B \Rightarrow A$ est vraie par contraposition (S31), et $A \Leftrightarrow B$ est vraie par introduction de $\Leftrightarrow$ (S13). La démonstration détaillée de S33 suivant ces indications est laissée au lecteur.

7.3 Disjonctions de cas contraires

Le schéma élémentaire de disjonction des cas (S6) est souvent appliqué à deux propositions A , B contradictoires, c'est-à-dire en prenant pour B la proposition $\text{non}A$. Car la proposition $(A \text{ ou } \text{non}A)$ est une tautologie d'après la loi du tiers exclu (S34). On parle alors d'un raisonnement par disjonction des cas contraires: A , $\text{non}A$.

S34. *Tiers exclu*

$$\frac{}{\Gamma \vdash A \text{ ou } \text{non}A.}$$

Ce schéma sans prémisses a été présenté au paragraphe 6.4. Sa démonstration est donnée dans la figure 6.12.

S35. *Disjonction de cas contraires*

$$\frac{\Gamma; A \vdash C \quad \Gamma; \text{non}A \vdash C}{\Gamma \vdash C} \quad \frac{\Gamma; A \vdash C \quad \Gamma; \text{non}A \vdash D}{\Gamma \vdash (C \text{ ou } D).}$$

La démonstration naturelle du premier de ces schémas est donnée ci-dessous. Elle comporte deux lignes (2° et 4°) qui ne servent qu'à la forme

naturelle. Aucune autre ligne n'en est déduite. En renonçant à cette forme, on peut supprimer ces lignes. Le corps de la démonstration est alors la séquence de jugements: 1° $\Gamma \vdash A$ ou $\text{non}A$ (S34), 2° $\Gamma; A \vdash C$ (base), 3° $\Gamma; \text{non}A \vdash C$ (base), 4° $\Gamma \vdash C$ (1°, 2°, 3°, disj cas).

Γ		
1°	A ou $\text{non}A$	S34
2°	A	hyp
3°	C	base
4°	$\text{non}A$	hyp
5°	C	base
6°	C	1°, 3°, 5°, disj cas.

Pour démontrer le deuxième schéma, on utilise le premier. Le corps de la démonstration minimale, de forme non naturelle, est la séquence de jugements 1° $\Gamma; A \vdash C$ (base); 2° $\Gamma; A \vdash C$ ou D (1°, ou-intro); 3° $\Gamma; \text{non}A \vdash D$ (base); 4° $\Gamma; \text{non}A \vdash C$ ou D (3°, ou-intro); 5° $\Gamma \vdash C$ ou D (2°, 4°, S35 (premier)).

S36.

$$\frac{\Gamma \vdash A \Rightarrow C \quad \Gamma \vdash \text{non}A \Rightarrow C}{\Gamma \vdash C.}$$

$$\frac{\Gamma \vdash A \Rightarrow C \quad \Gamma \vdash \text{non}A \Rightarrow D}{\Gamma \vdash (C \text{ ou } D).}$$

Deux schémas qui dérivent de manière évidente de S24 et du tiers exclu. Nous donnons la démonstration du premier:

Γ		
1°	$A \Rightarrow C$	base
2°	$\text{non}A \Rightarrow C$	base
3°	$(A \text{ ou } \text{non}A) \Rightarrow C$	1°, 2°, S24
4°	$A \text{ ou } \text{non}A$	S34
5°	C	4°, 3°, mod ponens.

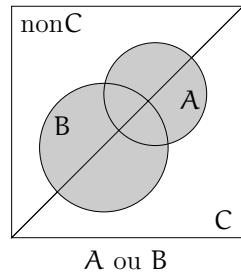
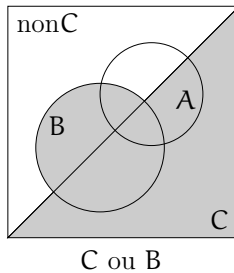
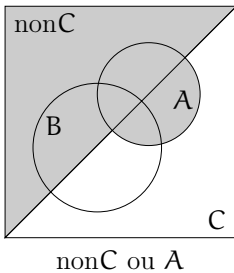


Figure 7.2 Illustration de S37.

Comme application de S36, nous démontrerons le schéma S37 ci-dessous, qui est illustré par la figure 7.2. On considère un carré dans lequel sont inscrits

deux cercles. Soit x un point se trouvant dans ce carré. Désignons par C (resp. $\text{non}C$) l'assertion « x est au-dessous (resp. au-dessus) de la diagonale »; par A (resp. B) l'assertion « x est dans le petit cercle (resp. le grand cercle) ». Si l'on sait que x est dans le petit cercle *ou* au-dessus de la diagonale ($\text{non}C$ ou A), et que x est dans le grand cercle *ou* au-dessous de la diagonale (C ou B), on peut en déduire que x est dans l'un des deux cercles (A ou B). On admet naturellement que (C ou $\text{non}C$) est vraie, donc que x n'est pas sur la diagonale.

S37.

$$\frac{\Gamma \vdash \text{non}C \text{ ou } A \quad \Gamma \vdash C \text{ ou } B}{\Gamma \vdash A \text{ ou } B.}$$

Démonstration:

	Γ	
1°	$\text{non}C \text{ ou } A$	<i>base</i>
2°	C	<i>hyp</i>
3°	A	1°, 2°, S18 (2ème)
4°	$C \Rightarrow A$	3°, $\Rightarrow$ -intro
5°	$C \text{ ou } B$	<i>base</i>
6°	$\text{non}C$	<i>hyp</i>
7°	B	5°, 6°, S18 (1er)
8°	$\text{non}C \Rightarrow B$	7°, $\Rightarrow$ -intro
9°	$A \text{ ou } B$	4°, 8°, S36 (2ème).

S38.

$$\frac{\Gamma \vdash A \text{ et } B}{\Gamma \vdash (C \text{ et } A) \text{ ou } (\text{non}C \text{ et } B).}$$

Démonstration:

	Γ	
1°	$A \text{ et } B$	<i>base</i>
2°	C	<i>hyp</i>
3°	A	1°, et-élim
4°	$C \text{ et } A$	2°, 3°, et-intro
5°	$\text{non}C$	<i>hyp</i>
6°	B	1°, et-élim
7°	$\text{non}C \text{ et } B$	5°, 6°, et-intro
8°	$(C \text{ et } A) \text{ ou } (\text{non}C \text{ et } B)$	4°, 7°, S35.

Exercices. Pour démontrer une proposition C par disjonction de cas contraires A , $\text{non}A$, on peut choisir arbitrairement la proposition A . Il faut la choisir judicieusement pour parvenir rapidement au but. En général le choix à faire

est évident. On demande de démontrer de cette manière que les jugements suivants d'un langage $\mathcal{L}$ sont des théorèmes de $\text{Lprop}(\mathcal{L})$.

- $\Gamma \vdash (\text{non}A \Rightarrow A) \Rightarrow A$.
- $\Gamma \vdash (C \Rightarrow A) \text{ ou } (A \Rightarrow B)$.
- $\Gamma \vdash (A \Rightarrow B) \text{ ou } (A \Rightarrow \text{non}B)$.
- $\Gamma \vdash A \text{ ou } (A \Rightarrow B)$.
- $\Gamma \vdash (\text{non}A \Rightarrow B) \text{ ou } (A \Rightarrow B)$.

7.4 Propriétés de l'équivalence

S39. *modus ponens (bis)*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash B} \qquad \frac{\Gamma \vdash B \quad \Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash A}$$

Si une proposition A ainsi qu'une proposition de la forme $A \Leftrightarrow B$ sont vraies dans un contexte, B est vraie dans ce contexte. Cela est clair, puisque $A \Rightarrow B$ est vraie dans ce contexte. La seconde déduction est semblable.

Γ	
1°	A <i>base</i>
2°	$A \Leftrightarrow B$ <i>base</i>
3°	$A \Rightarrow B$ 2°, $\Leftrightarrow$ -élim
4°	B 1°, 3°, mod ponens.

Dans les chapitres suivants, nous appliquerons souvent le schéma S39 sans le mentionner. Une séquence de trois lignes de la forme suivante dans une démonstration :

$$\begin{array}{ll} 1^\circ \Gamma \vdash A & \dots \\ 2^\circ \Gamma \vdash A \Leftrightarrow B & \mathbf{S_i} \\ 3^\circ \Gamma \vdash B & 1^\circ, 2^\circ, \text{S39} \end{array}$$

dans laquelle la deuxième ligne contient une équivalence bien connue, déduite de rien suivant un schéma sans prémisses $\mathbf{S_i}$, sera abrégée par

$$\begin{array}{ll} 1^\circ \Gamma \vdash A & \dots \\ 2^\circ \Gamma \vdash B & \mathbf{S_i}, \end{array}$$

en mentionnant seulement l'équivalence $\Gamma \vdash A \Leftrightarrow B$ par le nom du schéma $\mathbf{S_i}$, mais sans écrire cette équivalence.

S40. $\Leftrightarrow$ -introduction (bis)

$$\frac{\Gamma; A \vdash B \quad \Gamma; B \vdash A}{\Gamma \vdash A \Leftrightarrow B} \qquad \frac{\Gamma; A \vdash B \quad \Gamma; \text{non}A \vdash \text{non}B}{\Gamma \vdash A \Leftrightarrow B}$$

Le premier de ces schémas fournit la méthode la plus employée pour démontrer une équivalence $A \Leftrightarrow B$ sous des hypothèses Γ . Il suffit de démontrer B sous les hypothèses $\Gamma; A$ et de démontrer A sous les hypothèses $\Gamma; B$. Car alors $A \Rightarrow B$ et $B \Rightarrow A$ sont vraies sous les hypothèses Γ . Le deuxième schéma se démontre

de façon analogue par S33.

	Γ	
1°	A	hyp
2°	B	base
3°	$A \Rightarrow B$	2°, $\Rightarrow$ -intro
4°	B	hyp
5°	A	base
6°	$B \Rightarrow A$	5°, $\Rightarrow$ -intro
7°	$A \Leftrightarrow B$	3°, 6°, $\Leftrightarrow$ -intro.

S41.

$$\Gamma \vdash (A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A)).$$

La démonstration ci-dessous est une application du schéma S40. Si nous désignons $A \Leftrightarrow B$ par A' , et $((A \Rightarrow B) \text{ et } (B \Rightarrow A))$ par B' , la démonstration consiste à démontrer B' sous les hypothèses $\Gamma; A'$, puis A' sous les hypothèses $\Gamma; B'$, et à conclure suivant S40.

	Γ	
1°	$A \Leftrightarrow B$	hyp
2°	$A \Rightarrow B$	1°, $\Leftrightarrow$ -élim
3°	$B \Rightarrow A$	1°, $\Leftrightarrow$ -élim
4°	$(A \Rightarrow B) \text{ et } (B \Rightarrow A)$	2°, 3°, et-intro
5°	$(A \Rightarrow B) \text{ et } (B \Rightarrow A)$	hyp
6°	$A \Rightarrow B$	5°, et-élim
7°	$B \Rightarrow A$	5°, et-élim
8°	$A \Leftrightarrow B$	6°, 7°, $\Leftrightarrow$ -intro
9°	$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$	4°, 8°, S40.

S42. *Deux propositions vraies sont équivalentes*

$$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \Leftrightarrow B}.$$

	Γ	
Si deux propositions A, B sont vraies dans	1°	A base
un contexte, chacune des propositions $A \Rightarrow B$,	2°	B base
$B \Rightarrow A$ est vraie dans ce contexte d'après S21	3°	$A \Rightarrow B$ 2°, S21
(<i>tout implique le vrai</i>), d'où la conclusion.	4°	$B \Rightarrow A$ 1°, S21
	5°	$A \Leftrightarrow B$ 3°, 4°, $\Leftrightarrow$ -intro.

S43. *Deux propositions fausses sont équivalentes*

$$\frac{\Gamma \vdash \text{non}A \quad \Gamma \vdash \text{non}B}{\Gamma \vdash A \Leftrightarrow B}.$$

La démonstration est semblable à la précédente, mais elle utilise **S22** (le faux implique tout) à la place de **S21**. Les deux schémas suivants complètent **S42** et **S43**.

S44.

$$\frac{\Gamma \vdash A \quad \Gamma \vdash \text{non}B}{\Gamma \vdash \text{non}(A \Leftrightarrow B)}.$$

S45.

$$\frac{}{\Gamma \vdash \text{non}(A \Leftrightarrow \text{non}A)}.$$

Γ			Γ		
1°	A	base	1°	A $\Leftrightarrow$ nonA	hyp
2°	nonB	base	2°	A	hyp
3°	A $\Leftrightarrow$ B	hyp	3°	nonA	1°, 2°, S39
4°	B	1°, 3°, S39	4°	$\perp$	2°, 3°, $\perp$ -intro
5°	$\perp$	2°, 4°, $\perp$ -intro	5°	nonA	hyp
6°	non(A $\Leftrightarrow$ B)	5°, red abs J.	6°	A	1°, 5°, S39
			7°	$\perp$	5°, 6°, $\perp$ -intro
			8°	$\perp$	4°, 7°, S35
			9°	non(A $\Leftrightarrow$ nonA)	8°, red abs J

S46. *Réflexivité de l'équivalence*

$$\frac{}{\Gamma \vdash A \Leftrightarrow A}.$$

S47. *Symétrie de l'équivalence*

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash B \Leftrightarrow A}.$$

S48. *Transitivité de l'équivalence*

$$\frac{\Gamma \vdash A \Leftrightarrow B \quad \Gamma \vdash B \Leftrightarrow C}{\Gamma \vdash A \Leftrightarrow C}.$$

réflexivité

Γ		
1°	A $\Rightarrow$ A	S19
2°	A $\Leftrightarrow$ A	1°, 1°, $\Leftrightarrow$ -intro.

symétrie

Γ		
1°	A $\Leftrightarrow$ B	base
2°	A $\Rightarrow$ B	1°, $\Leftrightarrow$ -élim
3°	B $\Rightarrow$ A	1°, $\Leftrightarrow$ -élim
4°	B $\Leftrightarrow$ A	3°, 2°, $\Leftrightarrow$ -intro.

transitivité

Γ		
1°	A $\Leftrightarrow$ B	base
2°	B $\Leftrightarrow$ C	base
3°	A $\Rightarrow$ B	1°, $\Leftrightarrow$ -élim
4°	B $\Rightarrow$ C	2°, $\Leftrightarrow$ -élim
5°	A $\Rightarrow$ C	3°, 4°, S20
6°	C $\Rightarrow$ B	2°, $\Leftrightarrow$ -élim
7°	B $\Rightarrow$ A	1°, $\Leftrightarrow$ -élim
8°	C $\Rightarrow$ A	6°, 7°, S20
9°	A $\Leftrightarrow$ C	5°, 8°, $\Leftrightarrow$ -intro.

Chaînes d'équivalences. Une démonstration de la forme

$$\begin{array}{c|cl} \Gamma & & \\ 1^\circ & A_1 \Leftrightarrow A_2 & j_1 \\ 2^\circ & A_2 \Leftrightarrow A_3 & j_2 \\ & \vdots & \\ n^\circ & A_n \Leftrightarrow A_{n+1} & j_n \end{array} \quad (1)$$

mettant en jeu $n + 1$ propositions $A_1, \dots, A_{n+1}$ et comportant n lignes dont la k -ième (pour $k = 1, \dots, n$) est de la forme $\Gamma \vdash A_k \Leftrightarrow A_{k+1}$, est appelée une *chaîne d'équivalences* sous les hypothèses Γ . Nous avons désigné par $j_1, \dots, j_n$ les justifications qui doivent être données pour chaque ligne. Des n lignes de cette démonstration, on déduit

$$\Gamma \vdash A_1 \Leftrightarrow A_{n+1}$$

de façon évidente, par $n - 1$ applications de la transitivité de l'équivalence (S48): on tire d'abord $\Gamma \vdash A_1 \Leftrightarrow A_3$ des lignes 1° et 2° et l'on inscrit cette nouvelle ligne; puis on tire $\Gamma \vdash A_1 \Leftrightarrow A_4$ de cette nouvelle ligne et de 3° , et ainsi de suite. Nous abrègerons en général une démonstration de la forme (1) de la manière suivante:

$$\begin{array}{c|cl} \Gamma & & \\ 1^\circ & A_1 \Leftrightarrow A_2 & j_1 \\ 2^\circ & \Leftrightarrow A_3 & j_2 \\ & \vdots & \\ n^\circ & \Leftrightarrow A_{n+1} & j_n. \end{array} \quad (2)$$

Lorsqu'en outre chacune des lignes de la démonstration se déduit seulement de celle qui la précède immédiatement — c'est-à-dire que dans la justification de la $(k + 1)$ -ième ligne on ne se réfère pas à des lignes de numéro inférieur à k — alors on omet la numérotation, et l'on ne donne dans les justifications que les schémas de déduction utilisés. On omet souvent aussi le cadre et la liste d'hypothèses Γ lorsqu'elle est supposée connue. On écrit simplement

$$\begin{array}{cl} A_1 \Leftrightarrow A_2 & j_1 \\ \Leftrightarrow A_3 & j_2 \\ \vdots & \\ \Leftrightarrow A_{n+1} & j_n. \end{array} \quad (3)$$

Lorsque nous parlons de démontrer dans un contexte une proposition de la forme $A \Leftrightarrow B$ par *chaîne d'équivalences*, nous entendons par là construire une démonstration de la forme ci-dessus dans laquelle A_1 est la proposition A , et A_{n+1} est la proposition B . Cette forme de démonstration sera très souvent utilisée dans la suite. C'est la seconde manière standard de démontrer une équivalence — la première étant de suivre le premier schéma S40, comme dans

la démonstration de **S41** par exemple. Le premier schéma **S40** est une variante de **S13**, en vertu de **S7**. De même, le deuxième schéma **S40** est une variante de **S33**. C'est pourquoi nous parlerons de démontrer une équivalence $A \Leftrightarrow B$ par *double implication*, lorsque nous le ferons en suivant l'un ou l'autre de ces quatre schémas.

Chaînes d'implications et équivalences. Moins fréquentes que les chaînes d'équivalences (1) sont les *chaînes d'implications*, c'est-à-dire les démonstrations de la forme

$$\begin{array}{c|cc} \Gamma & & \\ 1^o & A_1 \Rightarrow A_2 & j_1 \\ 2^o & A_2 \Rightarrow A_3 & j_2 \\ & \vdots & \\ n^o & A_n \Rightarrow A_{n+1} & j_n. \end{array} \quad (4)$$

Par application de la *transitivité de l'implication* (**S20**), on déduit de ces lignes les $n - 1$ lignes $\Gamma \vdash A_1 \Rightarrow A_3$, $\Gamma \vdash A_1 \Rightarrow A_4$, $\dots$, $\Gamma \vdash A_1 \Rightarrow A_{n+1}$, cette dernière étant en général le but que l'on voulait atteindre.

Finalement, on rencontre assez fréquemment des chaînes qui sont un mélange d'implications et d'équivalences, c'est-à-dire des démonstrations de la forme

$$\begin{array}{c|cc} \Gamma & & \\ 1^o & A_1 \sqsubset_1 A_2 & j_1 \\ 2^o & A_2 \sqsubset_2 A_3 & j_2 \\ & \vdots & \\ n^o & A_n \sqsubset_n A_{n+1} & j_n \end{array} \quad (5)$$

dans lesquelles chacun des symboles $\sqsubset_i$ est un symbole d'équivalence ($\Leftrightarrow$) ou un symbole d'implication ($\Rightarrow$). De ces lignes on peut déduire les $n - 1$ lignes suivantes :

$$\Gamma \vdash A_1 \sim_3 A_3, \quad \Gamma \vdash A_1 \sim_4 A_4, \quad \dots \quad \Gamma \vdash A_1 \sim_{n+1} A_{n+1} \quad (6)$$

où chacun des symboles $\sim_k$ est un symbole d'équivalence ou d'implication. On se sert pour cela de la transitivité de l'équivalence (**S48**), de la transitivité de l'implication (**S20**), et des schémas **S49** ci-dessous. Dans (6), $\sim_k$ est le symbole $\Leftrightarrow$ si chacun des symboles $\sqsubset_1$ à $\sqsubset_k$ de la chaîne (5) est le symbole $\Leftrightarrow$; sinon, s'il y a au moins un symbole d'implication ($\Rightarrow$) parmi les symboles $\sqsubset_1$ à $\sqsubset_k$, le symbole $\sim_k$ de (6) est $\Rightarrow$.

Une chaîne d'implications et d'équivalences (5) est abrégée de la manière (2) ou (3), mais avec le symbole $\sqsubset_k$ ($\Leftrightarrow$ ou $\Rightarrow$) à la k -ième ligne.

S49.

$$\frac{\Gamma \vdash A \Rightarrow B \quad \Gamma \vdash B \Leftrightarrow C}{\Gamma \vdash A \Rightarrow C} \qquad \frac{\Gamma \vdash A \Leftrightarrow B \quad \Gamma \vdash B \Rightarrow C}{\Gamma \vdash A \Rightarrow C}.$$

Démonstrations:

Γ		Γ	
1°	$A \Rightarrow B$ <i>base</i>	1°	$A \Leftrightarrow B$ <i>base</i>
2°	$B \Leftrightarrow C$ <i>base</i>	2°	$A \Rightarrow B$ 1°, $\Leftrightarrow$ -élim
3°	$B \Rightarrow C$ 2°, $\Leftrightarrow$ -élim	3°	$B \Rightarrow C$ <i>base</i>
4°	$A \Rightarrow C$ 1°, 3°, S20.	4°	$A \Rightarrow C$ 2°, 3°, S20.

7.5 Substitutivité de l'équivalence

Les neuf formes de déductions décrites ci-dessous par S50.1 à S50.9, ont une structure commune. Dans chacune d'elles, la prémisse est une équivalence $A \Leftrightarrow B$ sous une liste d'hypothèses Γ . La conclusion est dans chaque cas une équivalence sous les mêmes hypothèses, et le membre de droite de cette équivalence s'obtient en prenant celui de gauche et en y remplaçant A par B . Les neuf cas correspondent à toutes les propositions que l'on peut construire en prenant une proposition A et en lui appliquant le connecteur logique « non », ou en prenant A et une seconde proposition C , et en combinant les deux au moyen d'un connecteur logique binaire (ou, et, $\Rightarrow$, $\Leftrightarrow$).

S50. *Substitutivité de l'équivalence*

1.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash \text{non}A \Leftrightarrow \text{non}B.}$$
2.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (A \text{ ou } C) \Leftrightarrow (B \text{ ou } C).}$$
3.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (C \text{ ou } A) \Leftrightarrow (C \text{ ou } B).}$$
4.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (A \text{ et } C) \Leftrightarrow (B \text{ et } C).}$$
5.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (C \text{ et } A) \Leftrightarrow (C \text{ et } B).}$$
6.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (A \Rightarrow C) \Leftrightarrow (B \Rightarrow C).}$$
7.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (C \Rightarrow A) \Leftrightarrow (C \Rightarrow B).}$$
8.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (A \Leftrightarrow C) \Leftrightarrow (B \Leftrightarrow C).}$$
9.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash (C \Leftrightarrow A) \Leftrightarrow (C \Leftrightarrow B).}$$

La démonstration de quatre de ces déductions est donnée dans la figure 7.3. Les autres sont analogues. Notamment, la démonstration de S50.4 est semblable à celle de S50.2 mais utilise S29 au lieu de S24.

Propositions équivalentes. On dit que deux propositions A , B sont *équivalentes* dans un contexte particulier, lorsque la proposition $A \Leftrightarrow B$ est vraie dans ce contexte. On dit aussi que A est *équivalente* à B dans ce contexte. La référence à un contexte, dans cette locution, est essentielle puisque la vérité de la proposition $A \Leftrightarrow B$ n'est définie que par rapport à un contexte. Néanmoins,

Γ		
1 ^o	$A \Leftrightarrow B$	<i>base</i>
2 ^o	$A \Rightarrow B$	1 ^o , $\Leftrightarrow$ -élim
3 ^o	$\text{non}B \Rightarrow \text{non}A$	2 ^o , S31
4 ^o	$B \Rightarrow A$	1 ^o , $\Leftrightarrow$ -élim
5 ^o	$\text{non}A \Rightarrow \text{non}B$	4 ^o , S31
6 ^o	$\text{non}A \Leftrightarrow \text{non}B$	5 ^o , 3 ^o , $\Leftrightarrow$ -intro.
Γ		
1 ^o	$A \Leftrightarrow B$	<i>base</i>
2 ^o	$A \Rightarrow B$	1 ^o , $\Leftrightarrow$ -élim
3 ^o	$C \Rightarrow C$	S19
4 ^o	$(A \text{ ou } C) \Rightarrow (B \text{ ou } C)$	2 ^o , 3 ^o , S24 (2ème)
5 ^o	$B \Rightarrow A$	1 ^o , $\Leftrightarrow$ -élim
6 ^o	$(B \text{ ou } C) \Rightarrow (A \text{ ou } C)$	5 ^o , 3 ^o , S24 (2ème)
7 ^o	$(A \text{ ou } C) \Leftrightarrow (B \text{ ou } C)$	4 ^o , 6 ^o , $\Leftrightarrow$ -intro.
Γ		
1 ^o	$A \Leftrightarrow B$	<i>base</i>
2 ^o	$B \Leftrightarrow A$	1 ^o , S47
3 ^o	$A \Rightarrow C$	hyp
4 ^o	$B \Rightarrow C$	2 ^o , 3 ^o , S49 (2ème)
5 ^o	$B \Rightarrow C$	hyp
6 ^o	$A \Rightarrow C$	1 ^o , 5 ^o , S49 (2ème)
7 ^o	$(A \Rightarrow C) \Leftrightarrow (B \Rightarrow C)$	4 ^o , 6 ^o , S40.
Γ		
1 ^o	$A \Leftrightarrow B$	<i>base</i>
2 ^o	$B \Leftrightarrow A$	1 ^o , S47
3 ^o	$A \Leftrightarrow C$	hyp
4 ^o	$B \Leftrightarrow C$	2 ^o , 3 ^o , S48
5 ^o	$B \Leftrightarrow C$	hyp
6 ^o	$A \Leftrightarrow C$	1 ^o , 5 ^o , S48
7 ^o	$(A \Leftrightarrow C) \Leftrightarrow (B \Leftrightarrow C)$	4 ^o , 6 ^o , S40.

Figure 7.3 Démonstrations des schémas de substitutivité de l'équivalence S50.1, S50.2, S50.6, S50.8.

lorsque ce contexte est connu, on a tendance à dire simplement que A et B sont équivalentes. Par exemple, le schéma sans prémisses

$$\overline{\Gamma \vdash A \Leftrightarrow \text{nonnon}A}$$

qui figure plus loin dans notre répertoire (S51) peut s'énoncer en disant que, dans un contexte quelconque, une proposition est équivalente à sa double négation. En utilisant cette manière de parler, on peut exprimer en une seule phrase les neuf schémas S50 de la manière suivante:

Substitutivité de l'équivalence. Si deux propositions A , B sont équivalentes dans un contexte, alors les propositions $\text{non}A$, $\text{non}B$ sont équivalentes dans ce contexte, et étant donné une troisième proposition C , les propositions A ou C , C ou A , A et C , C et A , $A \Rightarrow C$, $C \Rightarrow A$, $A \Leftrightarrow C$, $C \Leftrightarrow A$ sont équivalentes respectivement à B ou C , C ou B , B et C , C et B , $B \Rightarrow C$, $C \Rightarrow B$, $B \Leftrightarrow C$, $C \Leftrightarrow B$ dans ce contexte.

Exemple. Comme exemple d'application des schémas S50, nous allons donner une démonstration de la déduction suivante:

$$\frac{\Gamma \vdash C \Leftrightarrow C'}{\Gamma \vdash (A \text{ et } \text{non}(B \text{ ou } C)) \Leftrightarrow (A \text{ et } \text{non}(B \text{ ou } C'))} \quad (1)$$

Démonstration:

Γ		
1°	$C \Leftrightarrow C'$	<i>base</i>
2°	$(B \text{ ou } C) \Leftrightarrow (B \text{ ou } C')$	1°, S50.3
3°	$\text{non}(B \text{ ou } C) \Leftrightarrow \text{non}(B \text{ ou } C')$	2°, S50.1
4°	$(A \text{ et } \text{non}(B \text{ ou } C)) \Leftrightarrow (A \text{ et } \text{non}(B \text{ ou } C'))$	3°, S50.5.

Schéma général de substitutivité de l'équivalence. La déduction (1) peut s'écrire

$$\frac{\Gamma \vdash C \Leftrightarrow C'}{\Gamma \vdash R \Leftrightarrow R'} \quad (2)$$

si l'on désigne par R la proposition $(A \text{ et } \text{non}(B \text{ ou } C))$ et par R' la proposition $(A \text{ et } \text{non}(B \text{ ou } C'))$. L'important est de remarquer que l'on passe de R à R' en remplaçant C par C' .

Cet exemple suggère immédiatement que toute déduction de la forme (2), dans laquelle R' s'obtient en prenant R et en y remplaçant C par C' , admet une démonstration dans $\mathbf{Lprop}(\mathcal{L})$, en vertu des schémas S50. Les déductions de cette forme sont la généralisation de celles décrites par les schémas S50. Mais il faut préciser que C doit être une *sous-proposition* de R au sens de la logique propositionnelle. Cela veut dire que C doit être une partie de R qui est elle-même une proposition et que R est construite avec C et éventuellement d'autres propositions en utilisant exclusivement des *connecteurs logiques*. Il n'intervient pas de quantificateurs. Grâce à cette construction de R , on peut passer du jugement $\Gamma \vdash C \Leftrightarrow C'$ au jugement $\Gamma \vdash R \Leftrightarrow R'$ par une succession de déductions suivant S50, comme dans l'exemple ci-dessus.

Cette classe générale de déductions est décrite par le très important schéma de déduction suivant, dont il est fait constamment usage dans le raisonnement

mathématique.

Substitutivité générale de l'équivalence

$$\frac{\Gamma \vdash C \Leftrightarrow C'}{\Gamma \vdash R \Leftrightarrow R'}$$

où C est une sous-proposition de R et R' est la proposition obtenue en remplaçant C par C' dans R .

La notion de sous-proposition doit être prise pour le moment au sens de la logique propositionnelle, défini ci-dessus. Plus loin cette notion pourra être prise au sens plus large de la logique des prédicats, sous certaines conditions sur les quantificateurs utilisés dans la construction de R .

On peut énoncer ce schéma de déduction en termes de *propositions équivalentes*, au sens défini plus haut : si deux propositions C , C' sont équivalentes dans un contexte, alors deux propositions R , R' telles que C soit une sous-proposition de R et R' soit obtenue en remplaçant C par C' dans R sont équivalentes dans ce contexte. Cet énoncé un peu compliqué peut se tourner plus simplement comme suit :

Si l'on remplace une sous-proposition C d'une proposition R par une proposition C' qui est équivalente à C dans un contexte donné, on obtient une proposition R' qui est équivalente à R dans ce contexte.

Il arrive qu'on simplifie encore cet énoncé en disant sommairement : *on peut remplacer une sous-proposition d'une proposition R par une sous-proposition équivalente*. Ce qui est sous-entendu, c'est qu'on obtient de cette manière une proposition R' équivalente à R .

7.6 Schémas de déduction booléens

Liste des principaux schémas booléens. Presque tous les schémas de ce paragraphe, §51.1 à §52.16, sont sans prémisses. Pour condenser leur présentation, ci-dessous, nous utilisons une barre de déduction verticale « $\mid$ » à la place de la barre horizontale usuelle. Ainsi le schéma

$$\frac{}{\Gamma \vdash A \Leftrightarrow \text{nonnon}A}$$

est noté $\mid \Gamma \vdash A \Leftrightarrow \text{nonnon}A$.

À l'exception des schémas avec prémisses §51.2, tous les schémas ci-dessous décrivent des *tautologies*. Nous rappelons qu'une tautologie (§6.5) d'un langage $\mathcal{L}$ est une proposition A de $\mathcal{L}$ telle que le jugement $\vdash A$ soit un théorème de $\text{Lprop}(\mathcal{L})$. De plus, toutes les tautologies décrites dans ce paragraphe sont des *équivalences*, dont les deux membres sont construits au moyen de propositions A , B , C en utilisant exclusivement les connecteurs logiques *et*, *ou*, *non*, que nous appellerons les connecteurs *booléens*. La discussion détaillée de ces schémas et

de leur démonstration est donnée dans la suite de ce paragraphe. Nous ne faisons ici que les présenter tous ensemble.

On observe que les schémas S52.1 à S52.16 sont groupés par paires obéissant à la symétrie suivante appelée *dualité*: chacun des deux schémas d'une paire s'obtient en remplaçant les disjonctions et les conjonctions de l'autre respectivement par des conjonctions et des disjonctions. On dit que deux propositions qui se transforment ainsi l'une en l'autre par cet échange des deux types de connecteurs sont des propositions *duales*, ou que l'une est la *duale* de l'autre. Ce sujet sera traité en détail au paragraphe 7.7.

S51.1 $\mid \Gamma \vdash A \Leftrightarrow \text{nonnon}A.$

S51.2

$$\frac{\Gamma \vdash A \Leftrightarrow \text{non}B}{\Gamma \vdash \text{non}A \Leftrightarrow B} \qquad \frac{\Gamma \vdash \text{non}A \Leftrightarrow B}{\Gamma \vdash A \Leftrightarrow \text{non}B.}$$

De Morgan

S52.1 $\mid \Gamma \vdash \text{non}(A \text{ ou } B) \Leftrightarrow (\text{non}A \text{ et } \text{non}B).$

S52.2 $\mid \Gamma \vdash \text{non}(A \text{ et } B) \Leftrightarrow (\text{non}A \text{ ou } \text{non}B).$

S52.3 $\mid \Gamma \vdash (A \text{ ou } B) \Leftrightarrow \text{non}(\text{non}A \text{ et } \text{non}B).$

S52.4 $\mid \Gamma \vdash (A \text{ et } B) \Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B).$

Idempotence de la disjonction et de la conjonction

S52.5 $\mid \Gamma \vdash (A \text{ ou } A) \Leftrightarrow A.$

S52.6 $\mid \Gamma \vdash (A \text{ et } A) \Leftrightarrow A.$

Commutativité de la disjonction et de la conjonction

S52.7 $\mid \Gamma \vdash (A \text{ ou } B) \Leftrightarrow (B \text{ ou } A).$

S52.8 $\mid \Gamma \vdash (A \text{ et } B) \Leftrightarrow (B \text{ et } A).$

Associativité de la disjonction et de la conjonction

S52.9 $\mid \Gamma \vdash (A \text{ ou } (B \text{ ou } C)) \Leftrightarrow ((A \text{ ou } B) \text{ ou } C).$

S52.10 $\mid \Gamma \vdash (A \text{ et } (B \text{ et } C)) \Leftrightarrow ((A \text{ et } B) \text{ et } C).$

Distributivité de la conjonction rapport à la disjonction, et vice-versa

S52.11 $\mid \Gamma \vdash (A \text{ et } (B \text{ ou } C)) \Leftrightarrow ((A \text{ et } B) \text{ ou } (A \text{ et } C)).$

S52.12 $\mid \Gamma \vdash (A \text{ ou } (B \text{ et } C)) \Leftrightarrow ((A \text{ ou } B) \text{ et } (A \text{ ou } C)).$

Absorption

S52.13 $\mid \Gamma \vdash ((A \text{ et } B) \text{ ou } B) \Leftrightarrow B.$

S52.14 $\mid \Gamma \vdash ((A \text{ ou } B) \text{ et } B) \Leftrightarrow B.$

Neutralité

S52.15 $\mid \Gamma \vdash ((A \text{ et } \text{non}A) \text{ ou } B) \Leftrightarrow B.$

S52.16 $\mid \Gamma \vdash ((A \text{ ou } \text{non}A) \text{ et } B) \Leftrightarrow B.$

Γ	
1°	A hyp
2°	$\text{non}A$ hyp
3°	$\perp$ 1°, 2°, $\perp$ -intro
4°	$\text{nonnon}A$ 3°, red abs J
5°	$\text{nonnon}A$ hyp
6°	$\text{non}A$ hyp
7°	$\perp$ 6°, 5°, $\perp$ -intro
8°	A 7°, red abs K
9°	$A \Leftrightarrow \text{nonnon}A$ 4°, 8°, S40.

Figure 7.4 Démonstration de S51.1

Γ		Γ	
1°	$A \Leftrightarrow \text{non}B$ base	1°	$A \Leftrightarrow \text{non}B$ base
2°	$\text{non}A \Leftrightarrow \text{nonnon}B$ 1°, S50.1	2°	$\text{non}A \Leftrightarrow \text{nonnon}B$ 1° (S50.1)
3°	$\text{nonnon}B \Leftrightarrow B$ S51.1	3°	$\Leftrightarrow B$ S51.1.
4°	$\text{non}A \Leftrightarrow B$ 2°, 3°, S48.		
(a)		(b)	

Figure 7.5 (a) Démonstration complète du premier schéma S51.2. (b) Démonstration abrégée.

Utilisation tacite de la symétrie de l'équivalence. En utilisant le schéma de symétrie de l'équivalence S47, on obtient immédiatement, pour chacun des schémas de déduction ci-dessus, un schéma dérivé symétrique dans lequel on a simplement permuté les deux membres de l'équivalence, par exemple le schéma

$$\frac{}{\Gamma \vdash \text{nonnon}A \Leftrightarrow A} \quad (1)$$

symétrique de S51.1. Nous n'écrirons pas ces schémas, dont la démonstration est triviale, et nous considérons que le numéro S51.1 vaut aussi pour le schéma (1). Cette convention s'applique désormais à tous les schémas dont la conclusion est une équivalence.

Schémas de double négation. La démonstration du schéma S51.1 est donnée dans la figure 7.4. Celle du premier schéma S51.2 est donnée dans la figure 7.5 de deux manières. D'abord complètement, (a), en utilisant toutefois tacitement la symétrie de l'équivalence à la ligne 2°, comme nous venons de l'annoncer. Les lignes 2° et 3° constituent une chaîne d'équivalences (§7.4), et la ligne 4° est la conclusion tirée de cette chaîne par transitivité de l'équivalence (S48). Nous abrégeons donc cette démonstration sous la forme (b), comme nous l'avons décrit pour les chaînes d'équivalence au paragraphe 7.4. La conclusion tirée d'une telle chaîne (ligne 4°) par transitivité est généralement omise, car on la

voit bien sans qu'il soit besoin de l'écrire. Au lieu de S50.1, à la ligne 2°, on peut invoquer le schéma général de substitutivité de l'équivalence (§7.5) dont S50.1 est un cas particulier. Par ailleurs, ce schéma général est si souvent appliqué que nous omettrons de le mentionner, sauf exception. Donc la mention de S50.1 à la deuxième ligne peut être omise. La démonstration du second schéma S51.2, abrégée de cette manière, est la suivante :

Γ		
1°	$\text{non}A \Leftrightarrow B$	<i>base</i>
2°	$A \Leftrightarrow \text{nonnon}A$	S51.1
3°	$\Leftrightarrow \text{non}B$	1°.

Γ		
1°	$A \Rightarrow (A \text{ ou } B)$	S23
2°	$B \Rightarrow (A \text{ ou } B)$	S23
3°	$\text{non}(A \text{ ou } B) \Rightarrow \text{non}A$	1°, S31
4°	$\text{non}(A \text{ ou } B) \Rightarrow \text{non}B$	2°, S31
5°	$\text{non}(A \text{ ou } B) \Rightarrow (\text{non}A \text{ et } \text{non}B)$	3°, 4°, S29
6°	$\text{non}A \text{ et } \text{non}B$	hyp
7°	$A \text{ ou } B$	hyp
8°	$\text{non}A$	6°, et-élim
9°	B	7°, 8°, S18
10°	$\text{non}B$	6°, et-élim
11°	$\perp$	9°, 10°, $\perp$ -intro
12°	$\text{non}(A \text{ ou } B)$	11°, red abs J
13°	$(\text{non}A \text{ et } \text{non}B) \Rightarrow \text{non}(A \text{ ou } B)$	12°, $\Rightarrow$ -intro
14°	$\text{non}(A \text{ ou } B) \Leftrightarrow (\text{non}A \text{ et } \text{non}B)$	5°, 13°, $\Leftrightarrow$ -intro.

Figure 7.6 Démonstration de S52.1.

Schémas de De Morgan et dualité. Les schémas du groupe S52.1–4 portent le nom du mathématicien anglais Augustus De Morgan (1806–1871). La démonstration des deux premiers est donnée dans la figure 7.6. Nous allons maintenant utiliser le premier schéma de De Morgan (S52.1) pour établir le deuxième (S52.2), par une démonstration en forme de *chaîne d'équivalences* (§7.4), en appliquant le schéma général de *substitutivité de l'équivalence* (§7.5). Ceci est notre premier exemple d'une méthode de démonstration que nous emploierons très fréquemment dans la suite, aussi nous le commenterons en détail.

Γ		
	$\text{non}(A \text{ et } B) \Leftrightarrow \text{non}(\text{nonnon}A \text{ et } B)$	S51.1, subst equiv
	$\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } \text{nonnon}B)$	S51.1, subst equiv
	$\Leftrightarrow \text{nonnon}(\text{non}A \text{ ou } \text{non}B)$	S52.1, subst equiv
	$\Leftrightarrow (\text{non}A \text{ ou } \text{non}B)$	S51.1.

Désignons par R_1 à R_5 les cinq propositions de cette chaîne, la première étant $\text{non}(A \text{ et } B)$, la cinquième ($\text{non}A$ ou $\text{non}B$). Le principe de substitutivité de l'équivalence est appliqué trois fois. On passe de R_1 à R_2 en remplaçant la sous-proposition A de R_1 par la proposition $\text{nonnon}A$, qui lui est équivalente d'après S51.1; de R_2 à R_3 en remplaçant la sous-proposition B de R_2 par $\text{nonnon}B$; de R_3 à R_4 en remplaçant la sous-proposition ($\text{nonnon}A$ et $\text{nonnon}B$) de R_3 par la proposition $\text{non}(\text{non}A \text{ ou } \text{non}B)$, qui lui est équivalente d'après S52.1 — ou plus exactement d'après S52.1 et S47 (voir plus haut: utilisation tacite de la symétrie de l'équivalence). Et enfin l'on passe de R_4 à R_5 suivant S51.1. Il est permis de dire que dans cette dernière étape on applique le cas trivial du principe de substitutivité de l'équivalence, où l'on considère R_4 comme sous-proposition d'elle-même.

Il faut bien voir qu'une telle chaîne est inversible en vertu de la symétrie de l'équivalence (S47). On peut l'écrire dans ce sens:

$$\begin{aligned}
 (\text{non}A \text{ ou } \text{non}B) &\Leftrightarrow \text{nonnon}(\text{non}A \text{ ou } \text{non}B) && \text{S51.1} \\
 &\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } \text{nonnon}B) && \text{S52.1} \\
 &\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } B) && \text{S51.1} \\
 &\Leftrightarrow \text{non}(A \text{ et } B) && \text{S51.1.}
 \end{aligned}$$

Dorénavant nous omettrons généralement, comme ci-dessus, de signaler l'usage de la substitutivité de l'équivalence, qui est constant dans les démonstrations par chaînes d'équivalence. Nous omettrons souvent aussi le cadre dans lequel se place la chaîne et la désignation de la liste d'hypothèses correspondante, qui est généralement connue.

Les schémas S52.3 et S52.4 dérivent des trois schémas booléens précédents d'après les chaînes d'équivalences évidentes:

$$\begin{aligned}
 (A \text{ ou } B) &\Leftrightarrow \text{nonnon}(A \text{ ou } B) && \text{S51.1} \\
 &\Leftrightarrow \text{non}(\text{non}A \text{ et } \text{non}B) && \text{S52.1} \\
 (A \text{ et } B) &\Leftrightarrow \text{nonnon}(A \text{ et } B) && \text{S51.1} \\
 &\Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B) && \text{S52.2.}
 \end{aligned}$$

Idempotence, commutativité, associativité, distributivité de la disjonction et de la conjonction. Les propriétés de la disjonction et de la conjonction ainsi nommées s'expriment par les schémas de déduction S52.5 à S52.12. Ces schémas sont groupés par paires duales. L'un des deux schémas de chaque paire peut toujours se démontrer par une chaîne d'équivalence en utilisant l'autre schéma et ceux de De Morgan. C'est un procédé standard qui permet pratiquement de ne démontrer qu'un des deux schémas de chaque paire — en vertu d'un principe général de dualité que nous présenterons plus loin (§7.7). Les démonstrations de l'idempotence de la disjonction (S52.5) et de la commutativité de la disjonction (S52.7) sont données dans la figure 7.7. Celles de l'associativité de la disjonction (S52.9) et de la distributivité de la conjonction par rapport à la disjonction (S52.11) sont données dans la figure 7.8.

Nous montrons ici comment démontrer la commutativité de la conjonction (S52.8) d'après celle de la disjonction et les schémas de De Morgan. La méthode est standard — introduire une double négation et appliquer ces schémas :

$$\begin{aligned}
 (A \text{ et } B) &\Leftrightarrow \text{nonnon}(A \text{ et } B) && \text{S51.1} \\
 &\Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B) && \text{S52.2} \\
 &\Leftrightarrow \text{non}(\text{non}B \text{ ou } \text{non}A) && \text{S52.7} \\
 &\Leftrightarrow (B \text{ et } A) && \text{S52.4.}
 \end{aligned}$$

La démonstration de l'idempotence (S52.6) et de l'associativité (S52.10) de la conjonction, et de la distributivité de la disjonction par rapport à la conjonction (S52.12) de cette manière, en utilisant S52.5/9/11 et De Morgan, est laissée au lecteur.

Γ		
1°	$A \Rightarrow A$	S19
2°	$(A \text{ ou } A) \Rightarrow A$	1°, 1°, S24
3°	$A \Rightarrow (A \text{ ou } A)$	S23
4°	$(A \text{ ou } A) \Leftrightarrow A$	2°, 3°, $\Leftrightarrow$ -intro
Γ		
1°	$A \Rightarrow (B \text{ ou } A)$	S23
2°	$B \Rightarrow (B \text{ ou } A)$	S23
3°	$(A \text{ ou } B) \Rightarrow (B \text{ ou } A)$	1°, 2°, S24
4°	$B \Rightarrow (A \text{ ou } B)$	S23
5°	$A \Rightarrow (A \text{ ou } B)$	S23
6°	$(B \text{ ou } A) \Rightarrow (A \text{ ou } B)$	4°, 5°, S24
7°	$(A \text{ ou } B) \Leftrightarrow (B \text{ ou } A)$	3°, 6°, $\Leftrightarrow$ -intro

Figure 7.7 Démonstrations de S52.5 et S52.7.

Emploi tacite de la commutativité de ou/et. Nous combinerons souvent une déduction suivant un schéma S_i avec une déduction suivant l'un ou l'autre des schémas de commutativité S52.7–8, en nous référant seulement au schéma S_i . Par exemple, nous nous référerons simplement à S52.13, c'est-à-dire à :

$$\frac{}{\Gamma \vdash ((A \text{ et } B) \text{ ou } B) \Leftrightarrow B,} \quad (2)$$

lorsque nous ferons une déduction sans prémisses de la forme :

$$\frac{}{\Gamma \vdash (B \text{ ou } (B \text{ et } A)) \Leftrightarrow B.}$$

Car une telle déduction se démontre par une chaîne d'équivalences évidente

Γ		
1°	$A \Rightarrow (A \text{ ou } B)$	S23
2°	$(A \text{ ou } B) \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	S23
3°	$A \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	1°, 2°, S20
4°	$B \Rightarrow (A \text{ ou } B)$	S23
5°	$B \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	4°, 2°, S20
6°	$C \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	S23
7°	$(B \text{ ou } C) \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	5°, 6°, S24
8°	$(A \text{ ou } (B \text{ ou } C)) \Rightarrow ((A \text{ ou } B) \text{ ou } C)$	3°, 7°, S24
:	Démonstration analogue de :	
16°	$((A \text{ ou } B) \text{ ou } C) \Rightarrow (A \text{ ou } (B \text{ ou } C))$	...
17°	$(A \text{ ou } (B \text{ ou } C)) \Leftrightarrow ((A \text{ ou } B) \text{ ou } C)$	8°, 16°, $\Leftrightarrow$ -intro.

Γ		
1°	$A \text{ et } (B \text{ ou } C)$	hyp
2°	A	1°, et-élim
3°	$B \text{ ou } C$	1°, et-élim
4°	B	hyp
5°	$A \text{ et } B$	2°, 4°, et-intro
6°	$(A \text{ et } B) \text{ ou } (A \text{ et } C)$	5°, ou-intro
7°	C	hyp
8°	$A \text{ et } C$	2°, 7°, et-intro
9°	$(A \text{ et } B) \text{ ou } (A \text{ et } C)$	8°, ou-intro
10°	$(A \text{ et } B) \text{ ou } (A \text{ et } C)$	3°, 6°, 9°, disj cas
11°	$(A \text{ et } B) \text{ ou } (A \text{ et } C)$	hyp
12°	$A \text{ et } B$	hyp
13°	A	12°, et-élim
14°	B	12°, et-élim
15°	$B \text{ ou } C$	14°, ou-intro
16°	$A \text{ et } (B \text{ ou } C)$	13°, 15°, et-intro
17°	$A \text{ et } C$	hyp
18°	A	17°, et-élim
19°	C	17°, et-élim
20°	$B \text{ ou } C$	19°, ou-intro
21°	$A \text{ et } (B \text{ ou } C)$	18°, 20°, et-intro
22°	$A \text{ et } (B \text{ ou } C)$	11°, 16°, 21°, disj cas
23°	$(A \text{ et } (B \text{ ou } C)) \Leftrightarrow ((A \text{ et } B) \text{ ou } (A \text{ et } C))$	10°, 22°, S40.

Figure 7.8 Démonstrations de S52.9 et S52.11.

suivant le schéma (2) et ceux de commutativité mentionnés :

$$\begin{aligned}
 (B \text{ ou } (B \text{ et } A)) &\Leftrightarrow ((B \text{ et } A) \text{ ou } B) && \text{S52.7} \\
 &\Leftrightarrow ((A \text{ et } B) \text{ ou } B) && \text{S52.8} \\
 &\Leftrightarrow B && (2).
 \end{aligned}$$

Convention d'association de ou/et. Étant donné des propositions A, B, C , nous convenons dorénavant que l'expression sans parenthèses

$$A \text{ ou } B \text{ ou } C \tag{3}$$

représente la proposition

$$A \text{ ou } (B \text{ ou } C),$$

c'est-à-dire, en notation préfixée: $\vee A \vee B \vee C$. C'est la convention d'association à droite pour la disjonction. En vertu de l'associativité de la disjonction (S52.9), la proposition représentée par (3) est équivalente à $((A \text{ ou } B) \text{ ou } C)$ dans tout contexte logique, de sorte que nous pourrions tout aussi bien convenir que (3) représente cette dernière proposition (convention d'association à gauche). Le choix est sans importance en vertu du principe général de substitutivité de l'équivalence (§7.5), suivant lequel on peut toujours remplacer une sous-proposition de la forme $(A \text{ ou } (B \text{ ou } C))$ par $((A \text{ ou } B) \text{ ou } C)$, ou vice-versa. On étend cette convention à un nombre quelconque de propositions $A, B, C, D, \dots$ en convenant que

$$A \text{ ou } B \text{ ou } C \text{ ou } D \tag{4}$$

représente $A \text{ ou } (B \text{ ou } C \text{ ou } D)$, c'est-à-dire $A \text{ ou } (B \text{ ou } (C \text{ ou } D))$, et ainsi de suite. Le schéma d'associativité S52.9 permet de parenthéser une disjonction multiple telle que (4) de manière arbitraire. Par exemple les propositions

$$(A \text{ ou } B) \text{ ou } (C \text{ ou } D), \quad A \text{ ou } (B \text{ ou } C) \text{ ou } D,$$

sont équivalentes à (4). En outre, en vertu de la commutativité de la disjonction (S52.7), on obtient des propositions équivalentes à (4) en permutant les propositions $A, B, C, D, \dots$ de manière arbitraire.

Nous adoptons des conventions semblables pour la conjonction.

Commentaire sur les schémas de distributivité. Pour appliquer correctement et rapidement les schémas de distributivité S52.11–12, il est utile, comme moyen mnémotechnique, de les mettre chacun en relation avec la loi bien familière de distributivité de la multiplication par rapport à l'addition des nombres réels :

$$\begin{array}{l}
 (A \text{ et } (B \text{ ou } C)) \Leftrightarrow ((A \text{ et } B) \text{ ou } (A \text{ et } C)) \\
 (A \text{ ou } (B \text{ et } C)) \Leftrightarrow ((A \text{ ou } B) \text{ et } (A \text{ ou } C)) \\
 \hline
 (a \cdot (b + c)) = ((a \cdot b) + (a \cdot c))
 \end{array}$$

Mais cette analogie de forme ne doit pas faire oublier que le symbole d'égalité est un symbole *relationnel*, de chaque côté duquel on a des *termes*, $a \cdot (b + c)$,

Γ		
1°	$(A \text{ et } B) \Rightarrow B$	S28
2°	$B \Rightarrow B$	S19
3°	$((A \text{ et } B) \text{ ou } B) \Rightarrow B$	1°, 2°, S24
4°	$B \Rightarrow ((A \text{ et } B) \text{ ou } B)$	S23
5°	$((A \text{ et } B) \text{ ou } B) \Leftrightarrow B$	3°, 4°, $\Leftrightarrow$ -intro.
1°	$\text{non}(A \text{ et non}A)$	S17
2°	$(A \text{ et non}A) \Rightarrow B$	1°, S22
3°	$B \Rightarrow B$	S19
4°	$((A \text{ et non}A) \text{ ou } B) \Rightarrow B$	2°, 3°, S24
5°	$B \Rightarrow ((A \text{ et non}A) \text{ ou } B)$	S23
6°	$((A \text{ et non}A) \text{ ou } B) \Leftrightarrow B$	4°, 5°, $\Leftrightarrow$ -intro.

Figure 7.9 Démonstrations de S52.13 et S52.15.

$(a \cdot b) + (a \cdot c)$, alors que le symbole d'équivalence est un connecteur logique, de chaque côté duquel on a des *propositions*.

Les schémas S52.11–12 doivent être appelées plus précisément schémas de distributivité à *gauche* de la conjonction par rapport à la disjonction et vice-versa. Les schémas de distributivité à *droite*, c'est-à-dire

$$\begin{aligned} & \Gamma \vdash ((B \text{ ou } C) \text{ et } A) \Leftrightarrow ((B \text{ et } A) \text{ ou } (C \text{ et } A)), \\ & \Gamma \vdash ((B \text{ et } C) \text{ ou } A) \Leftrightarrow ((B \text{ ou } A) \text{ et } (C \text{ ou } A)), \end{aligned}$$

dérivent immédiatement de S52.11–12 par des chaînes d'équivalence utilisant la commutativité de la conjonction et de la disjonction. Nous les emploierons en nous référant aussi à S52.11–12.

Exercice. Soient A, B, C, D des propositions d'un langage $\mathcal{L}$. Montrer par des chaînes d'équivalences que les deux propositions suivantes sont des tautologies.

$$((A \text{ ou } B) \text{ et } (C \text{ ou } D)) \Leftrightarrow ((A \text{ et } C) \text{ ou } (A \text{ et } D) \text{ ou } (B \text{ et } C) \text{ ou } (B \text{ et } D)).$$

$$((A \text{ et } B) \text{ ou } (C \text{ et } D)) \Leftrightarrow ((A \text{ ou } C) \text{ et } (A \text{ ou } D) \text{ et } (B \text{ ou } C) \text{ et } (B \text{ ou } D)).$$

Autres schémas booléens. Les quatre schémas booléens qui restent sont ceux d'absorption (S52.13–14) et de neutralité (S52.15–16). Deux d'entre eux sont démontrés dans la figure 7.9. On dit qu'une proposition $(A \text{ et } B)$ est *absorbée* par B dans une disjonction (S52.13), alors que $(A \text{ ou } B)$ est absorbée par B dans une conjonction (S52.14). Une proposition $(A \text{ et non}A)$ (fausse, S17) est *neutre* pour la disjonction (S52.15), alors que $(A \text{ ou non}A)$ (vraie, S34) est neutre pour la conjonction (S52.16).

7.7 La règle de dualité

Exemple. Soient A, B, C des propositions d'un langage $\mathcal{L}$, et Γ une liste quelconque de propositions de $\mathcal{L}$. Les schémas de déduction de De Morgan permettent de démontrer aisément, par chaîne d'équivalence, le théorème suivant

de $\text{Lprop}(\mathcal{L})$:

$$\Gamma \vdash \text{non}((A \text{ ou } B) \text{ et } C) \Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } \text{non}C). \quad (1)$$

Démonstration:

$$\begin{array}{l|l} \Gamma & \\ \hline \text{non}((A \text{ ou } B) \text{ et } C) \Leftrightarrow (\text{non}(A \text{ ou } B) \text{ ou } \text{non}C) & \text{S52.2} \\ & \Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } \text{non}C) \quad \text{S52.1.} \end{array}$$

Il est important d'observer dans cette chaîne comment la négation, initialement tout à gauche de la proposition de départ, se propage à l'intérieur de cette proposition jusque devant chacune des propositions A , B , C , et comment, en même temps, le OU est changé en ET et le ET est changé en OU.

Cette transformation est représentée dans la figure 7.10, au moyen des arbres des propositions. Le NON initial est attaché à la racine de l'arbre (a). Il descend le long des branches (arbres (b) et (c)) jusqu'au dessus des sous-arbres A , B , C . Lorsque cette négation se propage au travers d'un nœud ET, celui-ci est changé en OU, et inversement. La transformation globale, $(a) \rightarrow (c)$, peut s'obtenir en effectuant les opérations suivantes simultanément:

- 1° changer chaque OU en ET et chaque ET en OU;
- 2° remplacer A , B , C par $\text{non}A$, $\text{non}B$, $\text{non}C$.

Plus précisément, si nous désignons par R la proposition $((A \text{ ou } B) \text{ et } C)$, encadrée en pointillé dans la figure (a), l'opération 1° s'applique à cette proposition R , et le résultat de cette opération est la proposition $((A \text{ et } B) \text{ ou } C)$, figure (d), appelée la *duale de R*. L'opération 2° doit être appliquée à la duale de R . On voit que d'après les schémas de déduction de De Morgan, la proposition $\text{non}R$ est équivalente dans tout contexte logique à la proposition obtenue en remplaçant A , B , C par $\text{non}A$, $\text{non}B$, $\text{non}C$ dans la duale de R .

Schémas de De Morgan généralisés. Cet exemple se généralise, et l'on en tire des schémas de déduction qu'on appelle les *schémas de De Morgan généralisés*. Mais pour en donner un énoncé précis, il est important de remarquer encore une ou deux choses sur l'exemple qui précède.

Premièrement, on considère une proposition R qui est construite au moyen de propositions $A, B, \dots$ uniquement au moyen de connecteurs logiques de conjonction et/ou de disjonction, car les schémas de De Morgan S52.1–2 concernent ces connecteurs. Deuxièmement, il ne suffit pas de parler de la duale de R , car cette dénomination est insuffisamment précise. Considérons en effet la proposition R de la figure 7.10(a). Les propositions A , B , C sont indéterminées — c'est pourquoi elles sont représentées par des triangles symbolisant un arbre indéterminé. Il se peut que ces propositions soient elles-mêmes construites au moyen d'autres propositions avec des connecteurs ET/OU. Dans ce cas, les transformations $(a) \rightarrow (b) \rightarrow (c)$ de la figure pourraient se poursuivre. Les négations pourraient descendre plus bas. Lorsqu'on parle de la *duale de R*, il faut donc savoir où l'on s'arrête dans ces transformations. C'est

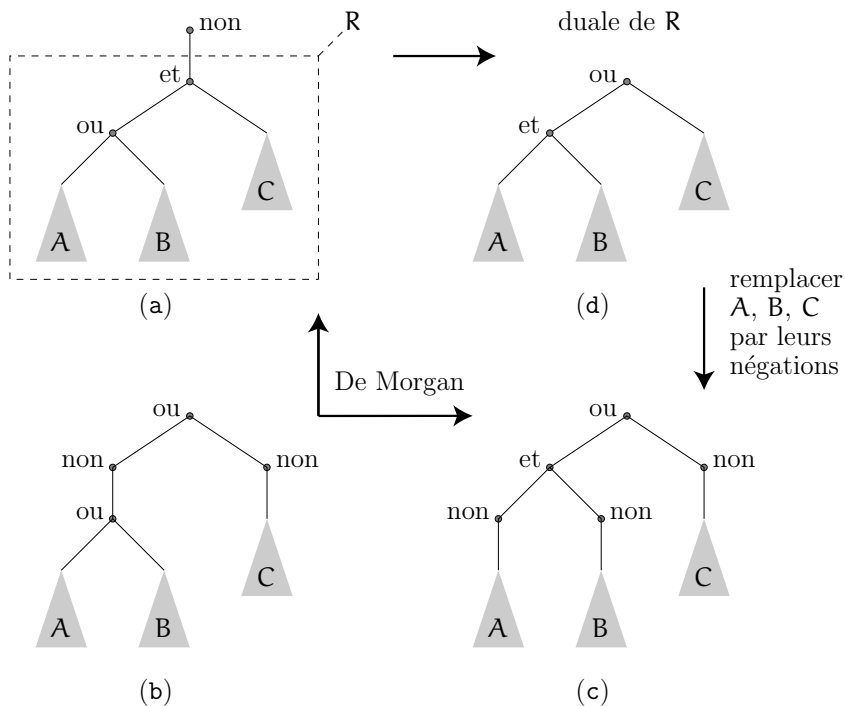


Figure 7.10 Illustration du schéma de De Morgan généralisé: la proposition $\text{non}R$ est équivalente à la proposition obtenue en remplaçant A, B, C par leurs négations dans la duale de R .

pourquoi nous dirons que la proposition de la figure 7.10(c) est la duale de R par rapport aux propositions A, B, C .

Pour formuler les schémas de De Morgan généralisés, nous utiliserons les notations suivantes. Si $A, B, \dots$ désignent des propositions d'un langage $\mathcal{L}$, la notation $R[A, B, \dots]$ désigne une proposition construite à partir de $A, B, \dots$ par des conjonctions et/ou des disjonctions uniquement, et $R^D[A, B, \dots]$ désigne la duale de R par rapport à $A, B, \dots$, c'est-à-dire la proposition obtenue en échangeant les connecteurs de disjonction et de conjonction dans la construction de R à partir de $A, B, \dots$. Enfin $R^D[\text{non}A, \text{non}B, \dots]$ et $R[\text{non}A, \text{non}B, \dots]$ désignent les propositions obtenues en remplaçant $A, B, \dots$ par $\text{non}A, \text{non}B, \dots$ dans R^D et dans R . Si n est le nombre de connecteurs (et, ou) utilisés pour construire $R[A, B, \dots]$ à partir de $A, B, \dots$ alors, en appliquant n fois les schémas de De Morgan S52.1–2, on peut démontrer l'équivalence

$$\text{non}R[A, B, \dots] \Leftrightarrow R^D[\text{non}A, \text{non}B, \dots] \quad (2)$$

dans n'importe quel contexte logique, comme nous l'avons fait dans l'exemple

précédent (1). Nous pouvons donc admettre le schéma de déduction général :

$$\frac{}{\Gamma \vdash \text{non}R[A, B, \dots] \Leftrightarrow R^D[\text{non}A, \text{non}B, \dots]} \quad (3)$$

Remarquons ensuite que la duale de la duale d'une proposition $R[A, B, \dots]$ est $R[A, B, \dots]$ elle-même. Donc, pouvant démontrer l'équivalence (2) pour n'importe quelle proposition $R[A, B, \dots]$, on peut aussi démontrer l'équivalence $\text{non}R^D[A, B, \dots] \Leftrightarrow (R^D)^D[\text{non}A, \text{non}B, \dots]$, c'est-à-dire

$$\text{non}R^D[A, B, \dots] \Leftrightarrow R[\text{non}A, \text{non}B, \dots].$$

Nous pouvons donc admettre le deuxième schéma de déduction général :

$$\frac{}{\Gamma \vdash \text{non}R^D[A, B, \dots] \Leftrightarrow R[\text{non}A, \text{non}B, \dots]} \quad (4)$$

Les schémas (3) et (4) sont les schémas de De Morgan généralisés.

Exercice. Démontrer les déductions (3) et (4) par des chaînes d'équivalences analogues à (1) dans le cas d'une proposition $R[A, B, C, D]$ de la forme

$$((A \text{ et } B) \text{ ou } C) \text{ et } D.$$

La règle de dualité. Les schémas de De Morgan généralisés (3) et (4) permettent d'expliquer la règle de dualité que nous avons observée dans notre répertoire de schémas de déduction booléens S52, à savoir que ces schémas vont par paires: chacun d'eux est accompagné de son dual. Ce que nous appelons deux schémas duals, ce sont deux schémas de déduction sans prémisses de la forme

$$S[A, B, \dots].$$

$$\frac{}{\Gamma \vdash R_1[A, B, \dots] \Leftrightarrow R_2[A, B, \dots]}$$

$$S^D[A, B, \dots].$$

$$\frac{}{\Gamma \vdash R_1^D[A, B, \dots] \Leftrightarrow R_2^D[A, B, \dots]}$$

où $R_1[A, B, \dots]$, $R_2[A, B, \dots]$ sont deux propositions construites à partir de propositions $A, B, \dots$ par des conjonctions et/ou des disjonctions uniquement. Là encore, lorsque nous parlons de schémas duals, nous devons préciser par rapport à quelles propositions $A, B, \dots$. Par exemple, les schémas de neutralité (S52.15–16)

$$\frac{}{\Gamma \vdash ((A \text{ et } \underbrace{\text{non}A}_C) \text{ ou } B) \Leftrightarrow B}$$

$$\frac{}{\Gamma \vdash ((\underbrace{A \text{ ou } \text{non}A}_C) \text{ et } B) \Leftrightarrow B}$$

sont duals par rapport à A, B, C où C est la proposition $\text{non}A$. Dans cet exemple, $R_2[A, B, C]$ est simplement B , et elle est identique à sa duale.

L'explication de la règle ou principe de dualité est que le schéma dual $S^D[A, B, \dots]$ d'un schéma $S[A, B, \dots]$ *dérive* de celui-ci en vertu des schémas de

De Morgan généralisés. En utilisant le schéma $S[A, B, \dots]$, nous pouvons écrire en effet la chaîne d'équivalences:

$$\begin{array}{l|l} \Gamma & \\ \hline R_1^D[A, B, \dots] \Leftrightarrow \text{non}(\text{non}R_1^D[A, B, \dots]) & S51.1 \\ & \Leftrightarrow \text{non}(R_1[\text{non}A, \text{non}B, \dots]) & \text{De Morgan gén-(4)} \\ & \Leftrightarrow \text{non}(R_2[\text{non}A, \text{non}B, \dots]) & \text{schéma } S[A, B, \dots] \\ & \Leftrightarrow \text{non}(\text{non}R_2^D[A, B, \dots]) & \text{De Morgan gén-(4)} \\ & \Leftrightarrow R_2^D[A, B, \dots] & S51.1. \end{array}$$

On applique ici le schéma $S[A, B, \dots]$ aux propositions $\text{non}A, \text{non}B, \dots$, ce qui est légitime, puisque ce schéma peut être appliqué à des propositions $A, B, \dots$ quelconques.

7.8 Formes booléennes de l'implication et de l'équivalence

Les démonstrations de tous les schémas de déduction présentés dans ce paragraphe sont rassemblées dans la figure 7.11. Les schémas S53 à S56 permettent d'exprimer une implication au moyen de connecteurs logiques de négation et de disjonction (S53) ou de négation et de conjonction (S55), et inversement, d'exprimer une disjonction ou une conjonction au moyen de connecteurs de négation et d'implication (S54, S56). Les schémas S59 permettent d'exprimer une équivalence avec des connecteurs booléens (non, et, ou) de deux manières différentes. Les schémas de contraposition et de distribution S57, S58 n'ont pas de rapport avec le thème principal de ce paragraphe. Ils sont présentés à cet endroit car leur démonstration est une application facile de S53.

Expression booléenne de l'implication. La base de tous les schémas de déduction de ce paragraphe est le schéma suivant.

S53.

$$\frac{}{\Gamma \vdash (A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)}.$$

En vertu de ce schéma et du schéma général de substitutivité de l'équivalence (§7.5), on pourra toujours remplacer une sous-proposition de la forme $A \Rightarrow B$ dans une proposition R par $(\text{non}A \text{ ou } B)$ et obtenir ainsi une proposition R' équivalente à R .

L'équivalence de $A \Rightarrow B$ et de $(\text{non}A \text{ ou } B)$ n'est pas toujours bien perçue intuitivement — et souvent même elle ne l'est pas du tout. Celui qui n'a pas étudié la logique formelle et qui se base seulement sur le sens qu'ont pour lui les mots « non », « ou », « si... alors » peut ne pas sentir que des propositions de la forme $A \Rightarrow B$ (si A alors B) et $(\text{non}A \text{ ou } B)$ sont logiquement équivalentes.

Nous voulons donc discuter cette question en considérant un exemple. La figure 7.12 représente une cible avec deux cercles concentriques A et B , le

Γ																												
(S53)	<table><tr><td>1°</td><td>$A \Rightarrow B$</td><td>hyp</td></tr><tr><td>2°</td><td>$\text{non}A \Rightarrow \text{non}A$</td><td>S19</td></tr><tr><td>3°</td><td>$B \text{ ou } \text{non}A$</td><td>1°, 2°, S36 (2ème)</td></tr><tr><td>4°</td><td>$\text{non}A \text{ ou } B$</td><td>3°, S52.7, S39</td></tr><tr><td>5°</td><td>$\text{non}A \text{ ou } B$</td><td>hyp</td></tr><tr><td>6°</td><td>A</td><td>hyp</td></tr><tr><td>7°</td><td>B</td><td>5°, 6°, S18 (2ème)</td></tr><tr><td>8°</td><td>$A \Rightarrow B$</td><td>7°, $\Rightarrow$-intro</td></tr><tr><td>9°</td><td>$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$</td><td>4°, 8°, S40.</td></tr></table>	1°	$A \Rightarrow B$	hyp	2°	$\text{non}A \Rightarrow \text{non}A$	S19	3°	$B \text{ ou } \text{non}A$	1°, 2°, S36 (2ème)	4°	$\text{non}A \text{ ou } B$	3°, S52.7, S39	5°	$\text{non}A \text{ ou } B$	hyp	6°	A	hyp	7°	B	5°, 6°, S18 (2ème)	8°	$A \Rightarrow B$	7°, $\Rightarrow$ -intro	9°	$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	4°, 8°, S40.
1°	$A \Rightarrow B$	hyp																										
2°	$\text{non}A \Rightarrow \text{non}A$	S19																										
3°	$B \text{ ou } \text{non}A$	1°, 2°, S36 (2ème)																										
4°	$\text{non}A \text{ ou } B$	3°, S52.7, S39																										
5°	$\text{non}A \text{ ou } B$	hyp																										
6°	A	hyp																										
7°	B	5°, 6°, S18 (2ème)																										
8°	$A \Rightarrow B$	7°, $\Rightarrow$ -intro																										
9°	$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	4°, 8°, S40.																										
(S54)	<table><tr><td>$(A \text{ ou } B) \Leftrightarrow (\text{nonnon}A \text{ ou } B)$</td><td>S51.1</td></tr><tr><td>$\Leftrightarrow (\text{non}A \Rightarrow B)$</td><td>S53.</td></tr></table>	$(A \text{ ou } B) \Leftrightarrow (\text{nonnon}A \text{ ou } B)$	S51.1	$\Leftrightarrow (\text{non}A \Rightarrow B)$	S53.																							
$(A \text{ ou } B) \Leftrightarrow (\text{nonnon}A \text{ ou } B)$	S51.1																											
$\Leftrightarrow (\text{non}A \Rightarrow B)$	S53.																											
(S55)	<table><tr><td>$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$</td><td>S53</td></tr><tr><td>$\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } \text{non}B)$</td><td>De Morgan (S52.3)</td></tr><tr><td>$\Leftrightarrow \text{non}(A \text{ et } \text{non}B)$</td><td>S51.1.</td></tr></table>	$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	S53	$\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } \text{non}B)$	De Morgan (S52.3)	$\Leftrightarrow \text{non}(A \text{ et } \text{non}B)$	S51.1.																					
$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	S53																											
$\Leftrightarrow \text{non}(\text{nonnon}A \text{ et } \text{non}B)$	De Morgan (S52.3)																											
$\Leftrightarrow \text{non}(A \text{ et } \text{non}B)$	S51.1.																											
(S56)	<table><tr><td>$(A \text{ et } B) \Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B)$</td><td>De Morgan (S52.4)</td></tr><tr><td>$\Leftrightarrow \text{non}(A \Rightarrow \text{non}B)$</td><td>S53.</td></tr></table>	$(A \text{ et } B) \Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B)$	De Morgan (S52.4)	$\Leftrightarrow \text{non}(A \Rightarrow \text{non}B)$	S53.																							
$(A \text{ et } B) \Leftrightarrow \text{non}(\text{non}A \text{ ou } \text{non}B)$	De Morgan (S52.4)																											
$\Leftrightarrow \text{non}(A \Rightarrow \text{non}B)$	S53.																											
(S57)	<table><tr><td>$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$</td><td>S53</td></tr><tr><td>$\Leftrightarrow (B \text{ ou } \text{non}A)$</td><td>S52.7</td></tr><tr><td>$\Leftrightarrow (\text{nonnon}B \text{ ou } \text{non}A)$</td><td>S51.1</td></tr><tr><td>$\Leftrightarrow (\text{non}B \Rightarrow \text{non}A)$</td><td>S53.</td></tr></table>	$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	S53	$\Leftrightarrow (B \text{ ou } \text{non}A)$	S52.7	$\Leftrightarrow (\text{nonnon}B \text{ ou } \text{non}A)$	S51.1	$\Leftrightarrow (\text{non}B \Rightarrow \text{non}A)$	S53.																			
$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	S53																											
$\Leftrightarrow (B \text{ ou } \text{non}A)$	S52.7																											
$\Leftrightarrow (\text{nonnon}B \text{ ou } \text{non}A)$	S51.1																											
$\Leftrightarrow (\text{non}B \Rightarrow \text{non}A)$	S53.																											
(S58.1)	<table><tr><td>$((A \text{ ou } B) \Rightarrow C) \Leftrightarrow (\text{non}(A \text{ ou } B) \text{ ou } C)$</td><td>S53</td></tr><tr><td>$\Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } C)$</td><td>De Morgan</td></tr><tr><td>$\Leftrightarrow ((\text{non}A \text{ ou } C) \text{ et } (\text{non}B \text{ ou } C))$</td><td>S52.12</td></tr><tr><td>$\Leftrightarrow ((A \Rightarrow C) \text{ et } (B \Rightarrow C))$</td><td>S53.</td></tr></table>	$((A \text{ ou } B) \Rightarrow C) \Leftrightarrow (\text{non}(A \text{ ou } B) \text{ ou } C)$	S53	$\Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } C)$	De Morgan	$\Leftrightarrow ((\text{non}A \text{ ou } C) \text{ et } (\text{non}B \text{ ou } C))$	S52.12	$\Leftrightarrow ((A \Rightarrow C) \text{ et } (B \Rightarrow C))$	S53.																			
$((A \text{ ou } B) \Rightarrow C) \Leftrightarrow (\text{non}(A \text{ ou } B) \text{ ou } C)$	S53																											
$\Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } C)$	De Morgan																											
$\Leftrightarrow ((\text{non}A \text{ ou } C) \text{ et } (\text{non}B \text{ ou } C))$	S52.12																											
$\Leftrightarrow ((A \Rightarrow C) \text{ et } (B \Rightarrow C))$	S53.																											
(S58.2)	<table><tr><td>$(C \Rightarrow (A \text{ et } B)) \Leftrightarrow (\text{non}C \text{ ou } (A \text{ et } B))$</td><td>S53</td></tr><tr><td>$\Leftrightarrow ((\text{non}C \text{ ou } A) \text{ et } (\text{non}C \text{ ou } B))$</td><td>S52.12</td></tr><tr><td>$\Leftrightarrow ((C \Rightarrow A) \text{ et } (C \Rightarrow B))$</td><td>S53.</td></tr></table>	$(C \Rightarrow (A \text{ et } B)) \Leftrightarrow (\text{non}C \text{ ou } (A \text{ et } B))$	S53	$\Leftrightarrow ((\text{non}C \text{ ou } A) \text{ et } (\text{non}C \text{ ou } B))$	S52.12	$\Leftrightarrow ((C \Rightarrow A) \text{ et } (C \Rightarrow B))$	S53.																					
$(C \Rightarrow (A \text{ et } B)) \Leftrightarrow (\text{non}C \text{ ou } (A \text{ et } B))$	S53																											
$\Leftrightarrow ((\text{non}C \text{ ou } A) \text{ et } (\text{non}C \text{ ou } B))$	S52.12																											
$\Leftrightarrow ((C \Rightarrow A) \text{ et } (C \Rightarrow B))$	S53.																											
(S59.1)	<table><tr><td>$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$</td><td>S41</td></tr><tr><td>$\Leftrightarrow ((\text{non}A \text{ ou } B) \text{ et } (\text{non}B \text{ ou } A))$</td><td>S53.</td></tr></table>	$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$	S41	$\Leftrightarrow ((\text{non}A \text{ ou } B) \text{ et } (\text{non}B \text{ ou } A))$	S53.																							
$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$	S41																											
$\Leftrightarrow ((\text{non}A \text{ ou } B) \text{ et } (\text{non}B \text{ ou } A))$	S53.																											
(S59.2)	<table><tr><td>$(A \Leftrightarrow B) \Leftrightarrow ((\neg A \text{ ou } B) \text{ et } (\neg B \text{ ou } A))$</td><td>S59.1</td></tr><tr><td>$\Leftrightarrow ((\neg A \text{ ou } B) \text{ et } \neg B) \text{ ou } ((\neg A \text{ ou } B) \text{ et } A)$</td><td>S52.11</td></tr><tr><td>$\Leftrightarrow (\neg B \text{ et } (\neg A \text{ ou } B)) \text{ ou } (A \text{ et } (\neg A \text{ ou } B))$</td><td>S52.8</td></tr><tr><td>$\Leftrightarrow ((\neg B \text{ et } \neg A) \text{ ou } (\neg B \text{ et } B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$</td><td>S52.11</td></tr><tr><td>$\Leftrightarrow ((B \text{ et } \neg B) \text{ ou } (\neg A \text{ et } \neg B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$</td><td>S52.7–8</td></tr><tr><td>$\Leftrightarrow ((\neg A \text{ et } \neg B) \text{ ou } (A \text{ et } B))$</td><td>S52.15.</td></tr></table>	$(A \Leftrightarrow B) \Leftrightarrow ((\neg A \text{ ou } B) \text{ et } (\neg B \text{ ou } A))$	S59.1	$\Leftrightarrow ((\neg A \text{ ou } B) \text{ et } \neg B) \text{ ou } ((\neg A \text{ ou } B) \text{ et } A)$	S52.11	$\Leftrightarrow (\neg B \text{ et } (\neg A \text{ ou } B)) \text{ ou } (A \text{ et } (\neg A \text{ ou } B))$	S52.8	$\Leftrightarrow ((\neg B \text{ et } \neg A) \text{ ou } (\neg B \text{ et } B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$	S52.11	$\Leftrightarrow ((B \text{ et } \neg B) \text{ ou } (\neg A \text{ et } \neg B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$	S52.7–8	$\Leftrightarrow ((\neg A \text{ et } \neg B) \text{ ou } (A \text{ et } B))$	S52.15.															
$(A \Leftrightarrow B) \Leftrightarrow ((\neg A \text{ ou } B) \text{ et } (\neg B \text{ ou } A))$	S59.1																											
$\Leftrightarrow ((\neg A \text{ ou } B) \text{ et } \neg B) \text{ ou } ((\neg A \text{ ou } B) \text{ et } A)$	S52.11																											
$\Leftrightarrow (\neg B \text{ et } (\neg A \text{ ou } B)) \text{ ou } (A \text{ et } (\neg A \text{ ou } B))$	S52.8																											
$\Leftrightarrow ((\neg B \text{ et } \neg A) \text{ ou } (\neg B \text{ et } B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$	S52.11																											
$\Leftrightarrow ((B \text{ et } \neg B) \text{ ou } (\neg A \text{ et } \neg B)) \text{ ou } ((A \text{ et } \neg A) \text{ ou } (A \text{ et } B))$	S52.7–8																											
$\Leftrightarrow ((\neg A \text{ et } \neg B) \text{ ou } (A \text{ et } B))$	S52.15.																											

Figure 7.11 Démonstrations des schémas S53 à S59.

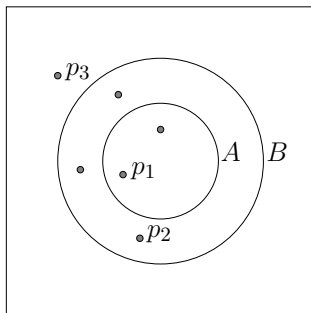


Figure 7.12 Propositions équivalentes:

- (a) si p est dans A alors p est dans B ;
- (b) p n'est pas dans A ou p est dans B .

premier étant à l'intérieur du second. La cible porte un certain nombre de marques d'impacts de projectiles, dont trois sont désignées par p_1, p_2, p_3 .

Chacun reconnaîtra sans doute la vérité de l'assertion suivante: si un point d'impact est à l'intérieur du cercle A , alors il est à l'intérieur du cercle B . On peut donner à cette assertion une tournure plus formelle en utilisant une variable individuelle, disons p , pour désigner un point d'impact *indéterminé*, et en affirmant: si p est à l'intérieur de A , alors p est à l'intérieur de B , ce que nous abrègerons par

$$p \text{ int } A \Rightarrow p \text{ int } B. \quad (1)$$

Là où le sens intuitif des mots «si... alors» peut se révéler insuffisant, c'est lorsqu'on *particularise* l'assertion (1) pour un point *déterminé*, c'est-à-dire lorsqu'on remplace p par la désignation d'un point connu, en écrivant par exemple: si p_3 est à l'intérieur du cercle A , alors p_3 est à l'intérieur du cercle B . Le lecteur peut très bien ne pas saisir le sens de cette phrase — vu la position du point p_3 — ni par conséquent savoir si elle est vraie ou fausse. Il en va ainsi des trois propositions

$$\begin{aligned} p_3 \text{ int } A &\Rightarrow p_3 \text{ int } B, \\ p_2 \text{ int } A &\Rightarrow p_2 \text{ int } B, \\ p_1 \text{ int } A &\Rightarrow p_1 \text{ int } B. \end{aligned} \quad (2)$$

Même la troisième sonne un peu bizarrement avec son conditionnel «si p_1 est à l'intérieur de A », puisque p_1 est à l'intérieur de A . Celui qui n'est pas familiarisé avec la logique formelle doit donc *apprendre* à reconnaître sans hésitation la vérité de ces trois propositions.

Revenons donc à l'assertion générale (1) et essayons de la transformer intuitivement en une forme équivalente permettant de l'interpréter pour chaque point p particulier. En fait, cette forme sera justement celle qui est fournie par S53, à savoir:

$$\text{non}(p \text{ int } A) \text{ ou } (p \text{ int } B).$$

Pour simplifier les expressions, nous remplacerons $\text{non}(p \text{ int } A)$ par la forme équivalente « $p \text{ ext } A$ », signifiant: p est à l'extérieur du cercle A .

Nous admettons que le sens de (1) peut être exprimé de manière complète comme ceci:

si p est à l'intérieur de A alors p est à l'intérieur de B ,
et si p est à l'extérieur de A , alors p peut être soit à
l'intérieur de B , soit à l'extérieur de B .

Mais ce dernier « alors » est trivialement vrai: p est toujours soit à l'intérieur soit à l'extérieur de B (on néglige l'épaisseur des cercles). On peut donc simplifier la description des choses en disant que pour un point d'impact p il y a deux cas possibles: p est à l'extérieur de A , ou alors p est simultanément à l'intérieur de A et à l'intérieur de B :

$$(p \text{ ext } A) \text{ ou } (p \text{ int } A \text{ et } p \text{ int } B). \quad (3)$$

Il nous paraît assez naturel d'admettre l'équivalence de (1) et (3). Cela permet de reconnaître aussitôt la vérité des assertions (2), en les mettant chacune sous la forme équivalente (3). Car pour chacun des trois points p_1, p_2, p_3 l'un au moins des membres de la disjonction (3) est vrai — on a: $p_3 \text{ ext } A, p_2 \text{ ext } A, (p_1 \text{ int } A \text{ et } p_1 \text{ int } B)$.

Il reste à reconnaître que l'on peut simplifier la proposition (3), à savoir qu'elle équivaut à

$$(p \text{ ext } A) \text{ ou } (p \text{ int } B). \quad (4)$$

De manière générale, pour des propositions A, B quelconques, on peut écrire la démonstration suivante dans n'importe quel contexte logique:

$$\begin{aligned} (\text{non}A \text{ ou } (A \text{ et } B)) &\Leftrightarrow ((\text{non}A \text{ ou } A) \text{ et } (\text{non}A \text{ ou } B)) & \text{S52.12} \\ &\Leftrightarrow (\text{non}A \text{ ou } B) & \text{S52.16.} \end{aligned}$$

L'équivalence de (1) et (4) peut être reconnue, dans cet exemple, par des considérations géométriques. Il est clair que la proposition (1) est vraie pour un point p quelconque parce que le cercle A est à l'intérieur du cercle B . En vertu de cette configuration, chaque point de la cible est soit à l'extérieur de A soit à l'intérieur de B (figure 7.13). Les domaines constitués par l'extérieur de A et l'intérieur de B recouvrent ensemble toute la cible. Ce ne serait pas le cas si le cercle A n'était pas à l'intérieur du cercle B .

Les considérations qui précèdent ne remplacent pas une démonstration formelle de S53 (figure 7.11). Elles ne visent qu'à montrer la compatibilité de ce schéma avec le sens intuitif de l'implication, et à développer un peu ce sens. Dorénavant, le lecteur est invité à interpréter toute implication $A \Rightarrow B$ comme ayant le même sens que $(\text{non}A \text{ ou } B)$, en particulier à interpréter (1) comme (4). Avec cette interprétation, les schémas

$$\begin{array}{ll} \text{S21} & \frac{\Gamma \vdash B}{\Gamma \vdash A \Rightarrow B} \qquad \text{S22} \quad \frac{\Gamma \vdash \text{non}A}{\Gamma \vdash A \Rightarrow B} \end{array}$$

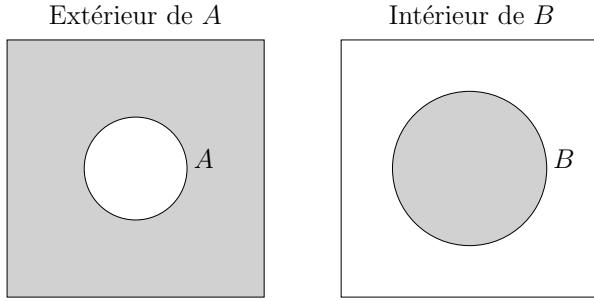


Figure 7.13 Propositions équivalentes:

- (a) si p est dans A alors p est dans B ;
- (b) p n'est pas dans A ou p est dans B .

ne sont qu'un cas particulier de l'introduction de ou (S5). On peut dire que, dans notre exemple, la proposition $p_3 \text{ int } A \Rightarrow p_3 \text{ int } B$ est vraie en vertu de S22; la proposition $p_2 \text{ int } A \Rightarrow p_2 \text{ int } B$ est vraie aussi bien en vertu de S21 que de S22; enfin $p_1 \text{ int } A \Rightarrow p_1 \text{ int } B$ est vraie en vertu de S21.

Expression d'une disjonction au moyen d'une implication. Le schéma suivant dérive immédiatement de S53 (voir figure 7.11):

S54.

$$\frac{}{\Gamma \vdash (A \text{ ou } B) \Leftrightarrow (\text{non}A \Rightarrow B)}.$$

En vertu de ce schéma, pour démontrer une proposition de la forme $(A \text{ ou } B)$ dans un contexte donné, il suffit de démontrer $\text{non}A \Rightarrow B$ et de conclure en invoquant S54 et S39. C'est *très souvent ainsi* que l'on procède pour démontrer une disjonction. Car pour démontrer $\text{non}A \Rightarrow B$, il suffit de démontrer B sous l'hypothèse supplémentaire $\text{non}A$, ce qui simplifie en général le problème de démonstration. Cette méthode est suffisamment importante pour que nous lui consacrons une figure (fig. 7.14).

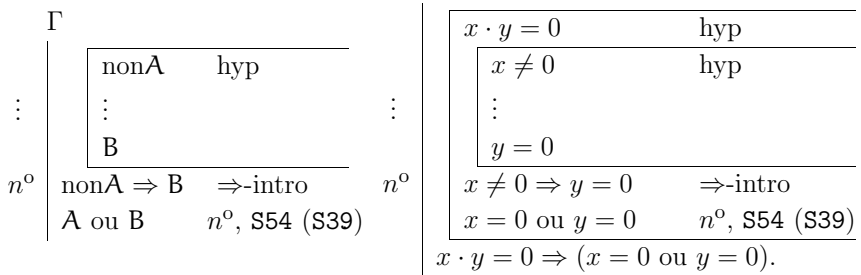


Figure 7.14 Méthode de démonstration d'une disjonction et exemple.

Cette méthode est illustrée par le plan de démonstration du théorème

$$\vdash (x \cdot y = 0) \Rightarrow (x = 0 \text{ ou } y = 0)$$

sur les nombres réels. La démarche naturelle est de faire l'hypothèse $x \cdot y = 0$, puis de montrer (sous cette hypothèse) que si $x \neq 0$ alors $y = 0$.

Négation d'une implication. Les schémas S55 suivants dérivent aussi immédiatement de S53. La démonstration du premier se trouve dans la figure 7.11. Le second dérive évidemment du premier, d'après la démonstration

$$\begin{aligned} \text{non}(A \Rightarrow B) &\Leftrightarrow \text{nonnon}(A \text{ et non}B) && \text{S55 (1er)} \\ &\Leftrightarrow (A \text{ et non}B) && \text{S51.1.} \end{aligned}$$

S55.

$$\Gamma \vdash (A \Rightarrow B) \Leftrightarrow \text{non}(A \text{ et non}B)$$

$$\Gamma \vdash \text{non}(A \Rightarrow B) \Leftrightarrow (A \text{ et non}B).$$

Le second de ces schémas est le plus fréquemment utilisé, car on a très souvent l'occasion de prendre la négation d'une implication. Par exemple (en anticipant sur le chapitre 10), la proposition d'inclusion d'un ensemble a dans un ensemble b , notée $a \subset b$, signifie que tout élément de a est élément de b , et cela s'exprime par le jugement

$$\vdash a \subset b \Leftrightarrow \forall x(x \in a \Rightarrow x \in b). \quad (5)$$

Intuitivement, un ensemble a n'est pas inclus dans un ensemble b , s'il n'est pas vrai que tout élément de a est élément de b , autrement dit s'il existe un élément de a qui n'est pas élément de b . On devrait donc pouvoir déduire de (5) le jugement:

$$\vdash a \not\subset b \Leftrightarrow \exists x(x \in a \text{ et } x \notin b). \quad (6)$$

C'est possible en effet, en trois pas de déduction, le deuxième relevant de la logique des prédicats:

$$\begin{aligned} \text{non}(a \subset b) &\Leftrightarrow \text{non}\forall x(x \in a \Rightarrow x \in b) && (5), \text{S50.1} \\ &\Leftrightarrow \exists x \text{non}(x \in a \Rightarrow x \in b) && \text{logique des prédicats} \\ &\Leftrightarrow \exists x(x \in a \text{ et } x \notin b) && \text{S55 (2ème).} \end{aligned}$$

Le schéma de déduction de logique des prédicats utilisé ici dit que la négation d'une proposition de la forme $\forall xR$ est équivalente à $\exists x \text{non}R$, ce qui est assez intuitif.

Expression d'une conjonction par une implication. Une conjonction peut s'exprimer au moyen de connecteurs d'implication et de négation, d'après le schéma suivant démontré dans la figure 7.11, de même qu'une disjonction d'après S54:

S56.

$$\Gamma \vdash (A \text{ et } B) \Leftrightarrow \text{non}(A \Rightarrow \text{non}B).$$

Contraposition. Le schéma de *contraposition (bis)* suivant est également démontré dans la figure 7.11. Il est plus fort que le schéma de contraposition S31, puisqu'il permet non seulement de déduire $\Gamma \vdash \text{non}B \Rightarrow \text{non}A$ de $\Gamma \vdash A \Rightarrow B$, en utilisant S39, mais aussi de faire la déduction inverse.

S57. *Contraposition (bis)*

$$\frac{}{\Gamma \vdash (A \Rightarrow B) \Leftrightarrow (\text{non}B \Rightarrow \text{non}A)}.$$

Distribution d'une implication. Nous appelons les deux schémas qui suivent les *schémas de distribution d'une implication par rapport à une conjonction, resp. par rapport à une disjonction*. Le second pourrait s'appeler d'ailleurs *distributivité à gauche* de l'implication par rapport à la conjonction. Mais pour le premier, il n'est pas indiqué de parler de distributivité (à droite), puisque la disjonction est changée en une conjonction lorsqu'on passe du membre de gauche de l'équivalence à celui de droite. C'est pourquoi nous parlons simplement d'un schéma de distribution. Ces schémas sont présentés dans ce paragraphe parce que leur démonstration (fig. 7.11) est une application directe, typique, de l'équivalence de $A \Rightarrow B$ avec $(\text{non}A \text{ ou } B)$.

S58.1.

$$\frac{}{\Gamma \vdash ((A \text{ ou } B) \Rightarrow C) \Leftrightarrow ((A \Rightarrow C) \text{ et } (B \Rightarrow C))}.$$

S58.2.

$$\frac{}{\Gamma \vdash (C \Rightarrow (A \text{ et } B)) \Leftrightarrow ((C \Rightarrow A) \text{ et } (C \Rightarrow B))}.$$

Formes booléennes de l'équivalence. Les deux derniers schémas de ce paragraphe fournissent deux manières d'exprimer une équivalence $A \Leftrightarrow B$ au moyen de connecteurs booléens (et, ou, non). La première est dite de forme *conjonctive*, car son connecteur principal est un ET. La seconde est dite de forme *disjonctive*. Son connecteur principal est un OU. Les démonstrations de ces schémas se trouvent dans la figure 7.11.

S59.

1.
$$\frac{}{\Gamma \vdash (A \Leftrightarrow B) \Leftrightarrow ((\text{non}A \text{ ou } B) \text{ et } (\text{non}B \text{ ou } A))}.$$
2.
$$\frac{}{\Gamma \vdash (A \Leftrightarrow B) \Leftrightarrow ((\text{non}A \text{ et } \text{non}B) \text{ ou } (A \text{ et } B))}.$$

Pour illustrer le deuxième de ces schémas, considérons une proposition d'égalité $a = b$ entre deux ensembles a et b . Par définition, deux ensembles a et b sont égaux si et seulement s'ils ont les mêmes éléments, et cela s'exprime par le jugement

$$\vdash a = b \Leftrightarrow \forall x(x \in a \Leftrightarrow x \in b). \quad (7)$$

Le second membre de cette équivalence peut s'exprimer en disant que, quel que soit l'objet x , x est élément de a si et seulement si x est élément de b .

D'après S59.2, on déduit de (7) le jugement

$$\vdash a = b \Leftrightarrow \forall x((x \notin a \text{ et } x \notin b) \text{ ou } (x \in a \text{ et } x \in b)). \quad (8)$$

Le membre de droite de cette autre équivalence signifie: un objet x quelconque n'appartient ni à l'un ni à l'autre des ensembles a , b ou alors il appartient aux deux à la fois. Intuitivement, c'est bien une autre manière d'exprimer que ces ensembles ont les mêmes éléments.

7.9 Simplifications conditionnelles

Les schémas de déduction rassemblés dans ce paragraphe, S60.1 à S60.8, possèdent tous *une* prémisses. Leur conclusion est chaque fois une équivalence dont le membre de droite est plus simple que celui de gauche. Ces schémas sont souvent utilisés dans les démonstrations en forme de chaînes d'équivalences, où ils permettent de simplifier les expressions. Mais ces simplifications sont *conditionnelles* dans le sens suivant: elles peuvent s'effectuer dans un contexte donné à *condition* que la proposition de la prémisse du schéma soit vraie dans ce contexte.

S60.

- | | |
|--|--|
| 1. $\frac{\Gamma \vdash A}{\Gamma \vdash (A \text{ et } B) \Leftrightarrow B}$ | 2. $\frac{\Gamma \vdash \text{non}A}{\Gamma \vdash (A \text{ ou } B) \Leftrightarrow B}$ |
| 3. $\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (A \text{ et } B) \Leftrightarrow A}$ | 4. $\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash (A \text{ ou } B) \Leftrightarrow B}$ |
| 5. $\frac{\Gamma \vdash A}{\Gamma \vdash (A \Rightarrow B) \Leftrightarrow B}$ | 6. $\frac{\Gamma \vdash \text{non}B}{\Gamma \vdash (A \Rightarrow B) \Leftrightarrow \text{non}A}$ |
| 7. $\frac{\Gamma \vdash A}{\Gamma \vdash (A \Leftrightarrow B) \Leftrightarrow B}$ | 8. $\frac{\Gamma \vdash \text{non}A}{\Gamma \vdash (A \Leftrightarrow B) \Leftrightarrow \text{non}B}$ |

Ces schémas sont démontrés dans la figure 7.15. Certaines de ces démonstrations sont des chaînes d'équivalences dans lesquelles des simplifications sont faites suivant ceux de ces schémas qui sont déjà démontrés. Le lecteur a ainsi un exemple de l'utilisation typique de ces schémas. Par exemple, dans la chaîne d'équivalences qui démontre S60.2, on fait une simplification évidente suivant S60.1.

7.10 Conditions nécessaires et suffisantes

Conditions nécessaires. Pour réussir l'examen d'entrée dans une université, par exemple, il est certainement *nécessaire* de savoir lire, il *faut* savoir lire. Mais cette condition est sans doute loin d'être *suffisante*. Les deux adjectifs soulignés — nécessaire, suffisant — ont un rapport avec le connecteur logique

(S60.1)	1°	A	<i>base</i>
	2°	$B \Rightarrow A$	1°, S21
	3°	$B \Rightarrow B$	S19
	4°	$B \Rightarrow (A \text{ et } B)$	2°, 3°, S29
	5°	$(A \text{ et } B) \Rightarrow B$	S28
	6°	$(A \text{ et } B) \Leftrightarrow B$	5°, 4°, $\Leftrightarrow$ -intro.
(S60.2)	1°	nonA	<i>base</i>
	2°	$(A \text{ ou } B) \Leftrightarrow \text{non}(\text{non}A \text{ et } \text{non}B)$	De Morgan
	3°	$\Leftrightarrow \text{non}(\text{non}B)$	1°, S60.1
	4°	$\Leftrightarrow B$	S51.1.
(S60.3)	1°	$A \Rightarrow B$	<i>base</i>
	2°	$(A \text{ et } B) \Rightarrow A$	S28
	3°	$A \Rightarrow A$	S19
	4°	$A \Rightarrow (A \text{ et } B)$	3°, 1°, S29
	5°	$(A \text{ et } B) \Leftrightarrow A$	2°, 4°, $\Leftrightarrow$ -intro.
(S60.4)	1°	$A \Rightarrow B$	<i>base</i>
	2°	$\text{non}B \Rightarrow \text{non}A$	1°, S31
	3°	$(A \text{ ou } B) \Leftrightarrow \text{non}(\text{non}B \text{ et } \text{non}A)$	De Morgan
	4°	$\Leftrightarrow \text{non}(\text{non}B)$	2°, S60.3
	5°	$\Leftrightarrow B$	S51.1.
(S60.5)	1°	A	<i>base</i>
	2°	nonnonA	1°, S51.1, S39
	3°	$(A \Rightarrow B) \Leftrightarrow (\text{non}A \text{ ou } B)$	S53
	4°	$\Leftrightarrow B$	2°, S60.2.
(S60.6)	1°	nonB	<i>base</i>
	2°	$(A \Rightarrow B) \Leftrightarrow \text{non}(A \text{ et } \text{non}B)$	S55
	3°	$\Leftrightarrow \text{non}(\text{non}B \text{ et } A)$	S52.8
	4°	$\Leftrightarrow \text{non}(A)$	1°, S60.1.
(S60.7)	1°	A	<i>base</i>
	2°	$B \Rightarrow A$	1°, S21
	3°	$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$	S41
	4°	$\Leftrightarrow (A \Rightarrow B)$	2°, S60.1
	5°	$\Leftrightarrow B$	1°, S60.5.
(S60.8)	1°	nonA	<i>base</i>
	2°	$A \Rightarrow B$	1°, S22
	3°	$(A \Leftrightarrow B) \Leftrightarrow ((A \Rightarrow B) \text{ et } (B \Rightarrow A))$	S41
	4°	$\Leftrightarrow (B \Rightarrow A)$	S60.1
	5°	$\Leftrightarrow (\text{non}A \Rightarrow \text{non}B)$	S57
	6°	$\Leftrightarrow \text{non}B$	1°, S60.5.

Figure 7.15 Démonstrations des schémas S60.1 à S60.8.

d'implication. Il importe de bien connaître ce rapport, vu l'usage constant de ces mots dans le langage mathématique.

Nous partirons de notre exemple — une condition nécessaire de réussite d'un examen — mais au lieu de la forme impersonnelle, nous écrirons, à propos d'une personne X :

$$\text{Pour que } \boxed{X \text{ puisse réussir l'examen}} \text{ il faut que } \boxed{X \text{ sache lire.}} \quad (1)$$

A B

Cette assertion combine deux propositions A, B qui sont essentiellement : « X peut réussir l'examen » et « X sait lire ». Le mode subjonctif des verbes n'est pas essentiel du point de vue logique. Deux autres formes équivalentes de cette assertion sont les suivantes :

$$\begin{aligned} \text{Une condition nécessaire pour que : } X \text{ puisse réussir l'examen (A)} & \quad (2) \\ \text{est que : } X \text{ sache lire. (B)} & \end{aligned}$$

$$\begin{aligned} X \text{ peut réussir l'examen (A) seulement si} & \quad (3) \\ X \text{ sait lire (B).} & \end{aligned}$$

Le bon usage de ces tournures découlera de leur traduction correcte au moyen de symboles logiques, est cette traduction est simplement :

$$\text{Si } X \text{ peut réussir l'examen (A), alors } X \text{ sait lire (B).} \quad (4)$$

Autrement dit, les assertions (1), (2), (3) ne sont qu'une autre manière d'énoncer : $A \Rightarrow B$. Le sens de cette implication n'est pas toujours reconnu immédiatement. On commet parfois l'erreur de traduire (1) par $B \Rightarrow A$, alors que c'est $A \Rightarrow B$ qui convient. Pour s'en convaincre, on pourra reconnaître que la nécessité — il faut — exprimée par (1) peut aussi s'exprimer par :

$$\begin{array}{ccc} \text{Si } \boxed{X \text{ ne sait pas lire}} & \text{alors} & \boxed{X \text{ ne peut pas réussir l'examen.}} \\ \text{nonB} & & \text{nonA} \end{array}$$

Or cette phrase, $\text{nonB} \Rightarrow \text{nonA}$, est équivalente à $A \Rightarrow B$ (S57).

Une autre manière de considérer l'assertion (1) est de penser : si je sais que X peut réussir l'examen, j'en déduis que X satisfait *nécessairement* — dans le sens de *inévitablement* — à la condition B. Dire que B est une condition nécessaire de A, c'est dire que cette condition est toujours ou inévitablement satisfaite lorsque A a lieu. C'est donc dire simplement que $A \Rightarrow B$.

On peut encore expliquer le sens de (1) en proclamant : ou bien X sait lire, ou bien X ne peut pas réussir l'examen. Or ceci est la proposition (B ou nonA), équivalente à $A \Rightarrow B$ (S54).

Ces commentaires ne sont pas une preuve formelle de l'équivalence des tournures (1), (2), (3), (4). Ils n'ont pour but que de faciliter l'acceptation de la *convention de langage* suivante.

Convention 1. Une proposition de la forme $A \Rightarrow B$ peut s'exprimer des quatre manières suivantes:

- (a) Si A , alors B .
- (b) Pour que A , il faut que B .
- (c) Une condition nécessaire pour que A est que B .
- (d) A seulement si B .

Exemple 1. Un paragraphe d'un livre d'Analyse¹ commence par ces mots:

3.1.4 Condition nécessaire pour qu'une série converge.

Si la série de terme général x_n est convergente, alors la suite (x_n) converge vers zéro.

Le théorème ainsi énoncé est de la forme $A \Rightarrow B$, avec les propositions

- A: la série de terme général x_n est convergente,
- B: la suite (x_n) converge vers zéro.

C'est pour comprendre le *titre du paragraphe* qu'il faut connaître la convention 1. Selon celle-ci, la condition nécessaire dont il est question dans le titre, c'est la proposition B.

Le théorème aurait pu être formulé des trois autres manières mentionnées dans cette convention: (b) Pour que la série de terme général x_n soit convergente, il faut que la suite (x_n) converge vers zéro. (c) Une condition nécessaire pour que la série de terme général x_n soit convergente est que la suite (x_n) converge vers zéro. (d) La série de terme général x_n est convergente seulement si la suite (x_n) converge vers zéro.

Conditions suffisantes. Lorsque l'argent fait le bonheur d'une personne X , qui ignore tous les vieux adages, on peut affirmer:

$$\text{Pour que } \boxed{X \text{ soit heureux}} \text{ il suffit que } \boxed{X \text{ ait de l'argent.}} \quad (5)$$

B
A

$$\begin{array}{ll} \text{Une condition suffisante pour que: } X \text{ soit heureux} & (B) \\ \text{est que: } X \text{ ait de l'argent.} & (A) \end{array} \quad (6)$$

$$\begin{array}{ll} X \text{ est heureux} & (B) \text{ si} \\ X \text{ a de l'argent} & (A). \end{array} \quad (7)$$

Ces trois affirmations ne signifient rien d'autre que $A \Rightarrow B$, c'est-à-dire: si X a de l'argent, alors X est heureux. Ceci est encore une convention de langage.

¹ Jacques Douchet et Bruno Zwahlen, *Calcul différentiel et intégral*, Vol. 1, Lausanne, Presses polytechniques et universitaires romandes, (1990).

Convention 2. Une proposition de la forme $A \Rightarrow B$ peut encore s'exprimer des trois manières suivantes:

- (e) Pour que B, il suffit que A.
- (f) Une condition suffisante pour que B est que A.
- (g) B *si* A.

Il revient donc au même de dire que B est une condition nécessaire pour que A, ou que A est une condition suffisante pour que B.

Exemple 2. Les auteurs du livre d'Analyse cité plus haut¹ pourraient, dans une prochaine édition de leur ouvrage, changer comme suit le titre de leur paragraphe 3.1.4, tout en laissant inchangé le contenu:

3.1.4 Condition suffisante pour qu'une suite converge vers zéro.

Si la série de terme général x_n est convergente, alors la suite (x_n) converge vers zéro.

Mais ils pourraient aussi changer du même coup le texte du paragraphe, en l'écrivant de l'une ou l'autre des trois manières citées dans la convention 2: (e) Pour que la suite (x_n) converge vers zéro, il suffit que la série de terme général x_n soit convergente. (f) Une condition suffisante pour que la suite (x_n) converge vers zéro est que la série de terme général x_n soit convergente. (g) La suite (x_n) converge vers zéro si la série de terme général x_n est convergente.

Conditions nécessaires et suffisantes. D'après S41, une proposition de la forme $A \Leftrightarrow B$ est toujours équivalente à

$$(A \Rightarrow B) \text{ et } (B \Rightarrow A). \quad (8)$$

En appliquant les conventions de langage qui précèdent, nous pouvons traduire (8) des trois manières suivantes:

Pour que A	et	pour que A
il faut que B		il suffit que B
Une condition nécessaire	et	une condition suffisante
pour que A est que B		pour que A est que B
A seulement si B	et	A si B

(9)

On abrège traditionnellement ces expressions en disant: pour que A, il faut et il suffit que B; une condition nécessaire et suffisante pour que A est que B; et enfin, en permutant les deux membres de (9):

$$A \text{ si et seulement si } B. \quad (10)$$

Exemple 3. Le livre d'Analyse cité précédemment contient le théorème et la démonstration suivants, que nous ne citons pas intégralement:

Pour qu'un sous-ensemble non vide S de $\mathbb{R}$ possède une borne supérieure, il faut et il suffit que S soit majoré.

DÉMONSTRATION. La condition est nécessaire, car si S admet une borne supérieure b , alors b est un majorant de S . D'où S est majoré. Supposons maintenant que S soit majoré. Alors... Fin de citation.

Pour discuter ce passage, désignons par H, A, B les propositions suivantes :

H : S est un sous-ensemble non vide de $\mathbb{R}$;

A : S possède une borne supérieure;

B : S est majoré.

Le théorème affirme l'équivalence $A \Leftrightarrow B$ sous l'hypothèse H :

$$H \vdash A \Leftrightarrow B.$$

La démonstration comporte deux parties (sous l'hypothèse H): la preuve de $A \Rightarrow B$ et celle de $B \Rightarrow A$. Ce sont les mots « La condition est nécessaire », au début de la démonstration, qui indiquent que l'on commence par la preuve de $A \Rightarrow B$. Il faut bien connaître les conventions de langage qui précèdent pour comprendre ce texte. L'énoncé du théorème utilise la tournure « il faut et il suffit », alors que la démonstration parle d'une « condition nécessaire ». On peut commencer par mettre l'énoncé sous la forme

Une condition nécessaire et suffisante pour que S possède une borne supérieure est que S soit majoré.

On voit bien alors que la condition dont il est question au début de la démonstration est B , et que la première partie de cette démonstration est celle de $A \Rightarrow B$.

Remarque finale. La maîtrise de ces expressions est une condition nécessaire à la compréhension des ouvrages de mathématique. « Mais pas suffisante » ajoutera le lecteur, avec raison et à-propos.

7.11 Exercices

Soient A, B, C, R, S des propositions d'un langage du premier ordre $\mathcal{L}$. Démontrer les théorèmes suivants de $\text{Lprop}(\mathcal{L})$.

- (a) $\vdash (A \Rightarrow B) \Leftrightarrow ((A \text{ ou } B) \Leftrightarrow B)$.
- (b) $\vdash (A \Rightarrow (B \Rightarrow C)) \Leftrightarrow ((A \Rightarrow B) \Rightarrow (A \Rightarrow C))$.
- (c) $\vdash ((A \text{ et } B) \Rightarrow C) \Leftrightarrow (A \Rightarrow (B \Rightarrow C))$.
- (d) $\vdash (A \text{ ou } B) \Leftrightarrow ((A \Rightarrow B) \Rightarrow B)$.
- (e) $\vdash (\text{non}A \Leftrightarrow B) \Leftrightarrow \text{non}(A \Leftrightarrow B)$.
- (f) $\vdash (A \Rightarrow (B \text{ ou } C)) \Leftrightarrow (B \text{ ou } (A \Rightarrow C))$.
- (g) $\vdash (A \text{ et } B) \Rightarrow ((R \text{ et } A) \text{ ou } (\text{non}R \text{ et } B))$.
- (h) $\vdash ((R \Rightarrow A) \text{ et } (\text{non}R \Rightarrow B)) \Leftrightarrow ((R \text{ et } A) \text{ ou } (\text{non}R \text{ et } B))$.
- (i) $\vdash [(A \Rightarrow B) \Rightarrow (R \Rightarrow C)] \Rightarrow [(A \Rightarrow (S \Rightarrow C)) \Rightarrow (R \Rightarrow (S \Rightarrow C))]$.
- (j) $\vdash (A \text{ ou } (B \Leftrightarrow C)) \Leftrightarrow ((A \text{ ou } B) \Leftrightarrow (A \text{ ou } C))$.
- (k) $\vdash (A \Rightarrow (B \Leftrightarrow C)) \Leftrightarrow ((A \text{ et } B) \Leftrightarrow (A \text{ et } C))$.
- (l) $\vdash (A \Leftrightarrow (B \Leftrightarrow C)) \Leftrightarrow ((A \Leftrightarrow B) \Leftrightarrow C)$.

Introduction à la théorie des ensembles

8.1 Le rôle de la théorie des ensembles et son traitement dans cet ouvrage

La théorie des ensembles est prise dans ce livre comme domaine d'application de la logique. Son traitement est réparti dans les chapitres qui suivent. Dès le début du chapitre 9, les schémas de déduction de la logique des prédicats seront illustrés par des exemples de déductions pris dans cette théorie. Nous espérons faciliter ainsi l'assimilation de ces schémas de déduction et d'autre part inculquer au lecteur des notions précises de théorie des ensembles et la faculté de raisonner formellement dans ce domaine. Dès que nous disposerons d'un répertoire suffisant de schémas de déduction dérivés de logique des prédicats (Chap. 10), nous consacrerons des paragraphes entiers au développement de la théorie des ensembles: présentation des axiomes et schémas de déduction spécifiques, démonstration de théorèmes.

Le présent chapitre n'est pas une « Introduction à la théorie des ensembles » dans le sens complet où l'on pourrait l'entendre. C'est tout le reste de l'ouvrage qui est une telle introduction — en plus d'un manuel de logique. Le présent chapitre ne contient qu'une brève présentation de l'univers de l'interprétation de cette théorie.

Des éléments de théorie des ensembles appartiennent de nos jours à l'enseignement des mathématiques de tous les niveaux. Le langage de la théorie des ensembles est utilisé dans tous les domaines. Mais lorsqu'on passe à une théorie formelle des ensembles, il y a un pas d'abstraction à franchir. Il faut considérer non plus des ensembles particuliers — l'ensemble des nombres réels, l'ensemble des cantons suisses, l'ensemble des atomes d'un corps, etc. Il faut considérer « les ensembles » de façon générale, en tant qu'objets de la pensée et en tant qu'*individus* d'un univers dans lequel il n'y a pas d'autre espèce d'individus que les ensembles. C'est des objets de cet univers que parle la théorie des ensembles.

La théorie des ensembles en tant qu'outil général. La théorie des ensembles est utilisée de deux manières en mathématique: (a) comme *outil* dans toutes les branches des mathématiques, et (b) comme *théorie fondamentale* à laquelle peuvent se ramener toutes les théories mathématiques particulières.

Le deuxième usage de la théorie des ensembles (b) est fait par ceux qui s'intéressent aux fondements des mathématiques, à la construction logique de tout l'édifice des mathématiques sur la base d'un nombre minimum d'axiomes et de schémas de déduction. Nous en dirons quelques mots plus bas.

Pour la grande majorité des mathématiciens, travaillant dans un domaine particulier, et des utilisateurs des mathématiques, ingénieurs et scientifiques qui modélisent mathématiquement les systèmes auxquels ils s'intéressent, la théorie des ensembles est un outil (a). Pour la plupart d'entre eux d'ailleurs cet outil n'est même pas considéré comme une « théorie », mais simplement comme un répertoire de quelques notions générales. Un auteur français renommé s'exprime en ces termes¹: « Pour l'immense majorité des mathématiciens, la partie de cette dernière [la théorie des ensembles] qu'ils utilisent se réduit à quelques notions élémentaires. »

Si nous prenons le terme de « notion » dans le sens de *concept* précis, faisant l'objet d'une définition en théorie des ensembles — la notion de fonction, la notion d'ensemble ordonné, la notion de borne supérieure, de relation d'équivalence, etc. — ces « quelques notions élémentaires » dont parle l'auteur cité sont en fait quelques dizaines. Nous en dénombrons en gros quarante. La moitié d'entre elles environ sont omniprésentes. Le langage de la théorie des ensembles a pénétré dans tous les domaines.

L'informaticien qui n'a jamais étudié ces notions ailleurs que dans ces « rappels » sommaires qui figurent souvent dans les manuels de mathématique pure ou appliquée (analyse, probabilités, recherche opérationnelle, algorithmique, théorie des langages et automates, etc.) ne les maîtrise pas suffisamment pour pouvoir les employer avec sûreté dans le cadre d'un langage de spécification formelle de logiciel.

La place logique de la théorie des ensembles. Toutes les mathématiques se ramènent ou peuvent se ramener à la théorie des ensembles. Le sens exact de cette assertion se précisera lorsque nous aurons présenté la notion d'extension définitionnelle d'une théorie (Chap. 12). Mais l'idée est très simple: on peut démontrer tous les théorèmes des mathématiques en partant des seuls axiomes et schémas de déduction de la théorie des ensembles et de *définitions* appropriées. La découverte de cette situation logique privilégiée de la théorie des ensembles eut lieu à la fin du XIX^e siècle lorsque Weierstrass et Dedekind donnèrent une « construction » des nombres réels à partir de l'ensemble des entiers naturels n'utilisant que des opérations ensemblistes. Lorsqu'on eut encore *défini* les entiers naturels eux-mêmes comme des ensembles construits à

¹ Jean Dieudonné, dans la *Préface à l'édition française* du traité d'algèbre de S. MacLane et G. Birkhoff, Paris, Gauthier-Villars, 1971.

partir de l'ensemble vide $\emptyset$, représentant le nombre 0, en n'utilisant que des opérations ensemblistes, on put dire qu'il n'y a dans toutes les mathématiques qu'une seule espèce d'objets — les ensembles — et que les axiomes de la théorie des ensembles permettent de démontrer tout ce qui est démontrable.

Les langages de spécification ensemblistes. La spécification formelle d'un système de traitement de l'information décrit de façon précise les propriétés qu'il doit avoir, sans être contraignante sur la manière de réaliser ces propriétés. Elle décrit *ce que* le système doit faire, sans dire *comment* il doit le faire. Cette abstraction la rend utile dans le processus de développement car elle permet de répondre avec sûreté aux questions sur le comportement du système sans avoir à extraire cette information d'une masse de code de programme détaillé, ou à spéculer sur le sens des phrases d'un document de spécification en prose imprécise.

Une spécification formelle peut servir de référence unique pour ceux qui déterminent les besoins exacts du client, ceux qui implémentent les programmes devant satisfaire ces besoins, ceux qui testent les résultats, et ceux qui écrivent les manuels d'instruction. Comme elle est indépendante du code de programme, elle peut être achevée dans la première période du processus de développement.

Une manière d'atteindre ces buts est de modéliser les données traitées par le système au moyen de *types de données mathématiques* abstraits qui ne sont pas orientés vers la représentation en machine, mais qui obéissent à des lois mathématiques permettant de raisonner effectivement sur le comportement qu'aura le système.

Ces types de données sont exactement ceux de la théorie des ensembles. Les mathématiques n'utilisent pas le terme de « type de donnée ». Celui-ci est une création des informaticiens et il est apparu dans leur vocabulaire sitôt qu'ils se sont rendus compte qu'il était judicieux de considérer les données de manière plus abstraite que sous la forme de leur représentation en machine. La manière la plus abstraite est celle des mathématiques, ou les « données » sont des systèmes d'ensembles, de fonctions, de relations, etc. qui se combinent suivant le répertoire d'opérations standard étudié en théorie des ensembles.

8.2 L'univers de la théorie des ensembles

Du point de vue sémantique, la théorie des ensembles parle des ensembles en général, de certaines opérations de construction d'ensembles à partir d'autres ensembles, de certaines relations entre ensembles, etc. On pourrait s'attendre toutefois à ce que l'univers de l'interprétation de cette théorie ne se compose pas seulement d'ensembles, mais aussi d'objets qui ne sont pas des ensembles.

La notion d'ensemble nous est en effet inculquée à l'école par des exemples d'ensembles dont les éléments sont des objets concrets ou abstraits qui ne sont pas eux-mêmes des ensembles, ou du moins que ne nous ne concevons pas

naturellement comme étant des ensembles. Lorsqu'on nous parle de l'*ensemble des élèves de la classe*, les éléments de cet ensemble sont des êtres humains, non des ensembles. Les éléments de l'*ensemble des lettres de l'alphabet* sont des symboles, non des ensembles. L'*ensemble des notes de la gamme* a pour éléments des sons musicaux, l'*ensemble des points d'une figure plane* sont des points, une espèce d'objet géométrique qui n'est pas un ensemble. Les nombres réels, éléments de l'*ensemble des nombres réels*, sont des entités abstraites que nous manipulons généralement sans nous poser de question sur leur nature, mais qui ne nous apparaissent certainement pas comme étant des ensembles.

Pour qu'il y ait des ensembles, il faut d'abord qu'il y ait des objets que nous puissions assembler ou grouper mentalement pour former des ensembles d'objets, construire des ensembles par la pensée. Ensuite, tous les ensembles que nous pouvons concevoir ainsi sont eux-mêmes des objets de la pensée qui peuvent être pris comme éléments d'autres ensembles, de niveau supérieur. Nous pouvons parler par exemple de l'ensemble des classes d'une école comme étant un ensemble d'ensembles d'élèves. Nous pouvons concevoir ainsi des ensembles d'ensembles puis des ensembles d'ensembles d'ensembles, et ainsi de suite sans limite. Mais au départ il nous faut, semble-t-il, d'autres espèces d'objets.

La théorie des ensembles ignore ces autres objets. Son univers ne se compose que d'ensembles. Car elle ne s'intéresse pas à ce que « sont » les ensembles, mais seulement à certaines relations typiques qu'il peut y avoir entre objets de cette espèce, et à certaines opérations que l'on peut effectuer sur ces objets. Tous les individus de son univers sont donc de la même nature abstraite: ce sont tous des ensembles.

La relation typique fondamentale qui peut exister entre deux objets de cet univers est la relation d'*appartenance*

$$a \in b \tag{1}$$

qui se lit: *a est élément de b*, ou *a appartient à b*. Lorsque nous lisons l'expression (1), nous interprétons automatiquement le symbole *b* comme désignant un ensemble. Mais pas automatiquement de la même manière le symbole *a*, au contraire. Nous pensons spontanément à des exemples comme ceux dont nous avons parlé ci-dessus, dans lesquels les éléments de l'ensemble *b* ne sont pas vus comme des ensembles. Mais *b* peut être un ensemble d'ensembles, et dans ce cas *a* est un ensemble. La théorie des ensembles ne considère que ce cas. Elle ne s'intéresse qu'aux relations entre ensembles et aux opérations sur les ensembles, et la proposition (1) exprime une relation qui peut certainement exister entre deux ensembles *a* et *b*.

Cette restriction de l'univers de la théorie aux seuls ensembles n'a pas empêché cette théorie d'absorber toutes les mathématiques, de sorte que cet univers exclusif est devenu celui de toutes les mathématiques. Le lecteur peut se dire que même lorsqu'il écrit $\sqrt{2} \in \mathbb{R}$, les *deux* termes de cette proposition, $\sqrt{2}$ et $\mathbb{R}$, désignent des ensembles, même s'il ne voit pas le premier comme

un ensemble. Il y a déjà un siècle que les mathématiciens savent définir les nombres réels comme des ensembles et déduire toutes leurs propriétés de cette définition. Il n'y a aucun besoin pratique de connaître cette définition. On peut même étudier l'analyse sans la connaître, si l'on se contente d'admettre comme axiomes un certain nombre de propriétés des nombres réels. Il suffit de savoir que cette définition existe. Il peut être intéressant aussi de savoir que tous les nombres (entiers naturels, entiers, rationnels, réels, complexes) sont construits en tant qu'ensembles à partir d'un unique ensemble de départ — l'ensemble *vide*, noté $\emptyset$, qui n'a aucun élément — au moyen des opérations standard de construction d'ensembles.

L'une de ces opérations consiste, par exemple, étant donné un objet quelconque x , à construire l'ensemble noté $\{x\}$, dont l'unique élément est l'objet x . Une manière possible de construire les entiers naturels $0, 1, 2, \dots$ en tant qu'ensembles est de les définir comme les ensembles $\emptyset, \{\emptyset\}, \{\{\emptyset\}\}$, etc. Toutes les propriétés des entiers naturels découlent alors des axiomes de la théorie des ensembles. La construction des nombres réels est beaucoup plus compliquée et ne peut être esquissée ainsi en quelques phrases.

Puisqu'il n'y a que des ensembles dans l'univers de la théorie des ensembles, des phrases telles que « x est un ensemble », ou « soit x un ensemble », sont en principe superflues, puisqu'un terme ne peut pas désigner autre chose qu'un ensemble, dans ce contexte. Et de fait, dans le langage formel de la théorie des ensembles, il *n'y a pas* de prédicat (symbole relationnel unaire) tel que « *est-un-ensemble* ». La phrase « x est un ensemble » n'appartient pas au langage de cette théorie. Dans ce langage, on peut dire quantité de choses sur un ensemble x , sauf qu'il est un ensemble — ce qui fait partie de notre *interprétation* du symbole x dans ce contexte. Ce genre de phrase ne figurera donc jamais dans une démonstration formelle. On l'utilise souvent cependant dans les démonstrations informelles, dans cet ouvrage comme ailleurs, lorsqu'on veut faire appel à l'intuition du lecteur.

Appartenance et inclusion. La théorie des ensembles ne définit pas le sens de la relation d'appartenance $a \in b$. Ce sens fait partie de notre *interprétation* du langage de cette théorie, et nous admettons qu'il est plus ou moins acquis par le lecteur. Pour le reste, ce sens se précisera de lui-même *d'après* le développement formel de la théorie.

La représentation graphique « naïve » de la notion d'ensemble et des relations d'appartenance est celle de la figure 8.1 (a), dans laquelle on entoure les éléments d'un ensemble par un « cercle ». Cette figure est censée représenter des objets $a, b, c, \dots$ (par les petits cercles); l'ensemble x ayant pour éléments les objets a, b, c , ce qu'on note $x = \{a, b, c\}$; l'ensemble $y = \{d, e\}$; l'ensemble z dont les éléments sont les ensembles x et y , $z = \{x, y\} = \{\{a, b, c\}, \{d, e\}\}$. Les éléments de l'ensemble s ne sont pas tous nommés. Mais la figure montre que l'objet u appartient à z et à s . Les « objets » $a, b, c, \dots$ sont représentés eux-mêmes par des petits cercles dans notre figure pour indiquer qu'ils sont eux-

mêmes des ensembles, suivant le point de vue de cette théorie. On ne s'intéresse simplement pas à leurs éléments. Lorsque nous parlerons ainsi « d'objets », en interprétant les expressions de théorie des ensembles, il faudra prendre ce mot dans le sens d'ensembles dont les éléments n'interviennent pas et sont ignorés dans la discussion.

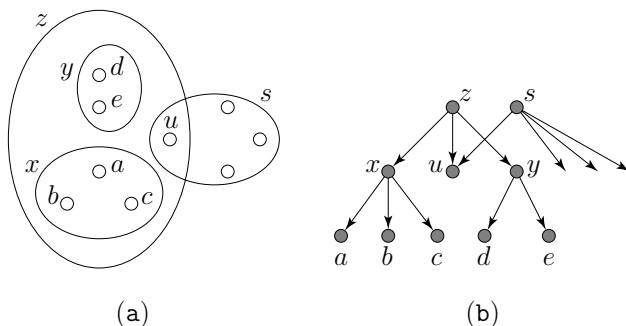


Figure 8.1 Représentation de relations d'appartenance.

Une meilleure représentation, pour interpréter les expressions de la théorie des ensembles, est celle de la figure 8.1 (b), qui correspond à sa voisine. Là, tous les ensembles sont représentés de la même manière par des « points », et les relations d'appartenance sont représentées par des flèches. Pour chaque élément d'un ensemble p , il y a une flèche allant du point p à cet élément.

On peut imaginer le diagramme de points et de flèches de la figure 8.1 comme étant un fragment d'un immense diagramme de ce genre représentant l'univers entier de la théorie des ensembles et c'est cette image de l'univers des ensembles que nous suggérons au lecteur de se faire. C'est un univers d'objets ou entités entre certains desquels il y a une relation dite d'appartenance que l'on note $x \in y$ et qui est représentée dans le diagramme par une flèche allant de y à x . Ce sont les propriétés de ce diagramme que décrit la théorie des ensembles. Car il n'est pas quelconque. Il a des propriétés qui sont énoncées dans les axiomes de la théorie.

Une relation d'inclusion

$$x \subset y$$

se lit: x est *inclus dans* y , ou x est un *sous-ensemble* de y , ou encore x est une *partie* de y . Elle signifie que tout élément de x est élément de y . Ceci fait l'objet de notre premier axiome de théorie des ensembles, que nous appellerons *axiome de l'inclusion*:

$$\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y)). \quad (2)$$

D'habitude, on parle plutôt d'une *définition* que d'un axiome. La proposition (2) est présentée comme une définition de l'inclusion. Mais nous verrons que les définitions, du point de vue de la logique, ne sont que des axiomes d'une forme particulière. Nous en restons donc pour le moment au terme d'axiome.

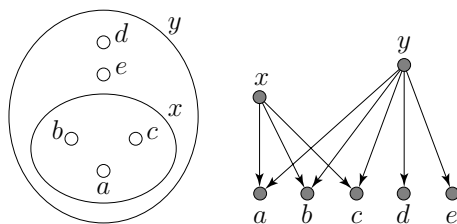


Figure 8.2 Relation d'inclusion $x \subset y$:
tout élément de x est élément de y .

La figure 8.2 représente deux ensembles x et y vérifiant la relation $x \subset y$: les éléments du premier sont a, b, c et sont aussi éléments du second. Il ne faut pas confondre les symboles d'appartenance ($\in$) et d'inclusion ($\subset$). Dans l'exemple de la figure 8.2 on a $x \subset y$, mais non($x \in y$), ce qu'on abrège $x \not\in y$: il n'y pas de flèche de y à x . On pourrait modifier l'exemple en *ajoutant* une flèche de y à x . Dans ce cas l'ensemble x serait à la fois élément de y ($x \in y$) et sous-ensemble de y ($x \subset y$). On écrirait: $x = \{a, b, c\}$ et $y = \{a, b, c, d, e, x\} = \{a, b, c, d, e, \{a, b, c\}\}$.

On pourrait se représenter le diagramme de l'univers des ensembles avec deux espèces de flèches: celles dont nous avons parlé jusqu'ici, marquant les relations d'appartenance, et des flèches d'une « autre couleur » pour marquer les relations d'inclusion. Une de ces autres flèches serait alors tracée entre y et x dans la figure 8.2.

La relation d'inclusion est *transitive*. On entend par là que si $x \subset y$ et $y \subset z$, alors on a nécessairement $x \subset z$. S'il y a une flèche d'inclusion de z à y dans le diagramme de l'univers des ensembles et une flèche d'inclusion de y à x , alors il y a nécessairement une flèche d'inclusion de z à x . C'est une conséquence logique de l'axiome (2) et nous en ferons la démonstration formelle ultérieurement. Mais cela est bien évident: si tout élément de x est élément de y et si tout élément de y est élément de z , alors tout élément de x est élément de z . Pour illustrer cela, le lecteur peut compléter la figure 8.2 avec un ensemble z tel que $y \subset z$, c'est-à-dire tel que a, b, c, d, e soient éléments de z .

Par contre, la relation d'appartenance *n'est pas* transitive. Dans l'exemple de la figure 8.1, on a $a \in x$, $x \in z$, mais non $a \in z$. Ceci n'est pas une preuve, bien entendu. C'est des axiomes de la théorie qu'il faut déduire l'existence de trois ensembles x, y, z tels que $x \in y$, $y \in z$ et $x \not\in z$. Nous ne pouvons le faire ici. Mais pour le lecteur qui est déjà un peu familiarisé avec cette théorie, il suffit de considérer les ensembles $x = \emptyset$, $y = \{x\}$, et $z = \{y\}$.

L'ensemble vide. Le symbole $\emptyset$ est une constante individuelle qui représente un ensemble sans éléments, que l'on dit *vide*. On l'appelle *l'ensemble vide*, et pas seulement *un* ensemble vide, parce qu'on peut démontrer qu'il est le *seul* ensemble sans élément. Dans le diagramme de l'univers des ensembles, il y

a un objet et un seul dont il ne part aucune flèche d'appartenance: l'objet $\emptyset$. Cela sera démontré au chapitre 10 lorsque nous disposerons des schémas de déduction nécessaires de la logique des prédicats. Pour le moment nous ne faisons que poser l'axiome déclarant que l'ensemble vide n'a pas d'éléments, ou *axiome de l'ensemble vide*:

$$\forall x(x \notin \emptyset). \quad (3)$$

Systèmes d'axiomes. Bien qu'il soit coutume de parler de « la » théorie des ensembles, il faut savoir qu'il n'y en a pas qu'une. Mais il n'entre pas dans le cadre de cet ouvrage d'enseigner les distinctions entre les variantes possibles. Pour les applications courantes, ces différences ont d'ailleurs peu d'importance, car seul le tronc commun de ces théories est utilisé. Celle que nous présentons dans cet ouvrage, et que nous appelons $\mathcal{T}_{\text{ENS}}$, correspond assez bien à ce tronc commun.

La raison des différences qui peuvent exister entre diverses variantes de théorie des ensembles est que la notion d'ensemble est une notion abstraite sur laquelle les esprits peuvent avoir des vues divergentes — à propos de certaines questions. Ce sont des questions qui ne se posent pas ou se posent rarement dans les applications, mais que l'on ne peut pas éviter de se poser lorsqu'on choisit un système d'axiomes.

L'une de ces questions est la suivante. Un ensemble x peut être élément d'un autre ensemble y , ce que l'on écrit $x \in y$. Existe-t-il un ensemble x qui soit élément de lui-même: $x \in x$? Autrement dit, peut-on assembler mentalement des objets et parmi ces objets avoir aussi celui qui est l'assemblage lui-même? Certaines versions de la théorie des ensembles condamnent cette idée au moyen d'un axiome qui entraîne la réponse *non* et entraîne la même réponse aux questions analogues: existe-t-il deux ensembles x et y tels que l'on ait $x \in y$ et $y \in x$? trois ensembles x, y, z tels que l'on ait $x \in y, y \in z$, et $z \in x$? etc. Mais ce sont des questions marginales et la réponse n'a pas d'incidence sur ce que nous appelons le tronc commun des théories des ensembles. Nous choisirons de ne *pas y répondre*, c'est-à-dire que le système d'axiomes retenu ne permettra de démontrer ni l'affirmative, ni la négative — du moins personne n'y est parvenu, à notre connaissance.

La théorie des ensembles est caractérisée non seulement par ses axiomes mais encore par des schémas de déduction élémentaires spécifiques. Ceux-ci ne seront présentés toutefois qu'au chapitre 14. L'ordre dans lequel nous abordons cette théorie n'est pas celui des exposés usuels dans lesquels toute la logique est considérée comme acquise, voir innée. Notre présentation de cette théorie en parallèle avec la logique nous impose un ordre particulier et un système d'axiomes et de déductions élémentaires qui n'est pas le plus économique possible. Mais cela n'a aucune importance dans l'optique de cet ouvrage et ne peut en avoir que pour des spécialistes.

Logique des prédicats

9.1 Dédutions élémentaires

Pour tout langage du premier ordre $\mathcal{L}$, nous allons définir une théorie appelée *logique des prédicats de langage $\mathcal{L}$* . Nous abrègerons cette dénomination par $\mathbf{Lpréd}(\mathcal{L})$. Le langage de cette théorie est le langage $\mathcal{L}$. De même que la logique propositionnelle de langage L , $\mathbf{Lprop}(\mathcal{L})$, la théorie $\mathbf{Lpréd}(\mathcal{L})$ est une théorie sans axiomes (cf. §6.1). La classe des déductions de cette théorie *contient* celle de la logique propositionnelle. Autrement dit, la théorie $\mathbf{Lpréd}(\mathcal{L})$ est une extension (§5.7) de la théorie $\mathbf{Lprop}(\mathcal{L})$.

De façon précise, les déductions élémentaires de la logique des prédicats de langage $\mathcal{L}$ sont décrites par les schémas de déduction S1–S14 (Fig. 6.1), et par les quatre nouveaux schémas S61–S64 de la figure 9.1, qui déterminent l’usage des quantificateurs.

Au chapitre 11 nous définirons la théorie $\mathbf{Lprédég}(\mathcal{L})$, ou *logique des prédicats avec égalité* de langage $\mathcal{L}$, qui sera une extension de la théorie $\mathbf{Lpréd}(\mathcal{L})$ obtenue en ajoutant deux schémas de déduction élémentaires relatifs au symbole d’égalité.

Nous appelons *théorie logique* de langage $\mathcal{L}$ toute théorie $\mathcal{T}$ de langage $\mathcal{L}$ dont la classe des déductions inclut celles de la logique des prédicats avec égalité $\mathbf{Lprédég}(\mathcal{L})$. Pratiquement toutes les théories mathématiques sont, dans ce sens, des théories logiques.

Prédicats. Avant de présenter ces schémas, nous voulons nous étendre un peu sur le sens du mot *prédicat*, et sur ce qui distingue la logique des prédicats de la logique propositionnelle.

Ce mot provient de la grammaire des langues naturelles, où il est associé au mot « sujet » dans la description des phrases simples. Celles-ci sont composées d’un *sujet* et d’un *prédicat*. Le sujet est ce dont on parle, le prédicat est ce que l’on dit du sujet. Ainsi lorsqu’on dit « la loi est dure », on applique au sujet « la loi » le prédicat « est dure ».

S61. *Particularisation*

$$\frac{\Gamma \vdash \forall x A}{\Gamma \vdash (T|x)A}$$

Si T est librement
substituable à x dans A .

S62. *Généralisation*

$$\frac{\Gamma \vdash A}{\Gamma \vdash \forall x A}$$

Si x ne figure pas
librement dans Γ .

S63. *Introduction de $\exists$*

$$\frac{\Gamma \vdash (T|x)A}{\Gamma \vdash \exists x A}$$

Si T est librement
substituable à x dans A .

S64. *Élimination de $\exists$*

$$\frac{\Gamma \vdash \exists x A \quad \Gamma; A \vdash B}{\Gamma \vdash B}$$

Si x ne figure librement
ni dans B ni dans Γ .

Figure 9.1 Schémas de déduction élémentaires de la logique des prédicats, dans lesquels A et B désignent des propositions, x une variable, et T un terme.

Si l'on transpose cette notion dans un langage du premier ordre, un prédicat n'est autre qu'un symbole relationnel unaire. Par exemple, dans notre introduction aux langages du premier ordre (§3.2), nous avons esquissé dans un langage du premier ordre la description d'une base de données sur une université. Les langages de ce genre contiennent de nombreux symboles relationnels unaires de la forme «est un ...», ou «est une ...». Nous avons écrit par exemple les propositions **est-un-prof**(x), **est-un-cours**(y). Ces symboles, **est-un-prof**, **est-un-cours**, sont des symboles relationnels unaires, donc des prédicats.

Cependant les langages du premier ordre ne comportent pas que des symboles relationnels unaires. Le sens du mot «prédicat» doit être élargi, dans le cadre de ces langages, pour désigner tous les symboles relationnels, d'arité quelconque. Au paragraphe 3.2, nous avons écrit par exemple la proposition **enseigne**(p, x) pour dire que le professeur p enseigne la matière x . Le prédicat **enseigne** s'applique donc à deux sujets. En grammaire on ne parle pas de deux sujets, mais plutôt d'un sujet et d'un complément direct.

Nous rappelons que les *propositions atomiques* d'un langage du premier ordre $\mathcal{L}$ (§4.1) sont les expressions de la forme

$$rT_1 \cdots T_n$$

obtenues en écrivant successivement un symbole relationnel r du langage $\mathcal{L}$ et n termes de ce langage, T_1 à T_n , distincts ou non, n étant l'arité de r . On peut donc dire qu'une proposition atomique est construite en appliquant un prédicat n -aire à n termes. Avec la notation préfixée il n'y a pas besoin de parenthèses, mais on écrit parfois tout de même $r(T_1, \dots, T_n)$ pour plus de

lisibilité. Par exemple, $\text{enseigne}(p, x)$ est écrit pour $\text{enseigne } p x$.

Finalement, la notion de prédicat est élargie parfois pour désigner tout ce qui reste d'une proposition *quelconque*, et pas seulement atomique, lorsqu'on lui enlève toutes ses variables libres. Considérons par exemple la proposition

$$\exists x(a + x = b), \quad (1)$$

dans le cadre de la théorie des entiers naturels. On peut la considérer comme étant construite en appliquant le prédicat

$$\exists x(_1 + x = _2),$$

aux sujets a et b . Les traits numérotés, dans le prédicat, marquent les emplacements pour les sujets potentiels. On sait que dans le cadre de la théorie des entiers naturels, ce prédicat est défini comme étant le prédicat *d'inégalité*, et l'on utilise le symbole $\leq$ pour l'abrégé, en posant comme *définition*, que $a \leq b$ équivaut à (1).

Logique propositionnelle et logique des prédicats. La logique propositionnelle ne s'occupe que des connecteurs logiques (et, ou, non, $\Rightarrow$, $\Leftrightarrow$). Elle ne « voit pas » le détail de ce qu'il y a dans une proposition entre les occurrences de ces connecteurs. Par exemple la proposition

$$\underbrace{(x \subset y \text{ et } \exists a(a \in x))}_{\text{A}} \Rightarrow \underbrace{\exists a(a \in y)}_{\text{C}}$$

est vue seulement comme $(A \text{ et } B) \Rightarrow C$. L'intérieur des propositions A, B, C est ignoré. Chacune d'elles n'est qu'une suite de symboles sans rôles particuliers. Il suffit de savoir distinguer ces trois suites, et reconnaître deux suites identiques lorsqu'elles se présentent.

La logique des prédicats traite non seulement de l'usage des connecteurs logiques, mais encore de celui des quantificateurs. Cependant lorsqu'on utilise des quantificateurs, on utilise aussi des variables individuelles. La logique des prédicats traite donc aussi de l'usage de ces variables, et notamment de la substitution. C'est ici qu'interviendront les notions de variables libres, de variables liées, et de termes librement substituables (Chap. 4), dont la maîtrise est indispensable pour la lecture de ce chapitre.

Les schémas de particularisation et d'introduction de $\exists$. La présentation des schémas de déduction élémentaires de la figure 9.1 que nous faisons dans ce paragraphe n'est qu'un premier contact, assez bref, avec le sujet. Une étude plus approfondie en sera faite dans la suite de ce chapitre, où nous verrons d'autres exemples et nous établirons les premiers schémas dérivés.

Il est indiqué de considérer ensemble les deux schémas symétriques S61 et S63 dits de *particularisation* et d'*introduction de $\exists$* . Ils nous paraissent aussi comme devant être les plus familiers au lecteur. Le premier est le schéma suivant le quel on déduit « Socrate est mortel » de $\forall x(x \text{ est mortel})$, parce que si l'on appelle A la proposition « x est mortel », alors la proposition « Socrate est

mortel » est obtenue en remplaçant la variable x dans A par le terme Socrate (Fig. 9.2). Le schéma affirme que ce qui est vrai pour *tout* objet est vrai pour un objet *particulier*, comme Socrate. D'où le nom de *particularisation*.

$$\frac{\Gamma \vdash \overbrace{\forall x(x \text{ est mortel})}^A}{\Gamma \vdash \underbrace{(\text{Socrate est mortel})}_{(\text{Socrate}|x)A}} \qquad \frac{\Gamma \vdash \overbrace{(\text{Socrate est philosophe})}^{(\text{Socrate}|x)A}}{\Gamma \vdash \underbrace{\exists x(x \text{ est philosophe})}_A}$$

Figure 9.2 Dédutions de la forme S61 (particularisation) et S63 ($\exists$ -introduction).

Le schéma S63 permet de déduire $\exists x(x \text{ est philosophe})$ de « Socrate est philosophe », en vertu de la même substitution qui transforme « x est philosophe » en « Socrate est philosophe » (Fig. 9.2). Si un objet particulier, comme Socrate, possède une certaine propriété, on ne peut certes pas en déduire que tout objet possède cette propriété, mais bien qu'il en existe un.

Les schémas de déduction de la logique des prédicats comportent souvent une *condition*, notée en principe à côté de la fraction, et qui est une restriction sur la forme des prémisses ou de la conclusion des déductions décrites par le schéma. Par exemple, les déductions de la classe définie par S61 sont les séquences de jugements de la forme

$$\frac{\Gamma \vdash \forall xA}{\Gamma \vdash (T|x)A}$$

où A est une proposition, x est une variable, et T est un terme *librement substituable à x dans A* . Si cette condition n'est pas satisfaite, la séquence de jugements n'est pas une déduction conforme au schéma S61.

La notion de terme librement substituable à une variable dans une proposition a été définie au paragraphe 4.6. Le lecteur est invité à la revoir, car c'est une condition qui est imposée en logique des prédicats à *toutes* les déductions où interviennent des substitutions.

Le schéma de généralisation. Le schéma de *généralisation* S62, appelé aussi *introduction de $\forall$* , est soumis à une condition qui intervient très fréquemment. Nous disons qu'une variable x *ne figure pas librement* dans une liste de propositions Γ , lorsque cette liste ne contient aucune proposition dans laquelle x figure librement. Il va de soi que cette condition est trivialement vérifiée lorsque la liste Γ est vide. Lorsque cette condition sur une variable x et sur les hypothèses d'un contexte est satisfaite, on dit aussi qu'aucune hypothèse n'est faite *sur x* dans ce contexte, ou encore que x , ou plutôt que l'objet désigné par x , est *quelconque* dans ce contexte. Car du point de vue sémantique, une proposition est une assertion sur les objets que désignent ses variables *libres* (§4.4, Variables libres et sémantique). L'idée du schéma de généralisation est

que ce qui est vrai pour un objet *quelconque* — au sens d'un objet sur lequel on n'a fait aucune hypothèse — est vrai pour tout objet.

Par exemple, lorsqu'on veut démontrer que tout nombre réel x possède une certaine propriété, par exemple la propriété $x^2 \geq 0$ (proposition **A**), c'est-à-dire lorsqu'on veut démontrer $\forall x(x^2 \geq 0)$, on considère *un seul* nombre réel x , mais *quelconque*, et l'on démontre que son carré est positif. Par « quelconque », on veut dire que l'on prend garde de ne faire sur ce nombre aucune hypothèse qui en ferait un cas particulier que l'on ne pourrait pas généraliser.

On dit qu'une proposition de la forme $\forall xA$ est la *généralisation sur x* de la proposition **A**. Lorsque **A** est vraie dans un contexte et lorsqu'aucune hypothèse n'est faite sur x dans ce contexte, on dit que l'on peut généraliser **A** sur x dans ce contexte.

Le schéma d'élimination de $\exists$. Il est un peu plus difficile de reconnaître dans le schéma d'*élimination de $\exists$* (**S64**) un mode de raisonnement habituel. Car bien que celui-ci soit très fréquent, il est de ceux qui s'effectuent le plus inconsciemment, et souvent par ailleurs en des termes logiquement incorrects. Nous nous étendrons donc un peu plus longuement sur l'exemple qui suit.

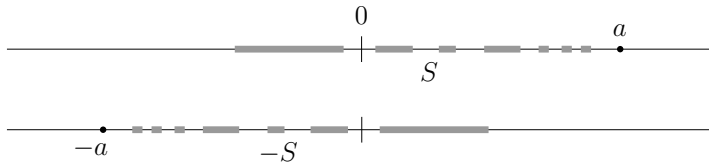


Figure 9.3 Si un ensemble S de nombres réels est majoré, l'ensemble $-S$ des nombres $-x$ tels que $x \in S$ est minoré.

Soit S un ensemble de nombres réels, $S \subset \mathbb{R}$, et supposons que cet ensemble soit *majoré*, c'est-à-dire qu'il existe un nombre réel a tel que tout nombre de S soit inférieur ou égal à a . Ces hypothèses, une fois formalisées, constituent une liste de propositions Γ . La théorie du contexte est celle des nombres réels. Dans ce contexte, la proposition suivante est vraie (par hypothèse) :

$$\exists a(a \text{ est majorant de } S). \quad (\exists xA)$$

Cette proposition correspond au $\exists xA$ du schéma **S63**. On veut démontrer, dans ce contexte, la proposition **B** suivante :

$$\text{l'ensemble } -S \text{ est minoré}, \quad (\mathbf{B})$$

où par $-S$ nous désignons l'ensemble des nombres $-x$ tels que $x \in S$, et par « $-S$ est minoré » nous voulons dire qu'il existe un nombre réel b qui est inférieur à tout nombre de cet ensemble. La démonstration formelle est esquissée dans la figure 9.4. Il s'agit de voir comment l'on passe de la ligne 2^o à la conclusion (dernière ligne). Le raisonnement informel typique que l'on peut tenir dans cet exemple est illustré par la figure 9.3. Il est le suivant :

Raisonnement informel. Soit a un majorant de S — il en existe un par hypothèse. Pour tout nombre x appartenant à l'ensemble S , on a $x \leq a$, donc $-a \leq -x$, de sorte que $-a$ est un minorant de $-S$. Donc l'ensemble $-S$ est minoré.

Nous ne détaillerons pas tous les éléments de ce raisonnement, mais seulement ceux qui concernent l'application de **S64**. Remarquons d'abord que le raisonnement commence par une *hypothèse*: soit a un majorant de S .

On fait donc l'hypothèse **A**. C'est pourquoi un cadre intérieur Γ' est tracé dans la figure 9.4. Γ' est la liste d'hypothèses $\Gamma; \mathbf{A}$. Ce qui n'est pas très clair, dans le raisonnement informel, c'est à quel endroit l'hypothèse **A** cesse d'être en vigueur. Car ce n'est pas dans le cadre Γ' que l'on veut démontrer **B**, mais bien dans le cadre Γ . Intuitivement, l'assertion **B** doit être vraie sous la seule hypothèse que l'ensemble S est majoré. L'hypothèse auxiliaire faite sur un nombre particulier a sert sans doute à la démonstration, mais on doit pouvoir établir la vérité de **B** *en dehors* du bloc de lignes de cette hypothèse. Or, dans le raisonnement informel, lorsqu'on parvient à l'assertion

$$-a \text{ est un minorant de } -S, \quad (2)$$

il est clair qu'on est encore dans le cadre Γ' , car cette proposition n'est vraie évidemment qu'en vertu de l'hypothèse faite sur a . De (2), on déduit que $-S$ est minoré, c'est-à-dire que

$$\exists b(b \text{ est un minorant de } -S). \quad (3)$$

Formellement (Fig. 9.4) cette déduction se fait suivant **S63**, et dans ce schéma, la prémisse et la conclusion ont la même liste d'hypothèses. Cette déduction se fait donc encore dans le cadre Γ' . Pour le lecteur qui n'est pas encore familiarisé avec **S63**, nous pouvons annoter cette déduction comme ceci, en désignant par $\mathbf{A}'$ la proposition « b est un minorant de $-S$ »:

$$\frac{\overbrace{\Gamma' \vdash -a \text{ est un minorant de } -S}^{(-a|b)\mathbf{A}'}}{\Gamma' \vdash \underbrace{\exists b(b \text{ est un minorant de } -S)}_{\mathbf{A}'}} \quad \exists\text{-intro}$$

C'est maintenant qu'intervient le nouveau schéma d'*élimination de $\exists$* (**S64**). Il permet de copier dans le cadre extérieur Γ la proposition **B** qui se trouve dans le cadre intérieur Γ' , et plus exactement, de faire la déduction

$$\frac{\Gamma \vdash \exists a(\overbrace{a \text{ est majorant de } S}^{\mathbf{A}}) \quad \Gamma; \overbrace{a \text{ est majorant de } S}^{\mathbf{A}} \vdash \overbrace{-S \text{ est minoré}}^{\mathbf{B}}}{\Gamma \vdash \underbrace{-S \text{ est minoré}}_{\mathbf{B}}}.$$

Cette déduction respecte les *conditions* du schéma **S64**. Premièrement, la

S64.

$$\frac{\Gamma \vdash \exists x A \quad \Gamma; A \vdash B}{\Gamma \vdash B}$$

Si x ne figure librement
ni dans B , ni dans Γ .

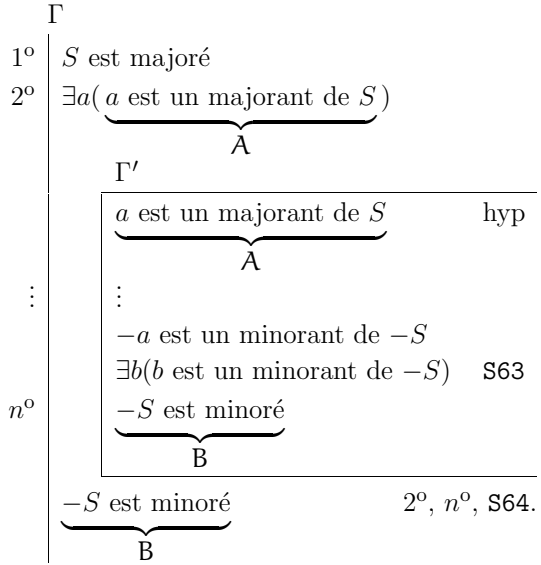


Figure 9.4 Exemple de raisonnement par élimination de $\exists$.

variable a (qui correspond au x du schéma) ne figure pas librement dans la proposition B (elle n'y figure même pas du tout). Deuxièmement, *aucune hypothèse* n'est faite sur a dans la liste Γ , c'est-à-dire que l'hypothèse « a est un majorant de S » est la *seule hypothèse* sur a en vigueur dans le cadre Γ' .

On peut décrire cette déduction intuitivement comme suit. On a démontré une proposition B (ligne n°), en considérant un objet particulier a , sur lequel on a fait l'hypothèse qu'il est un majorant de S . Mais cet objet joue un rôle purement accessoire pour les deux raisons suivantes: premièrement il n'est pas question de cet objet dans l'assertion B et deuxièmement aucune *autre* hypothèse n'est faite sur cet objet. Hormis le fait que a soit un majorant de S , aucune autre propriété particulière de a ne peut donc intervenir dans la preuve de B . En d'autres termes, a est un majorant *quelconque* de S . Si B est vraie c'est donc simplement parce qu'il *existe* un majorant de S — on peut en « prendre » un — et non en vertu du choix d'un majorant particulier. On se sent donc le droit de transférer la conclusion B du cadre Γ' dans le cadre Γ .

Cette répétition de la conclusion B est évidemment omise dans le raisonnement informel. On ne dit qu'une fois « donc $-S$ est minoré ». La fermeture du cadre Γ' est laissée « en l'air ». Bien souvent d'ailleurs, l'ouverture du cadre Γ' est aussi effacée dans le raisonnement informel. On passe directement de 2° aux conséquences de l'hypothèse « a est un majorant de S », sans déclarer cette hypothèse. On s'exprime par exemple ainsi:

il existe un nombre a qui est majorant de S (2°),
et alors (!) on a $-a \leq -x$ pour tout $x \in S$.

Les points d'interrogation et d'exclamation expriment évidemment le malaise du logicien qui lit pareil discours. Car en bonne logique, la proposition

$$-a \leq -x \text{ pour tout } x \in S,$$

qui parle d'un nombre particulier a , ne saurait se déduire de la proposition

$$\exists a(a \text{ est un majorant de } S),$$

dans laquelle il n'est question d'aucun nombre particulier — elle signifie seulement que S est majoré. Un quantificateur existentiel ne peut être éliminé qu'en faisant une *hypothèse*.

9.2 Convention d'un langage universel

La notion de théorie au sens de la logique (§5.3) est statique: une théorie est donnée par son langage, ses axiomes, et ses déductions élémentaires. Ces données sont fixées et *définissent* la théorie. Cependant, dans la pratique, les théories importantes sont en *développement permanent*. Leur langage s'enrichit constamment de nouveaux symboles, qui apparaissent dans les *définitions*. Nous examinerons plus loin (chap. 12) le mécanisme d'extension d'une théorie au moyen de définitions. Mais nous pouvons déjà en donner ici un exemple.

Exemple. Au paragraphe 5.3, nous avons défini la théorie appelée *Arithmétique du premier ordre*, en donnant son langage, ses axiomes, et en signalant que ses déductions élémentaires sont celles de la logique classique (logique des prédicats avec égalité) plus les déductions obéissant au schéma de récurrence. Les seuls symboles du langage donné sont la constante individuelle 0, le symbole fonctionnel unaire σ , (σx désigne le successeur du nombre x), les symboles fonctionnels binaires $+$ et $\cdot$, et le symbole relationnel binaire d'égalité. Les axiomes sont les six propositions

$$\begin{array}{ll} \forall x(\sigma x \neq 0) & \forall x(x \cdot 0 = 0) \\ \forall x \forall y(\sigma x = \sigma y \Rightarrow x = y) & \forall x \forall y(x + \sigma y = \sigma(x + y)) \\ \forall x(x + 0 = x) & \forall x \forall y(x \cdot \sigma y = x \cdot y + x). \end{array}$$

Mais l'égalité n'est pas la seule relation que l'on peut considérer entre deux entiers. Tôt ou tard on sera conduit à introduire de nouveaux symboles relationnels, le premier étant sans doute le symbole $\leq$, avec son sens usuel. Et puis il y aura tous les symboles relationnels qui consistent en des mots du langage naturel, comme par exemple le symbole relationnel binaire «divise» de la proposition « x divise y » ou « x est diviseur de y »; le symbole relationnel unaire «est premier» de la proposition « x est un nombre premier», etc.

Chacun de ces symboles fera l'objet d'une *définition*. Par exemple, pour le symbole d'inégalité, on posera comme définition:

$$\forall x \forall y(x \leq y \Leftrightarrow \exists a(x + a = y)).$$

Ceci est en fait un nouvel axiome, mais en vertu de sa forme et du fait que le symbole $\leq$ n'est pas utilisé dans les axiomes précédents, il possède des propriétés logiques particulières en raison desquelles on l'appelle une définition. Nous en parlerons au chapitre 12. Pour le moment nous voulons seulement mettre en évidence le besoin d'étendre le langage d'une théorie. Chaque définition s'accompagne d'une extension du langage.

Le langage universel $\mathcal{L}_\Omega$. Pour simplifier la suite de notre exposé, nous allons nous servir dorénavant d'un seul langage, *fixe*, et imaginer que nous n'en sommes pas les seuls utilisateurs, mais que ce langage est au contraire un langage « universel ». Nous allons préciser dans quel sens nous l'entendons. Ce langage doit être assez grand, car il doit contenir par avance suffisamment de symboles pour faire face à tous les besoins futurs de nouveaux symboles, comme nous venons de le décrire. Nous supposons donc donné un langage que nous appelons *le langage universel* et que nous désignons par $\mathcal{L}_\Omega$, possédant : une infinité de constantes individuelles $s_1, s_2, s_3, \dots$; une infinité de symboles fonctionnels unaires $f_1^1, f_2^1, f_3^1, \dots$; une infinité de symboles fonctionnels binaires $f_1^2, f_2^2, f_3^2, \dots$; et ainsi de suite, une infinité de symboles fonctionnels de chaque arité, et en outre une infinité de symboles relationnels $r_1^k, r_2^k, r_3^k, \dots$ pour chaque arité k .

Nous ne précisons pas quels sont ces symboles. Il nous suffit d'admettre qu'ils sont définis, ce qui est certainement possible. Notre hypothèse n'est pas irréaliste. Il suffit de prendre par exemple comme n -ième symbole fonctionnel d'arité k l'identificateur $f \circ \dots \circ \bullet \dots \bullet$, où le nombre de $\circ$ s est n et le nombre de $\bullet$ s est k . Le répertoire de symboles relationnels sera fourni par les identificateurs semblables commençant par r au lieu de f . Le répertoire de constantes individuelles se composera des identificateurs $c \circ \dots \circ$.

Convention. Nous convenons que désormais toutes les théories dont nous parlons dans la suite de cet ouvrage ont pour langage le langage universel $\mathcal{L}_\Omega$, que nous supposons bien défini, ce qui est possible comme nous venons de l'indiquer. C'est à cette convention que nous avons fait allusion plus haut en disant que nous allons nous servir dorénavant d'un langage fixe. Mais cela ne veut pas dire que nous allons réellement utiliser les symboles de ce langage. Nous utiliserons au contraire, dans tous nos exemples, des symboles *usuels*, par exemple les symboles $=, \in, \emptyset, \cup$, ou les symboles formés avec des mots comme le prédicat « est une fonction » dans la proposition « f est une fonction ».

Mais afin de pouvoir dire que nous utilisons toujours le seul et même langage $\mathcal{L}_\Omega$, nous admettrons que tous les symboles que nous employons, usuels ou non, *représentent* des symboles bien déterminés du langage universel $\mathcal{L}_\Omega$, suivant une table de correspondance contenant des indications telles que

Symbole usuel	Symbole de $\mathcal{L}_\Omega$
$=$	r_1^2
$\in$	r_2^2

$\emptyset$	$\mathbf{S}_{36}$
$\cup$	$\mathbf{f}_{157}^2$
« est une fonction »	$\mathbf{r}_{55}^1$
...	...

Ni le langage $\mathcal{L}_\Omega$, ni cette table de traduction ne sont réellement donnés. Mais ils pourraient l'être, et ce n'est pas une fiction irréaliste que d'admettre qu'ils le sont. Il pourrait y avoir un « organisme international de normalisation » chargé de traduire dans $\mathcal{L}_\Omega$ tous les langages formels utilisés dans ce monde. Le langage $\mathcal{L}_\Omega$ est suffisamment grand pour satisfaire aux besoins de tous les auteurs présents et futurs de textes formels, car à quelque date future que ce soit, tous ces auteurs n'auront jamais écrit qu'un nombre fini de pages, et n'auront donc utilisé qu'un nombre fini de symboles de ce langage.

Nous considérons ainsi tous les langages du premier ordre comme étant « plongés » dans le langage universel $\mathcal{L}_\Omega$, et lorsque nous parlons dans la suite d'un symbole, d'un terme, d'une proposition, d'une liste de propositions, d'un jugement, d'une séquence de jugements, etc. sans mentionner le langage dont il est question, il s'agit toujours du langage $\mathcal{L}_\Omega$.

Les théories $\mathbf{Lprop}$, $\mathbf{Lpréd}$, $\mathbf{Lprédég}$. Puisque nous ne considérons plus que des théories de langage $\mathcal{L}_\Omega$, nous n'avons plus besoin de mentionner explicitement le langage d'une théorie. En particulier, nous désignerons par $\mathbf{Lprop}$, $\mathbf{Lpréd}$, $\mathbf{Lprédég}$ les théories $\mathbf{Lprop}(\mathcal{L}_\Omega)$ (§6.1), $\mathbf{Lpréd}(\mathcal{L}_\Omega)$ et $\mathbf{Lprédég}(\mathcal{L}_\Omega)$ (§9.1). Lorsque nous parlons d'une *théorie logique*, il s'agit d'une théorie de langage $\mathcal{L}_\Omega$ qui est une extension (§5.7) de $\mathbf{Lprédég}$.

9.3 Particularisations

Dans ce paragraphe et dans les suivants, nous allons reprendre un à un les schémas de déduction élémentaires de la logique des prédicats (Fig. 9.1) sur les quantificateurs afin de permettre au lecteur de se familiariser avec ces formes de déduction, grâce à un certain nombre d'exemples. Quelques schémas dérivés seront introduits à cette occasion. L'élaboration d'un répertoire complet de schémas de déduction dérivés de logique des prédicats sera entreprise au chapitre 10. Il est important de bien maîtriser d'abord l'usage des schémas élémentaires $\mathbf{S61}$ – $\mathbf{S64}$. Nous consacrerons un paragraphe au rôle important du schéma de déduction *a fortiori* ($\mathbf{S2}$) en logique des prédicats.

Le schéma de particularisation. Les *conditions* qui font partie des schémas $\mathbf{S61}$ – $\mathbf{S64}$, et qui se retrouvent dans un grand nombre de schémas dérivés, sont la « difficulté » principale. La signification du schéma $\mathbf{S61}$:

$\Gamma \vdash \forall x A$	Si A est une proposition, x une variable,
$\Gamma \vdash (T x)A$	et T un terme librement substituable à x dans A

repose évidemment sur la définition précise de $(T|x)A$ (§4.5) et celle de

l'expression « T est librement substituable à x dans A » (§4.6). Nous supposons ces notions connues. Lorsqu'on applique **S61**, on dit que l'on *particularise* la proposition $\forall xA$ avec T . Par exemple (Fig. 9.2), lorsqu'on déduit « Socrate est mortel » de $\forall x(x \text{ est mortel})$, on dit que l'on particularise l'assertion générale $\forall x(x \text{ est mortel})$ avec le terme « Socrate », désignant un être particulier. On dit même: tout x est mortel, donc *en particulier* Socrate est mortel.

Avant d'examiner d'autres exemples de déductions de cette forme, nous allons présenter un schéma qui est un corollaire évident de celui-ci grâce à une démonstration triviale par introduction de $\Rightarrow$ (**S7**).

Corollaire de S61

	Γ	
$\frac{}{\Gamma \vdash \forall xA \Rightarrow (T x)A}$	1°	$\left \begin{array}{l} \forall xA \\ (T x)A \end{array} \right.$
Si T est librement substituable	2°	$\left \begin{array}{l} \text{hyp} \\ 1^\circ, \text{S61} \end{array} \right.$
à x dans A	3°	$\forall xA \Rightarrow (T x)A \quad 2^\circ, \Rightarrow\text{-intro.}$

Faisons d'abord une remarque générale à l'occasion de cette première démonstration de logique des prédicats. Lorsqu'un schéma dérivé, comme ce corollaire de **S61**, contient une ou plusieurs conditions (ici la condition: T est librement substituable à x dans A), nous supposons toujours, dans la démonstration qui est donnée, que les termes, les variables, les propositions, et la liste d'hypothèses Γ considérés dans cette démonstration *satisfont* aux conditions du schéma de déduction à démontrer. Dans cet exemple, il est admis à la ligne 2° que T désigne un terme librement substituable à x dans A , ce qui permet d'appliquer **S61**. La condition du schéma **S61** se retrouve dans le corollaire, parce que la démonstration de celui-ci contient une déduction suivant **S61**.

Ce corollaire est trivial, mais il faut se garder de croire qu'à partir d'un schéma de déduction de la forme

$$\frac{\Gamma \vdash \text{expression}_1}{\Gamma \vdash \text{expression}_2} \quad (\text{Conditions C}) \quad (1)$$

on peut toujours démontrer, par introduction de $\Rightarrow$, le schéma sans prémisses

$$\frac{}{\Gamma \vdash \text{expression}_1 \Rightarrow \text{expression}_2} \quad (\text{Conditions C}) \quad (2)$$

avec les *mêmes conditions* C . Nous verrons bientôt un exemple de schéma (1) pour lequel ce n'est pas possible: le schéma de généralisation **S62**.

Nous présentons maintenant quelques exemples de particularisations, c'est-à-dire de déductions suivant **S61**. Il s'agit de bien savoir reconnaître une substitution admissible $(T|x)A$, où A est une proposition, x une variable, T un terme *librement substituable* à x dans A .

Exemple 1. Du jugement $\Gamma \vdash \forall u(u = 0)$, dans lequel Γ est une liste d'hypothèses quelconque, on peut déduire $\Gamma \vdash 1 = 0$. Si l'on désigne en effet par A la proposition $u = 0$ et par T le terme 1 , qui est une constante individuelle, alors

la proposition $(T|u)A$ n'est autre que $1 = 0$. La condition du schéma **S61**, T est librement substituable à u dans A , est évidemment satisfaite. Elle l'est non seulement pour ce terme T particulier, mais aussi pour n'importe quel autre terme, puisque A ne contient pas de quantificateurs. Il ne peut pas y avoir de création de liens lorsqu'on substitue un terme T à u dans cette proposition.

Exemple 2. Le jugement

$$\vdash \forall a \underbrace{\exists b(b < a)}_A, \quad (3)$$

est un théorème de la théorie des nombres réels. Désignons par A la proposition $\exists b(b < a)$ qui suit le quantificateur $\forall a$ dans (3). D'après **S61**, si un terme T est librement substituable à la variable a dans la proposition A , le jugement $\vdash (T|a)A$, c'est-à-dire

$$\vdash \exists b(b < T) \quad (4)$$

est aussi un théorème de cette théorie. C'est le cas, par exemple pour le terme $\frac{x+1}{2}$. Le jugement

$$\vdash \exists b(b < \frac{x+1}{2})$$

se déduit de (3) suivant **S61**. De façon générale, un terme T est librement substituable à a dans la proposition A s'il ne contient pas d'occurrence libre de la variable b . Par contre, un terme T qui contient une occurrence libre de b n'est pas librement substituable à a dans A . Par exemple, le jugement

$$\vdash \exists b(b < \frac{b+1}{2})$$

ne se déduit pas de (3) suivant **S61**, parce que la condition du schéma n'est pas satisfaite. Il se trouve que ce jugement est quand même un théorème de la théorie en question, mais il ne se déduit pas de (3). Le lecteur n'aura pas de peine à trouver un terme T (contenant b) pour lequel l'assertion de (4) est fausse dans le contexte de la théorie des nombres réels.

Exemple 3. Soient T un terme quelconque, et Γ une liste quelconque de propositions. La séquence de jugements

$$\frac{\Gamma \vdash \forall x(x \notin \emptyset)}{\Gamma \vdash T \notin \emptyset} \quad (5)$$

est une déduction de la théorie **Lpréd** car elle est de la forme **S61**. Comme dans l'exemple 1, la condition du schéma est satisfaite puisque la proposition dans laquelle T est substitué à x est la proposition sans quantificateur $x \notin \emptyset$.

Sachant que la proposition $\forall x(x \notin \emptyset)$ est un axiome de la théorie $\mathcal{T}_{\text{ENS}}$ (§8.2) et que cette théorie est une théorie logique, donc que la déduction (5) est aussi une déduction de $\mathcal{T}_{\text{ENS}}$, nous pouvons faire sur cet exemple extrêmement simple un rappel de quelques notions fondamentales du chapitre 5 (théories, démonstrations, déductions, théorèmes).

Premièrement, le jugement $\Gamma \vdash \forall x(x \notin \emptyset)$ est un théorème de $\mathcal{T}_{\text{ENS}}$, puisque son assertion, $\forall x(x \notin \emptyset)$ est un axiome de cette théorie (§5.5, Règle 1). Deuxièmement, le jugement $\Gamma \vdash T \notin \emptyset$ est un théorème de $\mathcal{T}_{\text{ENS}}$, puisqu'il est la conclusion d'une déduction de $\mathcal{T}_{\text{ENS}}$ (5) dont la prémisse est un théorème de cette théorie (§5.5, Règle 3). Troisièmement, la séquence de jugements

$$\begin{array}{lll} 1^\circ & \Gamma \vdash \forall x(x \notin \emptyset) & \text{Axiome} \\ 2^\circ & \Gamma \vdash T \notin \emptyset & 1^\circ, \text{S61.} \end{array} \quad (6)$$

est le corps d'une démonstration de base vide de la théorie $\mathcal{T}_{\text{ENS}}$. Nous rappelons que chaque jugement du corps d'une démonstration d'une théorie $\mathcal{T}$ (§5.4, Déf. 1) vérifie l'une des conditions suivantes: (a) il figure dans la base de la démonstration; (b) il est un théorème de $\mathcal{T}$, autrement dit la conclusion d'une déduction sans prémisses de $\mathcal{T}$; (c) il se déduit de jugements qui le précèdent dans le corps de la démonstration suivant une déduction de la théorie. Dans la démonstration (6) de $\mathcal{T}_{\text{ENS}}$, le jugement de la première ligne est un théorème puisque c'est un axiome, ce qui est indiqué dans sa justification. La deuxième ligne se déduit de la première suivant une déduction de la théorie, puisque celle-ci est une théorie logique.

Un jugement de la forme $\Gamma \vdash T \notin \emptyset$, où Γ est une liste de propositions quelconque et T est un terme quelconque, est donc un théorème de $\mathcal{T}_{\text{ENS}}$. Nous pouvons donc noter comme schéma de déduction sans prémisses de $\mathcal{T}_{\text{ENS}}$:

$$\frac{}{\Gamma \vdash T \notin \emptyset.} \quad (7)$$

Par la suite, nous pouvons nous référer à ce schéma de déduction lorsque nous écrivons une ligne de cette forme dans une démonstration de $\mathcal{T}_{\text{ENS}}$. Enfin, nous pouvons donner quelques exemples de jugements qui sont des théorèmes de $\mathcal{T}_{\text{ENS}}$ de cette forme:

$$\vdash 1 \notin \emptyset, \quad \vdash \{a, b\} \notin \emptyset, \quad \vdash \sqrt{i+1} \notin \emptyset, \quad \vdash \emptyset \notin \emptyset.$$

Le schéma d'élimination de $\forall$. Notre premier schéma dérivé de logique des prédicats sera le suivant:

S65. *$\forall$ -élimination*

$$\frac{\Gamma \vdash \forall x A}{\Gamma \vdash A.}$$

Cette forme de déduction est en fait un cas particulier de la forme décrite par le schéma de particularisation S61. C'est le cas particulier où l'on prend pour terme T la variable x elle-même. La déduction

$$\frac{\Gamma \vdash \forall x A}{\Gamma \vdash (x|x)A}$$

est conforme en effet au schéma S61, car une variable x est trivialement librement substituable à elle-même dans une proposition A . La substitution ne

change rien à A et ne crée donc pas de lien. Et comme $(x|x)A$ est identique à A , la déduction est de la forme décrite par S65.

Le schéma d'élimination de quantificateurs universels S65 ne comporte *aucune condition*, contrairement au schéma d'introduction S62. C'est une déduction facile et très fréquente. Sous le même numéro S65, nous nous référerons aussi au schéma d'élimination de $\forall$ *généralisé*, permettant l'élimination simultanée de plusieurs quantificateurs :

$$\frac{\Gamma \vdash \forall x_1 \dots \forall x_n A}{\Gamma \vdash A.}$$

La démonstration d'une telle déduction n'est qu'une application répétée de S65. Par exemple, pour $n = 3$:

$$\begin{array}{lcl} \Gamma & & \\ 1^\circ & \left| \begin{array}{l} \forall x_1 \forall x_2 \forall x_3 A \end{array} \right. & \text{base} \\ 2^\circ & \left| \begin{array}{l} \forall x_2 \forall x_3 A \end{array} \right. & 1^\circ, \text{S65} \\ 3^\circ & \left| \begin{array}{l} \forall x_3 A \end{array} \right. & 2^\circ, \text{S65} \\ 4^\circ & \left| \begin{array}{l} A \end{array} \right. & 3^\circ, \text{S65.} \end{array}$$

Exemple 4. Dans la théorie $\mathcal{T}_{\text{ENS}}$, on peut écrire la démonstration ci-dessous, suivant laquelle le jugement

$$\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)$$

est un théorème de cette théorie. Le mot « nil » désigne traditionnellement, en informatique, une liste vide. Nous l'utiliserons dans la notation emboîtée des démonstrations pour signaler explicitement que la liste d'hypothèses associée au cadre extérieur est vide.

$$\begin{array}{lcl} \text{nil} & & \\ 1^\circ & \left| \begin{array}{l} \forall x \forall y (x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)) \end{array} \right. & \text{Axiome de l'inclusion} \\ 2^\circ & \left| \begin{array}{l} x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y) \end{array} \right. & 1^\circ, \text{S65.} \end{array}$$

9.4 Généralisations

Le schéma de généralisation S62.

$$\frac{\Gamma \vdash A}{\Gamma \vdash \forall x A} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array}$$

Nous rappelons que la phrase « x ne figure pas librement dans Γ » signifie (§9.1) que la liste Γ ne comporte pas d'hypothèse dans laquelle x figure librement, et que cette condition est satisfaite notamment lorsque cette liste est vide. Lorsque cette condition est satisfaite pour une variable x dans un contexte donné, on dit que x est *quelconque* dans ce contexte. Si une proposition A , exprimant une propriété de l'objet désigné par x , est vraie dans un tel contexte,

alors la proposition $\forall x A$ est vraie dans ce contexte, d'après S62. C'est ce qu'on exprime par la règle: ce qui est vrai pour un objet *quelconque* est vrai pour tout objet. Lorsqu'on effectue cette déduction pour une variable x et une proposition A , on dit que l'on *généralise* A sur x .

Une condition semblable à celle de S62 est imposée dans le schéma S64, d'élimination de $\exists$. Les schémas de déduction qui dériveront de ces deux schémas élémentaires comporteront la même condition, puisque leur démonstration contiendra des déductions suivant ces schémas élémentaires. Il est donc important, lorsqu'on manipule des quantificateurs dans les déductions, d'être attentif aux variables libres des hypothèses du contexte. Ces variables désignent des objets qui ne sont pas quelconques, puisque des hypothèses sont faites sur eux.

Exemple 1. Nous donnons ci-dessous une démonstration du théorème suivant de $\mathcal{T}_{\text{ENS}}$:

$$\vdash \forall a(a \neq \emptyset \Rightarrow \exists x(x \in a)).$$

Ce théorème affirme que tout ensemble différent de l'ensemble vide possède un élément. C'est une affirmation évidente pour quiconque a un peu l'habitude de raisonner informellement sur les ensembles. Un ensemble qui n'a pas d'éléments est égal à l'ensemble vide, donc, par contraposition, un ensemble qui n'est pas égal à l'ensemble vide possède un élément. Mais une démonstration *formelle* nécessite un axiome de $\mathcal{T}_{\text{ENS}}$ que nous n'avons pas encore énoncé. Le théorème contient un symbole d'égalité et ni les schémas de déduction de logique des prédicats, ni les deux axiomes de $\mathcal{T}_{\text{ENS}}$ introduits jusqu'ici (§8.2) ne contiennent ce symbole. Nous sommes donc obligés d'admettre sans démonstration la déduction faite ci-dessous à la ligne 2°.

nil		
1°	$a \neq \emptyset$	hyp
2°	$\exists x(x \in a)$	1°, déduction admise
3°	$a \neq \emptyset \Rightarrow \exists x(x \in a)$	2°, $\Rightarrow$ -intro
4°	$\forall a(a \neq \emptyset \Rightarrow \exists x(x \in a))$	3°, S62.

La déduction faite à la ligne 4°:

$$\frac{\vdash a \neq \emptyset \Rightarrow \exists x(x \in a)}{\vdash \forall a(a \neq \emptyset \Rightarrow \exists x(x \in a))}$$

est bien conforme au schéma S62, puisque la liste d'hypothèses de ces deux jugements est vide, donc a ne figure pas librement dans cette liste. Nous pouvons dire que dans cette démonstration a est *quelconque* dans les lignes 3° et 4° du bloc extérieur, auquel est associée la liste vide. Par contre, a n'est pas quelconque dans le bloc de lignes 1°, 2°, auquel est associée l'hypothèse $a \neq \emptyset$ dans laquelle a figure librement. On ne peut pas généraliser sur a dans le cadre

intérieur de la démonstration, par exemple faire la déduction erronée

nil	
1°	$a \neq \emptyset$ hyp
2°	$\exists x(x \in a)$ 1°
	$\forall a \exists x(x \in a)$ 2°, erreur !

Si la proposition $\exists x(x \in a)$ est vraie à la ligne 2°, c'est bien parce qu'on a fait l'hypothèse $a \neq \emptyset$. Cette proposition ne peut pas se généraliser sur a .

Exemple 2. Nous démontrons ci-dessous le théorème suivant de la théorie **Lpréd**:

$$\vdash \forall x(x \notin \emptyset) \Rightarrow \forall y(y \notin \emptyset). \quad (1)$$

Bien entendu, puisque la théorie $\mathcal{T}_{\text{ENS}}$ est une extension de **Lpréd**, ce jugement est aussi un théorème de $\mathcal{T}_{\text{ENS}}$, mais la démonstration qui suit est une démonstration de **Lpréd**. Elle ne fait intervenir aucun axiome de $\mathcal{T}_{\text{ENS}}$.

Les propositions $\forall x(x \notin \emptyset)$, $\forall y(y \notin \emptyset)$ sont formellement ou syntaxiquement *distinctes*. Mais sémantiquement elles ont le même sens: aucun objet n'est élément de l'ensemble vide. Ces deux propositions sont un exemple de propositions que l'on dit « identiques à un changement de variable liée près ». On peut démontrer dans **Lpréd** non seulement l'implication (1), mais encore l'implication réciproque, donc l'équivalence des deux propositions.

nil	
1°	$\forall x(x \notin \emptyset)$ hyp
2°	$y \notin \emptyset$ 1°, S61
3°	$\forall y(y \notin \emptyset)$ 2°, S62
4°	$\forall x(x \notin \emptyset) \Rightarrow \forall y(y \notin \emptyset)$ 3°, $\Rightarrow$ -intro.

La déduction

$$\frac{\forall x(x \in \emptyset) \vdash y \notin \emptyset}{\forall x(x \in \emptyset) \vdash \forall y(y \notin \emptyset)}$$

faite à la ligne 3° est conforme au schéma **S62**, car la variable y , sur laquelle se fait la généralisation, ne figure pas librement (elle ne figure pas du tout) dans l'hypothèse $\forall x(x \notin \emptyset)$ de ces deux jugements.

Exemple 3. Étant donné une proposition A et des variables x , y , on a dans **Lpréd** le théorème $\vdash \forall x \forall y A \Rightarrow \forall y \forall x A$. La démonstration est donnée dans la figure 9.5. Les déductions

$$\frac{\forall x \forall y A \vdash A}{\forall x \forall y A \vdash \forall x A} \qquad \frac{\forall x \forall y A \vdash \forall x A}{\forall x \forall y A \vdash \forall y \forall x A}$$

des lignes 4° et 5° sont conformes au schéma **S62**, car ni la variable x de la première généralisation, ni la variable y de la seconde, ne figurent librement dans l'hypothèse $\forall x \forall y A$ de ces jugements (§4.4, Règles sur les variables libres, (h) et (i)). Dans cet exemple, contrairement aux deux précédents, nous

avons affaire à une proposition *indéterminée* A , de même qu'à des variables indéterminées x , y . Nous ignorons quelles variables figurent librement dans la proposition $\forall x \forall y A$. La seule certitude que nous avons à ce sujet est que les variables x et y ne figurent pas librement dans cette proposition.

	nil	
1°	$\forall x \forall y A$	hyp
2°	$\forall y A$	1°, S65
3°	A	2°, S65
4°	$\forall x A$	3°, S62
5°	$\forall y \forall x A$	4°, S62
6°	$\forall x \forall y A \Rightarrow \forall y \forall x A$	4°, $\Rightarrow$ -intro.

Figure 9.5 Démonstration du théorème de l'exemple 3.

Cette démonstration présente un trait typique de nombreuses démonstrations de logique des prédicats. Elle commence par une élimination de quantificateurs et se termine par une introduction de quantificateurs. Entre ces deux phases, dans les cas plus compliqués, il y a d'autres déductions.

Exemple 4. Il ne faut pas confondre une déduction de la forme

$$\frac{\Gamma \vdash A}{\Gamma \vdash \forall x A} \quad (2)$$

avec la déduction sans prémisses

$$\frac{}{\Gamma \vdash A \Rightarrow \forall x A}. \quad (3)$$

Si x est une variable qui ne figure pas librement dans Γ , la déduction (2) est une déduction de **Lpréd** conforme au schéma de généralisation **S62**, mais la déduction (3) peut être *erronée*, c'est-à-dire n'est pas toujours admise comme déduction de **Lpréd**. Elle n'est admise, comme nous allons le voir, qu'à une condition supplémentaire extrêmement restrictive.

Il est important de bien saisir intuitivement la différence entre ces deux déductions. Pour cela, considérons un jugement $\Gamma \vdash A$ et une variable x particuliers, par exemple le jugement $\vdash x \geq 0$ et la variable x . Les déductions (2) et (3) sont dans ce cas

$$\frac{\vdash x \geq 0}{\vdash \forall x(x \geq 0)} \quad \frac{}{\vdash (x \geq 0) \Rightarrow \forall x(x \geq 0)}. \quad (4)$$

La première de ces déductions est une déduction de **Lpréd**, donc aussi une déduction de n'importe quelle théorie logique. En théorie des nombres réels, la prémisse $\vdash x > 0$ n'est pas un théorème : la proposition $x \geq 0$ ne peut être vraie que sous des hypothèses particulières. On ne peut donc rien démontrer avec

cette déduction. Mais bien dans la théorie des entiers naturels, où un nombre x quelconque est supérieur ou égal au nombre 0, ou dans la théorie d'une algèbre de Boole dont le plus petit élément est noté 0, etc. Dans toute théorie logique $\mathcal{T}$ où l'on peut démontrer $\vdash x \geq 0$, on peut démontrer $\vdash \forall x(x \geq 0)$, par **S62**. Si la seconde des déductions (4) était aussi une déduction de **Lpréd**, on aurait dans toute théorie logique le théorème

$$\vdash x \geq 0 \Rightarrow \forall x(x \geq 0). \quad (5)$$

Cela est immédiatement suspect si l'on pense aux nombres réels. Nous pouvons justifier cette suspicion, en montrant que

$$\frac{\vdash x \geq 0 \Rightarrow \forall x(x \geq 0) \quad \vdash 0 \geq 0}{\vdash \forall x(x \geq 0)} \quad (6)$$

est une déduction de **Lpréd**. Donc si (5) était un théorème de la théorie des nombres réels, comme $\vdash 0 \geq 0$ en est un, $\vdash \forall x(x \geq 0)$ en serait un troisième et alors cette théorie serait contradictoire car la négation de $\forall x(x \geq 0)$ est vraie. La démonstration de (6) est la suivante:

nil		
1°	$x \geq 0 \Rightarrow \forall x(x \geq 0)$	<i>base</i>
2°	$0 \geq 0$	<i>base</i>
3°	$\forall x(x \geq 0 \Rightarrow \forall x(x \geq 0))$	1°, S62
4°	$0 \geq 0 \Rightarrow \forall x(x \geq 0)$	3°, S61
5°	$\forall x(x \geq 0)$	1°, 4°, mod ponens.

Il est important de bien observer les liens de la proposition de la ligne 3° et la particularisation (**S61**) qui suit. Nous les mettons en évidence:

$$\frac{(3^\circ) \vdash \forall x(\boxed{x \geq 0} \Rightarrow \forall x(\boxed{x \geq 0}))}{(4^\circ) \vdash 0 \geq 0 \Rightarrow \forall x(\boxed{x \geq 0})} \text{S61} \quad (7)$$

Dans une déduction de schéma **S61** on passe de la prémisse à la conclusion en ôtant le quantificateur universel initial $\forall x$ de la prémisse $\forall xA$ et en remplaçant par un terme T les occurrences de x dans A qui étaient liées à ce quantificateur. Ce sont les occurrences libres de x dans A .

9.5 Rôle des déductions structurelles

Nous avons montré au paragraphe 6.5 que les séquences de jugements de la forme

$$\frac{\vdash B}{\Gamma \vdash B} \quad (1)$$

sont des déductions de la théorie **Lprop**. Si Γ est une liste de n propositions $A_1; \dots; A_n$, la déduction (1) se démontre simplement par n applications du

schéma **S2** (vérité a fortiori):

1°	$\vdash B$	<i>base</i>
2°	$A_1 \vdash B$	1°, a fortiori
3°	$A_1; A_2 \vdash B$	2°, a fortiori
...		

Les déductions de la forme (1) sont un cas particulier de toutes les déductions dans lesquelles on ne fait qu'agrandir la liste d'hypothèses d'un jugement, et que nous désignons toutes par *a fortiori* (§6.5).

Le schéma de déduction (1) est très utile en logique des prédicats, et plus généralement dans toute théorie logique. En vertu de ce schéma, lorsqu'on doit démontrer un théorème $\Gamma \vdash B$, il suffit de démontrer le théorème sans hypothèses $\vdash B$, puisque le premier s'en déduit suivant (1). Bien entendu ce n'est pas toujours possible. Mais lorsqu'on établit des *schémas de déduction* dérivés, on a affaire à des listes d'hypothèses indéterminées Γ qui ne peuvent pas intervenir dans la démonstration puisqu'on ne les connaît pas. L'avantage d'une liste d'hypothèses vide, est qu'une variable x ne figure pas librement dans cette liste, ce qui est souvent une condition d'application d'un schéma de déduction.

Exemple 1. Nous allons établir le schéma de déduction dérivé suivant de logique des prédicats:

$$\frac{}{\Gamma \vdash A \Rightarrow \forall x A} \quad \begin{array}{l} \text{si } x \text{ ne figure pas} \\ \text{librement dans } A. \end{array} \quad (2)$$

Nous avons déjà parlé des déductions de cette forme (§9.4, Exemple 4), en disant qu'elles étaient soumises à une condition très restrictive. Il ne faut pas confondre, en effet, la condition de (2) avec la condition « x ne figure pas librement dans Γ » du schéma de généralisation **S62**. Sont par exemple de la forme (2) les déductions sans prémisse suivantes:

$$\begin{array}{c} \frac{}{\Gamma \vdash 2 + 2 = 4 \Rightarrow \forall y(2 + 2 = 4)} \quad \frac{}{\Gamma \vdash x \in e \Rightarrow \forall a(x \in e)} \\[1em] \frac{}{\Gamma \vdash \exists x(x \in e) \Rightarrow \forall x \exists x(x \in e)} \quad \frac{}{\Gamma \vdash \forall x(x \notin \emptyset) \Rightarrow \forall x \forall x(x \notin \emptyset)}. \end{array}$$

Ce ne sont certes pas des déductions que l'on fait tous les jours. Mais le schéma de déduction (2) sera utilisé dans la démonstration de schémas couramment utilisés.

En vertu du schéma *a fortiori*, il nous suffit de démontrer les déductions de la forme (2) dans lesquelles la liste d'hypothèses Γ est vide. Cette démonstration est la suivante, où A représente une proposition dans laquelle x ne figure

pas librement :

	nil	
1°	A	hyp
2°	$\forall xA$	1°, S62
3°	$A \Rightarrow \forall xA$	2°, $\Rightarrow$ -intro

La déduction faite à la ligne 2° est

$$\frac{(1^\circ) \quad A \vdash A}{(2^\circ) \quad A \vdash \forall xA.} \quad (3)$$

Elle est conforme à S62, puisque x ne figure pas dans l'hypothèse A de ces deux jugements. Nous ne pourrions pas faire la même démonstration pour une liste d'hypothèses Γ *quelconque*, au lieu de la liste vide nil, car alors la déduction de

la deuxième ligne serait $\frac{\Gamma; A \vdash A}{\Gamma; A \vdash \forall xA}$ et une telle déduction n'est conforme

à S62 que si x ne figure pas librement dans la liste d'hypothèses $\Gamma; A$, c'est-à-dire ni dans Γ , ni dans A . Or le schéma (2) ne fait aucune restriction sur Γ . La seule façon de démontrer le théorème $\Gamma \vdash A \Rightarrow \forall xA$, pour une liste d'hypothèses Γ *quelconque* et une proposition A dans laquelle x ne figure pas librement, est de démontrer le théorème *sans hypothèses* $\vdash A \Rightarrow \forall xA$ et d'en déduire $\Gamma \vdash A \Rightarrow \forall xA$ par a fortiori (1).

Exemple 2. Si A est une proposition quelconque, x une variable, et Γ une liste de propositions quelconque, le jugement

$$\Gamma \vdash \forall x(A \text{ ou } \text{non}A) \quad (4)$$

est un théorème de **Lpréd**. Il suffit de le démontrer dans le cas particulier où Γ est vide :

	nil	
1°	$A \text{ ou } \text{non}A$	tiers exclu
2°	$\forall x(A \text{ ou } \text{non}A)$	1°, S62.

De façon générale, si $\vdash B$ est un théorème de **Lpréd**, en particulier si $\vdash B$ est un théorème de **Lprop** comme $\vdash A \text{ ou } \text{non}A$, $\Gamma \vdash \forall xB$ est un théorème de **Lpréd** qui se déduit du premier suivant S62 et a fortiori.

Exemple 3. Soient A et B des propositions et x une variable. Nous allons montrer que si Γ est une liste de propositions quelconque, la séquence de jugements

$$\frac{\Gamma \vdash \forall xA \quad \Gamma \vdash \forall x(A \Rightarrow C)}{\Gamma \vdash \forall xC} \quad (5)$$

est une déduction de **Lpréd**. Il se trouve cependant que la démonstration évidente de cette déduction par élimination du quantificateur $\forall x$, modus ponens et réintroduction du quantificateur (généralisation) ne peut se faire que

si x ne figure pas librement dans Γ (ligne 6°):

Γ		
1°	$\forall xA$	<i>base</i>
2°	A	1°, S65
3°	$\forall x(A \Rightarrow C)$	<i>base</i>
4°	$A \Rightarrow C$	3°, S65
5°	C	2°, 4°, mod ponens
6°	$\forall xC$	5°, S62.

Nous n'avons démontré ainsi que le schéma de déduction *restreint*

$$\frac{\Gamma \vdash \forall xA \quad \Gamma \vdash \forall x(A \Rightarrow C)}{\Gamma \vdash \forall xC} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array} \quad (6)$$

Mais il y a une manière standard de démontrer le schéma général (5) en utilisant le schéma particulier (6). Dans cette démonstration, la forme exacte des assertions des prémisses et de la conclusion de (5) n'intervient pas. Seul intervient le fait que la variable x ne figure pas librement dans les assertions des prémisses. Désignons ces assertions par A_1 et A_2 et désignons par B l'assertion $\forall xC$. Nous pouvons récrire le schéma (6) sous la forme

$$\frac{\Gamma \vdash A_1 \quad \Gamma \vdash A_2}{\Gamma \vdash B} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array} \quad (7)$$

Si Γ est une liste de propositions quelconque, dans laquelle x peut figurer librement, alors l'arbre suivant est le corps d'une démonstration de **Lpréd** de base $\Gamma \vdash A_1, \Gamma \vdash A_2$:

$$\frac{\Gamma \vdash A_1 \quad \Gamma \vdash A_2 \quad \frac{\frac{}{A_1; A_2 \vdash A_1} \text{hyp} \quad \frac{}{A_1; A_2 \vdash A_2} \text{hyp}}{A_1; A_2 \vdash B} \text{schéma restreint (7)}}{\Gamma \vdash B} \text{transitivité des jugements}$$

Le schéma restreint (7) est utilisé correctement puisque la variable x ne figure pas librement dans la liste de propositions A_1, A_2 . Les autres déductions de cette démonstration sont des déductions structurales de logique propositionnelle (§6.5).

Méthode de démonstration. On peut procéder comme dans l'exemple précédent chaque fois que l'on veut démontrer dans une théorie logique $\mathcal{T}$ un schéma de déduction de la forme

$$\frac{\Gamma \vdash A_1 \quad \dots \quad \Gamma \vdash A_n}{\Gamma \vdash B} \quad (8)$$

pour lequel il y a une variable x qui ne figure pas librement dans les propositions

$A_1, \dots, A_n$. Il suffit de démontrer le schéma restreint

$$\frac{\Gamma \vdash A_1 \quad \dots \quad \Gamma \vdash A_n}{\Gamma \vdash B.} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array} \quad (9)$$

De façon précise, nous pouvons énoncer la *règle* suivante.

Soient B une proposition, x une variable et $A_1, \dots, A_n$ des propositions dans lesquelles x ne figure pas librement. Pour montrer que toute séquence de jugements de la forme (8) est une déduction d'une théorie logique $\mathcal{T}$, il suffit de le montrer pour le cas particulier des séquences de jugements de cette forme telles que x ne figure pas librement dans Γ . Car alors

$$\frac{A_1; \dots; A_n \vdash A_1 \quad \dots \quad A_1; \dots; A_n \vdash A_n}{A_1; \dots; A_n \vdash B}$$

est une déduction de $\mathcal{T}$, de sorte qu'on peut démontrer (8) pour une liste Γ quelconque comme dans l'exemple précédent, en utilisant les schémas structuraux *hypothèse* généralisé et *transitivité des jugements* (§6.5):

$$\frac{}{A_1; \dots; A_n \vdash A_i} \quad \frac{\Gamma \vdash A_1 \quad \dots \quad \Gamma \vdash A_n \quad A_1; \dots; A_n \vdash B}{\Gamma \vdash B.}$$

9.6 Introduction de quantificateurs existentiels

Le schéma d'introduction de $\exists$ (S63).

$$\frac{\Gamma \vdash (T|x)A}{\Gamma \vdash \exists xA} \quad \begin{array}{l} \text{Si } T \text{ est librement} \\ \text{substituable à } x \text{ dans } A. \end{array}$$

Si, en substituant un terme particulier T à une variable x dans une proposition A , on obtient une proposition, $(T|x)A$, qui est vraie dans un contexte, alors on peut en déduire que la proposition $\exists xA$ est vraie dans ce contexte, à *condition* que T soit librement substituable à x dans A . En vertu de ce schéma de déduction, pour prouver qu'une proposition d'existence $\exists xA$ est vraie dans un contexte, il suffit de trouver un terme T tel que la proposition $(T|x)A$ soit vraie dans ce contexte et qui soit librement substituable à x dans A . Sémantiquement, un tel terme représente un objet qui « vérifie A », et lorsqu'on en a trouvé un, on sait qu'il en existe un.

À partir de ce schéma, comme pour S61 (§9.3), on peut démontrer immédiatement:

Corollaire de S63

$\frac{\Gamma \vdash (T x)A \Rightarrow \exists xA}{\text{Si } T \text{ est librement substituable à } x \text{ dans } A}$	Γ 1° 2° 3°	<table style="border: 1px solid black; margin: 0 auto; border-collapse: collapse;"> <tr> <td style="padding: 5px;">$(T x)A$</td> <td style="padding: 5px;">hyp</td> </tr> <tr> <td style="padding: 5px;">$\exists xA$</td> <td style="padding: 5px;">1°, S63</td> </tr> <tr> <td style="padding: 5px;">$(T x)A \Rightarrow \exists xA$</td> <td style="padding: 5px;">2°, $\Rightarrow$-intro.</td> </tr> </table>	$(T x)A$	hyp	$\exists xA$	1° , S63	$(T x)A \Rightarrow \exists xA$	2° , $\Rightarrow$ -intro.
$(T x)A$	hyp							
$\exists xA$	1° , S63							
$(T x)A \Rightarrow \exists xA$	2° , $\Rightarrow$ -intro.							

Exemple 1. La séquence de jugements

$$\frac{\vdash 1 > 0}{\vdash \exists x(x > 0)} \quad (1)$$

est une déduction de **Lprédég** conforme au schéma **S63**, car l'assertion du premier, $1 > 0$, est identique à $(1|x)(x > 0)$, et le terme 1 est librement substituable à x dans la proposition $x > 0$. Avec la même prémisse $\vdash 1 > 0$ on peut faire d'ailleurs une infinité d'autres déductions de cette forme, par exemple

$$\frac{\vdash 1 > 0}{\vdash \exists a(a > 0)} \quad \frac{\vdash 1 > 0}{\vdash \exists x'(x' > 0)} \quad \frac{\vdash 1 > 0}{\vdash \exists z''(z'' > 0)} \quad \dots$$

Car la proposition $1 > 0$ est identique à chacune des propositions $(1|a)(a > 0)$, $(1|x')(x' > 0)$, $(1|z'')(z'' > 0)$, etc.

Comparaison des schémas S61 et S63. Nous comparons ci-dessous la déduction (1), suivant **S63**, avec la *déduction symétrique* suivant le schéma de particularisation **S61**, en rappelant à côté les deux schémas :

$$\begin{array}{lll} \frac{\vdash \forall x(x > 0)}{\vdash 1 > 0} & (\text{S61}) & \frac{\Gamma \vdash \forall xA}{\Gamma \vdash (T|x)A.} \quad \begin{array}{l} T \text{ librement substituable} \\ \text{à } x \text{ dans } A. \end{array} \\ \\ \frac{\vdash 1 > 0}{\vdash \exists x(x > 0)} & (\text{S63}) & \frac{\Gamma \vdash (T|x)A}{\Gamma \vdash \exists xA.} \quad \begin{array}{l} T \text{ librement substituable} \\ \text{à } x \text{ dans } A. \end{array} \end{array}$$

Le même jugement, $\vdash 1 > 0$, est la conclusion de la première déduction et la prémisse de la seconde. La prémisse de la première déduction et la conclusion de la seconde sont identiques au changement $\exists/\forall$ près. Le but de cette discussion est de comparer la manière dont ces déductions sont décrites par les deux schémas. Le rôle d'un schéma de déduction (à une prémisse) est de décrire le rapport de forme qu'il y a entre la prémisse et la conclusion des déductions d'une certaine classe. La déduction, en tant qu'opération logique, va toujours dans le sens prémisse $\rightarrow$ conclusion. Mais pour décrire le rapport de forme qu'il y a entre la prémisse et la conclusion, on peut aussi bien choisir de caractériser la forme de la conclusion par rapport à celle de la prémisse, ou inversement, la forme de la prémisse par rapport à celle de la conclusion. C'est le rapport de l'une à l'autre qui importe. En tant que descriptions de déductions, les schémas **S61** et **S63** vont précisément en sens contraires. Le premier décrit la forme de la conclusion à partir de celle de la prémisse: l'assertion de la conclusion est décrite comme étant la proposition obtenue en ôtant le quantificateur de l'assertion de la prémisse et en remplaçant la variable correspondante par un terme T , le terme 1 dans l'exemple. Le schéma **S63** décrit au contraire la forme de la prémisse à partir de celle de la conclusion: une déduction conforme à ce schéma est une déduction dans laquelle l'assertion de la prémisse est identique à la proposition que l'on obtient en prenant l'assertion de la conclusion, en lui

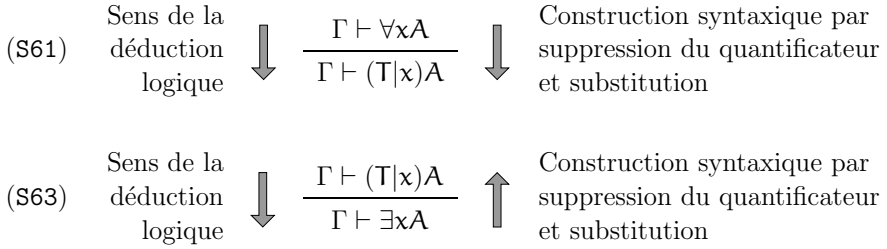


Figure 9.7 Comparaison des schémas
de déduction S61 et S63.

ôtant son quantificateur, et en remplaçant la variable correspondante par un terme T . Cette comparaison des deux schémas est résumée dans la figure 9.7.

Les déductions décrites par le schéma S63 peuvent aussi être décrites dans l'autre sens. Autrement dit, on peut remplacer ce schéma par un autre qui décrit la même classe de déductions, mais qui le fait en caractérisant la forme de la conclusion par rapport à celle de la prémisse. Cette description est plus naturelle, car il est sans doute plus naturel de considérer la conclusion d'une déduction comme étant construite par une certaine transformation de la prémisse que l'inverse. Le problème, avec cette description « naturelle », dans le cas de S63 et de quelques autres schémas dérivés, est qu'elle est beaucoup plus compliquée que la description de la prémisse par rapport à la conclusion, telle qu'elle est donnée par S63. Il n'est pas possible de s'en rendre compte d'après un seul exemple aussi simple que le précédent, mais nous le verrons plus loin.

Exemple 2. Nous allons voir maintenant un exemple qui montre l'importance de la condition « T est librement substituable à x » dans le schéma S63. Premièrement, la déduction

$$\frac{\vdash x < x + 1}{\vdash \exists y(x < y)} \quad (2)$$

est conforme à S63, car l'assertion de la prémisse est identique à la proposition obtenue en prenant l'assertion de la conclusion, en ôtant le quantificateur $\exists y$ et en remplaçant y par le terme $x + 1$. Autrement dit, la proposition $x < x + 1$ est identique à la proposition $(x + 1|y)(x < y)$. D'autre part, la condition de S63 est respectée: le terme $x + 1$ est librement substituable à y dans la proposition $x < y$. C'est un cas trivial, puisque cette proposition n'a pas de quantificateurs.

Considérons maintenant la séquence de jugements

$$\frac{\vdash \forall x(x < x + 1)}{\vdash \exists y \forall x(x < y)}. \quad (3)$$

Cette séquence n'est pas une déduction conforme à S63. Elle est d'ailleurs immédiatement suspecte si l'on pense aux nombres entiers (ou aux nombres

réels), pour lesquels la proposition $\forall x(x < x + 1)$ est vraie, mais la proposition $\exists y\forall x(x < y)$ est fausse: il n'existe pas de nombre y qui soit supérieur à tout nombre x . Pourtant, syntaxiquement, la première de ces propositions peut s'obtenir en prenant la seconde, en ôtant le quantificateur $\exists y$ et en remplaçant la variable y par un terme: le terme $x + 1$. Autrement dit,

$$\forall x(x < x + 1) \text{ est identique à } (x + 1|y)\forall x(x < y).$$

Mais la condition de **S63** n'est pas respectée. Le terme $x + 1$ n'est pas librement substituable à y dans la proposition $\forall x(x < y)$. La substitution crée un lien:

$$\forall x(\overline{\overline{x}} < y) \qquad \forall x(\overline{\overline{x}} < \overline{\overline{x}} + 1) .$$

Exemple 3. Considérons un théorème très simple de la théorie des ensembles, le théorème $\vdash \emptyset \notin \emptyset$, déjà mentionné précédemment (§9.3, fin de l'Exemple 3). Nous voulons énumérer tous les théorèmes d'existence (de $\mathcal{T}_{\text{ENS}}$), de la forme $\vdash \exists xA$, qui se déduisent de ce théorème suivant **S63**. Cet exemple a pour but de montrer la *multiplicité* des jugements que l'on peut déduire, suivant **S63**, d'un jugement particulier. Les trois déductions suivantes sont en effet conformes à **S63**:

$$\frac{\vdash \emptyset \notin \emptyset}{\vdash \exists x(x \notin \emptyset)} \qquad \frac{\vdash \emptyset \notin \emptyset}{\vdash \exists y(\emptyset \notin y)} \qquad \frac{\vdash \emptyset \notin \emptyset}{\vdash \exists z(z \notin z)}.$$

Dans les trois cas, la proposition de la prémisse peut se construire syntaxiquement à partir de celle de la conclusion par élimination du quantificateur existentiel et substitution du terme $\emptyset$ à la variable correspondante, substitution qui ne crée évidemment pas de lien. Du théorème $\vdash \emptyset \notin \emptyset$, on tire donc d'après **S63** les trois théorèmes d'existence:

$$(a) \vdash \exists x(x \notin \emptyset), \qquad (b) \vdash \exists y(\emptyset \notin y), \qquad (c) \vdash \exists z(z \notin z).$$

Le deuxième peut être utilisé encore comme prémisse d'une nouvelle déduction suivant **S63**:

$$\frac{\vdash \exists y(\emptyset \notin y)}{\vdash \exists x\exists y(x \notin y)}.$$

D'où le quatrième théorème: (d) $\vdash \exists x\exists y(x \notin y)$.

Ce sont les quatre théorèmes de $\mathcal{T}_{\text{ENS}}$ essentiellement distincts que l'on peut tirer de $\vdash \emptyset \notin \emptyset$ en n'utilisant que **S63**. Pour quiconque « possède le sens » des quantificateurs existentiels, ces quatre théorèmes d'existence sont des conséquences évidentes de $\vdash \emptyset \notin \emptyset$. Du fait que l'ensemble vide n'est pas élément de lui-même, on peut affirmer: (a) il existe un ensemble qui n'est pas élément de l'ensemble vide; (b) il existe un ensemble dont l'ensemble vide n'est pas élément; (c) il existe un ensemble qui n'est pas élément de lui-même; et enfin, (d) il existe deux ensembles tels que le premier n'est pas élément du deuxième — on ne dit pas que ces deux ensembles sont distincts: dans $\vdash \exists x\exists y(x \notin y)$, il n'est pas affirmé une existence d'objets x et y *distincts*.

Un autre théorème d'existence de $\mathcal{T}_{\text{ENS}}$ qui peut être regardé aussi comme une conséquence intuitivement évidente du théorème $\vdash \emptyset \notin \emptyset$ est le suivant : $\vdash$ il existe un ensemble qui n'est pas élément de lui-même et qui n'est pas élément de l'ensemble vide. Formellement :

$$\vdash \exists x(x \notin x \text{ et } x \notin \emptyset).$$

Démonstration. De $\vdash \emptyset \notin \emptyset$, par une déduction triviale de logique propositionnelle, on tire $\vdash (\emptyset \notin \emptyset \text{ et } \emptyset \notin \emptyset)$, et l'on applique ensuite **S63** :

$$\frac{\vdash (\emptyset \notin \emptyset \text{ et } \emptyset \notin \emptyset)}{\vdash \exists x(x \notin x \text{ et } x \notin \emptyset)}. \quad (4)$$

La proposition de la prémisse de cette déduction est en effet identique à $(\emptyset|x)\mathbf{A}$, où $\mathbf{A}$ est la proposition qui s'étend à droite du quantificateur $\exists x$ dans la conclusion, et le terme $\emptyset$ est trivialement librement substituable à x dans $\mathbf{A}$, puisqu'il ne contient pas de variable.

Avec la prémisse $\vdash (\emptyset \in \emptyset \text{ et } \emptyset \in \emptyset)$ et le quantificateur $\exists x$, on peut faire 16 déductions distinctes de schéma **S63**. Cela va de la déduction

$$\frac{\vdash (\emptyset \notin \emptyset \text{ et } \emptyset \notin \emptyset)}{\vdash \exists x(\emptyset \notin \emptyset \text{ et } \emptyset \notin \emptyset)} \quad (5)$$

à la déduction

$$\frac{\vdash (\emptyset \notin \emptyset \text{ et } \emptyset \notin \emptyset)}{\vdash \exists x(x \notin x \text{ et } x \notin x),}$$

en passant par la déduction (4). Le cas (5) est le cas limite du schéma **S63**, où la variable x ne figure pas librement dans la proposition $\mathbf{A}$ et donc $(\mathbf{T}|x)\mathbf{A}$ est identique à $\mathbf{A}$ (§4.5). Ces 16 déductions peuvent être décrites en disant que l'on passe de la prémisse à la conclusion en choisissant arbitrairement certaines occurrences du terme $\emptyset$ dans la prémisse, en les remplaçant par x , et en ajoutant $\exists x$ devant la proposition obtenue. C'est une description qui va dans le sens de la déduction (voir plus haut notre comparaison des schémas **S61** et **S63**). Mais c'est encore un exemple simple, dans lequel le choix de la variable (x) peut être quelconque, de même que le choix des occurrences du terme qui sont remplacées par x .

Exemple 4. Considérons les propositions $\mathbf{B}, \mathbf{A}_0, \dots, \mathbf{A}_7$ suivantes :

$$\begin{array}{llll} \mathbf{B} : & \forall x(x < x + 1) & \text{et} & a < x + 1 \quad \text{et} \quad x + 1 < b \\ \mathbf{A}_0 : & \forall x(x < x + 1) & \text{et} & a < x + 1 \quad \text{et} \quad x + 1 < b \\ \mathbf{A}_1 : & \forall x(x < x + 1) & \text{et} & a < y \quad \text{et} \quad x + 1 < b \\ \mathbf{A}_2 : & \forall x(x < x + 1) & \text{et} & a < x + 1 \quad \text{et} \quad y < b \\ \mathbf{A}_3 : & \forall x(x < x + 1) & \text{et} & a < y \quad \text{et} \quad y < b \\ \mathbf{A}_4 : & \forall x(x < y) & \text{et} & a < x + 1 \quad \text{et} \quad x + 1 < b \\ \mathbf{A}_5 : & \forall x(x < y) & \text{et} & a < y \quad \text{et} \quad x + 1 < b \\ \mathbf{A}_6 : & \forall x(x < y) & \text{et} & a < x + 1 \quad \text{et} \quad y < b \\ \mathbf{A}_7 : & \forall x(x < y) & \text{et} & a < y \quad \text{et} \quad y < b. \end{array}$$

Le lecteur peut vérifier que pour chacune des 8 propositions A_i , la proposition $(x + 1|y)A_i$ est identique à B. Mais le terme $x + 1$ n'est librement substituable à y que dans les quatre premières. Les quatre déductions

$$\frac{\vdash B}{\vdash \exists y A_0} \quad \frac{\vdash B}{\vdash \exists y A_1} \quad \frac{\vdash B}{\vdash \exists y A_2} \quad \frac{\vdash B}{\vdash \exists y A_3} \quad (6)$$

sont donc conformes à S63. Par contre les déductions analogues avec A_4 , A_5 , A_6 , A_7 ne sont pas conformes à S63.

On peut décrire les déductions (6) en disant que dans chacune d'elles on passe de la prémisse à la conclusion en choisissant certaines occurrences du terme $x + 1$, en les remplaçant par y et en écrivant $\exists y$ devant la proposition obtenue. Mais on ne peut pas choisir arbitrairement les occurrences de $x + 1$ que l'on remplace par y , puisque les propositions A_4 à A_7 ne conviennent pas. D'autre part, on peut prendre une autre variable que y pour remplacer les occurrences choisies de $x + 1$, et les possibilités de choix de cette variable *dépendent* du choix des occurrences de $x + 1$ que l'on remplace. Par exemple, la proposition

$$A'_3 : \forall x(x < x + 1) \text{ et } a < x \text{ et } x < b,$$

au lieu de A_3 , convient. La déduction $\frac{\vdash B}{\vdash \exists x A'_3}$ est conforme à S63. Seules les deux dernières occurrences de x dans A'_3 sont libres, et lorsqu'on les remplace par $x + 1$ on obtient B. En outre, le terme $x + 1$ est librement substituable à x dans A'_3 . Par contre la proposition

$$A'_1 : \forall x(x < x + 1) \text{ et } a < x \text{ et } x + 1 < b,$$

au lieu de A_1 , ne convient pas. La déduction $\frac{\vdash B}{\vdash \exists x A'_1}$ n'est pas conforme à S63. Lorsqu'on remplace les deux occurrences libres de x dans A'_1 par $x + 1$, on n'obtient pas B.

Le lecteur peut se rendre compte, d'après cet exemple, que la seule manière simple de décrire les déductions de schéma S63 est ce schéma lui-même, qui décrit la forme de la prémisse par rapport à celle de la conclusion, et non l'inverse.

En pratique, on opère bien cependant de la manière décrite dans cet exemple, à savoir que l'on construit la conclusion d'une déduction de schéma S63 à partir de sa prémisse comme suit: on choisit dans l'assertion B de la prémisse une ou plusieurs occurrences d'un certain terme T, on les remplace par une variable x appropriée, et l'on écrit $\exists x$ devant la proposition A obtenue. On *vérifie* ensuite que la déduction est bien conforme au schéma, à savoir que B est identique à $(T|x)A$ et que T est librement substituable à x dans A.

Exemple 5. On veut démontrer le théorème suivant de la théorie des entiers:

$$\vdash (x = 2a + 1) \Rightarrow (x^2 \text{ est impair}).$$

En vertu du schéma $\Rightarrow$ -introduction, il suffit de démontrer (x^2 est impair) sous l'hypothèse $x = 2a + 1$. Par définition de « est impair », la proposition (x^2 est impair) équivaut à

$$\exists b(x^2 = 2b + 1).$$

C'est celle-ci que nous démontrerons, sous l'hypothèse $x = 2a + 1$. Pour cela, d'après S63, il nous suffit de trouver un terme T tel que la proposition

$$(T|b)(x^2 = 2b + 1)$$

soit vraie, sous l'hypothèse mentionnée. Or, sous l'hypothèse $x = 2a + 1$, un tel terme est vite trouvé, puisqu'on a $x^2 = 4a^2 + 4a + 1$, donc

$$x^2 = 2(2a^2 + 2a) + 1.$$

On peut prendre pour T le terme $2a^2 + 2a$. La démonstration formelle complète se présente comme suit :

nil		
1°	$x = 2a + 1$	hyp
2°	$x^2 = 2(2a^2 + 2a) + 1$	1°, calcul
3°	$\exists b(x^2 = 2b + 1)$	2°, S63
4°	x^2 est impair	3°, définition de impair
5°	$(x = 2a + 1) \Rightarrow (x^2 \text{ est impair})$	4°, $\Rightarrow$ -intro.

Le contrôle de l'application correcte de S63 (ligne 3°) consiste encore une fois à vérifier que la proposition de la ligne 2° est de la forme $(T|b)A$, où A est la proposition de la ligne 3° privée de son quantificateur, et T est un terme librement substituable à b dans A . Cette dernière condition est d'ailleurs trivialement vérifiée puisque A est une proposition sans quantificateurs.

Introduction de quantificateur existentiel banale. Nous appelons introduction de quantificateur existentiel *banale* une déduction de la forme

S66.

$$\frac{\Gamma \vdash A}{\Gamma \vdash \exists x A}.$$

Une telle déduction est conforme au schéma S63, car elle peut s'écrire

$$\frac{\Gamma \vdash (x|x)A}{\Gamma \vdash \exists x A}.$$

Pour une proposition A et une variable x quelconques, en effet, la proposition $(x|x)A$ est identique à A , et la variable x est trivialement librement substituable à elle-même dans A . Cette forme de déduction est symétrique de celle du schéma S65 (élimination de $\forall$). Elle est suffisamment fréquente pour que nous lui donnions un numéro particulier et un nom. Observons que le schéma S66, tout comme son symétrique S65, est sans conditions.

Le schéma suivant est un corollaire de S66 par le fait que tout se déduit d'une contradiction (S16).

S67.

$\Gamma \vdash \text{non}\exists xA$		
$\Gamma \vdash A \Rightarrow B.$		
Γ		
1°	$\text{non}\exists xA$	<i>base</i>
2°	A	<i>hyp</i>
3°	$\exists xA$	2°, S66
4°	B	1°, 3°, S16
5°	$A \Rightarrow B$	4°, $\Rightarrow$ -intro.

Exemple 6. On démontre ci-dessous le théorème suivant de Lpréd:

$$\vdash \text{non}\exists x(x \in a) \Rightarrow \forall x(x \notin a).$$

La notation $x \notin a$ est une abréviation de $\text{non}(x \in a)$. Sans cette abréviation, le théorème présente une forme plus symétrique:

$$\vdash \text{non}\exists x(x \in a) \Rightarrow \forall x \text{non}(x \in a).$$

Il sera généralisé plus loin sous la forme d'un schéma de déduction sans prémisses de Lpréd, suivant lequel une proposition de la forme $\text{non}\exists xA$ est *équivalente*, dans tout contexte logique, à la proposition $\forall x \text{non}A$.

<i>nil</i>		
1°	$\text{non}\exists x(x \in a)$	<i>hyp</i>
2°	$x \in a$	<i>hyp</i>
3°	$\exists x(x \in a)$	2°, S66
4°	$\perp$	3°, 1°, $\perp$ -intro
5°	$x \notin a$	4°, red abs J
6°	$\forall x(x \notin a)$	5°, S62
7°	$\text{non}\exists x(x \in a) \Rightarrow \forall x(x \notin a)$	6°, $\Rightarrow$ -intro.

Exemple 7. Si un jugement $\Gamma \vdash A$ est un théorème de Lpréd et si x est une variable, le jugement $\Gamma \vdash \exists xA$ est un théorème de Lpréd qui se déduit du premier suivant S66. En particulier, tout jugement de la forme $\Gamma \vdash \exists xA$ dans lequel A est une tautologie est un théorème de Lpréd. Par exemple, si A est une proposition quelconque, le jugement $\Gamma \vdash \exists x(A \text{ ou } \text{non}A)$ est un théorème de Lpréd qui se démontre comme suit:

Γ		
1°	$A \text{ ou } \text{non}A$	tiers exclu
2°	$\exists x(A \text{ ou } \text{non}A)$	1°, S66.

On peut tirer de cet exemple une conclusion générale d'ordre sémantique sur l'univers de l'interprétation d'une théorie logique, c'est-à-dire sur la classe

des objets qui sont désignés par les termes de la théorie. La conclusion est que cet univers n'est jamais vide: il doit contenir au moins un objet. Nous parlons ici d'une théorie dont le langage est le langage universel $\mathcal{L}_\Omega$ (§9.2). Dans ce langage, il y a une infinité de symboles relationnels de chaque arité. Mais il suffit d'un seul, n'importe lequel, pour que nous puissions tirer cette conclusion. Choisissons arbitrairement le symbole relationnel d'égalité, et prenons pour A la proposition $p = p$, et pour Γ la liste d'hypothèses vide. Nous obtenons le théorème de **Lpréd**:

$$\vdash \exists p(p = p \text{ ou } p \neq p).$$

Ce jugement est donc un théorème de toute théorie logique. Or on ne voit pas comment l'interpréter dans un univers dans lequel il n'y aurait aucun objet. S'il existe, dans l'univers considéré, un objet qui est égal à lui-même ou qui est différent de lui-même, alors il existe un objet dans cet univers. Au lieu de $p = p$, nous pourrions prendre n'importe quelle proposition dans laquelle p figure librement, exprimant une propriété de l'objet désigné par p . Il existe un objet ayant cette propriété ou la propriété contraire — donc il existe un objet. La logique classique suppose que l'univers de l'interprétation n'est pas vide.

Combinaison de S65 et S66. Une combinaison évidente de déductions suivant S65 et S66 fournit la démonstration du schéma de déduction

S68

$$\frac{\Gamma \vdash \forall x A}{\Gamma \vdash \exists x A}.$$

Son interprétation est que si tout objet de l'univers du discours possède une certaine propriété, énoncée dans une proposition A , alors il existe un objet de cet univers qui possède cette propriété. Celui qui a l'habitude du raisonnement sur les ensembles voit ici que la logique classique fait l'hypothèse d'un univers d'interprétation non vide. Mais pour celui qui n'est pas encore familiarisé avec ce raisonnement, la chose se voit plus facilement dans l'exemple précédent.

9.7 Élimination de quantificateurs existentiels

Nous avons discuté assez longuement au paragraphe 9.1, le schéma de déduction élémentaire

S64. *Élimination de $\exists$*

$$\frac{\Gamma \vdash \exists x A \quad \Gamma; A \vdash B}{\Gamma \vdash B} \quad \begin{array}{l} \text{Si } x \text{ ne figure librement} \\ \text{ni dans } B \text{ ni dans } \Gamma. \end{array}$$

Nous poursuivons ici cette discussion avec d'autres exemples. Auparavant, nous donnons ci-dessous le plan général de l'utilisation du schéma de déduction S64 dans une démonstration. Lorsqu'une proposition de la forme $\exists x A$ est vraie sous des hypothèses Γ dans lesquelles x ne figure pas librement, si l'on doit

démontrer sous ces hypothèses une proposition B dans laquelle x ne figure pas librement, il suffit de démontrer B sous l'hypothèse supplémentaire A , car alors B est vraie sous les hypothèses Γ d'après S64.

	Γ	
1°	$\exists xA$	
2°	A hyp	x ne figure librement
$\vdots$	$\vdots$	ni dans B , ni dans Γ .
n°	B	
	$1^\circ, n^\circ, \text{S64}$	

Exemple 1. Une démonstration formelle du théorème

$\vdash$ si x est impair, alors x^2 est impair,

de la théorie des entiers, est donnée dans la figure 9.8. Les calculs des lignes 4° et 5° sont admis sans justification. Pour faciliter la discussion, les listes d'hypothèses correspondant aux divers cadres sont désignées par

Γ_0 : (liste vide)

Γ_1 : x est impair

Γ_2 : x est impair; $x = 2a + 1$.

Le raisonnement informel ordinaire correspondant à cette démonstration est le suivant: Supposons que x soit impair (1°). Par définition de « impair », il existe un entier a tel que $x = 2a + 1$ (2°). Soit a un tel entier — autrement dit posons $x = 2a + 1$ (3°). On a $x^2 = 2(2a^2 + 2a) + 1$ ($4^\circ, 5^\circ$). On voit d'après cela qu'il existe un entier b tel que $x^2 = 2b + 1$ (6°) (§9.6, Exemple 5). Donc x^2 est impair (7°).

Nous nous arrêtons ici, comme cela se fait dans une démonstration ordinaire, mais en fait nous n'en sommes qu'à la ligne 7° de la démonstration formelle complète. Ce qui se passe à ce stade, dans l'esprit de celui qui raisonne informellement, est qu'il « oublie » l'hypothèse $x = 2a + 1$ dont il s'est servi, et qu'il considère la conclusion « x^2 est impair » à laquelle il est parvenu comme une conséquence de la seule hypothèse $\exists a(x = 2a + 1)$. Cela revient bien à recopier « x^2 est impair » dans le cadre Γ_1 (8°).

Une autre façon de procéder, dans une texte informel, est d'omettre l'hypothèse $x = 2a + 1$ (3°), donc de ne pas ouvrir (mentalement) de cadre intérieur Γ_2 , et de faire comme si la proposition 4° , et par suite 5° et 6° , se *déduisaient* de 2° , ce qui est logiquement faux — mais le résultat final est juste: on parvient à « x^2 est impair » dans le cadre Γ_1 . L'application finale de $\Rightarrow$ -introduction (9°) est d'ordinaire sous-entendue.

Le schéma S64 est appliqué à la ligne 8° de la manière suivante:

$$\frac{(2^\circ) \Gamma_1 \vdash \exists a(x = 2a + 1) \quad (7^\circ) \Gamma_1; x = 2a + 1 \vdash x^2 \text{ est impair}}{(8^\circ) \Gamma_1 \vdash x^2 \text{ est impair.}}$$

	Γ_0 (nil)	
	Γ_1	
1°	x est impair	hyp
2°	$\exists a(\underbrace{x = 2a + 1}_A)$	1°, déf. de impair
	Γ_2	
3°	$\underbrace{x = 2a + 1}_A$	hyp
4°	$x^2 = 4a^2 + 4a + 1$	3°, calcul
5°	$x^2 = 2(2a^2 + 2a) + 1$	4°, calcul
6°	$\exists b(x^2 = 2b + 1)$	5°, S63
7°	$\underbrace{x^2 \text{ est impair}}_B$	6°, déf. de impair
8°	$\underbrace{x^2 \text{ est impair}}_B$	2°, 7°, S64
9°	x est impair $\Rightarrow x^2$ est impair	8°, $\Rightarrow$ -intro.

Figure 9.8 Démonstration du théorème de l'exemple 1.

Les conditions du schéma sont respectées: la variable a (variable du quantificateur existentiel éliminé) ne figure librement ni dans Γ_1 , ni dans la proposition « x^2 est impair » (proposition B).

Rappelons encore l'idée logique fondamentale de ce schéma. On est parvenu à démontrer une proposition B (x^2 est impair, ligne 7°) en considérant un objet auxiliaire a dont il n'est pas question dans B, et sur lequel on a fait l'unique hypothèse A (3°). On sait qu'il existe un tel objet (2°), et en faisant l'hypothèse A, on en « prend » un. Puisque c'est la seule hypothèse faite sur a , B ne découle en fait que de cette possibilité de « prendre » un objet a vérifiant A, donc de son existence seulement (2°). B est donc aussi vraie dans le cadre Γ_1 .

Exemple 2. Le théorème $\vdash \exists x(x \in a) \Rightarrow \exists y(y \in a)$ de **Lpréd** est démontré ci-dessous. Le schéma S64 est appliqué à la ligne 4° de cette démonstration de la manière suivante:

$$\frac{(1^\circ) \Gamma_1 \vdash \exists x(x \in a) \quad (3^\circ) \Gamma_1; x \in a \vdash \exists y(y \in a)}{(4^\circ) \Gamma_1 \vdash \exists y(y \in a)}.$$

Les conditions du schéma sont respectées: la variable x (variable du quantificateur existentiel éliminé) ne figure librement ni dans la liste d'hypothèses Γ_1 , comprenant la seule proposition $\exists x(x \in a)$, ni dans la proposition $\exists y(y \in a)$.

	Γ_0 (nil)	
	Γ_1	
1°	$\exists x(x \in a)$	hyp
2°	$x \in a$	hyp
3°	$\exists y(y \in a)$	2°, S63 car 2° est $(x y)(y \in a)$
4°	$\exists y(y \in a)$	1°, 3°, S64
5°	$\exists x(x \in a) \Rightarrow \exists y(y \in a)$	4°, $\Rightarrow$ -intro

On démontre le théorème réciproque, $\vdash \exists y(y \in a) \Rightarrow \exists x(x \in a)$, en permutant x et y dans la démonstration. D'où le théorème $\vdash \exists x(x \in a) \Leftrightarrow \exists y(y \in a)$. Les propositions $\exists x(x \in a)$ et $\exists y(y \in a)$ sont identiques à un changement de variable liée près (§9.4, Exemple 2). L'équivalence de telles propositions est un règle générale qui sera présentée ultérieurement.

Exemple 3. On démontre ci-dessous le théorème

$$\vdash \forall x(x \notin a) \Rightarrow \text{non}\exists x(x \in a)$$

de **Lpréd**, réciproque du théorème démontré au paragraphe 9.6 (Exemple 6). Le schéma S64 est appliqué à la ligne 6° de cette démonstration de la manière suivante:

$$\frac{(2^\circ) \Gamma_2 \vdash \exists x(x \in a) \quad (5^\circ) \Gamma_2; x \in a \vdash \perp}{(6^\circ) \Gamma_2 \vdash \perp.}$$

Les conditions du schéma sont respectées: la variable x , variable du quantificateur existentiel éliminé, ne figure librement ni dans la liste d'hypothèses Γ_2 , à savoir $\forall x(x \notin a)$; $\exists x(x \in a)$, ni dans la proposition $\perp$.

De ce théorème et de sa réciproque, on conclut à l'équivalence des propositions $\forall x(x \notin a)$, $\text{non}\exists x(x \in a)$. Nous verrons plus loin que l'équivalence de propositions la forme $\forall x \text{non}A$, $\text{non}\exists xA$ est une règle générale.

	Γ_0 (nil)	
	Γ_1	
1°	$\forall x(x \notin a)$	hyp
	Γ_2	
2°	$\exists x(x \in a)$	hyp
3°	$x \in a$	hyp
4°	$x \notin a$	1°, S65
5°	$\perp$	3°, 4°, $\perp$ -intro
6°	$\perp$	2°, 5°, S64
7°	$\text{non}\exists x(x \in a)$	6°, red abs J
8°	$\forall x(x \notin a) \Rightarrow \text{non}\exists x(x \in a)$	7°, $\Rightarrow$ -intro.

Exemple 4. On démontre ci-dessous le théorème suivant de la théorie des nombres réels:

$$\vdash \text{ si } x \text{ est rationnel, alors } x + 1 \text{ est rationnel.} \quad (1)$$

On rappelle qu'un nombre réel x est dit rationnel s'il existe deux entiers a et b tels que $b \neq 0$ et $x = a/b$. En notant au moyen du symbole « Q » préfixé le prédicat «est rationnel», et au moyen du symbole « Z » le prédicat «est entier», cette définition se formalise par

$$Qx \Leftrightarrow \exists a \exists b (Za \text{ et } Zb \text{ et } b \neq 0 \text{ et } x = a/b). \quad (2)$$

	Γ_0 (nil)	
	Γ_1	
1°	Qx	hyp
2°	$\exists a \exists b (Za \text{ et } Zb \text{ et } b \neq 0 \text{ et } x = a/b)$	1°, définition
	Γ_2	
3°	$\exists b (Za \text{ et } Zb \text{ et } b \neq 0 \text{ et } x = a/b)$	hyp
	Γ_3	
4°	$Za \text{ et } Zb \text{ et } b \neq 0 \text{ et } x = a/b$	hyp
5°	$Z(a + b) \text{ et } Zb \text{ et } b \neq 0 \text{ et } x + 1 = (a + b)/b$	4°
6°	$\exists q (Z(a + b) \text{ et } Zq \text{ et } q \neq 0 \text{ et } x + 1 = (a + b)/q)$	5°, S63
7°	$\exists p \exists q (Zp \text{ et } Zq \text{ et } q \neq 0 \text{ et } x + 1 = p/q)$	6°, S63
8°	$Q(x + 1)$	7°, définition
9°	$Q(x + 1)$	3°, 8°, S64
10°	$Q(x + 1)$	2°, 9°, S64
11°	$Qx \Rightarrow Q(x + 1)$	10°, $\Rightarrow$ -intro.

Les deux quantificateurs existentiels de la ligne 2° sont éliminés aux lignes 3° et 4°. Il y a deux applications correspondantes de S64 (9°, 10°), imbriquées l'une dans l'autre. Lorsqu'on est parvenu à la ligne 5°, dont nous *admettons* qu'elle se déduit de 4°, on sait que $x + 1$ est rationnel, puisque $x + 1$ est le quotient des deux entiers $a + b$, b . Mais pour utiliser *formellement* la définition de « $x + 1$ est rationnel», il faut obtenir pour $x + 1$ une proposition de la forme du membre de droite de (2). C'est ce qui est fait aux lignes 6° et 7°, où l'on introduit deux quantificateurs existentiels par deux applications de S63. Le lecteur peut vérifier que ces deux déductions sont correctes: la proposition 5° peut se construire syntaxiquement à partir de la proposition 6° par élimination du quantificateur $\exists q$ et substitution ($b|q$); la proposition 6° peut se construire à partir de 7° par suppression de $\exists p$ et substitution ($a + b|p$). Dans les deux cas, la substitution ne crée pas d'occurrence de variable liée — grâce au fait que nous avons utilisé les nouvelles variables p et q . Les deux déductions suivant S64 de cette démonstration sont reproduites dans la figure 9.10. Le lecteur vérifiera que les conditions du schéma sont respectées dans les deux cas: la variable b ne figure pas librement dans les hypothèses de la liste Γ_2 , ni dans la

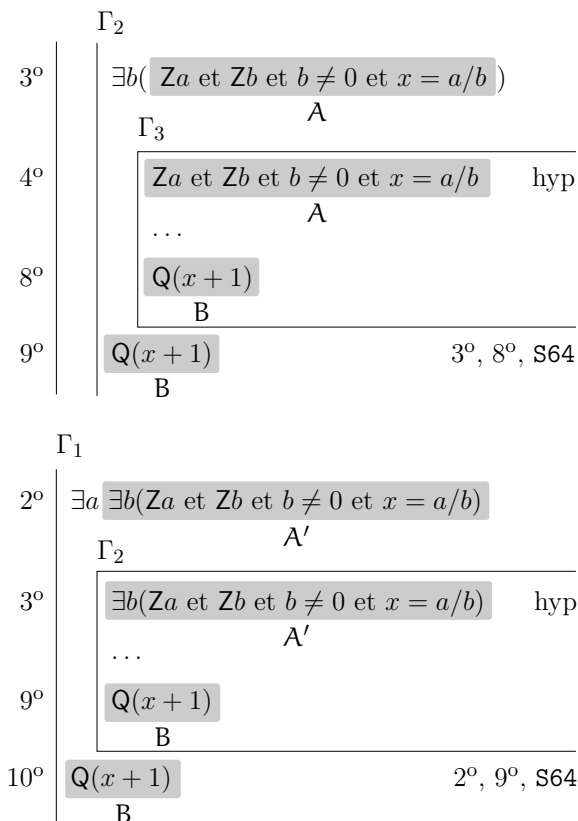


Figure 9.10 Applications de S64 dans la démonstration du théorème (1).

proposition $Q(x+1)$ désignée par B; la variable a ne figure librement ni dans l'hypothèse de Γ_1 , ni dans B.

Une démonstration ordinaire serait considérablement abrégée. De 2° , on sauterait directement à 5° , sans déclarer les hypothèses 3° et 4° , mais sans affirmer non plus que 5° se déduit de 2° — car on ne voudrait pas mentir tout de même! Puis on sauterait directement de 5° à la conclusion $Q(x+1)$, en se sentant mentalement dans le cadre Γ_1 donc en 10° , et l'on s'arrêterait à cet endroit.

Exemple 5. Une démonstration du schéma de déduction

$$\frac{\Gamma \vdash \exists x(A \text{ et } B)}{\Gamma \vdash \exists x(A \text{ et } \exists x B)} \quad (3)$$

est donnée ci-dessous, dans le cas particulier d'une liste d'hypothèses Γ vide. Le cas général (liste Γ quelconque) dérive par *a fortiori* du cas particulier comme nous l'avons expliqué au paragraphe 9.5. Le schéma S64 est appliqué à la ligne

8° de la manière suivante:

$$\frac{(1^\circ) \Gamma_1 \vdash \exists x(A \text{ et } B) \quad (7^\circ) \Gamma_1; A \text{ et } B \vdash \exists xA \text{ et } \exists xB}{(8^\circ) \Gamma_1 \vdash \exists xA \text{ et } \exists xB.}$$

Les conditions du schéma sont respectées: la variable x ne figure librement ni dans l'unique hypothèse de Γ_1 , ni dans la proposition $\exists xA$ et $\exists xB$.

	Γ_0 (nil)	
	Γ_1	
1°	$\exists x(A \text{ et } B)$	hyp
2°	$A \text{ et } B$	hyp
3°	A	2°, et-élim
4°	B	2°, et-élim
5°	$\exists xA$	3°, S66
6°	$\exists xB$	4°, S66
7°	$\exists xA \text{ et } \exists xB$	5°, 6°, et-intro
8°	$\exists xA \text{ et } \exists xB$	1°, 7°, S64
9°	$\exists x(A \text{ et } B) \Rightarrow (\exists xA \text{ et } \exists xB)$	8°, $\Rightarrow$ -intro.

Comme exemple de théorème de **Lpréd** de la forme décrite par (3), nous pouvons mentionner le jugement:

$$\vdash \exists x(x \in a \text{ et } x \in b) \Rightarrow (\exists x(x \in a) \text{ et } \exists x(x \in b)). \quad (4)$$

Ce théorème affirme que si deux ensembles a et b ont un élément en commun, c'est-à-dire s'il existe un objet qui appartient à a et à b , alors chacun de ces ensembles possède au moins un élément. Il est clair qu'on ne peut pas affirmer la réciproque. Deux ensembles a , b peuvent avoir chacun des éléments sans avoir pour autant un élément commun.

Il est instructif d'essayer de démontrer le théorème réciproque de (4), afin de voir pourquoi l'on n'y parvient pas. Cette tentative infructueuse est la démonstration inachevée suivante:

	Γ_0 (nil)	
	Γ_1	
1°	$\exists x(x \in a) \text{ et } \exists x(x \in b)$	hyp
2°	$\exists x(x \in a)$	1°, et-élim
	Γ_2	
3°	$x \in a$	hyp
4°	$\exists x(x \in b)$	1°, et-élim
5°	$x \in b$	hyp
6°	$x \in a \text{ et } x \in b$	3°, 5°, et-intro
7°	$\exists x(x \in a \text{ et } x \in b)$	7°, S66
8°	?	

On ne peut pas faire, à la ligne 8°, la déduction

$$\frac{(4^\circ) \Gamma_2 \vdash \exists x(x \in b) \quad (7^\circ) \Gamma_2; x \in b \vdash \exists x(x \in a \text{ et } x \in b)}{(8^\circ) \Gamma_2 \vdash \exists x(x \in a \text{ et } x \in b)}$$

car, l'une des conditions du schéma **S64** n'est pas vérifiée: la variable x du quantificateur éliminé à la ligne 5° *figure librement* dans une hypothèse de la liste Γ_2 , l'hypothèse $x \in a$ (3°). Il est clair qu'ayant fait les hypothèses 3° et 5° sur le *même* objet x , on parvient *sous ces hypothèses*, c'est-à-dire dans le cadre intérieur, à la conclusion que a et b ont un élément commun, mais en dehors de ce cadre, on ne peut pas tirer cette conclusion.

Exercice 1. Soient A, B, C des propositions, x une variable et Γ une liste de propositions. (a) Montrer que le jugement

$$\Gamma \vdash (\forall x A \text{ et } \exists x B) \Rightarrow \exists x(A \text{ et } B)$$

est un théorème de **Lpréd**. (b) Montrer que si x ne figure pas librement dans Γ , la séquence de jugements

$$\frac{\Gamma \vdash A \Rightarrow (B \Rightarrow C)}{\Gamma \vdash \exists x A \Rightarrow (\forall x B \Rightarrow \exists x C)}$$

est une déduction de **Lpréd**.

Exercice 2. Démontrer dans **Lpréd** la déduction

$$\frac{\vdash a \neq b \Rightarrow \forall x(x \neq a \text{ ou } x \neq b)}{\vdash a \neq b \Rightarrow \forall x \exists y(x \neq y)}.$$

Exercice 3. Soient A, B des propositions, x_1 et x_2 des variables et Γ une liste de propositions. Montrer que si x_1 et x_2 ne figurent librement ni dans Γ , ni dans B , on peut faire dans **Lpréd** la déduction suivante (élimination simultanée de deux quantificateurs existentiels):

$$\frac{\Gamma \vdash \exists x_1 \exists x_2 A \quad \Gamma; A \vdash B}{\Gamma \vdash B}.$$

Généraliser ce schéma de déduction pour n variables $x_1, \dots, x_n$ (élimination simultanée de n quantificateurs existentiels).

10

Déductions dérivées de logique des prédicats

10.1 Substitutivité de l'équivalence

Le but de ce paragraphe est de montrer que la règle de substitutivité de l'équivalence de la logique propositionnelle (§7.5) s'étend à la logique des prédicats. Cette propriété de l'équivalence repose sur les schémas de déduction dérivés suivants:

S69.

$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash \forall x A \Rightarrow \forall x B} \quad \frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash \exists x A \Rightarrow \exists x B} \quad \text{Si } x \text{ ne figure pas librement dans } \Gamma.$$

Dans les démonstrations ci-dessous, on suppose que Γ est une liste d'hypothèses dans lesquelles la variable x ne figure pas librement. Cette condition est donc aussi satisfaite pour les listes $\Gamma; \forall x A$ et $\Gamma; \exists x A$, ce qui permet d'appliquer S62 à la ligne 5°, respectivement S64 à la ligne 6°.

Γ			Γ		
1°	$A \Rightarrow B$	<i>base</i>	1°	$A \Rightarrow B$	<i>base</i>
2°	$\forall x A$	<i>hyp</i>	2°	$\exists x A$	<i>hyp</i>
3°	A	2°, S65	3°	A	<i>hyp</i>
4°	B	1°, 3°	4°	B	1°, 3°
5°	$\forall x B$	4°, S62	5°	$\exists x B$	4°, S66
6°	$\forall x A \Rightarrow \forall x B$	5°, $\Rightarrow$ -intro.	6°	$\exists x B$	2°, 5°, S64
			7°	$\exists x A \Rightarrow \exists x B$	6°, $\Rightarrow$ -intro.

Justification intuitive de ces schémas. Imaginons un univers d'individus, dans lequel on distingue diverses classes ou espèces d'individus: la classe A, la classe B, etc. , et écrivons « x est-A » pour dire que l'individu x appartient à la classe A. Les symboles *est-A*, *est-B* sont donc des symboles relationnels unaires (prédicats). Nous supposons ces classes définies par des propriétés caractéristiques de leurs individus. Certaines sont des sous-classes d'autres.

Il se peut qu'une certaine espèce A n'ait aucun individu, ou à l'opposé que A englobe tous les individus de l'univers. Si pour un individu x *quelconque*, c'est-à-dire sur lequel on ne fait aucune hypothèse, il est vrai que

$$x \text{ est-}A \Rightarrow x \text{ est-}B, \quad (1)$$

cela signifie qu'un individu quelconque ayant les propriétés caractéristiques de l'espèce A possède aussi celles de l'espèce B . Il est naturel d'en déduire que si l'espèce A englobe tous les individus, alors il en est de même de l'espèce B . Il est naturel aussi d'en déduire que s'il *existe* un individu de l'espèce A , alors il en existe un de l'espèce B . Ce sont précisément les propositions

$$\forall x(x \text{ est-}A) \Rightarrow \forall x(x \text{ est-}B), \quad \exists x(x \text{ est-}A) \Rightarrow \exists x(x \text{ est-}B),$$

qui se déduisent de (1) suivant S69.

Substitutivité de l'équivalence. Les déductions de la forme suivante se démontrent de manière évidente par élimination de $\Leftrightarrow$ (S14), S69 et introduction de $\Leftrightarrow$ (S13). La condition « x ne figure pas librement dans Γ » est valable pour les deux déductions.

S70.

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash \forall x A \Leftrightarrow \forall x B} \quad \frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash \exists x A \Leftrightarrow \exists x B} \quad \text{Si } x \text{ ne figure pas librement dans } \Gamma.$$

En tant que propriétés de l'équivalence, ces deux schémas sont à placer à côté des schémas S50.1–9 de la logique propositionnelle (§7.5). Nous rappelons que deux propositions A , B sont dites *équivalentes* dans un contexte particulier lorsque la proposition $A \Leftrightarrow B$ est vraie dans ce contexte (§7.5). Les schémas S50.1–9 et S70 peuvent s'énoncer tous ensemble de la manière suivante, qui complète l'énoncé du paragraphe 7.5:

S50.1–9 et S70. Si les propositions A , B sont équivalentes dans un contexte, alors les propositions $\text{non}A$, $\text{non}B$ sont équivalentes dans ce contexte, et étant donné une troisième proposition C , les propositions A ou C , C ou A , A et C , C et A , $A \Rightarrow C$, $C \Rightarrow A$, $A \Leftrightarrow C$, $C \Leftrightarrow A$ sont équivalentes respectivement à B ou C , C ou B , B et C , C et B , $B \Rightarrow C$, $C \Rightarrow B$, $B \Leftrightarrow C$, $C \Leftrightarrow B$ dans ce contexte. De plus, si x est une variable qui ne figure pas librement dans les hypothèses de ce contexte, les propositions $\forall x A$, $\exists x A$ sont équivalentes respectivement à $\forall x B$, $\exists x B$ dans celui-ci.

Nous avons vu au paragraphe 7.5 le schéma général de substitutivité de l'équivalence qui en découle, et qui dit en gros ceci : si l'on modifie une proposition R en remplaçant une *sous-proposition* C de R , c'est-à-dire une partie de R qui est elle-même une proposition, par une proposition C' qui est équivalente à C , on obtient une proposition R' qui est équivalente à R . Donc de $\Gamma \vdash C \Leftrightarrow C'$, on déduit $\Gamma \vdash R \Leftrightarrow R'$. Mais en raison de la condition du schéma S70, cette forme générale de déduction est soumise aussi à une condition, et le schéma général est formulé plus bas, après l'exemple qui suit.

Exemple. La figure 10.1 contient les arbres de deux propositions R , R' . La proposition R n'est autre que l'axiome de l'inclusion de $\mathcal{T}_{\text{ENS}}$ (§8.2). Une sous-proposition C de R est mise en évidence. De façon générale, les sous-propositions d'une proposition R sont faciles à identifier dans l'arbre de R : ce sont les parties de cet arbre formées par un nœud N portant un symbole relationnel, un connecteur logique, ou un quantificateur, et par tous les descendants de ce nœud N . La sous-proposition C considérée ici est donc la partie $a \in x \Rightarrow a \in y$ de R .

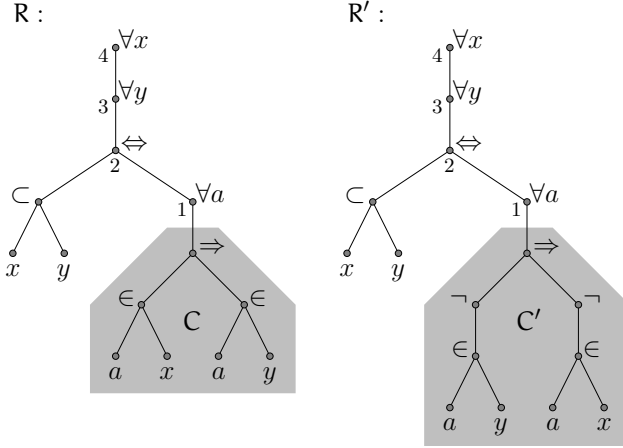


Figure 10.1 Illustration du schéma général de substitutivité de l'équivalence. De $\Gamma \vdash C \Leftrightarrow C'$ on déduit $\Gamma \vdash R \Leftrightarrow R'$ si les variables a, x, y ne figurent pas librement dans Γ .

Nous avons numéroté de 1 à 4 les nœuds qui *dominent* la sous-proposition C dans R . L'arbre de droite est celui de la proposition R' obtenue en remplaçant la sous-proposition C de R par la proposition $C' : a \notin y \Rightarrow a \notin x$. Les propositions C, C' sont équivalentes dans n'importe quel contexte (S57).

Désignons par $R_1, \dots, R_4$ (resp. $R'_1, \dots, R'_4$) les sous-propositions de R (resp. de R') correspondant aux nœuds 1, 2, 3, 4. Ce sont les propositions suivantes :

$$\begin{array}{ll} R_1 : \forall a C & R'_1 : \forall a C' \\ R_2 : (x \subset y) \Leftrightarrow R_1 & R'_2 : (x \subset y) \Leftrightarrow R'_1 \\ R_3 : \forall y R_2 & R'_3 : \forall y R'_2 \\ R_4 : \forall x R_3 & R'_4 : \forall x R'_3. \end{array}$$

De $\Gamma \vdash C \Leftrightarrow C'$, on déduit $\Gamma \vdash R_1 \Leftrightarrow R'_1$ suivant S70, à condition que la variable a ne figure librement dans aucune hypothèse de la liste Γ (condition de S70). Si tel est le cas, le schéma S50.9 permet d'en déduire $\Gamma \vdash R_2 \Leftrightarrow R'_2$. On en tire ensuite successivement $\Gamma \vdash R_3 \Leftrightarrow R'_3$, puis $\Gamma \vdash R_4 \Leftrightarrow R'_4$, suivant S70, à condition encore que les variables y et x ne figurent pas librement dans Γ .

Schéma général de substitutivité de l'équivalence. Une déduction suivant le schéma général de substitutivité de l'équivalence est une déduction de la forme

$$\frac{\Gamma \vdash C \Leftrightarrow C'}{\Gamma \vdash R \Leftrightarrow R'}$$

où C est une sous-proposition de R qui n'est pas dominée par un quantificateur dont la variable figure librement dans Γ , et R' est la proposition obtenue en remplaçant C par C' dans R . On peut énoncer ce schéma de la manière suivante (cf. §7.5):

Si une proposition R' est construite en remplaçant dans une proposition R une sous-proposition C par une proposition C' , alors de $\Gamma \vdash C \Leftrightarrow C'$ on peut déduire $\Gamma \vdash R \Leftrightarrow R'$, à condition qu'il n'y ait pas de quantificateur $\exists x$ ou $\forall x$ dominant C dans R dont la variable x figure librement dans Γ .

La combinaison d'une déduction par substitutivité de l'équivalence et d'une déduction par *a fortiori* (§9.5) permet souvent de satisfaire trivialement la condition d'une déduction par substitutivité de l'équivalence en effectuant celle-ci dans un contexte sans hypothèses:

$$\frac{\frac{\vdash C \Leftrightarrow C'}{\vdash R \Leftrightarrow R'} \text{ substitutivité de } \Leftrightarrow}{\Gamma \vdash R \Leftrightarrow R'} \text{ a fortiori}$$

Par exemple, pour les propositions R, R', C, C' de la figure 10.1, on démontre comme suit le théorème $\vdash R \Leftrightarrow R'$ dans $L_{\text{préd}}$ — la théorie des ensembles n'intervient pas:

$$\begin{array}{l|ll} \text{nil} & & \\ 1^\circ & C \Leftrightarrow C' & \text{S31} \\ 2^\circ & R \Leftrightarrow R' & 1^\circ, \text{ substitutivité de l'équivalence.} \end{array}$$

On peut en déduire $\Gamma \vdash R \Leftrightarrow R'$ par *a fortiori*, pour une liste d'hypothèses Γ quelconque, même si certaines des variables a, x, y des quantificateurs qui dominent C dans R figurent librement dans Γ .

Dans la suite, lorsque nous écrirons des démonstrations en forme de chaînes d'équivalences, nous le ferons souvent sous une liste d'hypothèses vide. Le lecteur se souviendra que la raison du choix de ce contexte est que l'on utilise constamment dans ces chaînes le schéma de substitutivité de l'équivalence et que sa condition est vérifiée trivialement dans un tel contexte.

10.2 Dualité en logique des prédicats

Négations et quantificateurs. Les schémas de déduction suivants peuvent être considérés comme étant les « schémas de De Morgan » pour les quantificateurs. Nous expliquerons plus bas la raison de cette dénomination.

S71.

- | | |
|--|--|
| 1. $\frac{}{\Gamma \vdash \text{non}\exists xA \Leftrightarrow \forall x\text{non}A.}$ | 2. $\frac{}{\Gamma \vdash \text{non}\forall xA \Leftrightarrow \exists x\text{non}A.}$ |
| 3. $\frac{}{\Gamma \vdash \exists xA \Leftrightarrow \text{non}\forall x\text{non}A.}$ | 4. $\frac{}{\Gamma \vdash \forall xA \Leftrightarrow \text{non}\exists x\text{non}A.}$ |

Les schémas 3 et 4 dérivent immédiatement des deux premiers suivant la logique propositionnelle (S51.2, S47). Nous démontrons ci-dessous le schéma 1 dans le cas particulier d'une liste d'hypothèses Γ vide (§9.5). Cela permet d'appliquer S62 à la ligne 6°, car la variable x ne figure pas librement dans l'hypothèse $\text{non}\exists xA$, et d'appliquer S64 à la ligne 12°, car x ne figure pas librement dans les hypothèses $\forall x\text{non}A$, $\exists xA$.

nil		
1°	non $\exists xA$	hyp
2°	A	hyp
3°	$\exists xA$	2°, S66
4°	$\perp$	1°, 3°, $\perp$ -intro
5°	nonA	4°, red abs J
6°	$\forall x\text{non}A$	5°, S62
7°	$\forall x\text{non}A$	hyp
8°	$\exists xA$	hyp
9°	A	hyp
10°	nonA	7°, S65
11°	$\perp$	9°, 10°
12°	$\perp$	8°, 11°, S64
13°	non $\exists xA$	12°, red abs J
14°	non $\exists xA \Leftrightarrow \forall x\text{non}A$	6°, 13°, $\Leftrightarrow$ -intro.

Le schéma S71.2 dérive de S71.1 par substitutivité de l'équivalence. On peut écrire en effet, sous une liste d'hypothèse vide, la démonstration en forme de chaîne d'équivalences suivante :

nil		
	$\exists x\text{non}A \Leftrightarrow \text{nonnon}\exists x\text{non}A$	S51
	$\Leftrightarrow \text{non}\forall x\text{nonnon}A$	S71.1
	$\Leftrightarrow \text{non}\forall xA$	S51.

Imaginons un univers dans lequel il n'y ait qu'un nombre fini n d'objets distincts, et choisissons n constantes individuelles $a_1, \dots, a_n$ du langage $\mathcal{L}_\Omega$ pour désigner ces individus. Supposons que l'on s'intéresse à une propriété que peuvent avoir les objets de cet univers, et que l'énoncé de cette propriété, pour un individu x s'exprime par une proposition P dans laquelle x figure librement. Pour simplifier la notation, désignons par $P(a_i)$ la proposition $(a_i|x)P$,

signifiant que l'individu $\mathbf{a}_i$ possède la propriété en question. Désignons encore la proposition P par $P(x)$. Vu le nombre fini d'individus de cet univers, nous pouvons traduire les propositions quantifiées $\forall xP(x)$ et $\exists xP(x)$ par des propositions sans quantificateurs, en admettant les déductions sans prémisses

$$\frac{\Gamma \vdash \forall xP(x) \Leftrightarrow (P(\mathbf{a}_1) \text{ et } \dots \text{ et } P(\mathbf{a}_n))}{\Gamma \vdash \exists xP(x) \Leftrightarrow (P(\mathbf{a}_1) \text{ ou } \dots \text{ ou } P(\mathbf{a}_n))}. \quad (1)$$

Il ne saurait être question de démontrer ces déductions. Elles reposent sur notre interprétation de n symboles distincts donnés comme désignant tous les individus d'un certain univers, et aussi sur notre interprétation des connecteurs logiques et/ou et des quantificateurs. Nous pouvons dire que ce sont les *dédutions élémentaires spécifiques* de la « théorie d'un univers à n individus », dont les individus sont désignés par les constantes $\mathbf{a}_i$. Dans cette théorie, on peut démontrer les théorèmes de la forme

$$\begin{aligned} \Gamma \vdash \text{non}\forall xP(x) &\Leftrightarrow \exists x\text{non}P(x) \\ \Gamma \vdash \text{non}\exists xP(x) &\Leftrightarrow \forall x\text{non}P(x) \end{aligned}$$

sans utiliser les schémas S71, mais en utilisant ceux de De Morgan (S52.1–2):

$$\left. \begin{array}{l} \Gamma \\ \vdash \end{array} \right\{ \begin{array}{ll} \text{non}\forall xP(x) \Leftrightarrow \text{non}(P(\mathbf{a}_1) \text{ et } \dots \text{ et } P(\mathbf{a}_n)) & \text{schémas (1)} \\ \Leftrightarrow (\text{non}P(\mathbf{a}_1) \text{ ou } \dots \text{ ou } \text{non}P(\mathbf{a}_n)) & \text{De Morgan} \\ \Leftrightarrow \exists x\text{non}P(x) & \text{schémas (1)}. \end{array} \right.$$

On voit pourquoi les schémas S71 peuvent être considérés comme des « schémas de De Morgan » relatifs aux quantificateurs. Une proposition de la forme $\forall xA$, peut-être interprétée dans un univers quelconque comme une sorte de conjonction généralisée. De même, la proposition $\exists xA$ peut être interprétée comme une disjonction généralisée. Certains auteurs utilisent d'ailleurs les notations

$$\bigwedge_x P(x), \quad \bigvee_x P(x)$$

au lieu de $\forall xP(x)$ et $\exists xP(x)$.

Exemple 1. En faisant une déduction suivant S71.2, nous démontrons le théorème suivant de $\mathcal{T}_{\text{ENS}}$, que nous pouvons appeler théorème de négation de l'inclusion:

$$\vdash x \not\subset y \Leftrightarrow \exists a(a \in x \text{ et } a \notin y).$$

Démonstration:

	nil	
1°	$\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y))$	axiome de l'inclusion
2°	$x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y)$	1°, S65
	$x \not\subset y \Leftrightarrow \text{non} \forall a (a \in x \Rightarrow a \in y)$	2°, S50.1
	$\Leftrightarrow \exists a \text{ non}(a \in x \Rightarrow a \in y)$	S71.2
	$\Leftrightarrow \exists a (a \in x \text{ et } a \notin y)$	S55.

Exemple 2. La combinaison des deux dernières déductions de la démonstration de l'exemple 1 (S71.2, S55) est assez fréquente. Elle permet de démontrer plus généralement le schéma de déduction sans prémisses suivant dont il est utile de se souvenir :

$$\Gamma \vdash \text{non} \forall x (A \Rightarrow B) \Leftrightarrow \exists x (A \text{ et non } B).$$

Exemple 3. Les propositions $\forall x (x \notin \emptyset)$ et $\text{non} \exists x (x \in \emptyset)$ sont équivalentes dans tout contexte d'après S71. Au lieu de choisir la première comme axiome de l'ensemble vide de $\mathcal{T}_{\text{ENS}}$, nous aurions pu choisir la seconde. Formellement, nous aurions défini par ce choix une *autre théorie*, disons $\mathcal{T}_{\text{ENS}}'$, puisqu'une théorie est *définie* par la donnée de ses axiomes et de ses déductions élémentaires (§5.3). Mais en vertu de S71, tout ce qui se déduit de l'un de ces deux axiomes peut aussi se déduire de l'autre, de sorte que les deux théories ont exactement les mêmes théorèmes. Le choix est sans importance.

Généralisation des schémas S71. Les schémas S71.1–2 peuvent être rassemblés en un seul sous la forme

$$\Gamma \vdash \text{non} Qx A \Leftrightarrow \overline{Q}x \text{non} A$$

en désignant par Qx l'un quelconque des quantificateurs $\forall x$, $\exists x$, et par $\overline{Q}x$ son *dual*: $\exists x$ est le dual de $\forall x$, $\forall x$ est le dual de $\exists x$. Ce schéma se généralise alors de la manière suivante, pour prendre la négation d'une proposition qui commence par n quantificateurs $Q_1 x_1 \cdots Q_n x_n$ quelconques :

$$\Gamma \vdash \text{non} Q_1 x_1 \cdots Q_n x_n A \Leftrightarrow \overline{Q}_1 x_1 \cdots \overline{Q}_n x_n \text{non} A.$$

En appliquant n fois l'un ou l'autre des schémas S71.1–2, et la substitutivité de l'équivalence, on peut écrire en effet la démonstration :

nil	
	$\text{non} Q_1 x_1 Q_2 x_2 \cdots A \Leftrightarrow \overline{Q}_1 x_1 \text{non} Q_2 x_2 \cdots A$
	$\Leftrightarrow \overline{Q}_1 x_1 \overline{Q}_2 x_2 \text{non} \cdots A$
	$\Leftrightarrow \dots$
	$\Leftrightarrow \overline{Q}_1 x_1 \cdots \overline{Q}_n x_n \text{non} A.$

Exemple 4. Dans le contexte de la théorie des nombres réels, on définit le prédicat « est surjective », pour une fonction $f : \mathbb{R} \rightarrow \mathbb{R}$, en posant comme définition:

$$f \text{ est surjective} \Leftrightarrow \forall y \exists x (y = f(x)). \quad (2)$$

Nous admettons ici pour simplifier que « $\forall y$ » signifie « pour tout *nombre réel* y » et que « $\exists x$ » signifie « il existe un *nombre réel* x ». Dans ce contexte, on peut écrire la démonstration suivante, utilisant le schéma de déduction S71 généralisé qui précède:

$$\begin{array}{l|l} \text{nil} & \\ \hline f \text{ n'est pas surjective} \Leftrightarrow \text{non} \forall y \exists x (y = f(x)) & (2), \text{S51.1} \\ & \Leftrightarrow \exists y \forall x \text{non}(y = f(x)) & \text{S71 généralisé} \\ & \Leftrightarrow \exists y \forall x (y \neq f(x)) & \text{abréviation.} \end{array}$$

Dualité. Les schémas S71.1–2 jouent un rôle analogue à celui des schémas de De Morgan (S52.1–2) en logique propositionnelle. Plus précisément, grâce à ces schémas, la règle de dualité (§7.7) peut être étendue à la logique des prédicats. On définit la *duale* d'une proposition $R[A, B, \dots]$ construite à partir de propositions $A, B, \dots$ au moyen exclusivement de connecteurs logiques de *conjonction* et de *disjonction* et de quantificateurs $\forall x, \exists x, \forall y, \exists y, \forall z, \exists z, \dots$. La proposition duale $R^D[A, B, \dots]$ de R (par rapport à $A, B, \dots$) s'obtient en remplaçant chaque disjonction par une conjonction et vice-versa, et chaque quantificateur par son dual, dans la construction de R . Par exemple:

$$\begin{array}{l} R[A, B, \dots] \text{ est la proposition } (A \text{ et } \exists x(B \text{ ou } C)), \\ R^D[A, B, \dots] \text{ est la proposition } (A \text{ ou } \forall x(B \text{ et } C)). \end{array} \quad (3)$$

Pour une proposition $R[A, B, \dots]$ et sa duale $R^D[A, B, \dots]$, on peut faire les déductions sans prémisses:

$$\begin{array}{l} 1. \frac{}{\Gamma \vdash \text{non} R[A, B, \dots] \Leftrightarrow R^D[\text{non} A, \text{non} B, \dots]}. \\ 2. \frac{}{\Gamma \vdash \text{non} R^D[A, B, \dots] \Leftrightarrow R[\text{non} A, \text{non} B, \dots]}. \end{array}$$

Elles se démontrent en appliquant les schémas de De Morgan et les schémas S71. La première par exemple, dans le cas des propositions (3), se démontre comme suit:

$$\begin{array}{l|l} \text{nil} & \\ \hline \text{non}(A \text{ et } \exists x(B \text{ ou } C)) \Leftrightarrow (\text{non} A \text{ ou } \text{non} \exists x(B \text{ ou } C)) & \text{De Morgan} \\ & \Leftrightarrow (\text{non} A \text{ ou } \forall x \text{non}(B \text{ ou } C)) & \text{S71} \\ & \Leftrightarrow (\text{non} A \text{ ou } \forall x(\text{non} B \text{ et } \text{non} C)) & \text{De Morgan.} \end{array}$$

Règle de dualité. La règle de dualité du paragraphe 7.7 s'étend ainsi à la logique des prédicats: chaque schéma de déduction de Lpréd de la forme

$$\frac{}{\Gamma \vdash R_1(A, B, \dots) \Leftrightarrow R_2(A, B, \dots)}$$

dans lequel $R_1(A, B, \dots)$ et $R_2(A, B, \dots)$ sont deux propositions construites à partir de propositions $A, B, \dots$ en n'utilisant que des connecteurs de disjonction et de conjonction et des quantificateurs s'accompagne du schéma *dual*:

$$\Gamma \vdash R_1^D(A, B, \dots) \Leftrightarrow R_2^D(A, B, \dots).$$

La preuve de cette règle est la même qu'au paragraphe 7.7.

10.3 Permutation et distribution de quantificateurs

Permutations de quantificateurs. Les schémas suivants sont les schémas de permutation de deux quantificateurs de *même espèce*, universels ou existentiels. C'est en raison de ces équivalences que l'on exprime souvent une proposition de la forme $\forall x \forall y A$ en disant « pour tout x et pour tout $y : A$ », ou encore « quels que soient x et $y : A$ ». Ces tournures laissent entendre que l'ordre dans lequel les variables x et y sont mentionnées n'est pas important. De même, $\exists x \exists y A$ s'exprime souvent par « il existe un x et un y tels que A ».

S72.

$$\Gamma \vdash \forall x \forall y A \Leftrightarrow \forall y \forall x A$$

$$\Gamma \vdash \exists x \exists y A \Leftrightarrow \exists y \exists x A.$$

Le premier de ces schémas a déjà été à moitié démontré au chapitre 9 (§9.4, Exemple 3), où nous avons démontré le théorème de **Lpréd**:

$$\vdash \forall x \forall y A \Rightarrow \forall y \forall x A,$$

pour une proposition A et des variables x, y quelconques. La démonstration du théorème réciproque, $\vdash \forall y \forall x A \Rightarrow \forall x \forall y A$, est identique, à la permutation près des variables x et y . D'où le théorème $\vdash \forall y \forall x A \Leftrightarrow \forall x \forall y A$ (par $\Leftrightarrow$ -intro). Le théorème $\Gamma \vdash \forall y \forall x A \Leftrightarrow \forall x \forall y A$, pour une liste d'hypothèses Γ quelconque, s'en déduit par a fortiori (§9.5).

La deuxième schéma S72 est le dual du premier. Sa démonstration utilise donc le premier, ainsi que le schéma S71 généralisé:

nil

$$\begin{array}{lcl} \exists x \exists y A \Leftrightarrow \text{nonnon} \exists x \exists y A & \text{S51} \\ \Leftrightarrow \text{non} \forall x \forall y \text{non} A & \text{S71} \\ \Leftrightarrow \text{non} \forall y \forall x \text{non} A & \text{S72 (premier)} \\ \Leftrightarrow \exists y \exists x \text{nonnon} A & \text{S71} \\ \Leftrightarrow \exists y \exists x A & \text{S51.} \end{array}$$

Deux quantificateurs d'espèces différentes ne peuvent pas être permutés de la même manière. Le schéma suivant ne permet d'écrire que des *implications*, non des équivalences.

S73.

$$\Gamma \vdash \exists x \forall y A \Rightarrow \forall y \exists x A.$$

Démonstration:

	nil	
1°	$\exists x \forall y A$	hyp
2°	$\forall y A$	hyp
3°	A	2°, S65
4°	$\exists x A$	3°, S66
5°	$\forall y \exists x A$	4°, S62
6°	$\forall y \exists x A$	1°, 5°, S64
7°	$\exists x \forall y A \Rightarrow \forall y \exists x A$	6°, $\Rightarrow$ -intro.

Exercice. Il faut prendre garde au sens de l'implication dans S73. C'est une erreur de raisonnement assez fréquente que de confondre cette implication avec la réciproque, $\forall y \exists x A \Rightarrow \exists x \forall y A$, qui est *fausse* en général. Pour apprendre à bien distinguer deux propositions de la forme $\exists x \forall y A$ et $\forall y \exists x A$, le lecteur est invité à considérer divers « graphes », ou diagrammes de points et de flèches tels que celui de la figure 10.2. Les points (ou sommets) du graphe sont pris comme individus de l'univers du discours, et l'on se sert du symbole $\rightarrow$ comme symbole relationnel binaire. La proposition $x \rightarrow y$ signifie qu'il y a une flèche allant du sommet x au sommet y . Lorsqu'il y a une telle flèche de x à y , on dit que le sommet y est un *successeur* du sommet x , ou encore que x est un *prédécesseur* de y . Un sommet peut être successeur (et prédécesseur) de lui-même: $x \rightarrow x$. On considère les quatre propositions suivantes, avec leur interprétation:

- (A₁) $\exists x \forall y (x \rightarrow y)$ Il existe un sommet qui est prédécesseur de chaque sommet;
- (A₂) $\forall y \exists x (x \rightarrow y)$ Tout sommet possède un prédécesseur;
- (A₃) $\exists y \forall x (x \rightarrow y)$ Il existe un sommet qui est successeur de chaque sommet;
- (A₄) $\forall x \exists y (x \rightarrow y)$ Tout sommet possède un successeur.

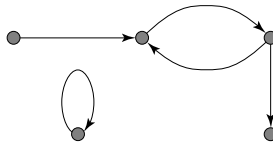


Figure 10.2

Le graphe de la figure 10.2 ne vérifie aucune de ces propositions. Pour chacun des graphes G_1 à G_4 de la figure 10.3, et pour chacune des propositions A_i , le lecteur dira laquelle des propositions A_i , non A_i le graphe vérifie. Il se convaincra aisément qu'un graphe qui vérifie A_1 (resp. A_3) vérifie *ipso facto* A_2 (resp. A_4), conformément à S73, et que conformément à ce schéma et à celui de contraposition (S31), un graphe qui vérifie non A_2 (resp. non A_4), vérifie *ipso*

facto $\text{non}A_1$ (resp. $\text{non}A_3$). Il observera par contre qu'un graphe peut vérifier A_2 (resp. A_4) et vérifier $\text{non}A_1$ (resp. $\text{non}A_2$).

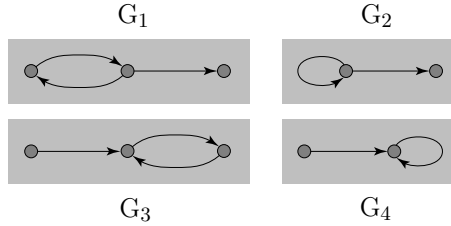


Figure 10.3

Schémas de distributivité de quantificateurs. Les deux schémas duals S74 ci-dessous sont les schémas de distributivité des quantificateurs universels par rapport à la conjonction et des quantificateurs existentiels par rapport à la disjonction. Leur démonstration est donnée dans la figure 10.4.

S74.

$$\frac{}{\Gamma \vdash \forall x(A \text{ et } B) \Leftrightarrow (\forall xA \text{ et } \forall xB)} \quad \frac{}{\Gamma \vdash \exists x(A \text{ ou } B) \Leftrightarrow (\exists xA \text{ ou } \exists xB)}.$$

Il faut faire attention à l'association des quantificateurs universels avec la conjonction, et des quantificateurs existentiels avec la disjonction dans ces schémas. Il n'y a pas de distributivité semblable des quantificateurs universels par rapport à la disjonction, ni des quantificateurs existentiels par rapport à la conjonction, mais seulement les schémas d'implication suivants:

S75.

$$\frac{}{\Gamma \vdash \exists x(A \text{ et } B) \Rightarrow (\exists xA \text{ et } \exists xB)} \quad \frac{}{\Gamma \vdash (\forall xA \text{ ou } \forall xB) \Rightarrow \forall x(A \text{ ou } B)}.$$

Les démonstrations se trouvent également dans la figure 10.4. Il faut prendre garde aussi au sens de ces implications. Les erreurs sur ce sens sont assez fréquentes.

Exemple. Pour pouvoir appliquer efficacement les schémas de déduction de la logique, il est important de les assimiler intuitivement. L'exemple qui suit devrait faciliter l'assimilation des règles ci-dessus. Il est question d'ensembles, mais de manière intuitive, non dans le cadre formel de la théorie $\mathcal{T}_{\text{ENS}}$.

Soient a et b deux ensembles, et considérons les propositions

$$\exists x(x \in a \text{ et } x \in b), \quad \exists x(x \in a) \text{ et } \exists x(x \in b).$$

La première signifie qu'il existe un objet appartenant aux deux ensembles, ou qui est élément des deux ensembles, autrement dit que ceux-ci ont un élément commun. La seconde signifie que chacun d'eux possède un élément (au moins).

nil		
1°	$\forall x(A \text{ et } B)$	hyp
2°	$A \text{ et } B$	1°, S65
3°	A	2°, et-élim
4°	B	2°, et-élim
5°	$\forall xA$	3°, S62
6°	$\forall xB$	4°, S62
7°	$\forall xA \text{ et } \forall xB$	5°, 6°, et-intro
8°	$\forall xA \text{ et } \forall xB$	hyp
9°	$\forall xA$	8°, et-élim
10°	$\forall xB$	8°, et-élim
11°	A	9°, S65
12°	B	10°, S65
13°	$A \text{ et } B$	11°, 12°, et-intro
14°	$\forall x(A \text{ et } B)$	13°, S62
15°	$\forall x(A \text{ et } B) \Leftrightarrow (\forall xA \text{ et } \forall xB)$	S40.
	$\exists x(A \text{ ou } B) \Leftrightarrow \text{non}\forall x\text{non}(A \text{ ou } B)$	S71
	$\Leftrightarrow \text{non}\forall x(\text{non}A \text{ et } \text{non}B)$	De Morgan
	$\Leftrightarrow \text{non}(\forall x\text{non}A \text{ et } \forall x\text{non}B)$	S74 (premier)
	$\Leftrightarrow \text{non}\forall x\text{non}A \text{ ou } \text{non}\forall x\text{non}B$	De Morgan
	$\Leftrightarrow \exists xA \text{ ou } \exists xB$	S71.
1°	$\exists x(A \text{ et } B)$	hyp
2°	$A \text{ et } B$	hyp
3°	A	2°, et-élim
4°	B	2°, et-élim
5°	$\exists xA$	3°, S66
6°	$\exists xB$	4°, S66
7°	$\exists xA \text{ et } \exists xB$	5°, 6°, et-intro
8°	$\exists xA \text{ et } \exists xB$	1°, 7°, S64
9°	$\exists x(A \text{ et } B) \Rightarrow (\exists xA \text{ et } \exists xB)$	8°, $\Rightarrow$ -intro.
1°	$\forall xA \text{ ou } \forall xB$	hyp
2°	$\forall xA$	hyp
3°	A	2°, S65
4°	$A \text{ ou } B$	3°, ou-intro
5°	$\forall xB$	hyp
6°	B	5°, S65
7°	$A \text{ ou } B$	6°, ou-intro
8°	$A \text{ ou } B$	4°, 7°, disj cas
9°	$\forall x(A \text{ ou } B)$	8°, S62
10°	$(\forall xA \text{ ou } \forall xB) \Rightarrow \forall x(A \text{ ou } B)$	9°, $\Rightarrow$ -intro.

Figure 10.4 Démonstrations des schémas S74 et S75.

Il faut bien percevoir le caractère local des variables liées. La seconde proposition a le même sens que $\exists x(x \in a)$ et $\exists x'(x' \in b)$. Si deux ensembles ont un élément commun, alors chacun d'eux possède un élément. Mais on ne peut pas affirmer la réciproque : si deux ensembles possèdent chacun un élément, ils n'ont pas nécessairement d'élément commun. Considérons ensuite les propositions

$$\forall x(x \notin a) \text{ ou } \forall x(x \notin b) \quad \forall x(x \notin a \text{ ou } x \notin b).$$

La première signifie que l'un (au moins) des deux ensembles a , b est sans éléments, autrement dit vide. La seconde signifie que les deux ensembles n'ont pas d'élément commun. Il est clair que si l'un des deux ensembles n'a pas d'éléments, ces ensembles n'ont pas d'élément commun. Mais on ne peut pas affirmer la réciproque. Si deux ensembles n'ont pas d'élément commun, cela n'implique pas que l'un ou l'autre soit vide.

Application: le théorème du barbier. Soit R un symbole relationnel binaire quelconque (du langage $\mathcal{L}_\Omega$). Nous démontrons plus loin dans ce paragraphe le théorème suivant de **Lpréd**:

$$\vdash \text{non} \exists y \forall x (x \not R y \Leftrightarrow x R y). \quad (1)$$

Nous utilisons la notation infixée pour le symbole R . La notation $x \not R y$ est l'abréviation habituelle de la proposition $\text{non}(x R y)$.

Avant de démontrer ce théorème, nous en présentons une interprétation amusante proposée par Bertrand Russel¹. L'univers de cette interprétation est l'ensemble des hommes d'un village, qui se divisent en deux classes : ceux qui se rasent eux-mêmes, et ceux qui ne se rasent pas ou se font raser par quelqu'un d'autre — un barbier par exemple. Suivant cette interprétation, la proposition $x R y$ signifie :

x se fait raser par y .

Le symbole relationnel R est donc le prédicat «se fait raser par». La proposition $x R x$ en particulier signifie que x se rase lui-même, et $x \not R x$ signifie que x ne se rase pas lui-même. Avant de démontrer le théorème (1), nous voulons le formuler en français, car il est bon de savoir mettre les expressions formelles telles que (1) en relation avec les phrases typiques que l'on rencontre à leur place dans les exposés moins formels. Une formulation possible de (1) en français est la suivante :

Il n'existe pas d'homme qui rase *tous* les hommes qui
ne se rasent pas eux-même et *seulement* des hommes
qui ne se rasent pas eux-mêmes. (2)

Pour comprendre cette formulation, il faut voir que l'assertion du théorème (1) est équivalente dans **Lpréd** à

$$\text{non} \exists y (\forall x (x \not R y \Rightarrow x R y) \text{ et } \forall x (x R y \Rightarrow x \not R x)). \quad (3)$$

¹ Bertrand Arthur William Russel (1872–1970), philosophe et logicien anglais. Prix Nobel de littérature 1950. Auteur avec A. N. Whitehead des *Principia Mathematica* (1910–1927), un traité de mathématiques dérivées formellement de la logique.

Cela est démontré par la chaîne d'équivalences suivante:

nil	
	$\text{non}\exists y\forall x(x \not\text{R} x \Leftrightarrow x \text{R} y)$
	$\Leftrightarrow \text{non}\exists y\forall x((x \not\text{R} x \Rightarrow x \text{R} y) \text{ et } (x \text{R} y \Rightarrow x \not\text{R} x))$ S41
	$\Leftrightarrow \text{non}\exists y(\forall x(x \not\text{R} x \Rightarrow x \text{R} y) \text{ et } \forall x(x \text{R} y \Rightarrow x \not\text{R} x))$ S74.

Les sous-propositions (A) $\forall x(x \not\text{R} x \Rightarrow x \text{R} y)$, et (B) $\forall x(x \text{R} y \Rightarrow x \not\text{R} x)$ de (3) sont traduites dans la figure 10.5, ainsi que leurs négations. La négation de A et celle de B sont équivalentes respectivement à $\exists x(x \not\text{R} x \text{ et } x \not\text{R} y)$ et $\exists x(x \text{R} y \text{ et } x \text{R} x)$ (cf. §10.2, Exemple 2). Cela devrait clarifier la traduction (2) du théorème (1).

A	$\forall x(x \not\text{R} x \Rightarrow x \text{R} y)$	tout homme qui ne se rase pas lui-même se fait raser par y y rase tous les hommes qui ne se rasent pas eux-mêmes
B	$\forall x(x \text{R} y \Rightarrow x \not\text{R} x)$	tout homme qui se fait raser par y ne se rase pas lui-même y rase seulement des hommes qui ne se rasent pas eux-mêmes
nonA	$\text{non}\forall x(x \not\text{R} x \Rightarrow x \text{R} y)$ $\exists x(x \not\text{R} x \text{ et } x \not\text{R} y)$	y ne rase pas tous les hommes qui ne se rasent pas eux-mêmes il existe quelqu'un qui ne se rase pas lui-même et qui n'est pas rasé par y
nonB	$\text{non}\forall x(x \text{R} y \Rightarrow x \not\text{R} x)$ $\exists x(x \text{R} y \text{ et } x \text{R} x)$	y ne rase pas seulement des hommes qui ne se rasent pas eux-mêmes y rase quelqu'un qui se rase lui-même

Figure 10.5

La démonstration du théorème (1) est la suivante:

nil	
1°	$\exists y\forall x(x \not\text{R} x \Leftrightarrow x \text{R} y)$ hyp
2°	$\forall x(x \not\text{R} x \Leftrightarrow x \text{R} y)$ hyp
3°	$y \not\text{R} y \Leftrightarrow y \text{R} y$ 2°, S61
4°	$\text{non}(y \not\text{R} y \Leftrightarrow y \text{R} y)$ S45
5°	$\perp$ 3°, 4°, $\perp$ -intro
6°	$\perp$ 1°, 5°, S64
7°	$\text{non}\exists y\forall x(x \not\text{R} x \Leftrightarrow x \text{R} y)$ 6°, red abs J.

L'idée clef de cette démonstration par réduction à l'absurde est la particularisation (S61) judicieuse effectuée à la ligne 3°. En éliminant le quantificateur

existentiel, on a fait l'hypothèse, à la ligne 2°, que y est un homme tel que, pour *tout* homme x : x ne se rase pas lui-même si et seulement si x se fait raser par y . Ce qui est vrai pour tout homme est vrai pour y , et en particulierisant avec y , on obtient immédiatement une proposition fausse, donc une contradiction.

On peut démontrer le théorème (1) d'une autre manière, en démontrant le théorème *équivalent* suivant :

$$\vdash \forall y(\exists x(x \not R x \text{ et } x \not R y) \text{ ou } \exists x(x R y \text{ et } x R x)). \quad (4)$$

Il faut s'assurer d'abord de l'équivalence des assertions de (1) et (4). Nous le faisons formellement ci-dessous. Mais cette équivalence peut être perçue d'abord intuitivement si l'on formule (4) en français :

Pour tout homme y il existe un homme qui ne se rase pas lui-même et qui ne se fait pas raser par y ou il existe un homme qui se rase lui-même et se fait aussi raser par y . (5)

La démonstration de l'équivalence des assertions de (1) et (4) est la suivante :

nil

non $\exists y \forall x (x \not R x \Leftrightarrow x R y)$	
$\Leftrightarrow \forall y \exists x \text{non}(x \not R x \Leftrightarrow x R y)$	S71
$\Leftrightarrow \forall y \exists x \text{non}((x \not R x \Rightarrow x R y) \text{ et } (x R y \Rightarrow x \not R x))$	S41
$\Leftrightarrow \forall y \exists x (\text{non}(x \not R x \Rightarrow x R y) \text{ ou } \text{non}(x R y \Rightarrow x \not R x))$	De Morgan
$\Leftrightarrow \forall y \exists x ((x \not R x \text{ et } x \not R y) \text{ ou } (x R y \text{ et } x R x))$	S55
$\Leftrightarrow \forall y (\exists x (x \not R x \text{ et } x \not R y) \text{ ou } \exists x (x R y \text{ et } x R x))$	S74.

Pour démontrer le théorème (4), on considère un homme y quelconque (Fig. 10.6) et l'on raisonne par disjonction des cas : y se rase lui-même, y ne se rase pas lui-même. Dans le premier cas il existe un homme qui se rase lui-même et qui est rasé par y ; dans le second, il existe un homme qui ne se rase pas lui-même et qui n'est pas rasé par y . Le lecteur peut comparer ces introductions de $\exists$ avec les dernières de l'exemple 3 du paragraphe 9.6.

nil

1°	$y R y$	hyp
2°	$y R y \text{ et } y R y$	1°, et-intro
3°	$\exists x (x R x \text{ et } x R y)$	2°, S63
4°	$y \not R y$	hyp
5°	$y \not R y \text{ et } y \not R y$	4°, et-intro
6°	$\exists x (x \not R x \text{ et } x \not R y)$	5°, S63
7°	$\exists x (x \not R x \text{ et } x \not R y) \text{ ou } \exists x (x R x \text{ et } x R y)$	6°, 3°, S35
8°	$\forall y \exists x (x \not R x \text{ et } x \not R y) \text{ ou } \exists x (x R x \text{ et } x R y)$	7°, S62.

Conséquence en théorie des ensembles. Le théorème (1) a joué un rôle historique dans le développement de la théorie des ensembles. On peut l'écrire en effet avec n'importe quel symbole relationnel binaire R . En prenant pour R

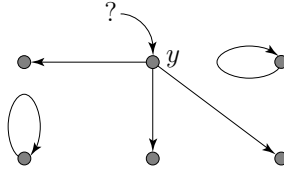


Figure 10.6 Graphe de la relation de rasage dans un village. Une flèche de a à b signifie que b est rasé par a . Ou bien y ne se rase pas lui-même, et alors il ne rase pas tous les hommes qui ne se rasent pas eux-mêmes. Ou bien y se rase lui-même, et alors il rase un homme qui se rase lui-même. Donc y ne peut pas à la fois : raser *tous* les hommes qui ne se rasent pas eux-mêmes et ne raser *que* des hommes qui ne se rasent pas eux-mêmes.

le symbole d'appartenance $\in$, le théorème devient

$$\text{non} \exists y \forall x (x \notin x \Leftrightarrow x \in y). \quad (6)$$

Il n'existe pas d'ensemble y tels que y ait comme éléments tous les ensembles qui ne sont pas éléments d'eux-mêmes et seulement les ensembles qui ne sont pas éléments d'eux-mêmes. C'est une propriété de l'univers des ensembles, et plus généralement de tout univers d'objets dans lequel on considère une relation $x R y$ entre objets. Ce qui se passe en théorie des ensembles, lorsqu'on choisit les axiomes et les schémas de déduction élémentaires de la théorie d'après l'intuition, c'est qu'on est tenté d'oublier ce théorème logique, et d'admettre des déductions qui entrent en contradiction avec celui-ci. Cela se produisit historiquement, et c'est à ce moment (1902) que Russel attira l'attention sur le théorème (6).

La tentation, en théorie des ensembles, est d'admettre que pour n'importe quelle propriété $P(x)$ d'un objet x , on peut parler de l'ensemble des objets qui possèdent cette propriété, c'est-à-dire croire qu'il existe un ensemble y ayant pour éléments tous les objets x qui ont la propriété $P(x)$ et seulement ces objets. Cet ensemble y vérifie la proposition

$$\forall x (P(x) \Leftrightarrow x \in y). \quad (7)$$

Or si l'on prend comme propriété $P(x)$ la propriété « x est un ensemble qui n'est pas élément de lui-même », on voit d'après (6) qu'il n'existe pas d'ensemble y qui vérifie (7). Le concept intuitif de l'ensemble des objets qui possèdent une certaine propriété ne peut pas être admis sans restriction. C'est ce dont on ne s'était pas aperçu avant l'intervention de Russel. Ce concept sera traité au chapitre 14.

Exercice. Soit $\mathcal{T}$ la théorie logique ayant pour axiomes les propositions

- (a) $\forall x (\text{non}(x \rightarrow x) \Leftrightarrow (x \rightarrow 0 \text{ ou } x \rightarrow 1)).$
- (b) $\forall a \forall b (a = b \Rightarrow \forall x (x \rightarrow a \Leftrightarrow x \rightarrow b)).$

Le symbole $\rightarrow$ est un symbole relationnel binaire et les symboles 0, 1 sont des constantes individuelles. Comme interprétation, on peut imaginer un univers de points reliés par des flèches (cf. Fig. 10.6). La proposition $x \rightarrow y$ signifie qu'il y a une flèche allant du point x au point y et les constantes 0 et 1 désignent deux points particuliers. Le premier axiome est que s'il n'y a pas de flèche allant d'un point x à lui-même, alors il y a une flèche allant de ce point x au point 0 ou au point 1. Le deuxième axiome est une propriété normale de l'égalité, si l'on interprète $a = b$ comme « a et b sont le même point ». Démontrer les théorèmes suivants de $\mathcal{T}$:

- (c) $\vdash 0 \neq 1$.
- (d) $\vdash \forall x \exists y (x \rightarrow y)$.
- (e) $\vdash \text{non}(0 \rightarrow 0)$.
- (f) $\vdash 0 \rightarrow 1$.

Pour (c) raisonner par l'absurde. En faisant l'hypothèse $0 = 1$ et en utilisant (b), obtenir une contradiction avec le théorème de Russel (1) pour le symbole relationnel $\rightarrow$. Pour (d) raisonner par disjonction des cas: $x \rightarrow x$, $\text{non}(x \rightarrow x)$. Pour (e) raisonner aussi par l'absurde. De l'hypothèse $0 \rightarrow 0$ on déduit une contradiction en particulierisant (a). Enfin (f) se déduit de (e) et d'une particularisation de (a), par la logique propositionnelle.

10.4 Changements de variables liées

Au paragraphe 9.4 (Exemple 2) nous avons démontré le théorème suivant de **Lpréd**: $\vdash \forall x (x \notin \emptyset) \Rightarrow \forall y (y \notin \emptyset)$. Le théorème réciproque, $\vdash \forall y (y \notin \emptyset) \Rightarrow \forall x (x \notin \emptyset)$ se démontre de la même manière, en permutant x et y . On a donc dans **Lpréd** le théorème

$$\vdash \forall x (x \notin \emptyset) \Leftrightarrow \forall y (y \notin \emptyset).$$

Les propositions $\forall x (x \notin \emptyset)$, $\forall y (y \notin \emptyset)$ sont équivalentes, par a fortiori, dans n'importe quel contexte logique. Rappelons que nous avons pris la première comme axiome de $\mathcal{T}_{\text{ENS}}$ (§8.2). Nous aurions pu prendre la seconde. Ces deux propositions ne sont pas identiques, mais on dit cependant qu'elles sont identiques à un *changement de variable liée près*.

Nous donnons dans ce paragraphe la définition syntaxique exacte d'un changement de variable liée dans une proposition, et nous démontrons les schémas de déduction suivant lesquels deux propositions qui se transforment l'une en l'autre par des changements de variables liées sont équivalentes dans tout contexte logique.

Nous rappelons la définition générale (§4.5) de $(T|x)A$ pour une proposition A , une variable x et un terme T : cette notation désigne la proposition obtenue en remplaçant *toutes les occurrences libres* de x par le terme T dans A ; s'il n'y en a pas, c'est-à-dire si x ne figure pas librement dans A , $(T|x)A$ désigne la proposition A . Dans ce paragraphe nous ne considérons que des substitutions

où le terme T est lui-même une variable y . Nous aurons donc affaire à la notation $(y|x)A$.

Exemples. Considérons la proposition A suivante:

$$\exists a(x + a = b). \quad (A)$$

Désignons par A' la proposition $(d|x)A$, c'est-à-dire

$$\exists a(d + a = b), \quad (A')$$

et observons que si l'on effectue dans A' la substitution inverse, c'est-à-dire si l'on prend la proposition $(x|d)A'$, on retrouve la proposition initiale A . Autrement dit la proposition $(x|d)((d|x)A)$ est identique à A . Pour cette raison, on dit que la variable d est *réversiblement substituable* à x dans A .

Nous voulons déterminer, d'abord dans cet exemple et ensuite de façon générale, quelles sont les variables réversiblement substituables à une variable donnée x dans une proposition donnée. Il y en a bien sûr une infinité — par exemple les variables d, d', d'', d''' , etc. sont toutes réversiblement substituables à x dans la proposition A ci-dessus. Il ne s'agit donc pas de les énumérer toutes, mais de les caractériser par des critères simples. Ceux-ci apparaissent immédiatement lorsqu'on examine la question inverse: quelles variables ne *sont pas* réversiblement substituables à une variable donnée dans une proposition donnée? Car celles-ci sont en nombre fini.

Dans la proposition A du présent exemple, il y a exactement deux variables qui ne sont pas réversiblement substituables à x , à savoir les variables a et b (Fig. 10.7). La proposition $(a|x)A$ est

$$\exists a(\overbrace{a + a} = b).$$

C'est une proposition B dans laquelle la variable a n'a pas d'occurrence libre, de sorte que $(x|a)B$ est identique à B , non à A . La proposition $(b|x)A$ est

$$\exists a(b + a = b)$$

et dans cette proposition C il y a *deux* occurrences libres de b , de sorte que $(x|b)C$ n'est pas identique à A , mais à $\exists a(x + a = x)$.

Il est clair que ce qui rend la variable a non réversiblement substituable à x dans A c'est le fait que a n'est *pas librement* substituable à x dans A . Pour la variable b , c'est le fait que b est distincte de x et figure elle-même librement dans A . Par contre, toutes les variables distinctes de a et de b sont réversiblement substituables à x dans A .

Notons encore que la variable x est trivialement réversiblement substituable à elle-même dans A , et qu'il en va ainsi de n'importe quelle variable x dans n'importe quelle proposition A , puisque $(x|x)A$ est identique à A .

Comme autre exemple, prenons pour A la proposition

$$(a \subset x \text{ et } \exists x(x \in a)) \Rightarrow \exists y(y \in x).$$

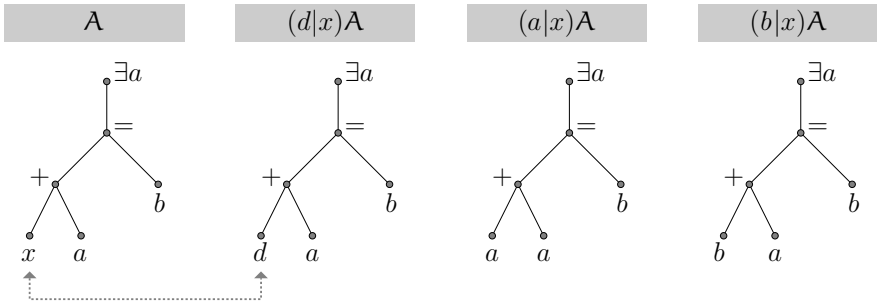


Figure 10.7 La variable d est réversiblement substituable à x dans la proposition A : A est identique à $(x|d)((d|x)A)$. Les variables a et b ne sont pas réversiblement substitubles à x dans A .

Les variables non réversiblement substitubles à x dans cette proposition A sont :

- (a) la variable y , qui n'est pas librement substituable à x ;
- (b) la variable a , qui est distincte de x et figure librement dans A .

Conditions de réversibilité d'une substitution. Les exemples qui précèdent suggèrent qu'étant donné une proposition A quelconque, et deux variables *distinctes* quelconques x et y , la variable y sera réversiblement substituable à x dans A , si et seulement si les deux conditions suivantes sont satisfaites :

- (a) y est librement substituable à x dans A ;
- (b) y ne figure pas librement dans A .

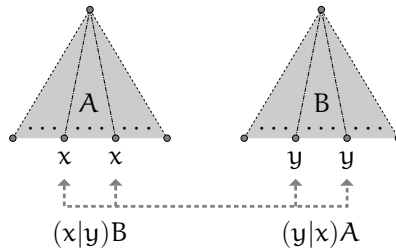


Figure 10.8 Si y (distincte de x) est librement substituable à x dans A et ne figure pas librement dans A , et si B est la proposition $(y|x)A$, alors A est identique à $(x|y)B$.

Pour montrer cela, nous nous appuierons sur les figures 10.8 à 10.10, dans lesquelles deux triangles symbolisent les arbres d'une proposition indéterminée A et de la proposition $(y|x)A$, désignée par B . On suppose que x et y sont deux variables *distinctes*. On a représenté schématiquement deux occurrences libres de x dans A , et les deux traits qui descendent de la racine de l'arbre à ces deux feuilles symbolisent les branches correspondantes de l'arbre.

Dans la figure 10.8, il est supposé (a) que y est librement substituable à x dans A . Cela signifie que sur les branches des occurrences libres de x dans A il n'y a pas de nœud portant le quantificateur $\forall y$ ou le quantificateur $\exists y$. Si l'on remplace les occurrences libres de x par y , on obtient donc une proposition B dans laquelle ces occurrences de y sont *libres*. D'autre part on suppose (b) que y ne figure pas librement dans A . Donc les seules occurrences libres de y dans B sont celles qui proviennent de la substitution effectuée. Lorsqu'on remplace toutes ces occurrences de y dans B par x , on retrouve la proposition A .

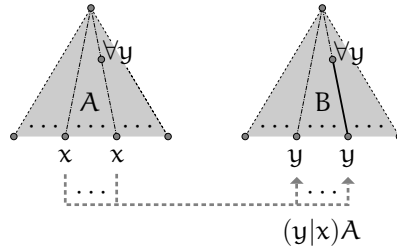


Figure 10.9 Si y (distincte de x) n'est pas librement substituable à x dans A , et si B est la proposition $(y|x)A$, alors A n'est pas identique à $(x|y)B$.

Dans la figure 10.9, on suppose que y ne satisfait pas à la condition (a), à savoir que y n'est pas librement substituable à x dans A . Cela signifie que dans l'arbre de A il y a une occurrence libre de x sur la branche de laquelle se trouve un nœud portant le quantificateur $\forall y$ ou $\exists y$. Lorsqu'on remplace toutes les occurrences libres de x par y , l'une au moins des occurrences de y qui en résultent dans B est une occurrence liée qui est désormais « inamovible », et subsiste lorsqu'on effectue la substitution inverse $(x|y)B$. Donc cette dernière proposition n'est pas identique à A .

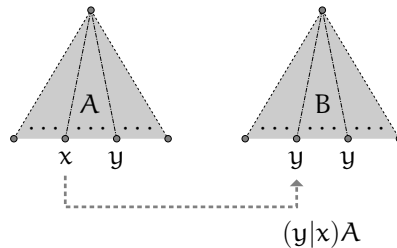


Figure 10.10 Si y (distincte de x) figure librement dans A , et si B est la proposition $(y|x)A$, alors A n'est pas identique à $(x|y)B$.

Enfin, dans la figure 10.10, on suppose que y ne satisfait pas à la condition (b), c'est-à-dire que y figure librement dans A . Dans ce cas il est clair que le nombre d'occurrences libres de y dans B n'est pas égal à celui des occurrences libres de x dans A , ne sorte que la proposition $(x|y)B$ n'est pas identique à A .

Relevons pour finir une propriété évidente, qui interviendra dans la démonstration du prochain schéma de déduction. Si une variable y est réversiblement substituable à une variable x dans une proposition A (cas de la figure 10.8), et si B désigne la proposition $(y|x)A$, alors x est réversiblement substituable à y dans B . Car alors $(y|x)((x|y)B)$ est identique à $(y|x)A$, donc à B . Si en outre x et y sont distinctes, alors elles satisfont vis-à-vis de B aux conditions symétriques de (a), (b):

- (a') x est librement substituable à y dans B ;
- (b') x ne figure pas librement dans B .

Schémas de déduction par changement de variable liée. Au début de ce paragraphe, nous avons mentionné le théorème

$$\vdash \forall x (x \notin \emptyset) \Leftrightarrow \forall y (y \notin \emptyset). \quad (1)$$

Ce théorème de **Lpréd** appartient à une classe de théorèmes de la même forme que l'on peut décrire par un schéma de déduction sans prémisses. Cette forme peut être mise en évidence dans cet exemple en désignant par A la proposition $x \notin \emptyset$. Alors la proposition $y \notin \emptyset$ peut se désigner par $(y|x)A$, et le théorème (1) peut être désigné par

$$\vdash \forall x A \Leftrightarrow \forall y (y|x)A.$$

On a le même théorème pour une proposition A quelconque et des variables x , y quelconques, à condition que y soit réversiblement substituable à x dans A . Si cette condition est satisfaite, on dit que la proposition $\forall y (y|x)A$ est construite par *changement de variable liée* dans la proposition $\forall x A$. Un changement de variable liée consiste en un changement de la variable d'un *quantificateur* (1) et du même changement pour toutes les occurrences de variables liées à ce quantificateur. Par $\forall y (y|x)A$, nous désignons évidemment la proposition $\forall y ((y|x)A)$.

S76.

$$\frac{}{\Gamma \vdash \forall x A \Leftrightarrow \forall y (y|x)A} \qquad \frac{}{\Gamma \vdash \exists x A \Leftrightarrow \exists y (y|x)A}.$$

Si y est réversiblement substituable à x dans A .

Démonstration. Le cas où les variables x et y sont identiques est trivial: les propositions $\forall x A$ et $\forall y (y|x)A$ sont alors identiques, donc équivalentes d'après S47, de même que les propositions $\exists x A$ et $\exists y (y|x)A$.

Supposons que x et y soient des variables distinctes et que y soit réversiblement substituable à x dans A . Désignons par B la proposition $(y|x)A$. La proposition $(x|y)B$ est identique à A , et les variables x et y vérifient par rapport à A et à B les conditions que nous rappelons:

- (a) y est librement substituable à x dans A ;
- (b) y ne figure pas librement dans A ;
- (a') x est librement substituable à y dans B ;

(b') x ne figure pas librement dans B .

En vertu de cela, les séquences de jugements de la figure 10.11 sont des démonstrations de **Lpréd**. Les conditions (a) et (a') garantissent que les déductions suivant **S61** aux lignes 2^o et 5^o respectent la condition de ce schéma. D'autre part (b) et (b') permettent les généralisations (**S62**) des lignes 3^o et 6^o, puisque y ne figure pas librement dans l'hypothèse $\forall xA$ et x ne figure pas librement dans l'hypothèse $\forall yB$. Nous laissons au lecteur le soin de contrôler que les conditions des schémas **S63** et **S64** sont aussi respectées dans la deuxième démonstration.

	nil	
1 ^o	$\forall xA$	hyp
2 ^o	$(y x)A$	1 ^o , S61 , (a)
3 ^o	$\forall y \underbrace{(y x)A}_B$	2 ^o , S62 , (b)
4 ^o	$\forall yB$	hyp
5 ^o	$(x y)B$	4 ^o , S61 , (a')
6 ^o	$\forall x \underbrace{(x y)B}_A$	5 ^o , S62 , (b')
7 ^o	$\forall xA \Leftrightarrow \forall y(y x)A$	3 ^o , 6 ^o , S40 .

1 ^o	$\exists xA$	hyp
2 ^o	$\underbrace{(x y)B}_A$	hyp
3 ^o	$\exists yB$	2 ^o , S63 , (a')
4 ^o	$\exists yB$	1 ^o , 3 ^o , S64 , (b')
5 ^o	$\exists yB$	hyp
6 ^o	$\underbrace{(y x)A}_B$	hyp
7 ^o	$\exists xA$	6 ^o , S63 , (a)
8 ^o	$\exists xA$	5 ^o , 7 ^o , S64 , (b)
9 ^o	$\exists xA \Leftrightarrow \exists y(y x)A$	4 ^o , 8 ^o , S40 .

Figure 10.11 Démonstrations des schémas de déduction **S76**

Changement de variable liée dans une sous-proposition. Nous appliquerons souvent le schéma de déduction **S76** pour transformer une proposition R en une proposition équivalente R' en effectuant un changement de variable liée dans une sous-proposition de R de la forme $\forall xA$ (resp. $\exists xA$). Cela veut dire remplacer une sous-proposition de R de cette forme par une proposition $\forall y(y|x)A$ (resp. $\exists y(y|x)A$), où y est une variable réversiblement substituable à x dans A .

Par exemple, la proposition

$$\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y)) \quad (2)$$

est équivalente dans tout contexte logique à

$$\forall x \forall y (x \subset y \Leftrightarrow \forall b (b \in x \Rightarrow b \in y)), \quad (3)$$

car on passe de la première à la seconde par un changement de variable liée dans la sous-proposition $\forall a (a \in x \Rightarrow a \in y)$, que l'on remplace par

$$\forall b ((b|a)(a \in x \Rightarrow a \in y)).$$

En d'autres termes, on effectue souvent d'un seul coup une déduction suivant **S76** et une déduction suivant le schéma général de substitutivité de l'équivalence (§ 10.1). Souvent il s'y ajoute encore une déduction par a fortiori, comme nous l'avons expliqué au paragraphe 10.1. C'est ce que nous faisons par exemple en disant que les propositions (2) et (3) sont équivalentes dans tout contexte, car le schéma général de substitutivité de l'équivalence permet de le dire seulement pour un contexte dans les hypothèses duquel les variables x et y ne figurent pas librement (ce sont les variables des quantificateurs qui dominent la sous-proposition de **R** considérée). Mais (2) et (3) sont bien équivalentes dans tout contexte, par a fortiori, car elles sont équivalentes dans tout contexte sans hypothèses.

Squelette d'une proposition. Nous présentons ici une manière commode de reconnaître que deux propositions **R**₁, **R**₂ sont identiques à des changements de variables liées près, en ce sens que l'une peut être transformée en l'autre par une succession d'opérations de changement de variable liée dans une sous-proposition. Considérons d'abord un changement de variable liée simple, par exemple le passage de la première à la seconde des propositions suivantes :

$$\begin{aligned} \forall a (a \in x \Rightarrow a \in y), \\ \forall b (b \in x \Rightarrow b \in y). \end{aligned}$$

Nous appelons *squelette* de ces deux propositions l'expression

$$\forall \square (\square \in x \Rightarrow \square \in y). \quad (4)$$

De façon générale, nous appelons *squelette d'une proposition* **R** l'expression que l'on obtient en effectuant les opérations suivantes, pour chaque sous-proposition de **R** de la forme **Qx****A** (où **Qx** est un quantificateur $\forall x$ ou $\exists x$) :

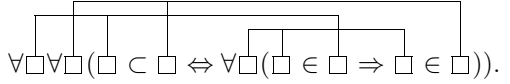
- (a) Remplacer cette sous-proposition par l'expression **Q** $\square(\square|x)$ **A**;
- (b) Relier par un trait tous les $\square$ s introduits dans l'opération (a).

La manière dont nous effectuons l'opération (b) est indiquée par l'exemple qui précède (4). La forme des traits qui servent de lien n'a pas d'importance. L'essentiel est que l'on voie clairement quels symboles $\square$ sont reliés entre eux, en particulier lorsque la proposition **R** dont on écrit le squelette comporte

plusieurs quantificateurs. Par exemple, le squelette de la proposition

$$\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y)) \quad (5)$$

est l'expression



$$\forall \square \forall \square (\square \subset \square \Leftrightarrow \forall \square (\square \in \square \Rightarrow \square \in \square)). \quad (6)$$

Nous utilisons dans ce graphisme la convention des schémas électriques selon laquelle il n'y a pas de contact entre deux lignes qui se croisent, à moins que l'on ne dessine un « point de soudure ». Cette expression ne comporte aucune variable, puisque la proposition (5) n'a pas de variables libres. La proposition suivante possède le même squelette (6) :

$$\forall a \forall u (a \subset u \Leftrightarrow \forall y (y \in a \Rightarrow y \in u)). \quad (7)$$

Nous allons voir que l'on peut transformer la proposition (5) en la proposition (7) par une succession de changements de variable liée (dans une sous-proposition). Les propositions (5) et (7) sont donc équivalentes dans tout contexte.

Les propositions qui ont le squelette (6) sont toutes les propositions de la forme

$$\forall x_1 \forall x_2 (x_1 \subset x_2 \Leftrightarrow \forall x_3 (x_3 \in x_1 \Rightarrow x_3 \in x_2)), \quad (R)$$

où x_1, x_2, x_3 sont trois variables quelconques mais *distinctes*. Considérons deux propositions de cette forme: **R** et

$$\forall y_1 \forall y_2 (y_1 \subset y_2 \Leftrightarrow \forall y_3 (y_3 \in y_1 \Rightarrow y_3 \in y_2)). \quad (R')$$

Nous supposons donc que les variables x_1, x_2, x_3 sont distinctes, de même que les variables y_1, y_2, y_3 . Mais les premières ne sont pas nécessairement distinctes des secondes: voir l'exemple des propositions (5) et (7). Choisissons trois variables z_1, z_2, z_3 distinctes entre elles et distinctes des variables x_i et y_i ($i = 1, 2, 3$). La figure 10.12 montre comment passer de **R** à **R'** par une succession de changements de variable liée. Si nous désignons par R_1 à R_7 les propositions de cette suite, nous voyons qu'en vertu du choix des variables le passage de R_i à R_{i+1} s'effectue chaque fois en remplaçant une variable de R_i par une variable qui ne figure pas du tout dans R_i (ni librement, ni dans un quantificateur). La condition de **S76** est donc vérifiée à chaque pas. Par exemple, lorsqu'on passe de la deuxième à la troisième de ces propositions, on remplace la sous-proposition

$$\underbrace{\forall x_2 (z_1 \subset x_2 \Leftrightarrow \forall x_3 (x_3 \in z_1 \Rightarrow x_3 \in x_2))}_A$$

par la proposition

$$\underbrace{\forall z_2 (z_1 \subset z_2 \Leftrightarrow \forall x_3 (x_3 \in z_1 \Rightarrow x_3 \in z_2))}_{(z_2|x_2)A}.$$

La variable z_2 est réversiblement substituable à x_2 dans la proposition **A** soulignée. Les conditions de réversibilité (a), (b) sont satisfaites: z_2 est librement substituable à x_2 dans **A** parce que z_2 est distincte de x_3 ; z_2 ne figure pas librement dans **A** parce que z_2 est distincte de z_1 .

$$\begin{array}{c}
\forall x_1 \forall x_2 (x_1 \subset x_2 \Leftrightarrow \forall x_3 (x_3 \in x_1 \Rightarrow x_3 \in x_2)) \\
\downarrow \quad \downarrow \quad \quad \quad \downarrow \\
\forall z_1 \forall x_2 (z_1 \subset x_2 \Leftrightarrow \forall x_3 (x_3 \in z_1 \Rightarrow x_3 \in x_2)) \\
\downarrow \quad \quad \quad \downarrow \quad \quad \quad \downarrow \\
\forall z_1 \forall z_2 (z_1 \subset z_2 \Leftrightarrow \forall x_3 (x_3 \in z_1 \Rightarrow x_3 \in z_2)) \\
\quad \quad \quad \downarrow \quad \downarrow \quad \quad \quad \downarrow \\
\forall z_1 \forall z_2 (z_1 \subset z_2 \Leftrightarrow \forall z_3 (z_3 \in z_1 \Rightarrow z_3 \in z_2)) \\
\downarrow \quad \downarrow \quad \quad \quad \downarrow \\
\forall y_1 \forall z_2 (y_1 \subset z_2 \Leftrightarrow \forall z_3 (z_3 \in y_1 \Rightarrow z_3 \in z_2)) \\
\downarrow \quad \downarrow \quad \quad \quad \downarrow \\
\forall y_1 \forall y_2 (y_1 \subset y_2 \Leftrightarrow \forall z_3 (z_3 \in y_1 \Rightarrow z_3 \in y_2)) \\
\quad \quad \quad \downarrow \quad \downarrow \quad \quad \quad \downarrow \\
\forall y_1 \forall y_2 (y_1 \subset y_2 \Leftrightarrow \forall y_3 (y_3 \in y_1 \Rightarrow y_3 \in y_2))
\end{array}$$

Figure 10.12 Transformation d'une proposition en une proposition équivalente de même squelette par une série de changements de variable liée. Les variables z_1, z_2, z_3 sont supposées distinctes entre elles et distinctes des variables x_1, x_2, x_3 et y_1, y_2, y_3 .

Schéma général de déduction par changement de variables liées. L'exemple qui précède montre clairement que deux propositions de même squelette sont équivalentes dans tout contexte logique. Nous pouvons noter ceci comme schéma de déduction général sans prémisses de **Lpréd**:

$$\frac{}{\Gamma \vdash R \Leftrightarrow R'} \quad \text{Si } R \text{ et } R' \text{ ont le même squelette.}$$

Ce schéma permet de faire d'un seul coup plusieurs changements de variables liées. Il suffit de s'assurer que ces changements conservent le squelette de la proposition à laquelle on les applique.

10.5 Particularisations simultanées

Substitutions simultanées. Soient **A** une expression (terme ou proposition), $x_1, \dots, x_n$ des variables *distinctes* et $T_1, \dots, T_n$ des termes quelconques. Nous désignons par

$$(T_1|x_1, \dots, T_n|x_n)A \tag{1}$$

l'expression obtenue en remplaçant *simultanément* dans **A**: toutes les occurrences libres de x_1 par T_1 ; toutes les occurrences libres de x_2 par T_2 , etc.

Par exemple, si A désigne la proposition $a \in x \Rightarrow a \in y$ et si T_1 et T_2 désignent les termes $x \cap y$, $x \cup \emptyset$, la proposition $(T_1|x, T_2|y)A$ est donnée dans le tableau suivant :

A	T_1	T_2	$(T_1 x, T_2 y)A$
$a \in x \Rightarrow a \in y$	$x \cap y$	$x \cup \emptyset$	$a \in x \cap y \Rightarrow a \in x \cup \emptyset$

La *simultanéité* des substitutions désignées par la notation (1) est importante. Dans l'exemple ci-dessus, la proposition $(T_1|x)((T_2|y)A)$, construite en remplaçant d'abord y par T_2 dans A , puis en remplaçant x par T_1 dans la proposition obtenue, est différente de la proposition $(T_1|x, T_2|y)A$:

A	T_1	T_2	$(T_2 y)A$
$a \in x \Rightarrow a \in y$	$x \cap y$	$x \cup \emptyset$	$a \in x \Rightarrow a \in x \cup \emptyset$

A	T_1	T_2	$(T_1 x)((T_2 y)A)$
$a \in x \Rightarrow a \in y$	$x \cap y$	$x \cup \emptyset$	$a \in x \cap y \Rightarrow a \in (x \cap y) \cup \emptyset$

La proposition $(T_2|y)((T_1|x)A)$ serait encore différente. Il est clair que ces différences proviennent du fait que les variables x et y auxquelles les termes T_1 et T_2 sont substitués *figurent* elles-mêmes (librement) dans ces termes. Lorsqu'on substitue T_2 à y dans A , on introduit une occurrence libre de x , et celle-ci est remplacée par T_1 si l'on effectue *ensuite* la substitution $(T_1|x)$. Par contre, si l'on choisit deux termes T'_1 et T'_2 dans lesquels les variables x et y ne *figurent pas librement*, les propositions

$$(T'_1|x, T'_2|y)A, \quad (T'_1|x)((T'_2|y)A), \quad (T'_2|y)((T'_1|x)A)$$

sont identiques. C'est un cas particulier de la propriété générale suivante.

Propriété. Soient $T_1, \dots, T_n$ des termes et $x_1, \dots, x_n$ des variables distinctes. Si ces variables ne figurent librement dans aucun des termes $T_1, \dots, T_n$, l'expression $(T_1|x_1, \dots, T_n|x_n)A$, pour un terme ou une proposition A quelconque, est identique à l'expression

$$(T_n|x_n) \cdots (T_2|x_2)((T_1|x_1)A) \quad (2)$$

obtenue en appliquant successivement l'opération $(T_1|x_1)$ à A , puis $(T_2|x_2)$ à l'expression $(T_1|x_1)A$, et ainsi de suite.

Nous admettons cette propriété comme évidente, mais nous indiquerons plus bas une manière de la prouver, à partir des règles de substitution ci-dessous qui constituent une définition opérationnelle précise de l'opération de substitution $(T_1|x_1, \dots, T_n|x_n)$.

Règles de substitution. La définition de la notation (1) peut être donnée de manière précise sous la forme des règles de substitution suivantes, semblables à celles du paragraphe 4.5 concernant le cas particulier $n = 1$.

(a) $(T_1|x_1, \dots, T_n|x_n)x_1$ est identique à T_1 , $(T_1|x_1, \dots, T_n|x_n)x_2$ est identique à T_2 , etc.

(b) Si $\mathbf{a}$ est une variable individuelle *distincte* de tous les $\mathbf{x}_i$, ou si $\mathbf{a}$ est une constante individuelle, $(T_1|x_1, \dots, T_n|x_n)\mathbf{a}$ est identique à $\mathbf{a}$.

(c) Si $\mathbf{A}$ est une expression composée, de la forme $s\mathbf{A}_1 \cdots \mathbf{A}_p$ dans laquelle s est un symbole fonctionnel ou un symbole relationnel ou un connecteur logique, l'expression $(T_1|x_1, \dots, T_n|x_n)\mathbf{A}$ est identique à

$$s((T_1|x_1, \dots, T_n|x_n)\mathbf{A}_1) \cdots ((T_1|x_1, \dots, T_n|x_n)\mathbf{A}_p).$$

Autrement dit, l'opération de substitution $(T_1|x_1, \dots, T_n|x_n)$ est appliquée à chacune des sous-expressions $\mathbf{A}_i$ subordonnées au symbole principal s de l'expression $\mathbf{A}$. Cas particulier: $(T_1|x_1, \dots, T_n|x_n)(\mathbf{A}_1 \text{ ou } \mathbf{A}_2)$ est identique à

$$((T_1|x_1, \dots, T_n|x_n)\mathbf{A}_1) \text{ ou } ((T_1|x_1, \dots, T_n|x_n)\mathbf{A}_2),$$

et de même pour les connecteurs et , $\Rightarrow$, $\Leftrightarrow$. Pour le connecteur de négation, $(T_1|x_1, \dots, T_n|x_n)(\text{non}\mathbf{A})$ est identique à $\text{non}((T_1|x_1, \dots, T_n|x_n)\mathbf{A})$.

(d) Si $\mathbf{A}$ est une proposition quelconque et si $\mathbf{y}$ est une variable *distincte* de $\mathbf{x}_1, \dots, \mathbf{x}_n$, les expressions $(T_1|x_1, \dots, T_n|x_n)(\forall \mathbf{y}\mathbf{A})$, $(T_1|x_1, \dots, T_n|x_n)(\exists \mathbf{y}\mathbf{A})$ sont identiques respectivement à

$$\forall \mathbf{y}((T_1|x_1, \dots, T_n|x_n)\mathbf{A}), \quad \exists \mathbf{y}((T_1|x_1, \dots, T_n|x_n)\mathbf{A}).$$

(e) Si $\mathbf{A}$ est une proposition quelconque, et si $\mathbf{x}_i$ est l'une quelconque des variables $\mathbf{x}_1, \dots, \mathbf{x}_n$, l'expression $(T_1|x_1, \dots, T_n|x_n)(\forall \mathbf{x}_i\mathbf{A})$ est identique à

$$(T_1|x_1, \dots, T_{i-1}|x_{i-1}, T_{i+1}|x_{i+1}, T_n|x_n)(\forall \mathbf{x}_i\mathbf{A}),$$

donc identique, d'après (d), à $\forall \mathbf{x}_i((T_1|x_1, \dots, T_{i-1}|x_{i-1}, T_{i+1}|x_{i+1}, T_n|x_n)\mathbf{A})$. On a la même règle pour $(T_1|x_1, \dots, T_n|x_n)(\exists \mathbf{x}_i\mathbf{A})$.

Application. Les règles qui précèdent *déterminent* exactement l'expression $(T_1|x_1, \dots, T_n|x_n)\mathbf{A}$ pour n'importe quelle expression $\mathbf{A}$, étant donné des termes $T_1, \dots, T_n$ et des variables distinctes $\mathbf{x}_1, \dots, \mathbf{x}_n$. Par exemple, si T_1 et T_2 sont des termes quelconques, l'expression

$$(T_1|x, T_2|y)(\forall a(\exists x(a \in x \cap y) \Rightarrow a \in y))$$

est déterminée comme suit d'après les règles mentionnées à droite. Pour la lisibilité, nous utilisons des paires de points au lieu de parenthèses. Chaque opérateur de substitution doit être appliqué à la sous-expression entre deux points qui est immédiatement à sa droite. Chaque ligne représente la même expression que celle qui la précède.

$$(T_1|x, T_2|y).\forall a(\exists x(a \in x \cap y) \Rightarrow a \in y).$$

$$\forall a(T_1|x, T_2|y).(\exists x(a \in x \cap y) \Rightarrow a \in y). \quad (\text{d})$$

$$\forall a((T_1|x, T_2|y).\exists x(a \in x \cap y). \Rightarrow (T_1|x, T_2|y).a \in y.) \quad (\text{c) pour } \Rightarrow$$

$$\forall a(\exists x(T_2|y).(a \in x \cap y). \Rightarrow (T_1|x, T_2|y).a \in (T_1|x, T_2|y).y.) \quad (\text{e), (c) pour } \in$$

$$\forall a(\exists x((T_2|y).a \in (T_2|y).x \cap y.) \Rightarrow a \in T_2) \quad (\text{c) pour } \in, (\text{b}), (\text{a})$$

$$\forall a(\exists x(a \in (T_2|y).x \cap (T_2|y).y.) \Rightarrow a \in T_2) \quad (\text{b}), (\text{c) pour } \cap$$

$$\forall a(\exists x(a \in x \cap T_2) \Rightarrow a \in T_2) \quad (\text{b}), (\text{a}).$$

Les règles de substitution (a) à (e) permettent de vérifier la propriété énoncée plus haut, à savoir l'identité des expressions

$$(T_1|x_1, \dots, T_n|x_n)A, \quad (T_n|x_n)(\dots (T_2|x_2)((T_1|x_1)A)\dots)$$

lorsque les variables $x_1, \dots, x_n$ ne figurent pas librement dans les termes $T_1, \dots, T_n$. Désignons la seconde expression par

$$(T_n|x_n) \cdots (T_1|x_1)A,$$

en omettant les parenthèses d'association à droite. Il suffit de montrer que l'opérateur de substitution séquentielle $(T_n|x_n) \cdots (T_1|x_1)$ satisfait aux *mêmes règles* (a) à (e) que l'opérateur de substitution simultanée $(T_1|x_1, \dots, T_n|x_n)$ lorsque les variables $x_1, \dots, x_n$ ne figurent pas librement dans les termes T_i . Par exemple:

(a) $(T_n|x_n) \cdots (T_1|x_1)x_1$ est identique à $(T_n|x_n) \cdots (T_2|x_2)T_1$, donc à T_1 si les variables $x_2, \dots, x_n$ ne figurent pas librement dans T_1 . Nous laissons au lecteur le soin de vérifier les autres règles.

Schéma des particularisations simultanées. Le schéma de déduction suivant est une généralisation du schéma de particularisation S61, qu'il englobe comme cas particulier ($n = 1$).

S77.

$$\frac{\Gamma \vdash \forall x_1 \cdots \forall x_n A}{\Gamma \vdash (T_1|x_1, \dots, T_n|x_n)A}$$

Si $x_1, \dots, x_n$ sont des variables distinctes, et si T_i est librement substituable à x_i dans A (pour $i = 1, \dots, n$).

Nous verrons plus bas comment une telle déduction peut se démontrer par une succession de n déductions suivant S61, précédées d'un changement de variables liées. Nous allons le voir d'abord dans un exemple.

Exemple. La déduction suivante est une déduction de Lpréd conforme à S77. Pour la discussion, nous désignerons par T_1 et T_2 les termes $x \cap y$ et $x \cup \emptyset$. Ces termes sont librement substituables respectivement à x et à y dans la proposition A indiquée, parce qu'ils ne contiennent pas d'occurrence de la variable a .

$$\frac{\Gamma \vdash \forall x \forall y \overbrace{(x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y))}^A}{\Gamma \vdash \underbrace{x \cap y \subset x \cup \emptyset \Leftrightarrow \forall a (a \in x \cap y \Rightarrow a \in x \cup \emptyset)}_{(x \cap y|x, x \cup \emptyset|y)A}} \quad (3)$$

Démonstration:

Γ		
1°	$\forall x \forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y))$	base
2°	$\forall x' \forall y' (x' \subset y' \Leftrightarrow \forall a (a \in x' \Rightarrow a \in y'))$	1°, chgmt var liées
3°	$\forall y' (T_1 \subset y' \Leftrightarrow \forall a (a \in T_1 \Rightarrow a \in y'))$	2°, S61
4°	$T_1 \subset T_2 \Leftrightarrow \forall a (a \in T_1 \Rightarrow a \in T_2)$	3°, S61.

Les propositions des lignes 1^o et 2^o sont identiques à des changements de variables liées près. Elles ont le même squelette et sont donc équivalentes d'après le schéma général de déduction par changement de variables liées (§10.4). On peut particulariser la proposition 2^o avec le terme $T_1 (x \cap y)$ (ligne 3^o) car ce terme est librement substituable à x' dans la proposition

$$\forall y' (x' \subset y' \Leftrightarrow \forall a (a \in x' \Rightarrow a \in y')).$$

Remarquons que sans le changement de variables liées 2^o cette déduction ne serait pas possible, car le terme T_1 n'est pas librement substituable à x dans la proposition $\forall y (x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y))$. On particularise ensuite la proposition de la ligne 3^o avec T_2 . À ce stade, les variables auxiliaires x' , y' ont complètement disparu. La proposition de la ligne 4^o est identique à celle qu'on obtient en prenant la proposition de la première ligne, en lui ôtant ses deux quantificateurs $\forall x \forall y$ et en remplaçant *simultanément* toutes les occurrences libres de x par T_1 et toutes celles de y par T_2 . C'est précisément la déduction décrite par S77.

Démonstration de S77. Il nous reste à montrer que l'exemple précédent se généralise, c'est-à-dire que l'on peut démontrer de manière analogue toute déduction de la forme décrite par S77. Pour ce faire, nous avons représenté une déduction de cette forme dans la figure 10.13, avec $n = 3$ pour fixer les idées. On imagine l'arbre d'une proposition de la forme $\forall x_1 \forall x_2 \forall x_3 A$, où A est une proposition quelconque et x_1, x_2, x_3 sont des variables *distinctes*. Une seule occurrence libre de chacune de ces variables x_i dans A est représentée, mais elle symbolise toutes les occurrences libres de x_i dans A . Les traits qui vont de la racine de l'arbre de A à ces trois feuilles symbolisent les branches correspondantes de l'arbre de A . Dans la déduction suivant S77, les quantificateurs $\forall x_1, \forall x_2, \forall x_3$ sont enlevés et les occurrences libres de x_1, x_2, x_3 dans A sont remplacées simultanément par trois termes T_1, T_2, T_3 . Les trois arbres de ces termes se substituent donc aux nœuds des occurrences libres de x_1, x_2, x_3 dans A . La proposition obtenue est celle qu'on désigne par $(T_1|x_1, T_2|x_2, T_3|x_3)A$. Il est supposé que les termes T_1, T_2, T_3 sont librement substituables à x_1, x_2, x_3 (respectivement) dans A . Cela signifie que sur la branche de l'arbre de A désignée par b_i , représentant les branches de toutes les occurrences libres de x_i dans A , il n'y a pas de nœud portant un quantificateur $\forall y$ ou $\exists y$ dont la variable y figure librement dans le terme T_i .

La même forme de représentation est utilisée dans la figure 10.14 pour la *démonstration* de cette déduction. La prémisse $\forall x_1 \forall x_2 \forall x_3 A$ de la déduction n'est pas reproduite dans cette figure pour des raisons de place. On part directement d'une proposition

$$\forall y_1 \forall y_2 \forall y_3 \underbrace{(y_1|x_1, y_2|x_2, y_3|x_3)}_{A'} A$$

où y_1, y_2, y_3 sont trois variables choisies selon les critères suivants: elles sont

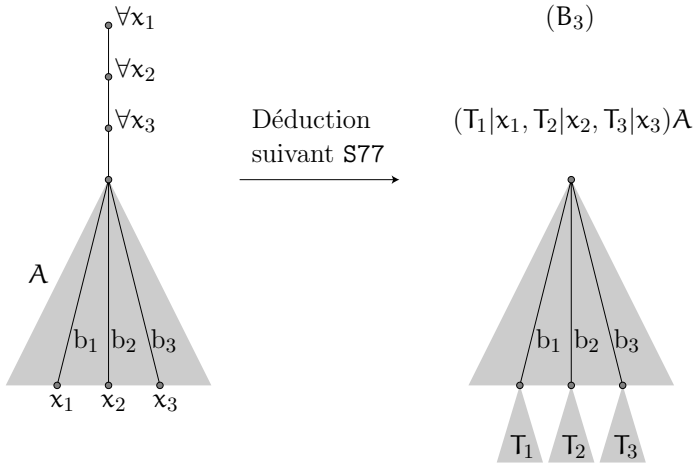


Figure 10.13 Dédution par particularisations simultanées.

distinctes et ne figurent d'aucune manière dans la proposition $\forall x_1 \forall x_2 \forall x_3 A$, ni dans les termes T_1, T_2, T_3 . Cela garantit évidemment que les propositions $\forall x_1 \forall x_2 \forall x_3 A$, $\forall y_1 \forall y_2 \forall y_3 A'$ ont le même squelette et sont donc équivalentes. Cela garantit aussi que la proposition $(T_3|y_3)((T_2|y_2)((T_1|y_1)A'))$ est identique à $(T_1|y_1, T_2|y_2, T_3|y_3)A'$, d'après la propriété que nous avons vue plus haut, et que cette dernière est identique à $(T_1|x_1, T_2|x_2, T_3|x_3)A$. On peut appliquer trois fois S61 comme indiqué dans la figure, grâce au fait que y_1, y_2, y_3 ne figurent pas dans les termes T_i et que sur la branche b_i il n'y a pas de quantificateur dont la variable figure librement dans T_i . La proposition B_3 est la même que dans la figure 10.13.

Le lecteur aura remarqué que l'on pourrait choisir y_1, y_2, y_3 selon des critères moins stricts. L'essentiel est que les propositions $\forall x_1 \forall x_2 \forall x_3 A$,

$$\forall y_1 \forall y_2 \forall y_3 \underbrace{(y_1|x_1, y_2|x_2, y_3|x_3)A}_{A'}$$

aient le même squelette et que: y_2, y_3 ne figurent pas librement dans T_1 ; y_3 ne figure pas librement dans T_2 . Le choix de variables y_1, y_2, y_3 qui ne figurent d'aucune manière (ni librement ni dans un quantificateur) dans la proposition $\forall x_1 \forall x_2 \forall x_3 A$ est simplement une façon particulièrement évidente de satisfaire à la première condition. C'est ce que nous avons fait dans l'exemple qui précède en remplaçant x et y par x' et y' . Nous aurions pu nous contenter de changer y en y' , comme le lecteur peut le vérifier.

Introductions simultanées de quantificateurs existentiels. Le schéma de déduction suivant dérive de S77 d'après la démonstration par réduction à l'absurde donnée ici. On suppose que $T_1, \dots, T_n$ sont des termes librement substituables à $x_1, \dots, x_n$ (respectivement) dans A . Alors il est clair que ces termes vérifient la même condition vis-à-vis de la proposition $\text{non}A$, ce qui

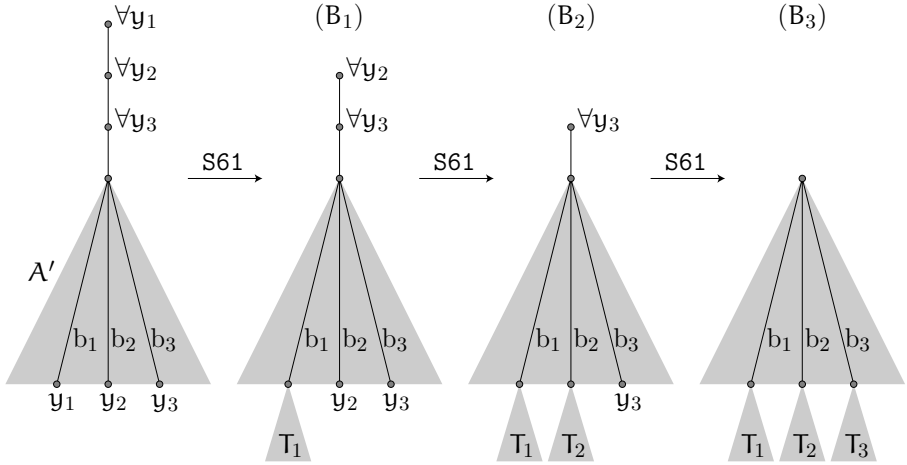


Figure 10.14 Trois particularisations successives, possibles si T_i ($i = 1, 2, 3$) est librement substituable à y_i dans A' et si y_1, y_2, y_3 ne figurent pas librement dans T_1, T_2, T_3 .

permet la déduction de la ligne 4°. La proposition de la ligne 5° est identique à celle de la ligne 4°, d'après les règles de substitution que nous avons vues dans ce paragraphe. Il s'agit ici de la règle (c), pour le connecteur logique de négation.

S78.

$$\frac{\Gamma \vdash (T_1|x_1, \dots, T_n|x_n)A}{\Gamma \vdash \exists x_1 \dots \exists x_n A}$$

Si $x_1, \dots, x_n$ sont des variables distinctes et si T_i ($i = 1, \dots, n$) est librement substituable à x_i dans A .

Γ		
1°	$(T_1 x_1, \dots, T_n x_n)A$	<i>base</i>
2°	$\text{non}\exists x_1 \dots \exists x_n A$	hyp
3°	$\forall x_1 \dots \forall x_n \text{non} A$	2°, S71
4°	$(T_1 x_1, \dots, T_n x_n)\text{non} A$	3°, S77
5°	$\text{non}(T_1 x_1, \dots, T_n x_n)A$	identique à 4°
6°	$\perp$	1°, 5°, $\perp$ -intro
7°	$\exists x_1 \dots \exists x_n A$	6°, S12.

10.6 Exemplifications

Considérons encore une proposition A , des variables distinctes $x_1, \dots, x_n$ et des termes $T_1, \dots, T_n$ que nous supposons librement substituables à $x_1, \dots, x_n$ (respectivement) dans A . Si A est vraie dans un contexte, et si les variables $x_1, \dots, x_n$ ne figurent librement dans aucune des hypothèses de celui-

ci, alors $\forall x_1 \cdots \forall x_n A$ est vraie dans ce contexte (S62), donc la proposition $(T_1|x_1, \dots, T_n|x_n)A$ est vraie aussi dans ce contexte d'après S77. C'est ce qu'énonce le schéma de déduction suivant—dont nous venons de décrire la démonstration.

S79 Exemplification

$$\frac{\Gamma \vdash A}{\Gamma \vdash (T_1|x_1, \dots, T_n|x_n)A}$$

Si $x_1, \dots, x_n$ sont des variables distinctes qui ne figurent pas librement dans Γ , et si T_i est librement substituable à x_i dans A (pour $i = 1, \dots, n$).

Nous dirons qu'un jugement $\Gamma \vdash (T_1|x_1, \dots, T_n|x_n)A$, déduit ainsi d'un jugement $\Gamma \vdash A$ par des substitutions simultanées de termes à des variables distinctes qui ne figurent pas librement dans les hypothèses de ces jugements est déduit par *exemplification* (en anglais *instantiation*) de la proposition A sous les hypothèses Γ . Ce genre de déduction est l'un des plus fréquents dans le raisonnement mathématique. Souvent en outre ces déductions sont accompagnées de changements de variables liées comme nous allons le voir.

Exemple 1. La déduction suivante de $L_{\text{préd}}$ est une déduction par exemplification de la proposition $(x \subset y \text{ et } y \subset z) \Rightarrow x \subset z$ sous la liste d'hypothèses vide. Les termes désignés par T_1, T_2, T_3 sont respectivement : $x \cap Y, x, A \times B$. Ces termes sont trivialement librement substituables à x, y, z dans la proposition A considérée, puisque celle-ci est sans quantificateurs.

$$\frac{\overbrace{\vdash (x \subset y \text{ et } y \subset z) \Rightarrow x \subset z}^A}{\underbrace{\vdash (y \cap X \subset x \text{ et } x \subset A \times B) \Rightarrow y \cap X \subset A \times B}_{(T_1|x, T_2|y, T_3|z)A}}$$

Nous verrons (§10.7) que la prémisse de cette déduction est un théorème de $\mathcal{T}_{\text{ENS}}$, appelé théorème de *transitivité de l'inclusion*. La conclusion de la déduction est donc aussi un théorème de cette théorie, et l'on peut en écrire une infinité de la même forme en choisissant arbitrairement trois termes T_1, T_2, T_3 . L'exemplification est une déduction si fréquente qu'on s'abstient en général de la mentionner. Ainsi, dans une démonstration de $\mathcal{T}_{\text{ENS}}$, lorsqu'une ligne est déduite du théorème de transitivité de l'inclusion par exemplification, on se contente de mentionner ce théorème comme justification, et l'on omet de mentionner S79.

Exemple 2. Une exemplification doit être précédée souvent d'un changement de variables liées. Considérons par exemple le théorème de $\mathcal{T}_{\text{ENS}}$:

$$\vdash x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y). \quad (1)$$

Nous ne pouvons pas l'exemplifier avec les termes T_1, T_2 suivants : $a, a \cup b$, car ces termes ne sont pas librement substituables à x et y dans (1). Mais nous pouvons faire un changement de variable liée dans (1) permettant ensuite

l'exemplification. Nous pouvons combiner par exemple les deux déductions suivantes :

$$\frac{\frac{\frac{\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)}{\vdash x \subset y \Leftrightarrow \forall z(z \in x \Rightarrow z \in y)} \text{ Chgmt var liée}}{\vdash a \subset a \cup b \Leftrightarrow \forall z(z \in a \Rightarrow z \in a \cup b)} \text{ S79} \quad (2)$$

Si nous formulons le premier et le troisième de ces jugements en français sans utiliser de variable liée, nous obtenons une déduction qui apparaît comme une simple exemplification :

$$\frac{\begin{array}{l} \vdash x \text{ est inclus dans } y \text{ si et seulement si} \\ \text{tout élément de } x \text{ est élément de } y \end{array}}{\begin{array}{l} \vdash a \text{ est inclus dans } a \cup b \text{ si et seulement si} \\ \text{tout élément de } a \text{ est élément de } a \cup b. \end{array}}$$

Cette combinaison d'une déduction par changement de variables liées et d'une déduction par exemplification est aussi tellement fréquente que l'on s'abstient de la mentionner explicitement. Lorsqu'on déduit

$$\vdash a \subset a \cup b \Leftrightarrow \forall z(z \in a \Rightarrow z \in a \cup b)$$

de (1), on se contente en général de mentionner (1) et l'on omet aussi bien la mention de S79 que celle du changement de variable liée, car ces déductions sont évidentes.

Exemple 3. L'exemple 2 n'est pas le cas le plus général de la combinaison d'une exemplification avec des changements de variables liées. Dans le cas général on effectue un changement de variables liées avant l'exemplification et un deuxième changement de variables liées après celle-ci. Un exemple est fourni par la déduction

$$\frac{\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)}{\vdash \{a, b\} \subset X \Leftrightarrow \forall x(x \in \{a, b\} \Rightarrow x \in X)}. \quad (3)$$

Celle-ci est perçue comme une simple exemplification lorsque, au lieu des deux propositions de ces jugements, on considère leurs squelettes :

$$\begin{array}{c} x \subset y \Leftrightarrow \forall \square (\square \in x \Rightarrow \square \in y) \\ \\ \{a, b\} \subset X \Leftrightarrow \forall \square (\square \in \{a, b\} \Rightarrow \square \in X). \end{array}$$

Mais il ne suffit pas d'un seul changement de variable liée et d'une exemplification pour passer de la prémisse de (3) à la conclusion. Il faut effectuer un premier changement de variable liée permettant la substitution des termes $\{a, b\}$ et X aux variables x et y . Mais à la place de la variable liée a on ne peut pas mettre la variable x puisque celle-ci figure librement dans la proposition $\forall a(a \in x \Rightarrow a \in y)$. On choisit donc une autre variable, z par exemple, et l'on

effectue les trois déductions suivantes :

$$\frac{\frac{\frac{\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)}{\vdash x \subset y \Leftrightarrow \forall z(z \in x \Rightarrow z \in y)} \text{ Chgmt var liée}}{\vdash \{a, b\} \subset X \Leftrightarrow \forall z(z \in \{a, b\} \Rightarrow z \in X)} \text{ Exemplification (S79)}}{\vdash \{a, b\} \subset X \Leftrightarrow \forall x(x \in \{a, b\} \Rightarrow x \in X)} \text{ Chgmt var liée}$$

Cette combinaison d'une exemplification et de deux changements de variables liées est assez fréquente, et l'on omet généralement de la mentionner. Lorsqu'on déduit ainsi dans $\mathcal{T}_{\text{ENS}}$ le théorème

$$\vdash \{a, b\} \subset X \Leftrightarrow \forall x(x \in \{a, b\} \Rightarrow x \in X)$$

du théorème (1), on ne mentionne que celui-ci, en omettant généralement de mentionner l'emploi de S79 et des changements de variables liées.

Schéma général de déduction par exemplification et changements de variables liées. L'exemple précédent, et notamment l'observation des squelettes des deux propositions de la déduction (3), suggère le schéma de déduction général suivant :

$$\frac{\Gamma \vdash A}{\Gamma \vdash B} \quad \begin{array}{l} \text{Si le squelette de B peut s'obtenir en remplaçant dans} \\ \text{celui de A des variables distinctes } x_1, \dots, x_n \text{ qui ne figu-} \\ \text{rent pas librement dans } \Gamma \text{ par des termes } T_1, \dots, T_n. \end{array}$$

Une démonstration d'une telle déduction peut toujours se construire comme nous l'avons fait dans l'exemple précédent. Supposons en effet que A et B soient deux propositions satisfaisant à la condition de ce schéma. On peut toujours, par des changements de variables liées dans A, construire une proposition A' de même squelette que A dans laquelle les termes T_i soient librement substituables aux variables x_i . Or une autre manière d'exprimer que les termes T_i sont librement substituables aux variables x_i dans A' est de dire que le squelette de la proposition

$$(T_1|x_1, \dots, T_n|x_n)A'$$

est identique à celui que l'on obtient en prenant le squelette de A' et en y remplaçant les variables x_i par les termes T_i . Comme le squelette de A' est le même que celui de A, celui de $(T_1|x_1, \dots, T_n|x_n)A'$ est le même que celui de B. On peut donc faire les trois déductions :

$$\frac{\frac{\frac{\Gamma \vdash A}{\Gamma \vdash A'} \text{ Chgmt var liées}}{\Gamma \vdash (T_1|x_1, \dots, T_n|x_n)A'} \text{ S79}}{\Gamma \vdash B} \text{ Chgmt var liées}$$

10.7 Application à la théorie des ensembles

La théorie des ensembles est une théorie logique (§5.3), c'est-à-dire une théorie admettant toutes les déductions de la logique des prédicats avec égalité. Elle admet en outre des déductions spécifiques que nous verrons plus loin (Chap. 13–14). Les schémas de déduction de la logique des prédicats permettent déjà de démontrer une partie du répertoire standard de théorèmes élémentaires de cette théorie. C'est le but de ce paragraphe. Pour la suite de ce répertoire, nous aurons besoin des schémas de déduction de la logique des prédicats avec égalité (Chap. 11), et plus loin, des schémas de déduction spécifiques de $\mathcal{T}_{\text{ENS}}$.

Nous appelons *théorie ensembliste* toute théorie $\mathcal{T}$ qui est une extension de la théorie $\mathcal{T}_{\text{ENS}}$ (§5.7) et *contexte ensembliste* un contexte logique (théorie + hypothèses) dont la théorie est une théorie ensembliste. Dans les paragraphes spécialement consacrés à la théorie des ensembles, lorsque nous parlons d'un « théorème » tout court, sans mentionner de théorie, il s'agit toujours d'un théorème de $\mathcal{T}_{\text{ENS}}$.

Tous les théorèmes et schémas de déduction de $\mathcal{T}_{\text{ENS}}$ seront numérotés dorénavant **E1**, **E2**, etc. La plupart de leurs démonstrations sont rassemblées dans l'annexe du livre (Chap. 16). Un certain nombre de démonstrations, faisant l'objet de discussions particulières, sont données dans le texte. L'exercice principal pour le lecteur est de démontrer lui-même le plus grand nombre possible de ces théorèmes et de ne consulter les démonstrations du livre que pour vérifier son travail. La plupart d'entre elles ne sont données que comme solution de cet exercice.

Nous commençons par numéroter nos deux premiers axiomes de $\mathcal{T}_{\text{ENS}}$ (§8.2). Nous écrivons le premier sans les quantificateurs universels initiaux $\forall x \forall y$, puisque ceux-ci peuvent être introduits suivant **S62**.

E1. *Axiome de l'inclusion:* $\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)$.

E2. *Axiome de l'ensemble vide:* $\vdash \forall x(x \notin \emptyset)$.

Un axiome d'une théorie $\mathcal{T}$ est une *proposition* (§5.3) et si **A** est un axiome de $\mathcal{T}$, le jugement $\vdash \mathbf{A}$, et plus généralement tout jugement $\Gamma \vdash \mathbf{A}$, est un *théorème* de $\mathcal{T}$ (§5.5, Règle 1). Ce que nous désignons par les numéros **E1** et **E2**, ce sont en fait les théorèmes $\vdash \mathbf{A}$ correspondant à ces deux axiomes (cf. §5.5, Axiomes en tant que théorèmes). Il en ira de même pour les autres axiomes de $\mathcal{T}_{\text{ENS}}$ ainsi numérotés.

Schémas de déduction sur l'inclusion. La proposition $x \subset y$ est appelée une *relation d'inclusion* entre l'ensemble x et l'ensemble y . C'est une relation très importante entre ensembles et nous allons formuler quelques schémas de déduction constamment utilisés relatifs à cette relation.

E3.

Si T_1 et T_2 sont des termes et z est une variable qui ne figure librement ni dans T_1 , ni dans T_2 .

$$\Gamma \vdash T_1 \subset T_2 \Leftrightarrow \forall z(z \in T_1 \Rightarrow z \in T_2)$$

Par exemple, d'après ce schéma, le jugement suivant est un théorème de $\mathcal{T}_{\text{ENS}}$:

$$\vdash \{a, b\} \subset X \Leftrightarrow \forall x(x \in \{a, b\} \Rightarrow x \in X). \quad (1)$$

Nous avons d'ailleurs montré précédemment (§10.6) comment il se déduit de E1 par exemplification et changements de variable liée. Dans cet exemple, les termes T_1 et T_2 du schéma sont $\{a, b\}$ et X , et la variable z est x . Elle ne figure librement ni dans T_1 , ni dans T_2 .

La démonstration d'une déduction de la forme E3 procède par changements de variable liée et exemplification (§10.6). Étant donné deux termes T_1 et T_2 , on choisit une variable auxiliaire z' qui ne figure librement ni dans l'un ni dans l'autre, et qui soit *distincte des variables x et y* , celles qui figurent dans l'axiome de l'inclusion. On peut alors écrire la démonstration

$$\begin{array}{lcl} \text{nil} & & \\ 1^\circ & \left| \begin{array}{l} x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y) \\ x \subset y \Leftrightarrow \forall z'(z' \in x \Rightarrow z' \in y) \\ T_1 \subset T_2 \Leftrightarrow \forall z'(z' \in T_1 \Rightarrow z' \in T_2) \\ \Leftrightarrow \forall z(z \in T_1 \Rightarrow z \in T_2) \end{array} \right. & \begin{array}{l} \text{E1} \\ 1^\circ, \text{S76} \\ 2^\circ, \text{S79} \\ \text{S76.} \end{array} \end{array} \quad (2)$$

Le schéma de déduction E3 est ainsi démontré dans le cas particulier d'une liste d'hypothèses Γ vide. Le cas général s'en déduit par a fortiori (§9.5). L'exemplification de la ligne 3° ne peut se faire que sous une liste d'hypothèses dans laquelle x et y ne figurent pas librement (S79), condition satisfaite par la liste vide.

La démonstration (2) est une démonstration de $\mathcal{T}_{\text{ENS}}$, de base vide. Ce n'est pas une démonstration de **Lpréd**, car E1 n'est pas un théorème de **Lpréd**. Elle devient une démonstration de **Lpréd** si nous prenons la première ligne comme *base*. Mais alors la déduction de **Lpréd** ainsi démontrée n'est pas la déduction sans prémisses E3 (avec Γ vide), mais la déduction avec prémisses

$$\frac{\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)}{\vdash T_1 \subset T_2 \Leftrightarrow \forall z(z \in T_1 \Rightarrow z \in T_2)}.$$

Dans les raisonnements ensemblistes, la plupart des déductions qui utilisent indirectement l'axiome de l'inclusion sont de la forme décrite par les schémas E4 et E5 ci-dessous, lesquels dérivent immédiatement de E3. Le premier (E4) caractérise ce que l'on peut *déduire* d'une inclusion $T_1 \subset T_2$, étant donné un terme E quelconque. Le second (E5) décrit les manières standard de *démontrer* une proposition d'inclusion $T_1 \subset T_2$.

E4. Conséquences d'une inclusion

$$1. \quad \frac{\Gamma \vdash T_1 \subset T_2}{\Gamma \vdash E \in T_1 \Rightarrow E \in T_2}.$$

$$2. \frac{\Gamma \vdash T_1 \subset T_2 \quad \Gamma \vdash E \in T_1}{\Gamma \vdash E \in T_2.} \qquad 3. \frac{\Gamma \vdash T_1 \subset T_2 \quad \Gamma \vdash E \notin T_2}{\Gamma \vdash E \notin T_1.}$$

Démonstration. De $\Gamma \vdash T_1 \subset T_2$, en prenant une variable z qui ne figure librement ni dans T_1 ni dans T_2 , on déduit $\Gamma \vdash \forall z(z \in T_1 \Rightarrow z \in T_2)$ d'après E3, ce dont on déduit $\Gamma \vdash E \in T_1 \Rightarrow E \in T_2$ en particulierisant avec E (S61). Le deuxième et le troisième schéma dérivent du premier suivant la logique propositionnelle.

E5. *Démonstration d'une inclusion*

$$\left. \begin{array}{l} 1. \frac{\Gamma \vdash z \in T_1 \Rightarrow z \in T_2}{\Gamma \vdash T_1 \subset T_2} \\ 2. \frac{\Gamma; z \in T_1 \vdash z \in T_2}{\Gamma \vdash T_1 \subset T_2} \end{array} \right\} \begin{array}{l} \text{Si } T_1 \text{ et } T_2 \text{ sont des termes} \\ \text{et } z \text{ est une variable qui ne} \\ \text{figure librement ni dans } \Gamma, \\ \text{ni dans } T_1, \text{ ni dans } T_2. \end{array}$$

Démonstration. En supposant que z vérifie ces conditions, on déduit de $\Gamma \vdash z \in T_1 \Rightarrow z \in T_2$, par généralisation (S62), $\Gamma \vdash \forall z(z \in T_1 \Rightarrow z \in T_2)$, d'où $\Gamma \vdash T_1 \subset T_2$ d'après E3. Le deuxième schéma dérive du premier suivant la logique propositionnelle ($\Rightarrow$ -intro).

Méthode de démonstration. Pour démontrer une proposition d'inclusion $T_1 \subset T_2$ dans un contexte ensembliste, on procède d'après E5 comme suit: on choisit une variable z qui ne figure librement ni dans T_1 , ni dans T_2 , ni dans les hypothèses du contexte, puis on démontre $z \in T_1 \Rightarrow z \in T_2$, ou l'on démontre $z \in T_2$ sous l'hypothèse supplémentaire $z \in T_1$.

$$\begin{array}{c} \Gamma \\ \vdots \\ n^o \left| \begin{array}{l} z \in T_1 \Rightarrow z \in T_2 \\ T_1 \subset T_2 \end{array} \right. \quad n^o, \text{ E5.1.} \end{array} \qquad \begin{array}{c} \Gamma \\ \vdots \\ n^o \left| \begin{array}{l} \boxed{z \in T_1 \quad \text{hyp}} \\ \vdots \\ z \in T_2 \end{array} \right. \quad n^o, \text{ E5.2.} \end{array}$$

Comme illustration de cette méthode, nous allons démontrer les deux théorèmes qui suivent.

E6. $\vdash x \subset x$.

E7. $\vdash \emptyset \subset x$.

Les deux démonstrations suivantes de E6 utilisent les deux formes de E5.

$$\begin{array}{c} \text{nil} \\ 1^o \left| \begin{array}{l} a \in x \Rightarrow a \in x \\ x \subset x \end{array} \right. \quad \begin{array}{l} \text{S19} \\ 1^o, \text{ E5.1.} \end{array} \end{array} \qquad \begin{array}{c} \text{nil} \\ 1^o \left| \begin{array}{l} \boxed{a \in x \quad \text{hyp}} \\ x \subset x \end{array} \right. \quad \begin{array}{l} 1^o, \text{ E5.2.} \end{array} \end{array}$$

Pour E7 on a la démonstration :

1 ^o	$\forall x(x \notin \emptyset)$	E2
2 ^o	$a \notin \emptyset$	1 ^o , S61
3 ^o	$a \in \emptyset \Rightarrow a \in x$	2 ^o , S22
4 ^o	$\emptyset \subset x$	3 ^o , E5.1.

Ex « $x \in \emptyset$ » quodlibet. La proposition $a \in \emptyset \Rightarrow a \in x$ de la troisième ligne de la démonstration qui précède est vraie d'après S22 parce que $a \in \emptyset$ est *fausse* (ligne 2^o). Cet exemple se généralise par le schéma de déduction sans prémisses suivant de $\mathcal{T}_{\text{ENS}}$, où $R(x)$ désigne une proposition quelconque.

E8.

$\Gamma \vdash \forall x(x \in \emptyset \Rightarrow R(x))$		nil
1 ^o	$x \notin \emptyset$	E2
2 ^o	$x \in \emptyset \Rightarrow R(x)$	1 ^o , S22
3 ^o	$\forall x(x \in \emptyset \Rightarrow R(x))$	2 ^o , S62.

Il nous arrivera parfois de désigner une proposition indéterminée par une lettre $A, B, C, R, \dots$ à côté de laquelle une variable est notée entre parenthèses, comme ici $R(x)$. Cette « décoration » de la lettre R avec (x) n'a aucune signification technique particulière. Elle ne signifie même pas que la variable x figure librement dans R . Elle a néanmoins pour but de suggérer que l'on s'intéresse particulièrement au cas où la variable x figure librement dans R et que R est interprétée comme l'énoncé d'une propriété de l'objet désigné par x . Considéré de cette manière, le schéma de déduction E8 peut s'énoncer ainsi : un élément de l'ensemble vide possède n'importe quelle propriété. Cette affirmation est vraie parce qu'il n'y a pas d'élément de l'ensemble vide !

Exemple. Le schéma de déduction E8 est appliqué en mathématique plus souvent qu'on ne pourrait le croire au premier abord. Cela provient du fait que l'on caractérise souvent des ensembles S par le fait que tous leurs éléments ont une certaine propriété. Considérons par exemple l'ensemble $\mathbb{N}$ des entiers naturels, et soit $b \in \mathbb{N}$. On dit qu'un sous-ensemble S de $\mathbb{N}$ est *majoré* par b si l'on a

$$\forall x(x \in S \Rightarrow x \leq b). \quad (3)$$

La proposition (3) caractérise les sous-ensembles S de $\mathbb{N}$ qui sont dits *majorés par b* . D'après E8, l'ensemble $\emptyset$ est majoré par b . On peut d'ailleurs remplacer b par 0 et affirmer que l'ensemble vide est majoré par 0 dans $\mathbb{N}$, ou encore que 0 est la borne supérieure de l'ensemble $\emptyset$ dans $\mathbb{N}$.

Transitivité de l'inclusion. Le théorème E9 ci-dessous de $\mathcal{T}_{\text{ENS}}$ est appelé le théorème de transitivité de l'inclusion. De façon générale on dit qu'une relation entre deux objets x et y notée $x R y$ au moyen d'un symbole relationnel binaire R est transitive, dans un contexte donné, si la proposition

$$\forall x \forall y \forall z ((x R y \text{ et } y R z) \Rightarrow x R z) \quad (4)$$

est vraie dans ce contexte. Le lecteur connaît les exemples de relations transitives que sont les relations d'égalité ($x = y$) et d'inégalité ($x \leq y$) entre nombres réels. La transitivité de la relation d'inclusion dans un contexte ensembliste quelconque se déduit du théorème E9 par généralisation et a fortiori. Lorsqu'on dit qu'une relation $x R y$ est *non transitive*, dans un contexte donné, on entend par là que la *négation* de (4) est vraie dans ce contexte, donc, d'après S71 et S56, que la proposition

$$\exists x \exists y \exists z ((x R y \text{ et } y R z) \text{ et } x \not R z)$$

est vraie dans ce contexte. Nous verrons plus loin que la relation d'appartenance $x \in y$ est non transitive dans un contexte ensembliste.

E9. $\vdash (x \subset y \text{ et } y \subset z) \Rightarrow x \subset z$.

	nil	
1 ^o	$x \subset y \text{ et } y \subset z$	hyp
2 ^o	$a \in x$	hyp
3 ^o	$x \subset y$	1 ^o , et-élim
4 ^o	$a \in y$	2 ^o , 3 ^o , E4
5 ^o	$y \subset z$	1 ^o , et-élim
6 ^o	$a \in z$	4 ^o , 5 ^o , E4
7 ^o	$x \subset z$	6 ^o , E5
8 ^o	$(x \subset y \text{ et } y \subset z) \Rightarrow x \subset z$	7 ^o , $\Rightarrow$ -intro.

Exercice. Démontrer le théorème

$$\vdash x \subset y \Leftrightarrow \forall z (z \subset x \Rightarrow z \subset y). \quad (5)$$

L'axiome d'extensionnalité. Nous introduisons ici un nouvel axiome de $\mathcal{T}_{\text{ENS}}$ qui *définit* le sens d'une relation d'égalité $x = y$ entre deux ensembles comme étant celui de la conjonction des inclusions $x \subset y$ et $y \subset x$.

E10. *Axiome d'extensionnalité:* $\vdash x = y \Leftrightarrow (x \subset y \text{ et } y \subset x)$.

Cet axiome ne permet pas de démontrer toutes les propriétés de l'égalité d'ensembles, mais seulement certaines d'entre elles. Deux ensembles x et y sont égaux, d'après cet axiome, si et seulement s'ils ont les *mêmes éléments*, puisque $x \subset y$ signifie que tout élément de x est élément de y , et $y \subset x$ signifie que tout élément de y est élément de x . Cette signification de l'égalité d'ensembles est exprimée formellement d'ailleurs par le théorème E16 ci-après.

Mais l'égalité $x = y$ de deux objets, que ce soient des ensembles ou des objets d'une autre espèce, est interprétée de façon générale comme le fait que x et y sont un seul et même objet, et par suite, qu'ils ont exactement les mêmes propriétés. Ce sont les schémas de déduction de la logique des prédicats avec égalité qui permettent de démontrer pour y toutes les propriétés de x et vice-versa. L'axiome d'extensionnalité à lui seul ne permet pas de le faire pour toutes les propriétés d'un ensemble x . Par exemple, si nous interprétons une égalité

d'ensembles $x = y$ comme étant l'identité de ces deux objets, et si nous savons que x est lui-même élément d'un ensemble u , nous en déduisons que y est élément de u , puisque x et y sont le même objet. Or c'est une déduction que nous ne pourrions pas faire avant de disposer de la logique des prédicats avec égalité (Chap. 11).

Le schéma de déduction suivant, dans lequel E , T_1 , T_2 désignent des termes quelconques, peut s'exprimer en disant que l'on peut remplacer un terme T_1 à droite du signe $\in$ par un terme T_2 égal à T_1 .

E11.

$$\frac{\Gamma \vdash T_1 = T_2}{\Gamma \vdash E \in T_1 \Leftrightarrow E \in T_2}.$$

Démonstration:

	Γ	
1°	$T_1 = T_2$	<i>base</i>
2°	$T_1 \subset T_2$ et $T_2 \subset T_1$	1°, E10 (exemplifié), S39
3°	$E \in T_1$	hyp
4°	$T_1 \subset T_2$	2°, et-élim
5°	$E \in T_2$	3°, 4°, E4.2
6°	$E \in T_2$	hyp
7°	$T_2 \subset T_1$	2°, et-élim
8°	$E \in T_1$	6°, 7°, E4.2
9°	$E \in T_1 \Leftrightarrow E \in T_2$	5°, 8°, S40.

L'axiome d'extensionnalité permet déjà de démontrer un certain nombre de théorèmes utiles. Il permet de démontrer notamment les théorèmes de réflexivité, de symétrie et de transitivité de l'égalité d'ensembles:

$$\begin{aligned} &\vdash x = x \\ &\vdash x = y \Rightarrow y = x \\ &\vdash (x = y \text{ et } y = z) \Rightarrow x = z. \end{aligned} \tag{6}$$

Cependant ce sont aussi des théorèmes de logique des prédicats avec égalité et nous ne les démontrerons pas ici. C'est par ailleurs un exercice très facile d'application de l'axiome E10, des théorèmes précédents et de la logique propositionnelle.

Théorèmes et schéma de déduction sur l'ensemble vide. Les théorèmes suivants sont démontrés dans la figure 10.15. Deux démonstrations sont données pour le premier théorème E13, illustrant les deux manières standard de démontrer une équivalence. Premièrement en démontrant deux implications et en appliquant S40. Deuxièmement — et souvent plus simplement — par une chaîne d'équivalences, au besoin en démontrant au préalable un ou deux théorèmes auxiliaires utilisés dans cette chaîne. Dans celle qui est donnée pour E13 on utilise le théorème $\vdash a \notin \emptyset$, déduit par exemplification de E2.

$$\text{E12. } \vdash x \subset \emptyset \Leftrightarrow x = \emptyset.$$

$$\begin{aligned} \text{E13. } & \vdash x = \emptyset \Leftrightarrow \forall a(a \notin x) \\ & \vdash x \neq \emptyset \Leftrightarrow \exists a(a \in x). \end{aligned}$$

$$\begin{aligned} \text{E14. } & \vdash (x \subset y \text{ et } y = \emptyset) \Rightarrow x = \emptyset \\ & \vdash (x \subset y \text{ et } x \neq \emptyset) \Rightarrow y \neq \emptyset. \end{aligned}$$

Informellement, la démonstration donnée dans la figure 10.15 pour le premier théorème E14 peut se lire comme ceci: supposons que $x \subset y$ et $y = \emptyset$; alors on a $y \subset \emptyset$, donc $x \subset \emptyset$ par transitivité de l'inclusion et finalement $x = \emptyset$ (E12).

Pour démontrer qu'une égalité de la forme $T = \emptyset$, où T est un terme, est vraie dans un contexte donné, on procède souvent de manière à pouvoir appliquer le schéma de déduction E15 ci-dessous. Pour cela, on choisit une variable z qui ne figure librement ni dans T ni dans les hypothèses du contexte, et l'on démontre une contradiction, donc $\perp$, sous l'hypothèse supplémentaire $z \in T$. Autrement dit, pour montrer qu'un ensemble est vide, il suffit de montrer qu'un objet *quelconque* n'est pas élément de cet ensemble.

$$\text{E15. } \frac{\Gamma; z \in T \vdash \perp}{\Gamma \vdash T = \emptyset} \quad \begin{array}{l} \text{Si } z \text{ ne figure librement ni} \\ \text{dans } \Gamma, \text{ ni dans } T. \end{array}$$

Ce schéma de déduction est aussi démontré dans la figure 10.15. La déduction

$$\frac{\Gamma \vdash x = \emptyset \Leftrightarrow \forall a(a \notin x)}{\Gamma \vdash T = \emptyset \Leftrightarrow \forall z(z \notin T)}$$

effectuée à la ligne 5° est une déduction par exemplification (S79) et changements éventuels de variables liées comme nous l'avons décrit au paragraphe 10.6. On doit d'abord, si nécessaire, changer la variable liée a en une variable z' distincte de x et ne figurant pas librement dans T de manière à pouvoir substituer T à x ; effectuer ensuite cette substitution; changer la variable liée z' en z si nécessaire (z' distincte de z).

Démonstrations d'égalités d'ensembles. L'axiome d'extensionnalité (E10) et celui de l'inclusion (E1) ont pour conséquence le théorème suivant. On l'exprime souvent en disant que deux ensembles sont égaux si et seulement s'ils ont les mêmes éléments.

$$\text{E16. } \vdash x = y \Leftrightarrow \forall a(a \in x \Leftrightarrow a \in y).$$

La démonstration est la chaîne d'équivalences suivante, dans un cadre sans hypothèses:

$$\begin{aligned} x = y & \Leftrightarrow (x \subset y \text{ et } y \subset x) & \text{E10} \\ & \Leftrightarrow \forall a(a \in x \Rightarrow a \in y) \text{ et } \forall a(a \in y \Rightarrow a \in x) & \text{E3} \\ & \Leftrightarrow \forall a((a \in x \Rightarrow a \in y) \text{ et } (a \in y \Rightarrow a \in x)) & \text{S74} \\ & \Leftrightarrow \forall a(a \in x \Leftrightarrow a \in y) & \text{S41.} \end{aligned}$$

	nil		
(E12)	1°	$x \subset \emptyset \Leftrightarrow (x \subset \emptyset \text{ et } \emptyset \subset x)$	E7, S60.1
	2°	$\Leftrightarrow x = \emptyset$	E10 (exemplifié).
(E13.1)	1°	$x = \emptyset$	hyp
	2°	$x \subset \emptyset$	1°, E12 (et S39)
	3°	$a \notin \emptyset$	E2 (exemplifié)
	4°	$a \notin x$	2°, 3°, E4
	5°	$\forall a(a \notin x)$	4°, S62
	6°	$\forall a(a \notin x)$	hyp
	7°	$a \notin x$	6°, S65
	8°	$a \in x \Rightarrow a \in \emptyset$	7°, S22
	9°	$x \subset \emptyset$	8°, E5
	10°	$x = \emptyset \Leftrightarrow \forall a(a \notin x)$	5°, 9°, S40.
(E13.1)		$x = \emptyset \Leftrightarrow x \subset \emptyset$	E12
		$\Leftrightarrow \forall a(a \in x \Rightarrow a \in \emptyset)$	E3
		$\Leftrightarrow \forall a(a \notin x \text{ ou } a \in \emptyset)$	S53
		$\Leftrightarrow \forall a(a \notin x)$	S60.2, E2 (exemplifié).
(E13.2)		$x \neq \emptyset \Leftrightarrow \text{non} \forall a(a \notin x)$	E13 (premier), S51
		$\Leftrightarrow \exists a(a \in x)$	S71.
(E14)	1°	$x \subset y \text{ et } y = \emptyset$	hyp
	2°	$x \subset y \text{ et } y \subset \emptyset$	1°, E12
	3°	$x \subset \emptyset$	2°, E, mod ponens
	4°	$x = \emptyset$	3°, E12
	5°	$(x \subset y \text{ et } y = \emptyset) \Rightarrow x = \emptyset$	$\Rightarrow$ -intro.
	1°	$x \subset y \text{ et } x \neq \emptyset$	hyp
	2°	$y = \emptyset$	hyp
	3°	$x \subset y \text{ et } y = \emptyset$	1°, 2°
	4°	$x = \emptyset$	3°, E14 (premier)
	5°	$\perp$	1°, 4°
	6°	$y \neq \emptyset$	5°, red abs J
	7°	$(x \subset y \text{ et } x \neq \emptyset) \Rightarrow y \neq \emptyset$	6°, $\Rightarrow$ -intro.
	Γ		
(E15)	1°	$z \in \mathbb{T}$	hyp
	2°	$\perp$	base
	3°	$z \notin \mathbb{T}$	2° red abs J
	4°	$\forall z(z \notin \mathbb{T})$	3°, S62
	5°	$\mathbb{T} = \emptyset \Leftrightarrow \forall z(z \notin \mathbb{T})$	E13 (Exempl + chgmt var liée)
	6°	$\mathbb{T} = \emptyset$	4°, 5°, S39.

Figure 10.15

Si T_1 et T_2 sont des termes quelconques, et si z est une variable qui ne figure librement ni dans l'un ni dans l'autre, alors on peut faire la déduction

$$\frac{\vdash x = y \Leftrightarrow \forall a(a \in x \Leftrightarrow a \in y)}{\vdash T_1 = T_2 \Leftrightarrow \forall z(z \in T_1 \Leftrightarrow z \in T_2)}$$

par exemplification et changements éventuels de variable liée. Cela permet de démontrer les déductions sans prémisses de la forme E17.1 ci-dessous (cf. démonstration de E3). La plupart des démonstrations d'égalité de deux ensembles se font suivant l'un des deux schémas de déduction E17.2-3, qui dérivent de façon évidente de E17.1 par la logique propositionnelle.

E17. Si T_1 et T_2 sont des termes et si z est une variable qui ne figure librement ni dans T_1 , ni dans T_2 :

$$1. \frac{}{\Gamma \vdash T_1 = T_2 \Leftrightarrow \forall z(z \in T_1 \Leftrightarrow z \in T_2)}$$

Si en outre z ne figure pas librement dans Γ :

$$2. \frac{\Gamma \vdash z \in T_1 \Leftrightarrow z \in T_2}{\Gamma \vdash T_1 = T_2} \quad 3. \frac{\Gamma; z \in T_1 \vdash z \in T_2 \quad \Gamma; z \in T_2 \vdash z \in T_1}{\Gamma \vdash T_1 = T_2}.$$

Pour démontrer une égalité $T_1 = T_2$ dans un contexte ensembliste suivant ces schémas, on choisit une variable z qui ne figure pas librement dans les termes T_1 , T_2 ni dans les hypothèses du contexte et l'on démontre

$$z \in T_1 \Leftrightarrow z \in T_2$$

par une chaîne d'équivalences, ou alors on démontre $z \in T_2$ sous l'hypothèse $z \in T_1$, ainsi que $z \in T_1$ sous l'hypothèse $z \in T_2$. Lorsqu'on y est parvenu, on s'arrête le plus souvent. La conclusion $T_1 = T_2$ tirée de ces résultats suivant E17.1 ou E17.2 est tacite. Nous représentons ici ces plans de démonstration standard d'une égalité d'ensembles. Les deux dernières lignes du premier et la dernière ligne du second sont généralement omises.

n^o	Γ $z \in T_1 \Leftrightarrow \dots$ $\Leftrightarrow \dots$ $\vdots$ $\Leftrightarrow \dots$ $\Leftrightarrow z \in T_2$ $z \in T_1 \Leftrightarrow z \in T_2$ $T_1 = T_2$	S48 n^o , E17.2.	k^o	Γ $z \in T_1$ hyp $\dots$ $z \in T_2$ $z \in T_2$ hyp $\dots$ $z \in T_1$ $T_1 = T_2$	n^o	k^o , n^o , E17.3.
-------	--	-----------------------	-------	---	-------	------------------------

Réunion et intersection de deux ensembles. Le symbole $\cup$, en théorie des ensembles, est un symbole fonctionnel binaire. L'ensemble désigné par $a \cup b$ est appelé la *réunion* des ensembles a et b . Cet ensemble est défini par l'axiome suivant:

E18. *Axiome de la réunion de deux ensembles*

$$\vdash x \in a \cup b \Leftrightarrow (x \in a \text{ ou } x \in b).$$

La démonstration des théorèmes E19–E27 est donnée dans l'annexe du livre (Chap. 16). Elle relève essentiellement de la logique propositionnelle. Ayant à démontrer des égalités ou des inclusions d'ensembles, on applique les méthodes de démonstration décrites dans ce paragraphe à propos des schémas de déduction E17 et E5.

E19. $\vdash a \cup a = a.$

E20. $\vdash a \cup b = b \cup a.$

E21. $\vdash a \cup (b \cup c) = (a \cup b) \cup c.$

E22. $\vdash a \cup \emptyset = a.$

E23. $\vdash a \subset a \cup b \quad \vdash b \subset a \cup b.$

E24. $\vdash (a \subset c \text{ et } b \subset c) \Leftrightarrow a \cup b \subset c.$

E25. $\vdash a \subset b \Leftrightarrow a \cup b = b.$

E26. $\vdash (a \subset a' \text{ et } b \subset b') \Rightarrow (a \cup b \subset a' \cup b').$

E27. $\vdash a \cup b = \emptyset \Leftrightarrow (a = \emptyset \text{ et } b = \emptyset).$

Le symbole $\cap$ est aussi un symbole fonctionnel binaire et l'ensemble désigné par $a \cap b$ est appelé *l'intersection* des ensembles a et b . Il est défini par l'axiome suivant :

E28. *Axiome de l'intersection de deux ensembles*

$$\vdash x \in a \cap b \Leftrightarrow (x \in a \text{ et } x \in b).$$

Les théorèmes E29–E36 sont symétriques des théorèmes E19–E26. Compte tenu des axiomes de la réunion et de l'intersection de deux ensembles, les théorèmes de distributivité E37, E38 sont une conséquence particulière des schémas de logique propositionnelle S52.11–12, avec lesquels on démontre les équivalences

$$\vdash x \in a \cap (b \cup c) \Leftrightarrow x \in (a \cap b) \cup (a \cap c)$$

$$\vdash x \in a \cup (b \cap c) \Leftrightarrow x \in (a \cup b) \cap (a \cup c)$$

Les démonstrations sont données dans l'annexe (Chap. 16) sauf pour E39.

E29. $\vdash a \cap a = a.$

E30. $\vdash a \cap b = b \cap a.$

E31. $\vdash a \cap (b \cap c) = (a \cap b) \cap c.$

E32. $\vdash a \cap \emptyset = \emptyset.$

E33. $\vdash a \cap b \subset a \quad \vdash a \cap b \subset b.$

E34. $\vdash (c \subset a \text{ et } c \subset b) \Leftrightarrow c \subset a \cap b.$

E35. $\vdash a \subset b \Leftrightarrow a \cap b = a.$

E36. $\vdash (a \subset a' \text{ et } b \subset b') \Rightarrow (a \cap b \subset a' \cap b').$

E37. $\vdash a \cap (b \cup c) = (a \cap b) \cup (a \cap c).$

E38. $\vdash a \cup (b \cap c) = (a \cup b) \cap (a \cup c).$

E39. $\vdash a \cap b = \emptyset \Leftrightarrow \forall x(x \in a \Rightarrow x \notin b).$

Ce dernier théorème se démontre par la chaîne d'équivalences suivante, dans un cadre sans hypothèses :

$$\begin{aligned}
 a \cap b = \emptyset &\Leftrightarrow \forall x (x \notin a \cap b) && \text{E13} \\
 &\Leftrightarrow \forall x (x \notin a \text{ ou } x \notin b) && \text{E28, De Morgan} \\
 &\Leftrightarrow \forall x (x \in a \Rightarrow x \notin b) && \text{S53.}
 \end{aligned}$$

Ensemble complémentaire d'un ensemble. Étant donné deux ensembles a et b , on désigne par $\complement ab$ l'ensemble dont les éléments sont les éléments de a qui n'appartiennent pas à b . Cet ensemble est appelé *l'ensemble complémentaire de b par rapport à a* . Le symbole $\complement$ est un symbole fonctionnel binaire. Pour améliorer la lisibilité, nous écrirons $\complement_a(b)$ au lieu de $\complement ab$. D'autres auteurs utilisent la notation $a \setminus b$, voire $a - b$. L'ensemble $\complement_a(b)$ est défini par l'axiome suivant :

E40. *Axiome de l'ensemble complémentaire*

$$\vdash x \in \complement_a(b) \Leftrightarrow (x \in a \text{ et } x \notin b).$$

La démonstration des théorèmes suivants est donnée dans l'annexe.

$$\text{E41. } \vdash \complement_a(b) \subset a.$$

$$\text{E42. } \vdash \complement_a(b) = a \Leftrightarrow a \cap b = \emptyset.$$

$$\text{E43. } \vdash \complement_a(b) = \emptyset \Leftrightarrow a \subset b.$$

$$\text{E44. } \vdash \complement_a(\emptyset) = a \quad \vdash \complement_a(a) = \emptyset.$$

$$\text{E45. } \vdash \complement_{\emptyset}(b) = \emptyset.$$

$$\text{E46. } \vdash \complement_X(a \cup b) = \complement_X(a) \cap \complement_X(b).$$

$$\text{E47. } \vdash \complement_X(a \cap b) = \complement_X(a) \cup \complement_X(b).$$

$$\text{E48. } \vdash a \subset b \Rightarrow \complement_X(b) \subset \complement_X(a).$$

$$a \subset X; b \subset X \vdash a \subset b \Leftrightarrow \complement_X(b) \subset \complement_X(a).$$

$$\text{E49. } \vdash \complement_X(\complement_X(a)) = X \cap a.$$

$$\text{E50. } \vdash a \subset X \Leftrightarrow \complement_X(\complement_X(a)) = a.$$

Le second des théorèmes E48 est notre premier exemple de théorème de $\mathcal{T}_{\text{ENS}}$ avec des hypothèses, à savoir des propositions à gauche du signe $\vdash$, les propositions $a \subset X$, $b \subset X$. La démonstration consiste naturellement à faire ces hypothèses, et à démontrer $a \subset b \Leftrightarrow \complement_X(b) \subset \complement_X(a)$ sous ces hypothèses. Pour se familiariser avec la correspondance entre démonstrations formelles et démonstrations usuelles, le lecteur est invité à comparer notre démonstration de ce théorème avec le texte suivant qu'on pourrait trouver dans un exposé traditionnel de théorie des ensembles :

Supposons que $a \subset X$ et $b \subset X$. En vertu du premier théorème E48, il nous suffit de démontrer que $\complement_X(b) \subset \complement_X(a) \Rightarrow a \subset b$. Supposons donc que $\complement_X(b) \subset \complement_X(a)$ et montrons que $a \subset b$. Soit $x \in a$. Nous devons montrer que $x \in b$. Raisonnons par l'absurde, en supposant que $x \notin b$. Comme $x \in a$ et $a \subset X$, nous avons $x \in X$, donc $x \in \complement_X(b)$. Comme $\complement_X(b) \subset \complement_X(a)$, il vient $x \in \complement_X(a)$, donc $x \notin a$, ce qui contredit notre hypothèse $x \in a$.

Utilité de la logique des prédicats avec égalité. Le théorème E50 est un exemple de théorème pour lequel la logique des prédicats avec égalité (Chap. 11) nous permettra d'écrire une démonstration plus simple et plus évidente que celle que nous pouvons donner au stade actuel. Au vu du théorème E49, on est naturellement tenté de démontrer E50 en écrivant, dans un cadre sans hypothèses, la chaîne d'équivalences suivante:

$$\begin{aligned}\mathbb{C}_X(\mathbb{C}_X(a)) = a &\Leftrightarrow X \cap a = a && \text{E49} \\ &\Leftrightarrow a \cap X = a && \text{E30} \\ &\Leftrightarrow a \subset X && \text{E35.}\end{aligned}$$

Puisque $\mathbb{C}_X(\mathbb{C}_X(a)) = X \cap a$ (E49), ces deux termes désignent le même ensemble, donc la proposition $\mathbb{C}_X(\mathbb{C}_X(a)) = a$ doit être équivalente à $X \cap a = a$. C'est une déduction correcte, mais elle n'est conforme à aucun des schémas de déduction que nous avons établis jusqu'ici. Ce remplacement d'un terme par un terme égal dans une proposition est précisément une déduction de logique des prédicats avec égalité. Pour le moment, rien ne nous y autorise. Nous sommes obligés d'appliquer systématiquement le schéma E11 comme suit:

$$\begin{aligned}\mathbb{C}_X(\mathbb{C}_X(a)) = a &\Leftrightarrow \forall x(x \in \mathbb{C}_X(\mathbb{C}_X(a)) \Leftrightarrow x \in a) && \text{E17.1} \\ &\Leftrightarrow \forall x(x \in X \cap a \Leftrightarrow x \in a) && \text{E49, E11} \\ &\Leftrightarrow \forall x(x \in a \cap X \Leftrightarrow x \in a) && \text{E30, E11} \\ &\Leftrightarrow a \cap X = a && \text{E17.1} \\ &\Leftrightarrow a \subset X && \text{E35.}\end{aligned}$$

Le schéma E11 permet aussi de remplacer un terme par un terme égal dans une proposition, mais seulement lorsque ce terme est à droite d'un symbole d'appartenance ($\in$).

10.8 Formes prénexes

Déplacements de quantificateurs. Les schémas S71 et S74 permettent certains déplacements de quantificateurs dans les propositions, opération que l'on effectue fréquemment dans les chaînes d'équivalences. Le but de ce paragraphe est de montrer que les schémas de la logique des prédicats permettent toujours de transformer une proposition en une proposition équivalente dans laquelle tous les quantificateurs se trouvent au début de la proposition. Nous établissons d'abord, dans ce but, quelques schémas dérivés supplémentaires.

S80. *Quantificateurs superflus*

$$\begin{array}{ccc} \hline \Gamma \vdash \forall x B \Leftrightarrow B & \Gamma \vdash \exists x B \Leftrightarrow B & \text{Si } x \text{ ne figure pas librement} \\ & & \text{dans } B. \end{array}$$

Lorsqu'une variable x ne figure pas librement dans une proposition B , on

peut écrire en effet les démonstrations suivantes :

nil			nil		
1°	$\forall xB \Rightarrow B$	S65			
2°	B	hyp		$\exists xB \Leftrightarrow \text{non}\forall x\text{non}B$	S71
3°	$\forall xB$	2°, S62		$\Leftrightarrow \text{nonnon}B$	S80 (premier)
4°	$B \Rightarrow \forall xB$	3°, $\Rightarrow$ -intro.		$\Leftrightarrow B.$	

Le premier des schémas S80 a déjà été « à moitié » présenté au paragraphe 9.5 (Exemple 5) et nous avons déjà signalé à cet endroit que l'on met rarement un quantificateur Qx ($\forall x$ ou $\exists x$) devant une proposition B dans laquelle la variable x ne figure pas librement. Les schémas S80 permettent de dire qu'un tel quantificateur est superflu. On peut l'enlever ou l'ajouter à volonté. En particulier, dans une proposition de la forme Q_1xQ_2xA où Q_1x et Q_2x sont deux quantificateurs avec la même variable x , le premier de ces quantificateurs est superflu au sens de S80.

S81. Localisation/délocalisation de quantificateurs

1.	$\frac{}{\Gamma \vdash \forall x(A \text{ ou } B) \Leftrightarrow (\forall xA \text{ ou } B)}$	} Si x ne figure pas librement dans B .
2.	$\frac{}{\Gamma \vdash \exists x(A \text{ et } B) \Leftrightarrow (\exists xA \text{ et } B)}$	
3.	$\frac{}{\Gamma \vdash \forall x(A \text{ et } B) \Leftrightarrow (\forall xA \text{ et } B)}$	
4.	$\frac{}{\Gamma \vdash \exists x(A \text{ ou } B) \Leftrightarrow (\exists xA \text{ ou } B)}$	

La démonstration des deux premiers de ces schémas est donnée ci-dessous. Le troisième et le quatrième dérivent de façon évidente de S74 et de S80 par substitutivité de l'équivalence.

nil			nil		
1°	$\forall x(A \text{ ou } B)$	hyp	1°	$\exists xA \text{ et } B$	hyp
2°	$\text{non}B$	hyp	2°	$\exists xA$	1°, et-élim
3°	$A \text{ ou } B$	1°, S65	3°	A	hyp
4°	A	2°, 3°, S18	4°	B	1°, et-élim
5°	$\forall xA$	4°, S62	5°	$A \text{ et } B$	3°, 4°
6°	$\text{non}B \Rightarrow \forall xA$	5°, $\Rightarrow$ -intro	6°	$\exists x(A \text{ et } B)$	5°, S66
7°	$\forall xA \text{ ou } B$	6°, S54	7°	$\exists x(A \text{ et } B)$	2°, 6°, S64
9°	$\forall xA \text{ ou } B$	hyp	8°	$\exists x(A \text{ et } B)$	hyp
10°	$\forall xA \text{ ou } \forall xB$	9°, S80	9°	$\exists xA \text{ et } \exists xB$	8°, S75
11°	$\forall x(A \text{ ou } B)$	10°, S75.	10°	$\exists xA \text{ et } B$	9°, S80.

On peut résumer S81 par le schéma de déduction :

$\frac{}{\Gamma \vdash Qx(A \text{ et/ou } B) \Leftrightarrow (QxA \text{ et/ou } B)}$	Si x ne figure pas librement dans B .
--	---

Qx désigne un quantificateur $\forall x$ ou $\exists x$. On applique ce schéma en prenant bien entendu le même connecteur de disjonction ou de conjonction des deux côtés de l'équivalence.

Formes prénexes. On appelle *proposition de forme prénex* une proposition de la forme $Q_1x_1 \cdots Q_nx_nA$, où les Q_ix_i sont des quantificateurs ($\forall x_i$ ou $\exists x_i$) et A est une proposition sans quantificateurs. C'est donc une proposition dans laquelle tous les quantificateurs se trouvent « en facteur » à gauche. Par exemple, la proposition $\forall x \exists y (y \neq x)$ est de forme prénex. Par contre, l'axiome de l'inclusion, $x \subset y \Leftrightarrow \forall a (a \in x \Rightarrow a \in y)$, ne l'est pas.

Une propriété générale de la théorie **Lpréd** est que toute proposition peut *se mettre sous forme prénex* dans cette théorie. On entend par là que pour toute proposition R on peut construire une proposition R_p de forme prénex telle que

$$\vdash R \Leftrightarrow R_p$$

soit un théorème de **Lpréd**. L'intérêt de cette transformation d'une proposition quelconque R est que pour démontrer l'équivalence de deux propositions données R et R' , il est souvent commode de montrer qu'elles peuvent se mettre sous la *même forme prénex*.

Les schémas de déduction qui permettent la mise sous forme prénex d'une proposition sont premièrement les schémas de logique propositionnelle **S53** et **S59**, qui permettent de transformer toute proposition en une proposition équivalente n'utilisant que les connecteurs logiques de négation, disjonction et conjonction. Deuxièmement, ce sont les schémas de déduction de logique des prédicats rappelés ci-dessous, à savoir **S71**, **S74**, **S81**, **S76**:

$$\begin{array}{c} \hline \Gamma \vdash \text{non} \forall x A \Leftrightarrow \exists x \text{non} A \\ \hline \Gamma \vdash \forall x (A \text{ et } B) \Leftrightarrow (\forall x A \text{ et } \forall x B) \quad \Gamma \vdash \text{non} \forall x A \Leftrightarrow \exists x \text{non} A. \\ \hline \Gamma \vdash \exists x (A \text{ ou } B) \Leftrightarrow (\exists x A \text{ ou } \exists x B). \\ \hline \Gamma \vdash Qx(A \text{ et/ou } B) \Leftrightarrow (QxA \text{ et/ou } B) \\ \text{Si } x \text{ ne figure pas librement dans } B. \\ \hline \Gamma \vdash \forall x A \Leftrightarrow \forall y(y|x)A \quad \Gamma \vdash \exists x A \Leftrightarrow \exists y(y|x)A \\ \text{si } y \text{ est réversiblement substituable à } x \text{ dans } A. \\ \hline \end{array}$$

La mise sous forme prénex d'une proposition R est une démonstration en forme de chaîne d'équivalence

$$\begin{array}{l} \text{nil} \\ \vdots \\ R \Leftrightarrow R_1 \\ \Leftrightarrow R_2 \\ \vdots \\ \Leftrightarrow R_n \end{array}$$

dans laquelle, à chaque ligne, un ou plusieurs quantificateurs sont déplacés vers la gauche par des déductions suivant l'un des schémas mentionnés ci-dessus. Bien entendu, il ne s'agit pas de conserver nécessairement les *mêmes* quantificateurs. Chaque application de S71 change un quantificateur en son dual. Des changements de variable liées sont nécessaires pour mettre sous forme prénexe des propositions de la forme

$$\exists xA \text{ et } B, \quad \forall xA \text{ ou } B \quad (1)$$

dans lesquelles x figure librement dans B , et que l'on ne peut donc pas transformer directement suivant S81. Dans ce cas il suffit de faire un changement de variable liée dans $\exists xA$, $\forall xA$ en prenant une variable y réversiblement substituable à x dans A et ne figurant pas librement dans B . Les propositions (1) sont alors équivalentes aux propositions

$$\exists y(y|x)A \text{ et } B, \quad \forall y(y|x)A \text{ ou } B \quad (2)$$

que l'on peut transformer en $\exists y((y|x)A \text{ et } B)$, $\forall y((y|x)A \text{ ou } B)$ suivant S81. Enfin, des propositions de la forme

$$\exists xA \text{ et } \exists xB, \quad \forall xA \text{ ou } \forall xB$$

dans lesquelles x figure librement et dans A , et dans B , peuvent se transformer d'abord suivant S81 en

$$\exists x(A \text{ et } \exists xB), \quad \forall x(A \text{ ou } \forall xB) \quad (3)$$

puisque x ne figure pas librement dans $\exists xB$, $\forall xB$, puis en

$$\exists x(A \text{ et } \exists y(y|x)B), \quad \forall x(A \text{ et } \forall y(y|x)B)$$

par changement de variable liée dans $\exists xB$, $\forall xB$ avec une variable y qui ne figure pas librement dans A , et enfin en

$$\exists x \exists y(A \text{ et } (y|x)B), \quad \forall x \forall y(A \text{ et } (y|x)B)$$

par une nouvelle application de S81 (et de la commutativité de l'équivalence).

Exemple 1. Dans la chaîne d'équivalences ci-dessous, il est montré que la proposition

$$x > 0 \Rightarrow \exists a(a > 0 \text{ et } \forall y(y < a \Rightarrow y < x))$$

est équivalente (dans **Lpréd**) à la proposition de forme prénexe

$$\exists a \forall y [x > 0 \Rightarrow (a > 0 \text{ et } (y < a \Rightarrow y < x))].$$

Une autre forme prénexe est déjà atteinte, d'ailleurs, à l'avant-dernière ligne de la démonstration.

nil

$$\left[\begin{array}{l} x > 0 \Rightarrow \exists a(a > 0 \text{ et } \forall y(y < a \Rightarrow y < x)) \\ \Leftrightarrow \text{non}(x > 0) \text{ ou } \exists a(a > 0 \text{ et } \forall y(y < a \Rightarrow y < x)) \\ \Leftrightarrow \exists a[\text{non}(x > 0) \text{ ou } (a > 0 \text{ et } \forall y(y < a \Rightarrow y < x))] \end{array} \right. \quad \begin{array}{l} \text{S53} \\ \text{S81.4} \end{array}$$

$\Leftrightarrow \exists a [\text{non}(x > 0) \text{ ou } \forall y (a > 0 \text{ et } (y < a \Rightarrow y < x))]$	S81.3
$\Leftrightarrow \exists a \forall y [\text{non}(x > 0) \text{ ou } (a > 0 \text{ et } (y < a \Rightarrow y < x))]$	S81.1
$\Leftrightarrow \exists a \forall y [x > 0 \Rightarrow (a > 0 \text{ et } (y < a \Rightarrow y < x))]$	S53.

Exemple 2. La proposition $\forall y(x \subset y) \Rightarrow \text{non} \exists y(y \in x)$, qui est vraie dans un contexte ensembliste, peut se mettre dans **Lpréd** sous la forme prénexe $\exists y \forall a(x \subset y \Rightarrow a \notin x)$, de la manière suivante :

nil	
$(\forall y(x \subset y) \Rightarrow \text{non} \exists y(y \in x))$	
$\Leftrightarrow \text{non} \forall y(x \subset y) \text{ ou } \text{non} \exists y(y \in x)$	S53
$\Leftrightarrow \exists y(x \not\subset y) \text{ ou } \forall y(y \notin x)$	S71
$\Leftrightarrow \exists y(x \not\subset y \text{ ou } \forall y(y \notin x))$	S81
$\Leftrightarrow \exists y(x \not\subset y \text{ ou } \forall a(a \notin x))$	S76
$\Leftrightarrow \exists y \forall a(x \not\subset y \text{ ou } a \notin x)$	S81
$\Leftrightarrow \exists y \forall a(x \subset y \Rightarrow a \notin x)$	S53.

Exercices. Mettre sous forme prénexe les propositions suivantes. Les symboles **N** et **P** (resp. **R** et **S**) sont des symboles relationnels unaires (resp. binaires) préfixés.

- (a) $\forall x(x \notin a) \text{ ou } \exists x(x \in a)$.
- (b) $\forall x(\exists y Rxy \Rightarrow \text{non} \exists y Sxy)$.
- (c) $\forall x[\mathbf{N}x \text{ ou } \text{non} \exists y(\mathbf{P}y \text{ et } Rxy)] \text{ ou } \text{non}[\mathbf{P}y \text{ et } \text{non} \forall x \mathbf{N}x]$.

10.9 Quantificateurs typés

Considérons la définition usuelle d'une suite convergente de nombres réels. On dit qu'une suite (x_n) de nombres réels converge vers un nombre réel a si pour tout nombre réel $\varepsilon > 0$ il existe un entier n tel que, pour tout entier $m \geq n$, on ait $|x_m - a| < \varepsilon$. Cette propriété d'une suite (x_n) par rapport à un nombre a ne peut pas se formaliser simplement par l'expression

$$\forall \varepsilon \exists n \forall m (|x_m - a| < \varepsilon),$$

car il est question dans cette proposition de nombres réels ε tels que $\varepsilon > 0$; de nombres n tels que $n \in \mathbb{N}$ et de nombres m tels que $m \in \mathbb{N}$ et $m \geq n$. Nous allons indiquer ces conditions imposées à ε , n , m en notant la proposition comme suit :

$$\left(\frac{\forall \varepsilon}{\varepsilon \in \mathbb{R} \text{ et } \varepsilon > 0} \right) \left(\frac{\exists n}{n \in \mathbb{N}} \right) \left(\frac{\forall m}{m \in \mathbb{N} \text{ et } m \geq n} \right) (|x_m - a| < \varepsilon). \quad (1)$$

Lire: pour tout ε [tel que $\varepsilon \in \mathbb{R}$ et $\varepsilon > 0$], il existe un n [tel que $n \in \mathbb{N}$], tel que, pour tout $m \dots$ Les « fractions » entre parenthèses dans (1) sont appelées des *quantificateurs typés*. Un quantificateur typé est un quantificateur usuel auquel est associé une proposition **R**, que l'on notera sous une barre de fraction ou

de l'autre manière introduite ci-dessous. Ces notations sont des *abréviations*, définies par la convention de notation suivante.

Notation. Si A et R sont des propositions et si x est une variable individuelle, les propositions

$$\forall x(R \Rightarrow A), \quad \exists x(R \text{ et } A)$$

sont abrégées respectivement par les expressions

$$\left(\frac{\forall x}{R}\right)A, \quad \left(\frac{\exists x}{R}\right)A$$

ou encore par

$$(\forall x : R)A, \quad (\exists x : R)A.$$

L'adjectif « typé » que nous utilisons à propos de cette notation vient de ce qu'un *type* — en un sens voisin de celui qu'a ce mot en informatique — est associé à la variable d'un quantificateur, ce type étant donné par une proposition R . L'expression (1), par exemple, pourrait se lire: pour tout ε de type réel strictement positif, il existe un n de type $\mathbb{N}$, tel que pour tout m de type entier naturel supérieur à m , on ait $|x_m - a| < \varepsilon$.

Abréviations particulières. Lorsque la proposition R est de la forme $x \ R \ T$, où R est un symbole relationnel binaire et T est un terme, les expressions

$$(\forall x : x \ R \ T)A, \quad (\exists x : x \ R \ T)A$$

sont souvent abrégées encore sous la forme

$$(\forall x \ R \ T)A, \quad (\exists x \ R \ T)A.$$

Par exemple, les propositions de la forme $(\forall x : x \in E)A$, $(\exists x : x \in E)A$, exprimant que tous les éléments de l'ensemble E ont une certaine propriété (exprimée par A), respectivement qu'il existe un élément de E ayant cette propriété, s'abrègent couramment

$$(\forall x \in E)A, \quad (\exists x \in E)A.$$

Ces expressions représentent donc les propositions

$$\forall x(x \in E \Rightarrow A), \quad \exists x(x \in E \text{ et } A).$$

Par exemple, le théorème $\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)$ de $\mathcal{T}_{\text{ENS}}$ peut être noté de cette manière

$$\vdash x \subset y \Leftrightarrow (\forall a \in x)(a \in y).$$

Schémas de déduction. Nous ne ferons que peu usage de ces abréviations dans la suite de ce livre. Mais les quantificateurs typés sont omniprésents dans le langage mathématique usuel, comme le suggère notre exemple initial de la propriété de convergence d'une suite de nombres réels, et il est utile de connaître les quelques schémas de déduction qui suivent. Les symboles A , B , R désignent comme d'habitude des propositions quelconques, et x , y désignent

des variables quelconques. Ce sont des schémas sans prémisses, notés avec une barre de déduction verticale à gauche (cf. §7.6). Le lecteur observera que ces schémas sont d'une forme analogue à celle des schémas **S71**, **S74**, **S72**, **S73** sur les quantificateurs simples. Les permutations de quantificateurs typés suivant les schémas (e), (f), (g) sont soumises cependant à la condition indiquée ci-dessous.

- (a) $\mid \Gamma \vdash \text{non}(\forall x: R)A \Leftrightarrow (\exists x: R)\text{non}A.$
- (b) $\mid \Gamma \vdash \text{non}(\exists x: R)A \Leftrightarrow (\forall x: R)\text{non}A.$
- (c) $\mid \Gamma \vdash (\forall x: R)(A \text{ et } B) \Leftrightarrow ((\forall x: R)A \text{ et } (\forall x: R)B).$
- (d) $\mid \Gamma \vdash (\exists x: R)(A \text{ ou } B) \Leftrightarrow ((\exists x: R)A \text{ ou } (\exists x: R)B).$

Si x ne figure pas librement dans S et y ne figure pas librement dans R :

- (e) $\mid \Gamma \vdash (\forall x: R)(\forall y: S)A \Leftrightarrow (\forall y: S)(\forall x: R)A.$
- (f) $\mid \Gamma \vdash (\exists x: R)(\exists y: S)A \Leftrightarrow (\exists y: S)(\exists x: R)A.$
- (g) $\mid \Gamma \vdash (\exists x: R)(\forall y: S)A \Rightarrow (\forall y: S)(\exists x: R)A.$

Tous ces schémas se démontrent facilement par une chaîne d'équivalence dans un cadre sans hypothèses. Cette démonstration est donnée ci-dessous pour (e). Les autres sont laissées au lecteur.

nil

$(\forall x: R)(\forall y: S)A$	
$\Leftrightarrow \forall x(R \Rightarrow \forall y(S \Rightarrow A))$	par abréviation
$\Leftrightarrow \forall x(\text{non}R \text{ ou } \forall y(\text{non}S \text{ ou } A))$	S53
$\Leftrightarrow \forall x\forall y(\text{non}R \text{ ou } (\text{non}S \text{ ou } A))$	S81
$\Leftrightarrow \forall y\forall x(\text{non}S \text{ ou } (\text{non}R \text{ ou } A))$	S72 , S52
$\Leftrightarrow \forall y(\text{non}S \text{ ou } \forall x(\text{non}R \text{ ou } A))$	S81
$\Leftrightarrow \forall y(S \Rightarrow \forall x(R \Rightarrow A))$	S53
$\Leftrightarrow (\forall y: S)(\forall x: R)A$	par abréviation.

Exemple. D'après les schémas (a) et (b) ci-dessus, la négation de la proposition (1), autrement dit l'assertion « la suite (x_n) ne converge pas vers a », est équivalente à

$$\left(\frac{\exists \varepsilon}{\varepsilon \in \mathbb{R} \text{ et } \varepsilon > 0} \right) \left(\frac{\forall n}{n \in \mathbb{N}} \right) \left(\frac{\exists m}{m \in \mathbb{N} \text{ et } m \geq n} \right) (|x_m - a| \geq \varepsilon).$$

Nous admettons naturellement que $|x_m - a| \geq \varepsilon$ équivaut à $\text{non}(|x_m - a| < \varepsilon)$.

11

Logique des prédicats avec égalité

11.1 Dédutions élémentaires

Étant donné un langage du premier ordre $\mathcal{L}$ dont le répertoire de symboles relationnels binaires contient le symbole d'égalité « = », nous appelons *logique des prédicats avec égalité* de langage $\mathcal{L}$, ou brièvement **Lprédég**($\mathcal{L}$), la théorie de langage $\mathcal{L}$ sans axiomes dont les déductions élémentaires sont celles de la théorie **Lpréd**($\mathcal{L}$) (§9.1) ainsi que les déductions décrites par les deux nouveaux schémas de la figure 11.1.

De façon précise, les déductions élémentaires de la théorie **Lprédég**($\mathcal{L}$) sont décrites par les schémas S1–S14 (Fig. 6.1), S61–S64 (Fig. 9.1) et les schémas S82, S83 (Fig. 11.1). La théorie **Lprédég**($\mathcal{L}$) est donc une extension de la théorie **Lpréd**($\mathcal{L}$), laquelle est une extension de **Lprop**($\mathcal{L}$) (§6.1).

Par **Lprédég**, nous désignons la théorie **Lprédég**($\mathcal{L}_\Omega$), où $\mathcal{L}_\Omega$ est le langage universel que nous supposons fixé une fois pour toutes (§9.2), et dont nous admettons évidemment qu'il contient le symbole relationnel binaire « = ». Nous appelons *théorie logique* toute théorie de langage $\mathcal{L}_\Omega$ qui est une extension de **Lprédég**, et *contexte logique* tout contexte (théorie + hypothèses) dont la théorie est une théorie logique. La théorie $\mathcal{T}_{\text{ENS}}$ est une théorie logique, et il en est de même de toutes les théories mathématiques importantes, si on les traduit dans le langage $\mathcal{L}_\Omega$.

Le schéma de réflexivité de l'égalité (S82) ne nécessite pas de commentaire. Le schéma de substitutivité (S83) est celui qui est appliqué à tout moment dans le raisonnement mathématique lorsqu'on remplace un terme par un terme égal dans une proposition. Cependant cette opération n'est pas décrite de cette manière. Le schéma ne parle pas de remplacer le terme T par le terme U dans une proposition. Il décrit l'assertion de l'une des prémisses de la déduction et l'assertion de la conclusion comme étant deux propositions construites à partir d'une même proposition A en substituant dans l'une le terme T , dans

S82. *Réflexivité de l'égalité*

$$\frac{}{\Gamma \vdash T = T.}$$

S83. *Substitutivité de l'égalité*

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash (T|x)A}{\Gamma \vdash (U|x)A} \quad \text{Si } T \text{ et } U \text{ sont librement substituables à } x \text{ dans } A.$$

Figure 11.1 Schémas de déduction élémentaires de la logique des prédicats avec égalité. T et U désignent des termes, A une proposition, x une variable.

l'autre le terme U , à toutes les occurrences libres d'une même variable x . Ceci n'est qu'une manière de décrire la relation qu'il y a entre la forme de cette prémisse et celle de la conclusion. On peut aussi décrire cette relation en parlant de remplacer *certaines* occurrences du terme T de la prémisse par le terme U . Mais il est relativement compliqué de décrire ainsi de façon *précise* la transformation qui s'opère entre la prémisse et la conclusion. Les exemples qui suivent permettront au lecteur de se familiariser avec la manière dont le schéma S83 décrit cette transformation.

Exemple 1. On peut faire dans une théorie logique, suivant S83, la déduction

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash x \in a}{\Gamma \vdash y \in a} \quad (1)$$

parce que les propositions $x \in a$, $y \in a$ peuvent être considérées comme étant obtenues en substituant x , resp. y , à une *même* variable, dans une *même* proposition, par exemple à la variable u dans la proposition $u \in a$. Autrement dit, cette déduction est conforme à S83 parce qu'elle peut s'écrire

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash (x|u)(u \in a)}{\Gamma \vdash (y|u)(u \in a)}$$

et parce que les termes x et y sont librement substituables à u dans la proposition $u \in a$. Mais on peut tout aussi bien justifier la déduction (1) en disant qu'elle peut s'écrire

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash (x|w)(w \in a)}{\Gamma \vdash (y|w)(w \in a)}.$$

Il y a une infinité de manières de considérer les propositions $x \in a$, $y \in a$ comme étant obtenues par substitution de x , resp. y , à une même variable dans une proposition A . Toutes les propositions A de la forme $z \in a$, où z est une variable distincte de a , font l'affaire.

Une déduction conforme à S83, par exemple la déduction (1), peut être

décrite comme étant de la forme

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash R}{\Gamma \vdash R'}, \quad (2)$$

où R et R' sont deux propositions pour lesquelles *on peut trouver* une proposition A et une variable x telles que R et R' soient identiques respectivement à $(T|x)A$ et $(U|x)A$, ces termes étant librement substituables à x dans A . Pratiquement bien entendu, lorsqu'on fait la déduction (1), on considère que l'on remplace x par y dans la proposition $x \in a$, vu l'égalité $x = y$. Dans tous les exemples, d'ailleurs, c'est de cette manière que l'on considère pratiquement la déduction. Le schéma S83 joue le rôle d'un *test* auquel doit se soumettre la déduction.

Exemple 2. La déduction

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash x + 1 = 2x}{\Gamma \vdash x + 1 = 2y} \quad (3)$$

est conforme à S83 parce que les propositions $x + 1 = 2x$, $x + 1 = 2y$ peuvent s'obtenir en substituant x , respectivement y , à une même variable dans une même proposition, par exemple à la variable u dans la proposition $x + 1 = 2u$. La déduction peut s'écrire

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash (x|u)(x + 1 = 2u)}{\Gamma \vdash (y|u)(x + 1 = 2u)}.$$

Pratiquement, la déduction (3) est vue comme le remplacement *d'une occurrence* de x par y dans l'assertion de la deuxième prémisse. On aurait pu faire une autre déduction de même schéma en remplaçant *l'autre* occurrence de x , et une troisième déduction en remplaçant les deux occurrences de x . Autrement dit, avec les mêmes prémisses, on peut faire les deux autres déductions suivantes, conformes à S83:

$$\frac{\Gamma \vdash x = y \quad \Gamma \vdash x + 1 = 2x}{\Gamma \vdash y + 1 = 2x} \quad \frac{\Gamma \vdash x = y \quad \Gamma \vdash x + 1 = 2x}{\Gamma \vdash y + 1 = 2y}.$$

Dans la première, les propositions R et R' (2) peuvent être considérées comme $(x|u)(u + 1 = 2x)$ et $(y|u)(u + 1 = 2x)$; dans la seconde, comme $(x|u)(u + 1 = 2u)$ et $(y|u)(u + 1 = 2u)$.

Exemple 3. Une déduction *erronée*, par rapport à S83, serait la suivante:

$$\frac{\Gamma \vdash a \cap b = X \quad \Gamma \vdash \forall b(a \cap b \subset b)}{\Gamma \vdash \forall b(X \subset b)} \quad (4)$$

Les termes T et U de cette déduction sont $a \cap b$ et X . Les propositions R , R' (2) sont $\forall b(a \cap b \subset b)$ et $\forall b(X \subset b)$. On peut certes trouver une proposition A et une variable x telles que R et R' soient identiques à $(a \cap b|x)A$ et $(X|x)A$,

par exemple la proposition $\forall b(z \subset b)$ et la variable z . Mais le terme $a \cap b$ n'est pas librement substituable à z dans cette proposition et la condition de S83 n'est pas respectée.

La déduction (4) serait d'ailleurs pour le moins bizarre en théorie des ensembles car $\vdash \forall b(a \cap b \subset b)$ est un théorème de $\mathcal{T}_{\text{ENS}}$ (E33, S62) et d'autre part, de $\vdash \forall b(X \subset b)$, on déduit $\vdash X \subset \emptyset$ par particularisation, donc $\vdash X = \emptyset$ (E12). On pourrait ainsi déduire $\vdash X = \emptyset$ de $\vdash a \cap b = X$.

En enlevant les quantificateurs $\forall b$ de (4), on obtient par contre une déduction conforme à S83:

$$\frac{\Gamma \vdash a \cap b = X \quad \Gamma \vdash a \cap b \subset b}{\Gamma \vdash X \subset b.} \quad (5)$$

Pratiquement, cette déduction est vue comme le remplacement du terme $a \cap b$ de la deuxième prémisse par le terme X . La déduction (4) est erronée parce que la variable b du terme $a \cap b$ de la deuxième prémisse, qui est remplacé par X dans la conclusion, est liée au quantificateur $\forall b$.

Exemple 4. Désignons par T et U les termes $a \cap b$ et $x \cup y$, et considérons la proposition R suivante dans laquelle il y a trois occurrences du terme T .

$$\mathsf{R}: \quad a \neq a \cap b \Rightarrow (\exists x(x \neq a \cap b) \text{ et } \forall b(a \cap b \subset b)).$$

L'arbre de cette proposition est représenté dans la figure 11.2 et les trois occurrences du terme T ($a \cap b$) y sont mises en évidence. La seule déduction

$$\frac{\Gamma \vdash \mathsf{T} = \mathsf{U} \quad \Gamma \vdash \mathsf{R}}{\Gamma \vdash \mathsf{R}_1}$$

conforme à S83 que l'on peut faire avec ces termes T , U et cette proposition R , est celle où R_1 est la proposition obtenue en remplaçant la première occurrence de T dans R par U . La deuxième occurrence de T n'est pas remplaçable par U , parce qu'elle est dominée par le quantificateur $\exists x$ dont la variable figure librement dans U . La troisième occurrence de T n'est remplaçable ni par U , ni par un autre terme, parce qu'elle est dominée par le quantificateur $\forall b$ dont la variable figure librement dans le terme T .

Examinons ces trois cas par rapport au schéma S83, en considérant les trois propositions $\mathsf{A}_1, \mathsf{A}_2, \mathsf{A}_3$ représentées dans la figure. Premièrement, la proposition R est identique à $(\mathsf{T}|w)\mathsf{A}_1$ et si l'on remplace la première occurrence de T dans R par U , la proposition R_1 obtenue est identique à $(\mathsf{U}|w)\mathsf{A}_1$. La déduction est donc dans ce cas

$$\frac{\Gamma \vdash \mathsf{T} = \mathsf{U} \quad \Gamma \vdash (\mathsf{T}|w)\mathsf{A}_1}{\Gamma \vdash (\mathsf{U}|w)\mathsf{A}_1.}$$

Elle est conforme à S83, car chacun des termes T , U est librement substituable à w dans A_1 . On le voit directement dans l'arbre de R à ceci que la première occurrence de T n'est dominée par aucun quantificateur dont la variable figure librement dans l'un ou l'autre des termes T , U .

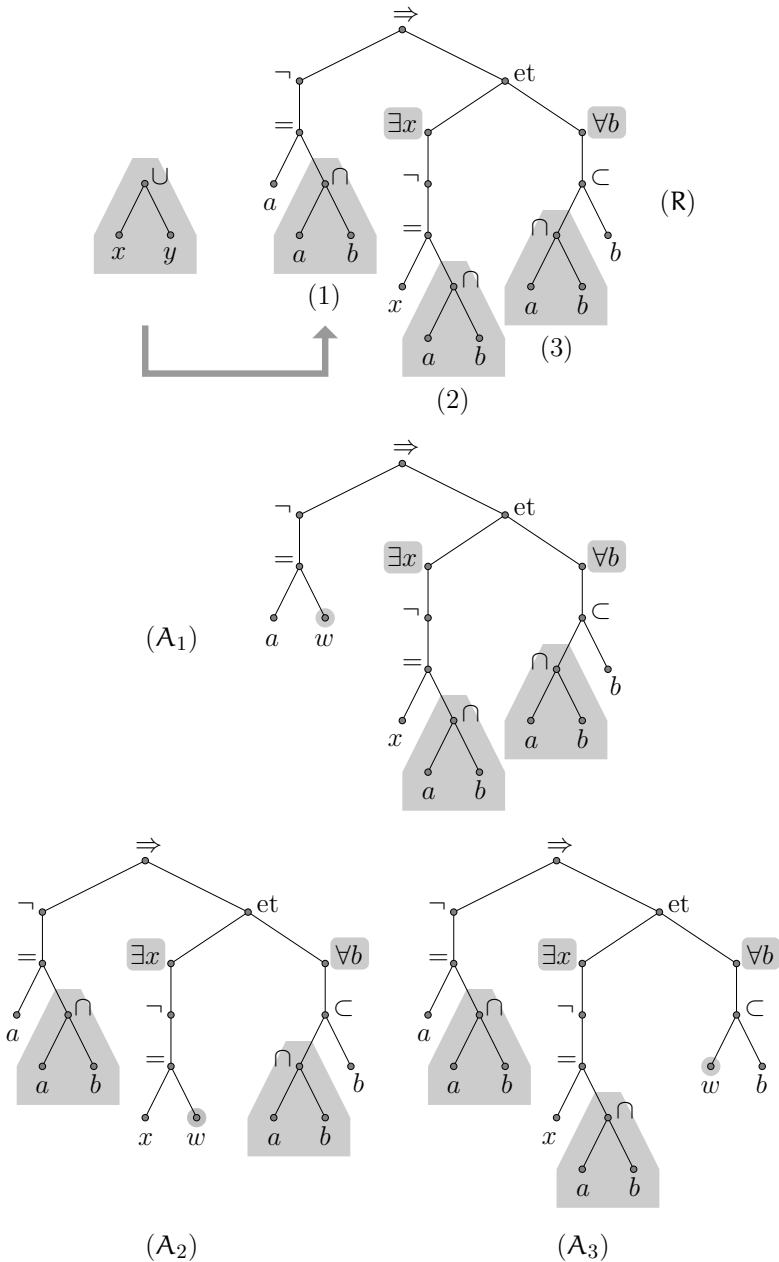


Figure 11.2 Si les propositions $a \cap b = x \cup y$ et R sont vraies, la proposition R_1 obtenue en remplaçant la première occurrence du terme $a \cap b$ dans R par le terme $x \cup y$ est vraie d'après S83, parce que cette occurrence de $a \cap b$ dans R n'est pas dominée par un quantificateur dont la variable figure librement dans l'un ou l'autre de ces termes.

Deuxièmement, la déduction

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash R}{\Gamma \vdash R_2}$$

où R_2 est la proposition obtenue en remplaçant la deuxième occurrence de T dans R par U est incorrecte, car c'est la déduction

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash (T|w)A_2}{\Gamma \vdash (U|w)A_2}$$

et le terme U n'est pas librement substituable à w dans la proposition A_2 . On le voit directement dans l'arbre de R au fait que la deuxième occurrence de T est dominée par le quantificateur $\exists x$ dont la variable figure librement dans U .

Troisièmement, la déduction

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash R}{\Gamma \vdash R_3}$$

où R_3 est la proposition obtenue en remplaçant la troisième occurrence de T dans R par U est incorrecte, car c'est la déduction

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash (T|w)A_3}{\Gamma \vdash (U|w)A_3}$$

et le terme T n'est pas librement substituable à w dans la proposition A_3 . On le voit directement dans l'arbre de R au fait que la troisième occurrence de T est dominée par le quantificateur $\forall b$ dont la variable figure librement dans T .

Autre formulation du schéma de substitutivité de l'égalité. L'exemple qui précède montre qu'on peut formuler le schéma de substitutivité de l'égalité S83 en disant qu'il permet les déductions de la forme

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash R}{\Gamma \vdash R'}$$

où R' est une proposition obtenue en remplaçant une ou plusieurs occurrences du terme T par U dans R , à condition que ces occurrences de T dans R ne soient pas dominées par un quantificateur dont la variable figure librement dans T ou dans U .

11.2 Schémas de déduction dérivés

Symétrie et transitivité de l'égalité. Dans les schémas dérivés suivants, on désigne par T , U , V des termes quelconques.

S84. *Symétrie de l'égalité*

$$\frac{\Gamma \vdash T = U}{\Gamma \vdash U = T} \quad \frac{}{\Gamma \vdash T = U \Leftrightarrow U = T.}$$

S85. *Transitivité de l'égalité*

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash U = V}{\Gamma \vdash T = V} \qquad \frac{}{\Gamma \vdash (T = U \text{ et } U = V) \Rightarrow T = V.}$$

Ces deux schémas sont donnés chacun sous deux formes. La seconde, sans prémisses, dérive évidemment de la première suivant la logique propositionnelle. Nous donnerons les démonstrations pour la première forme. Le schéma de symétrie de l'égalité se démontre en trois lignes :

Γ	$1^\circ \mid T = U \quad \textit{base}$ $2^\circ \mid T = T \quad \text{S82}$ $3^\circ \mid U = T \quad 1^\circ, 2^\circ, \text{S83}$	Γ	$T = U$ R R'	Γ	$T = U$ $(T x)(x = T)$ $(U x)(x = T).$
----------	--	----------	------------------------	----------	--

C'est une déduction par substitutivité de l'égalité. On passe de la proposition R ($T = T$) de la deuxième ligne à la proposition R' ($U = T$) de la troisième en remplaçant une occurrence de T par U . Pour comparer cette déduction avec le schéma S83, on considère la proposition $x = T$, où x est une variable qui ne figure pas librement dans T , et l'on voit que si l'on désigne cette proposition par A , les propositions R et R' sont identiques à $(T|x)A$ et $(U|x)A$, et les termes T, U sont librement substituables à x dans A .

La démonstration du schéma de transitivité de l'égalité fait une déduction semblable, précédée d'une déduction par symétrie de l'égalité. On passe ici de la proposition R ($U = V$) à la proposition à R' ($T = V$) en remplaçant une occurrence de U par T .

Γ	$1^\circ \mid T = U \quad \textit{base}$ $2^\circ \mid U = T \quad 1^\circ, \text{S84}$ $3^\circ \mid U = V \quad \textit{base}$ $4^\circ \mid T = V \quad 2^\circ, 3^\circ, \text{S83}$	Γ	$\dots$ $U = T$ R R'	Γ	$\dots$ $U = T$ $(U x)(x = V)$ $(T x)(x = V).$
----------	---	----------	-----------------------------------	----------	---

Schémas dérivés de substitutivité de l'égalité. Les schémas de déduction suivants, S86 et S87, dérivent des schémas élémentaires de substitutivité (S83) et de symétrie (S84) de l'égalité suivant la logique propositionnelle. Leur démonstration est laissée au lecteur. Ce sont des déductions très fréquentes dans les chaînes d'équivalences.

S86.

$$\frac{\Gamma \vdash T = U}{\Gamma \vdash (T|x)A \Leftrightarrow (U|x)A} \qquad \text{Si } T \text{ et } U \text{ sont librement substituables à } x \text{ dans } A.$$

D'après notre discussion de S83 (§11.1), ces déductions peuvent être décrites comme étant de la forme

$$\frac{\Gamma \vdash T = U}{\Gamma \vdash R \Leftrightarrow R'}$$

où R et R' sont des propositions et R' peut s'obtenir en remplaçant par U une ou plusieurs occurrences du terme T dans R qui ne sont pas dominées par des quantificateurs dont la variable figure librement dans T ou dans U .

S87. Si T et U sont librement substituables à x dans A :

1.
$$\frac{}{\Gamma \vdash (T = U \text{ et } (T|x)A) \Leftrightarrow (T = U \text{ et } (U|x)A)}.$$
2.
$$\frac{}{\Gamma \vdash (T = U \Rightarrow (T|x)A) \Leftrightarrow (T = U \Rightarrow (U|x)A)}.$$
3.
$$\frac{\Gamma \vdash (B \Rightarrow T = U)}{\Gamma \vdash (B \text{ et } (T|x)A) \Leftrightarrow (B \text{ et } (U|x)A)}.$$

Comme dans le cas de **S86**, ces déductions peuvent être décrites comme étant de la forme suivante, où R' est une proposition obtenue en remplaçant par U des occurrences de T dans R qui ne sont pas dominées par des quantificateurs dont la variable figure librement dans T ou dans U :

$$\frac{\frac{}{\Gamma \vdash (T = U \text{ et } R) \Leftrightarrow (T = U \text{ et } R')}}{\Gamma \vdash (T = U \Rightarrow R) \Leftrightarrow (T = U \Rightarrow R')} \quad \frac{\Gamma \vdash B \Rightarrow T = U}{\Gamma \vdash (B \text{ et } R) \Leftrightarrow (B \text{ et } R')}.$$

Exemple 1. Comme application de ces schémas, nous allons démontrer le théorème **E50** (§10.7) de $\mathcal{T}_{\text{ENS}}$:

$$\vdash a \subset X \Leftrightarrow \mathbb{C}_X(\mathbb{C}_X(a)) = a,$$

de la manière la plus simple possible, compte tenu des théorèmes antérieurs

$$\mathbf{E35.} \quad \vdash a \subset b \Leftrightarrow a \cap b = a;$$

$$\mathbf{E30.} \quad \vdash a \cap b = b \cap a;$$

$$\mathbf{E49.} \quad \vdash \mathbb{C}_X(\mathbb{C}_X(a)) = X \cap a.$$

Les deux premières lignes de la démonstration ne servent que de rappel. Elles sont normalement omises et au lieu de mentionner 1° à la ligne 4° on mentionne **E30**; au lieu de mentionner 2° à la ligne 5° on mentionne **E49**.

nil		
1°	$a \cap X = X \cap a$	E30 (exemplifié)
2°	$X \cap a = \mathbb{C}_X(\mathbb{C}_X(a))$	E49 , S84
3°	$a \subset X \Leftrightarrow a \cap X = a$	E35 (exemplifié)
4°	$\Leftrightarrow X \cap a = a$	1°, S86
5°	$\Leftrightarrow \mathbb{C}_X(\mathbb{C}_X(a)) = a$	2°, S86 .

Pour comparer cette démonstration avec le schéma **S86**, désignons par T , U , V les termes $a \cap X$, $X \cap a$, $\mathbb{C}_X(\mathbb{C}_X(a))$ et par R , R' , R'' les propositions $a \cap X = a$,

$X \cap a = a$, $\mathbb{C}_X(\mathbb{C}_X(a)) = a$. La démonstration peut s'écrire

nil	nil
1° $T = U$	1° $T = U$
2° $U = V$	2° $U = V$
3° $a \subset X \Leftrightarrow R$	3° $a \subset X \Leftrightarrow (T y)(y = a)$
4° $\Leftrightarrow R'$	4° $\Leftrightarrow (U y)(y = a)$
5° $\Leftrightarrow R''$	5° $\Leftrightarrow (V y)(y = a)$.

La proposition $y = a$ joue le rôle de la proposition A du schéma. On passe de R à R' en remplaçant une occurrence du terme T par le terme U , et de R' à R'' en remplaçant une occurrence du terme U par le terme V . Les propositions R, R', R'' sont identiques respectivement à $(T|y)A$, $(U|y)A$, $(V|y)A$.

Substitutivité de l'égalité dans un terme. Soient T, U, V des termes et x une variable. Considérons les propositions R et R' suivantes, dont les arbres sont représentés symboliquement dans la figure 11.3:

$$\begin{aligned} R : & \quad (T|x)V = (T|x)V, \\ R' : & \quad (T|x)V = (U|x)V. \end{aligned} \tag{1}$$

On suppose que les termes T et U sont librement substituables à x dans V . La figure ne représente qu'une occurrence libre de x dans V ainsi que la branche de l'arbre de V qui lui correspond, mais celle-ci symbolise toutes les occurrences libres de x dans V . Nous supposons donc que ces occurrences de x dans V ne sont pas dominées par des quantificateurs dont la variable figure librement dans le terme T ou dans le terme U . La proposition R est vraie dans tout contexte, d'après S82. Si l'égalité $T = U$ est vraie dans un contexte, alors la proposition R' est vraie dans ce contexte d'après S83, car on voit que cette proposition R' peut se construire en remplaçant certaines occurrences du terme T dans R par le terme U , et que ces occurrences de T dans R ne sont pas dominées par des quantificateurs dont la variable figure librement dans T ou dans U . Nous pouvons donc noter le schéma de déduction suivant:

S88.

$$\frac{\Gamma \vdash T = U}{\Gamma \vdash (T|x)V = (U|x)V}$$

Si T, U, V sont des termes
et si T et U sont librement
substituables à x dans V .

Nous pouvons encore vérifier que pour les propositions R, R' ci-dessus (1) la déduction

$$\frac{\Gamma \vdash T = U \quad \Gamma \vdash R}{\Gamma \vdash R'}$$

est bien conforme à S83, en montrant qu'il existe une proposition A et une variable y telles que R et R' soient identiques respectivement à $(T|y)A$ et $(U|y)A$ et dans laquelle T et U soient librement substituables à y . Or il suffit de prendre pour A la proposition $(T|x)V = (y|x)V$ (Fig. 11.3), où y est une

variable qui ne figure dans aucun des termes T , U , V et qui est librement substituable à x dans V . Car

$$\begin{aligned} R \text{ est identique à : } & (T|y) \overbrace{((T|x)V = (y|x)V)}^A, \\ R' \text{ est identique à : } & (U|y) ((T|x)V = (y|x)V), \end{aligned}$$

et les termes T et U sont librement substitubles à y dans A .

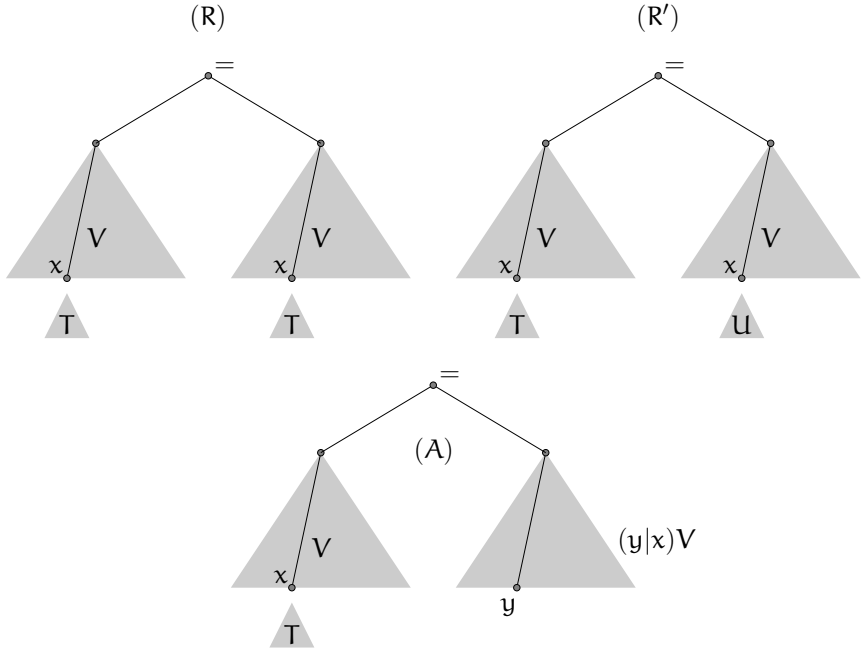


Figure 11.3 Justification de S88.

Substitutivité de l'égalité et de l'équivalence combinées. Dans les démonstrations, le schéma de substitutivité de l'égalité S86 s'emploie souvent de manière combinée avec le schéma général de substitutivité de l'équivalence (§10.1). Considérons par exemple les propositions R et R' suivantes :

$$\begin{aligned} R : & \forall a \forall b \forall x (x \in \{a, b\} \cap X \Rightarrow x \in X) \\ R' : & \forall a \forall b \forall x (x \in (\{a\} \cup \{b\}) \cap X \Rightarrow x \in X). \end{aligned}$$

Les termes $\{a, b\}$, $\{a\}$ sont des expressions bien connues de théorie des ensembles, désignant l'ensemble dont les éléments sont a et b , et l'ensemble ayant a pour unique élément. Les axiomes de $\mathcal{T}_{\text{ENS}}$ sur ces notations seront présentés au paragraphe suivant. Nous admettons ici le théorème intuitivement évident :

$$\vdash \{a, b\} = \{a\} \cup \{b\}. \quad (2)$$

Désignons le terme $\{a, b\}$ par T et le terme $\{a\} \cup \{b\}$ par U , et remarquons que R' peut s'obtenir syntaxiquement en remplaçant l'occurrence du terme T dans

R (il n'y en a qu'une) par U . Mais cette occurrence de T dans R est dominée par les quantificateurs $\forall a, \forall b$, dont les variables a, b figurent librement dans les termes T et U . Une autre manière de voir les choses est de considérer R et R' comme étant les propositions $(T|z)A, (U|z)A$, où A est la proposition

$$\forall a \forall b \forall x (x \in z \cap X \Rightarrow x \in X).$$

Les termes T et U ne sont pas librement substituables à z dans A .

Les arbres des propositions A, R, R' sont représentés dans la figure 11.4. Notons que dans le terme $\{a, b\}$ les symboles $\{, \}$ forment ensemble *un* symbole fonctionnel binaire, représentant l'opération de construction d'un ensemble avec deux objets. Dans le terme $\{a\}$, les symboles $\{ \}$ forment un symbole fonctionnel unaire, représentant l'opération de construction d'un ensemble à un seul élément.

Le schéma S86 à lui seul ne permet donc pas de déduire

$$\vdash R \Leftrightarrow R' \quad (3)$$

de (2). Cependant, si nous supprimons les quantificateurs $\forall a, \forall b, \forall x$ dans R et R' , c'est-à-dire si nous considérons les propositions R_1 et R'_1 suivantes :

$$\begin{aligned} R_1 : & \quad x \in \{a, b\} \cap X \Rightarrow x \in X \\ R'_1 : & \quad x \in (\{a\} \cup \{b\}) \cap X \Rightarrow x \in X, \end{aligned}$$

alors nous pouvons déduire de (2) le théorème $\vdash R_1 \Leftrightarrow R'_1$ par S86, puisque l'occurrence de T dans R_1 n'est plus dominée par les quantificateurs $\forall a, \forall b$. Autrement dit, si nous appelons A_1 la proposition $x \in z \cap X \Rightarrow x \in X$, les termes T et U sont librement substituables à z dans A_1 . Or, de $\vdash R_1 \Leftrightarrow R'_1$ on peut déduire $\vdash R \Leftrightarrow R'$ par S70. On peut donc bien déduire (3) de (2), mais pas uniquement par substitutivité de l'égalité. C'est la combinaison d'une déduction par substitutivité de l'égalité et d'une déduction par substitutivité de l'équivalence. Cette combinaison est très fréquente.

Par *a fortiori*, on peut déduire $\Gamma \vdash R \Leftrightarrow R'$ de $\vdash R \Leftrightarrow R'$ pour une liste d'hypothèses Γ quelconque. Les propositions R et R' sont donc équivalentes dans tout contexte ensembliste en vertu du théorème (2).

Le schéma d'existence triviale. Nous appelons ainsi le schéma de déduction sans prémisses suivant, dans lequel T désigne un terme et x une variable :

S89.

$$\frac{}{\Gamma \vdash \exists x (x = T)} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } T. \end{array}$$

Une déduction de cette forme se démontre en effet comme suit. Si x ne figure pas librement dans T , la proposition $T = T$ est identique à $(T|x)(x = T)$. De $\Gamma \vdash T = T$ (S82), on peut donc déduire $\Gamma \vdash \exists x (x = T)$:

$$\begin{array}{c} \Gamma \\ 1^\circ \left| \begin{array}{l} T = T \\ 2^\circ \left| \exists x (x = T) \end{array} \right. \end{array} \quad \begin{array}{l} \text{S82} \\ 1^\circ, \text{S63.} \end{array}$$

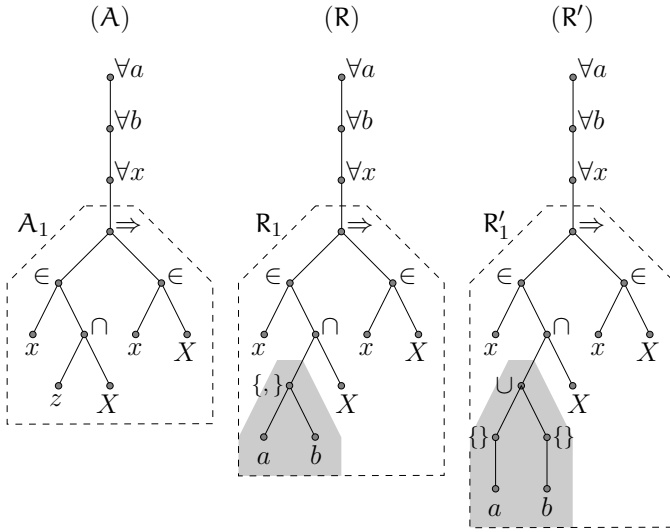


Figure 11.4 $\vdash R_1 \Leftrightarrow R'_1$ se déduit de $\vdash \{a, b\} = \{a\} \cup \{b\}$ par substitutivité de l'égalité (S86); $\vdash R \Leftrightarrow R'$ se déduit de $\vdash R_1 \Leftrightarrow R'_1$ par substitutivité de l'équivalence (§10.1).

Interprété informellement, ce raisonnement est le suivant. L'objet désigné par le terme T est égal à lui-même (S82). Donc il existe un objet qui est égal à l'objet désigné par T .

Par exemple, le jugement $\vdash \exists b(b = a + 1)$ est un théorème de **Lprédég**, d'après S89, parce que la variable b ne figure pas librement dans le terme $a + 1$. On ne peut pas en dire autant de $\vdash \exists a(a = a + 1)$.

Hypothèses de désignation. Une pratique courante dans le raisonnement mathématique est de dire « posons $x = T$ », lorsque T est un terme compliqué que l'on veut éviter d'écrire plusieurs fois. Ayant posé $x = T$, on se sert ensuite de la simple variable x chaque fois que l'on veut parler de l'objet désigné par le terme T . Bien entendu, on ne prend pas pour x n'importe quelle variable, mais une variable qui ne figure pas librement dans T et sur laquelle on n'a pas encore fait d'hypothèses dans le contexte. Lorsqu'on « pose » $x = T$, on fait alors une hypothèse sur x et ce que l'on démontre dans ce cadre n'est pas nécessairement vrai en dehors de cette hypothèse. Cependant, si l'on démontre sous cette hypothèse une proposition B dans laquelle x ne figure pas librement, cette proposition est vraie indépendamment de l'hypothèse. Cette déduction est décrite par le schéma suivant.

S90.

$$\frac{\Gamma; x = T \vdash B}{\Gamma \vdash B} \quad \begin{array}{l} \text{Si } x \text{ ne figure librement ni dans} \\ \Gamma, \text{ ni dans } T, \text{ ni dans } B. \end{array}$$

	Γ	
1 ^o	$\exists x(x = T)$	S89
2 ^o	$x = T$	hyp
3 ^o	B	base
4 ^o	B	1 ^o , 3 ^o , S64.

Une hypothèse de la forme $x = T$, où x ne figure librement ni dans T ni dans les hypothèses précédentes du contexte, ne sert qu'à désigner par x un objet désigné aussi par le terme T . Nous l'appelons une *hypothèse de désignation*. Dans le raisonnement courant, c'est un genre d'hypothèse qui est à peine perçu comme tel. Dans un raisonnement formel, une déduction suivant S90 se présente de la manière suivante.

	Γ	
	$x = T$	hyp
	$\vdots$	
n^o	B	
	B	n^o , S90.

Dans le raisonnement informel ordinaire, le cadre intérieur et la répétition finale de la proposition B sont omis. Mais on peut dire qu'ils sont logiquement présents.

Il arrive souvent dans le raisonnement mathématique qu'une hypothèse de désignation soit faite dans le but de démontrer un théorème d'existence, c'est-à-dire un théorème de la forme $\Gamma \vdash \exists xA$ où A est une proposition exprimant une propriété d'un objet désigné par x . On veut montrer qu'il existe un objet ayant cette propriété. Or on soupçonne qu'un objet connu, désigné par un terme T relativement complexe, possède justement la propriété en question. Dans cette situation, on fait l'hypothèse de désignation $x = T$ et l'on essaie de démontrer A . Si l'on y parvient, on peut conclure d'après le schéma de déduction suivant.

S91.

$\Gamma; x = T \vdash A$	Si x ne figure librement
$\Gamma \vdash \exists xA$	ni dans Γ , ni dans T .

Démonstration:

	Γ	
1 ^o	$x = T$	hyp
2 ^o	A	base
3 ^o	$\exists xA$	2 ^o , S66
4 ^o	$\exists xA$	3 ^o , S90.

Élimination et introduction de quantificateurs avec une égalité. Nous présentons ici, sous le numéro S92, deux schémas de déduction qui sont parmi les plus fréquemment employés dans les démonstrations en forme de chaîne d'équivalence, où ils permettent d'éliminer ou d'introduire des quantificateurs. S92.

$$\frac{}{\Gamma \vdash \exists x(x = T \text{ et } A) \Leftrightarrow (T|x)A} \quad \frac{}{\Gamma \vdash \forall x(x = T \Rightarrow A) \Leftrightarrow (T|x)A}$$

Si x ne figure pas librement dans T et si
 T est librement substituable à x dans A .

Exemples. Les deux théorèmes suivants de **Lprédégal**, déduits d'après ces schémas, devraient être intuitivement évidents:

$$\vdash \exists x(x = a \text{ et } x \in b) \Leftrightarrow a \in b; \quad \vdash \forall x(x = a \Rightarrow x \in b) \Leftrightarrow a \in b. \quad (4)$$

Le premier dit ceci: s'il existe un objet qui est égal à l'objet a et qui est élément de l'ensemble b , alors l'objet a est élément de l'ensemble b ; réciproquement, si l'objet a est élément de l'ensemble b , alors il existe un objet qui est élément de b et qui est égal à a . Le deuxième: si tout objet qui est égal à a est élément de b , alors a est élément de b , et réciproquement. Dans ces exemples, la proposition A du schéma est $x \in b$, et le terme T est la variable a .

Démonstration. Si une proposition A , une variable x et un terme T satisfont aux conditions du schéma, on peut écrire la chaîne d'équivalences ci-dessous dans un cadre sans hypothèses. L'application de S81 à la troisième ligne est correcte parce que les conditions mentionnées font que x ne figure pas librement dans la proposition $(T|x)A$.

$$\begin{aligned} \exists x(x = T \text{ et } A) &\Leftrightarrow \exists x(x = T \text{ et } (x|x)A) \\ &\Leftrightarrow \exists x(x = T \text{ et } (T|x)A) && \text{S87.1} \\ &\Leftrightarrow \exists x(x = T) \text{ et } (T|x)A && \text{S81} \\ &\Leftrightarrow (T|x)A && \text{S89, S60.1.} \end{aligned}$$

L'écriture d'une démonstration analogue pour le deuxième schéma, suivant S87.2, S53, S81, S71, et S60.2, est laissée au lecteur.

Lorsqu'on remplace une proposition de la forme $\exists x(x = T \text{ et } A)$ ou de la forme $\forall x(x = T \Rightarrow A)$ par $(T|x)A$ en vertu de S92, on dit que l'on élimine le quantificateur $\exists x$ ou $\forall x$ avec l'égalité $x = T$. On dit qu'on introduit ce quantificateur avec l'égalité $x = T$ lorsqu'on fait l'inverse.

11.3 Application à la théorie des ensembles

Ensembles à deux éléments. En théorie des ensembles, l'expression $\{a, b\}$ est un terme. Son interprétation est qu'il désigne un ensemble dont les éléments sont deux objets nommés a et b . On parle généralement d'un *ensemble à deux éléments*, mais la notation ne suppose pas que les objets désignés par a et b

soient distincts. S'ils ne le sont pas, il n'y a en fait qu'un seul élément dans cet ensemble.

La syntaxe du terme $\{a, b\}$ doit être considérée de la manière suivante: les signes $\{, \}$ forment un symbole fonctionnel binaire, qui est combiné avec les termes a et b . En notation préfixée on pourrait écrire $\{, \}ab$ ou utiliser un autre symbole préfixé, par exemple Δ , et écrire Δab . L'arbre de ce terme est ainsi



et son interprétation est exprimée par l'axiome suivant:

E51. *Axiome de l'ensemble à deux éléments*

$$\vdash x \in \{a, b\} \Leftrightarrow (x = a \text{ ou } x = b).$$

Les théorèmes et schémas de déduction E52–E58 ci-dessous sont démontrés dans l'annexe du livre (Chap. 16).

E52. $\vdash a \in \{a, b\} \quad \vdash b \in \{a, b\}.$

E53. $\vdash \{a, b\} \neq \emptyset.$

E54. $\vdash \forall a \exists x (a \in x).$

E55. $\vdash \{a, b\} = \{b, a\}.$

E56. *Schémas de déduction:*

$$1. \frac{}{\Gamma \vdash \exists x (x \in \{a, b\} \text{ et } R) \Leftrightarrow ((a|x)R \text{ ou } (b|x)R)}$$

$$2. \frac{}{\Gamma \vdash \forall x (x \in \{a, b\} \Rightarrow R) \Leftrightarrow ((a|x)R \text{ et } (b|x)R)}$$

Si R est une proposition dans laquelle les variables a et b sont librement substituables à x .

E57. $\vdash \{a, b\} \subset X \Leftrightarrow (a \in X \text{ et } b \in X).$

E58. $\vdash \{a, b\} = \{a, c\} \Leftrightarrow b = c.$

Du théorème **E53** on déduit par introduction de $\exists$ qu'il existe un ensemble non vide, et en outre qu'il existe deux ensembles différents, c'est-à-dire les théorèmes:

$$\begin{aligned} &\vdash \exists x (x \neq \emptyset); \\ &\vdash \exists y \exists x (x \neq y). \end{aligned}$$

Dans l'univers des ensembles, d'après les axiomes posés jusqu'ici, il y a donc au moins deux objets distincts. Nous verrons un peu plus loin qu'il y en a plus. Le théorème **E54** contient une autre information sur cet univers, à savoir sur les relations d'appartenance entre objets. Si nous nous représentons comme au paragraphe 8.2 les objets de cet univers (les ensembles) par des « points » et la relation d'appartenance $x \in y$ entre deux objets par une flèche allant du

point y au point x , nous voyons que dans ce système de points et de flèches, il existe pour chaque point a un objet x et une flèche allant de x à a (E54). Tout point est l'arrivée d'une flèche. Par contre, il y a un point qui n'est le départ d'aucune flèche, à savoir le point qui représente l'objet $\emptyset$, puisqu'on a le théorème $\vdash \text{non} \exists x(x \in \emptyset)$.

Ensembles à un seul élément. L'expression $\{a\}$ en théorie des ensembles est un terme, interprété comme un ensemble dont le seul élément est l'objet désigné par a . On parle aussi de *l'ensemble réduit au seul élément a* . Syntaxiquement, les signes $\{ \}$ forment un symbole fonctionnel unaire et en notation préfixée on pourrait écrire $\{ \}a$ au lieu de $\{a\}$. On pourrait prendre comme axiome, traduisant l'interprétation de ce terme, la proposition du jugement

$$\vdash x \in \{a\} \Leftrightarrow x = a. \quad (1)$$

Mais on préfère, chaque fois que cela est possible, poser des axiomes qui sont des *définitions*. Le sujet des définitions sera traité au chapitre 12. Nous pouvons déjà dire ici qu'une définition est un axiome d'une forme particulière. Il peut être vu comme l'introduction d'une nouvelle notation qui sert seulement à abrégé certaines expressions. Dans le cas de l'ensemble $\{a\}$, nous posons

E59. *Axiome de l'ensemble à un seul élément*

$$\vdash \{a\} = \{a, a\}.$$

La notation $\{a\}$ peut être ainsi considérée comme une abréviation de $\{a, a\}$, et toutes les propriétés des ensembles à un seul élément en découlent, en particulier le théorème (1). La plupart des théorèmes suivants sont démontrés dans l'annexe.

E60. Théorème principal $\vdash x \in \{a\} \Leftrightarrow x = a$.

E61. $\vdash a \in \{a\}$.

E62. $\vdash \{a\} \neq \emptyset$.

E63. $\vdash \{a\} \subset X \Leftrightarrow a \in X$.

E64. $\vdash \{a\} \cap X \neq \emptyset \Leftrightarrow a \in X$.

E65. $\vdash \{a\} \subset \{b\} \Leftrightarrow a = b$.

E66. $\vdash \{a\} = \{b\} \Leftrightarrow a = b$.

E67. $\vdash a = b \Leftrightarrow \{a, b\} = \{a\}$.

E68. $\vdash \{a, b\} = \{a\} \cup \{b\}$.

E69. $\vdash \{a\} \cap \{b\} \neq \emptyset \Leftrightarrow a = b$.

E70. $\vdash X \subset \{a\} \Leftrightarrow (X = \emptyset \text{ ou } X = \{a\})$.

La démonstration de E70 donnée dans l'annexe est la démonstration facile mais « laborieuse ». Avec un peu de pratique, on peut remplacer ce genre de démonstration par des chaînes d'équivalence plus simples et élégantes. Le lecteur est invité à étudier la suivante pour E70, qui utilise plusieurs techniques standard de ce genre de démonstration. Ces démonstrations ne sont courtes que si l'on dispose des théorèmes auxiliaires (lemmes) adéquats et ce n'est pas toujours le cas. On doit parfois démontrer séparément les lemmes nécessaires.

Mais il s'agit très souvent de théorèmes évidents. Dans l'exemple présent ce sont les deux théorèmes

$$\vdash x = \emptyset \Rightarrow x \subset y \quad (2)$$

$$\vdash (x \subset y \text{ et } a \in x) \Rightarrow a \in y \quad (3)$$

qui découlent évidemment de E7 et E4. Les techniques standard illustrées par la démonstration ci-dessous sont : l'introduction d'une « disjonction de cas contraires » suivant S52.16 (ligne 1°); la simplification d'une conjonction (3°) suivant S60.3, ou l'inverse, l'introduction d'une conjonction suivant le même schéma (5°) — c'est là que sont utilisés les lemmes (2) et (3) —; la mise en facteur d'un quantificateur (4°) suivant S81; l'élimination d'un quantificateur avec une égalité (7°) suivant S92.

1°	$X \subset \{a\} \Leftrightarrow (X \subset \{a\} \text{ et } (X = \emptyset \text{ ou } X \neq \emptyset))$	S52.16
2°	$\Leftrightarrow (X \subset \{a\} \text{ et } X = \emptyset) \text{ ou } (X \subset \{a\} \text{ et } X \neq \emptyset)$	S52.11
3°	$\Leftrightarrow X = \emptyset \text{ ou } (X \subset \{a\} \text{ et } \exists x(x \in X))$	(2), S60.3
4°	$\Leftrightarrow X = \emptyset \text{ ou } \exists x(X \subset \{a\} \text{ et } x \in X)$	S81
5°	$\Leftrightarrow X = \emptyset \text{ ou } \exists x(X \subset \{a\} \text{ et } x \in X \text{ et } x \in \{a\})$	(3), S60.3
6°	$\Leftrightarrow X = \emptyset \text{ ou } \exists x(X \subset \{a\} \text{ et } x \in X \text{ et } x = a)$	E60
7°	$\Leftrightarrow X = \emptyset \text{ ou } (X \subset \{a\} \text{ et } a \in X)$	S92
8°	$\Leftrightarrow X = \emptyset \text{ ou } (X \subset \{a\} \text{ et } \{a\} \subset X)$	E63
9°	$\Leftrightarrow X = \emptyset \text{ ou } X = \{a\}$	E10.

Ensembles à trois éléments et plus. Au moyen d'ensembles à deux éléments et d'ensembles à un seul élément, on peut définir les ensembles à trois éléments, à quatre éléments, etc. par des réunions. De façon précise, on pose :

$$\begin{aligned} \vdash \{a, b, c\} &= \{a, b\} \cup \{c\} \\ \vdash \{a, b, c, d\} &= \{a, b, c\} \cup \{d\} \\ &\text{etc.} \end{aligned}$$

Ce sont des *axiomes* — plus précisément des définitions — auxquels nous ne donnerons pas de numéro. Syntaxiquement, le terme $\{a, b, c\}$ se compose du symbole fonctionnel ternaire $\{, , \}$ appliqué aux trois termes a, b, c . En notation préfixée on pourrait écrire $\{, , \}abc$. Bien entendu, dans la dénomination « ensemble à trois éléments » l'adjectif numérique « trois » se réfère seulement à la syntaxe du terme $\{a, b, c\}$, car les objets désignés par a, b, c ne sont pas nécessairement distincts. Les théorèmes E18, E51, E60 entraînent

$$\begin{aligned} \vdash \{a, b, c\} &= \{a\} \cup \{b\} \cup \{c\} \\ \vdash \{a, b, c, d\} &= \{a\} \cup \{b\} \cup \{c\} \cup \{d\} \\ &\text{etc.} \\ \vdash x \in \{a, b, c\} &\Leftrightarrow (x = a \text{ ou } x = b \text{ ou } x = c) \\ \vdash x \in \{a, b, c, d\} &\Leftrightarrow (x = a \text{ ou } x = b \text{ ou } x = c \text{ ou } x = d) \\ &\text{etc.} \end{aligned}$$

Dans ces propositions, des parenthèses sont omises naturellement à droite des signes $=$ et $\Leftrightarrow$ en vertu de l'associativité de la disjonction (S52.9) et de la réunion (E21).

L'ensemble des parties d'un ensemble. La proposition $x \subset a$ se lit: x est inclus dans a , ou x est un sous-ensemble de a , ou encore x est une *partie* de l'ensemble a . On désigne par $\mathcal{P}(a)$ l'ensemble des parties ou sous-ensembles de l'ensemble a . Autrement dit, étant donné un ensemble a , on admet l'existence d'un ensemble unique, désigné par $\mathcal{P}(a)$, ayant pour éléments les sous-ensembles de a . Cela revient à admettre l'axiome suivant:

E71. *Axiome de l'ensemble des parties d'un ensemble*

$$\vdash x \in \mathcal{P}(a) \Leftrightarrow x \subset a.$$

Un ensemble x est élément de l'ensemble $\mathcal{P}(a)$ si et seulement s'il est inclus dans a . Syntaxiquement, le symbole $\mathcal{P}$ est un symbole fonctionnel unaire et l'on peut écrire naturellement $\mathcal{P}a$ sans parenthèses autour de a . Celles-ci sont cependant traditionnelles. Les théorèmes suivants, à l'exception de E76, sont démontrés dans l'annexe.

E72. $\vdash a \in \mathcal{P}(a)$ et $\emptyset \in \mathcal{P}(a)$.

E73. $\vdash \mathcal{P}(a) \neq \emptyset$.

E74. $\vdash \mathcal{P}(\emptyset) = \{\emptyset\}$.

E75. $\vdash \mathcal{P}(\{x\}) = \{\emptyset, \{x\}\}$.

E76. $\vdash \mathcal{P}(\{x, y\}) = \{\emptyset, \{x\}, \{y\}, \{x, y\}\}$.

E77. $\vdash a \subset b \Leftrightarrow \mathcal{P}(a) \subset \mathcal{P}(b)$.

La démonstration de E76 est facile mais laborieuse. Nous n'en donnerons que le plan. On raisonne essentiellement par disjonction des cas. On démontre

$$z \in \mathcal{P}(\{x, y\}) \Rightarrow z \in \{\emptyset, \{x\}, \{y\}, \{x, y\}\} \quad (4)$$

en faisant l'hypothèse $z \in \mathcal{P}(\{x, y\})$ et en démontrant $z \in \{\emptyset, \{x\}, \{y\}, \{x, y\}\}$ par disjonction de quatre cas, en vertu de la proposition vraie suivante:

$$((x \notin z \text{ et } y \notin z) \text{ ou } (x \in z \text{ et } y \notin z)) \text{ ou } ((x \notin z \text{ et } y \in z) \text{ ou } (x \in z \text{ et } y \in z)).$$

Dans le premier cas on démontre $z = \emptyset$; dans le deuxième, $z = \{x\}$; dans le troisième, $z = \{y\}$; dans le quatrième, $z = \{x, y\}$. Dans les quatre cas, on a donc $z \in \{\emptyset, \{x\}, \{y\}, \{x, y\}\}$. L'implication réciproque de (4) se démontre aussi par disjonction de quatre cas: $z = \emptyset$, $z = \{x\}$, $z = \{y\}$, $z = \{x, y\}$, puisque l'hypothèse $z \in \{\emptyset, \{x\}, \{y\}, \{x, y\}\}$ équivaut à la disjonction de ces quatre propositions.

Exercice 1. (a) Le mathématicien allemand Ernst Zermelo a proposé (1908) de définir les entiers naturels $0, 1, 2, \dots$ comme étant les ensembles

$$\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \{\{\{\emptyset\}\}\}, \dots \quad (5)$$

Cette proposition n'est évidemment raisonnable que si deux quelconques de ces ensembles sont distincts. Montrer que c'est bien le cas, en vertu de E62 et

E66. Plus précisément, indiquer comment peut se démontrer le théorème

$$\vdash \underbrace{\{\dots\{\emptyset\}\dots\}}_m \neq \underbrace{\{\dots\{\emptyset\}\dots\}}_{n+m} \quad (6)$$

dans lequel le premier symbole $\emptyset$ est noté entre m parenthèses $\{\}$ ($m \geq 0$), et le deuxième entre $n + m$ de ces parenthèses ($n \neq 0$).

(b) Cela montre que dans l'univers de la théorie des ensembles, tel qu'il est décrit par les axiomes que nous avons posés jusqu'ici, il y a un nombre infini d'objets (ensembles). Mais cela ne signifie pas qu'il y ait un objet qui soit un *ensemble infini*, c'est-à-dire un ensemble ayant une infinité d'éléments. Tous les objets de la liste (5) sont des ensembles à un seul élément. Pour qu'il y ait un ensemble infini dans l'univers des ensembles, il faut le décréter, c'est-à-dire poser un axiome correspondant. Une manière de le faire est d'affirmer l'existence d'un ensemble A tel que tous les objets de la liste (5) soient éléments de A . Le problème est de formuler cet axiome en n'utilisant que les notions de théorie des ensembles introduites jusqu'ici. Il ne peut être question de dénombrer des parenthèses $\{\}$. Par exemple on ne peut pas écrire

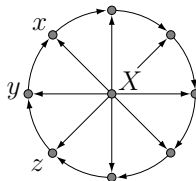
$$\exists A \forall n (n \in \mathbb{N} \Rightarrow \underbrace{\{\dots\{\emptyset\}\dots\}}_n \in A),$$

car cela utilise des notions ($\mathbb{N}$) qui ne font pas partie de notre théorie $\mathcal{T}_{\text{ENS}}$ au stade actuel. On ne peut pas non plus écrire une expression de longueur infinie telle que $\exists A (\emptyset \in A \text{ et } \{\emptyset\} \in A \text{ et } \{\{\emptyset\}\} \in A \text{ et } \dots)$. Mais on peut formuler un axiome ayant le même sens.

Exercice 2. Les axiomes de $\mathcal{T}_{\text{ENS}}$ posés dans cet ouvrage ne permettent de démontrer ni l'existence, ni la non existence d'un ensemble qui soit élément de lui-même. Du moins personne, à notre connaissance, n'est parvenu avec *ces* axiomes à démontrer l'un ou l'autre des théorèmes

$$\vdash \exists x (x \in x), \quad \vdash \text{non} \exists x (x \in x).$$

Bien entendu l'existence d'un ensemble qui *n'est pas* élément de lui-même est démontrée par le théorème $\vdash \emptyset \notin \emptyset$. La même question, d'existence ou non, se pose pour deux ensembles x, y tels que l'on ait $x \in y$ et $y \in x$; puis pour trois ensembles x, y, z tels que l'on ait $x \in y$ et $y \in z$ et $z \in x$. etc. Le dessin ci-dessous représente sous forme d'un diagramme une telle chaîne fermée de relations d'appartenances: $x \in y$ et $y \in z$ et $z \in \dots \in x$.



Les huit points sur le cercle représentent des ensembles et le point central représente l'ensemble à huit éléments $X = \{x, y, z, \dots\}$. Les flèches repré-

sentent les relations d'appartenance: $x \in y$, $y \in z$, etc. et les relations d'appartenance $x \in X$, $y \in X$, $z \in X$, etc. Une propriété remarquable de l'ensemble X est que pour chacun de ses éléments a , on a $X \cap a \neq \emptyset$. En effet, $X \cap x \neq \emptyset$ parce que $y \in X$ et $y \in x$; de même, $X \cap y \neq \emptyset$ parce que $z \in X$ et $z \in y$, etc. Ni l'existence, ni la non existence d'un tel ensemble X n'est démontrable avec nos axiomes. Sa non existence est souvent prise comme un axiome, appelé *l'axiome de fondation*:

$$\vdash \text{non} \exists X (X \neq \emptyset \text{ et } \forall a (a \in X \Rightarrow X \cap a \neq \emptyset)).$$

En utilisant cet axiome, démontrer les théorèmes

$$\begin{aligned} &\vdash \text{non}(x \in x) \\ &\vdash \text{non}(x \in y \text{ et } y \in x) \\ &\vdash \text{non}(x \in y \text{ et } y \in z \text{ et } z \in x). \end{aligned}$$

11.4 Quantificateurs d'unicité

Les quantificateurs Hx , $\exists!x$. Lorsqu'on lit une proposition de la forme $\exists xA$, on s'exprime souvent en disant: «il existe *au moins* un x tel que A ». On veut ainsi éviter tout malentendu, toute confusion avec la proposition: il existe un x *et un seul* tel que A . Cette dernière proposition sera notée $\exists!xA$, au moyen du quantificateur modifié $\exists!x$. C'est une notation assez peu employée dans le langage mathématique courant où l'on préfère en général l'usage des mots «il existe un et un seul». Mais dans un paragraphe comme celui-ci, spécialement consacré à ce genre d'expression, un abréviation s'impose. Celle-ci est assez fréquente dans les ouvrages de logique.

Une autre proposition dont le sens est distinct de celui de $\exists xA$ est l'assertion «il existe *au plus* un x tel que A ». Nous la noterons HxA , avec le quantificateur spécial Hx . Cette notation, contrairement à la précédente, est spécifique au présent ouvrage. Il est peu probable que le lecteur la rencontre ailleurs. Aucun symbole particulier ne semble s'être imposé chez les logiciens pour abréger l'expression «il existe au plus un».

Précisons encore le sens de ces divers quantificateurs au moyen d'un exemple. Lorsqu'on écrit $\exists x(x \in a)$, cela signifie que l'ensemble a possède *au minimum* un élément, c'est-à-dire qu'il possède un élément, ou plusieurs, ou une infinité. Lorsqu'on affirme qu'il existe *au plus* un x tel que $x \in a$, ce que nous écrivons $Hx(x \in a)$, on veut dire que l'ensemble a possède *au maximum* un élément, c'est-à-dire qu'il est vide (zéro élément) ou possède un seul élément. Enfin, dire qu'il existe *un x et un seul* tel que $x \in a$, ce que nous notons $\exists!x(x \in a)$, signifie que a possède *exactement* un élément, ni plus, ni moins. Cette assertion équivaut donc à la conjonction de « a possède au moins un élément» et de « a possède au plus un élément». Nous aurons le théorème logique: $\vdash \exists!x(x \in a) \Leftrightarrow (\exists x(x \in a) \text{ et } Hx(x \in a))$.

L'introduction de ces nouveaux quantificateurs, que nous appelons *quantificateurs d'unicité*, est une *extension* de la syntaxe des langages du premier ordre. Dans la définition de ces langages (Chap. 3–4), et notamment dans les diagrammes syntaxiques de la figure 4.1, il n'est question que des quantificateurs $\forall x$ et $\exists x$ associés à chaque variable individuelle x . Nous devons donc procéder ici à une extension de cette syntaxe en introduisant, pour chaque variable individuelle x , deux nouveaux quantificateurs Hx , $\exists!x$ et en ajoutant la règle syntaxique suivante (§4.1): si A est une proposition, les expressions HxA , $\exists!xA$ sont des propositions. Celles-ci se lisent

HxA : il existe au plus un x tel que A
 $\exists!xA$: il existe un x et un seul tel que A .

La correction correspondante des diagrammes syntaxiques de la figure 4.1 est évidente. Mais cette extension doit s'accompagner d'une révision de la définition des occurrences *libres* et des occurrences *liées* de variables dans une proposition (§4.4) et de toutes les notions qui reposent sur celles-ci, par exemple celle de terme *librement substituable* à une variable dans une proposition. Car il est entendu que les nouveaux quantificateurs ont le même effet de liaison des variables que les quantificateurs ordinaires $\forall x$, $\exists x$. Nous n'avons pas besoin de récrire toutes ces définitions. Il nous suffit de dire que les quantificateurs Hx , $\exists!x$ sont traités syntaxiquement comme les quantificateurs $\forall x$, $\exists x$.

Cette extension de la syntaxe des langages du premier ordre, en particulier du langage universel $\mathcal{L}_\Omega$, est notre premier pas dans le sens d'une adaptation directe au langage mathématique usuel. Les nouveaux quantificateurs correspondent à des expressions usuelles. Nous introduirons au chapitre 13 la classe générale des « opérateurs » du langage mathématique usuel qui n'est pas couverte par la syntaxe du chapitre 4.

L'usage des nouveaux quantificateurs Hx , $\exists!x$ obéit à des schémas de déduction bien déterminés qui sont le sujet principal de ce paragraphe. Nous posons d'abord deux schémas de déduction *élémentaires*. Cela veut dire que nous procédons à une *extension* de la théorie **Lprédég**_{al}. Nous pourrions introduire ici une nouvelle dénomination telle que *Logique des prédicats avec égalité et quantificateurs d'unicité*. Cependant nous nous en abstiendrons, et nous conserverons malgré l'extension le nom de logique des prédicats avec égalité (**Lprédég**_{al}). La raison de cette conservation du nom de la théorie est que les schémas de déduction élémentaires que nous lui ajoutons peuvent être considérés comme étant « seulement » des *définitions* (Chap. 12).

S93. *Définition de HxA*

$$\Gamma \vdash HxA \Leftrightarrow \forall x \forall y [(A \text{ et } (y|x)A) \Rightarrow x = y]$$

Si x et y sont des variables distinctes et si y est réversiblement substituable à x dans A .

S94. Définition de $\exists!x A$

$$\overline{\Gamma \vdash \exists!x A \Leftrightarrow (\exists x A \text{ et } Hx A)}.$$

Le deuxième schéma se comprend facilement d'après ce qui précède. La compréhension intuitive du premier (S93) peut être plus difficile. Il faut penser à la proposition A comme étant l'énoncé d'une certaine propriété d'un objet désigné par x . Si y est librement substituable à x dans A (ce qui est le cas en particulier si y est réversiblement substituable à x dans A), la proposition $(y|x)A$ exprime que l'objet désigné par y possède cette propriété. La proposition

$$\forall x \forall y [(A \text{ et } (y|x)A) \Rightarrow x = y],$$

peut s'interpréter comme ceci: si deux objets quelconques possèdent la propriété qui est exprimée par A pour x , alors ces objets sont égaux. Ceci est une affirmation *conditionnelle* dans laquelle on ne dit pas qu'il existe des objets ayant la propriété en question, mais seulement qu'il n'en existe pas deux *distincts*. Elle est en effet équivalente, d'après S71 et S56, à

$$\text{non} \exists x \exists y [A \text{ et } (y|x)A \text{ et } x \neq y].$$

D'ailleurs, s'il n'existe pas d'objet x ayant la propriété exprimée par A , c'est-à-dire si la proposition $\text{non} \exists x A$ est vraie dans un contexte, la proposition $Hx A$ doit être vraie dans ce contexte puisqu'elle est censée exprimer qu'il existe « 0 ou 1 objets » ayant cette propriété. Or c'est bien ce que l'on peut déduire de S93. De façon précise, on peut déduire $\Gamma \vdash Hx A$ de $\Gamma \vdash \text{non} \exists x A$ suivant S93 et S67.

Exemple. Les théorèmes suivants de $\mathcal{T}_{\text{ENS}}$ sont démontrés dans la figure 11.7:

$$\vdash Hx(x \in a) \Leftrightarrow \exists b(a \subset \{b\}); \quad (1)$$

$$\vdash \exists!x(x \in a) \Leftrightarrow \exists b(a = \{b\}). \quad (2)$$

D'une part c'est un exemple d'application de S93 et S94 et d'autre part ces théorèmes sont une confirmation de la justesse de ces schémas par rapport au sens intuitif des mots « il existe au plus un », « il existe un et un seul ». Car $a \subset \{b\}$ équivaut d'après E70 à $a = \emptyset$ ou $a = \{b\}$, et un ensemble a qui vérifie cette disjonction possède bien 0 ou 1 élément (intuitivement parlant). Un ensemble a qui vérifie une égalité $a = \{b\}$ possède bien un élément et un seul.

« **Au plus un** » signifie **zéro ou un**. Le reste de ce paragraphe est consacré à quelques schémas de déduction dérivés sur les quantificateurs d'unicité. Le titre du présent article est le nom que l'on peut donner au premier de ces schémas:

S95.

$$\overline{\Gamma \vdash Hx A \Leftrightarrow (\text{non} \exists x A \text{ ou } \exists!x A)}$$

$$\overline{\Gamma \vdash Hx A \Leftrightarrow (\exists x A \Rightarrow \exists!x A)}.$$

nil

1°	$Hx(x \in a)$	hyp
2°	$\forall x \forall y ((x \in a \text{ et } y \in a) \Rightarrow x = y).$	1°, S93
3°	$a = \emptyset$	hyp
4°	$a \subset \{b\}$	E7
5°	$\exists b(a \subset \{b\})$	4°, S66
6°	$a \neq \emptyset$	hyp
7°	$\exists b(b \in a)$	6°, E13
8°	$b \in a$	hyp
9°	$x \in a$	hyp
10°	$x \in a \text{ et } b \in a$	8°, 9°
11°	$(x \in a \text{ et } b \in a) \Rightarrow x = b$	2°, S77
12°	$x = b$	10°, 11°
13°	$x \in \{b\}$	12°, E60
14°	$a \subset \{b\}$	13°, E5
15°	$\exists b(a \subset \{b\})$	14°, S66
16°	$\exists b(a \subset \{b\})$	7°, 15°, S64
17°	$\exists b(a \subset \{b\})$	5°, 16°, S35
18°	$\exists b(a \subset \{b\})$	hyp
19°	$a \subset \{b\}$	hyp
20°	$x \in a \text{ et } y \in a$	hyp
21°	$x \in \{b\} \text{ et } y \in \{b\}$	19°, 20°, E4
22°	$x = b \text{ et } y = b$	21°, E60
23°	$x = y$	22°, S85
24°	$(x \in a \text{ et } y \in a) \Rightarrow x = y$	23°, $\Rightarrow$ -intro
25°	$(x \in a \text{ et } y \in a) \Rightarrow x = y$	18°, 24°, S64
26°	$\forall x \forall y [(x \in a \text{ et } y \in a) \Rightarrow x = y]$	25°, S62
27°	$Hx(x \in a)$	26°, S93
28°	$Hx(x \in a) \Leftrightarrow \exists b(a \subset \{b\})$	17°, 27°, S40.
	$\exists! x(x \in a) \Leftrightarrow (\exists x(x \in a) \text{ et } Hx(x \in a))$	S94
	$\Leftrightarrow a \neq \emptyset \text{ et } \exists b(a \subset \{b\})$	E13, théorème (1)
	$\Leftrightarrow \exists b(a \neq \emptyset \text{ et } a \subset \{b\})$	S81
	$\Leftrightarrow \exists b(a \neq \emptyset \text{ et } (a = \emptyset \text{ ou } a = \{b\}))$	E70
	$\Leftrightarrow \exists b((a \neq \emptyset \text{ et } a = \emptyset) \text{ ou } (a \neq \emptyset \text{ et } a = \{b\}))$	S52.11
	$\Leftrightarrow \exists b(a \neq \emptyset \text{ et } a = \{b\})$	S52.15
	$\Leftrightarrow \exists b(\{b\} \neq \emptyset \text{ et } a = \{b\})$	S87.1
	$\Leftrightarrow \exists b(a = \{b\})$	E62, S60.1.

Figure 11.7 Démonstration des théorèmes de $\mathcal{T}_{\text{ENS}}$ (1) $\vdash Hx(x \in a) \Leftrightarrow \exists b(a \subset \{b\})$;(2) $\vdash \exists! x(x \in a) \Leftrightarrow \exists b(a = \{b\})$.

La démonstration peut se décomposer en trois parties, démontrant les jugements

- (a) $\Gamma \vdash HxA \Rightarrow (\text{non}\exists xA \text{ ou } \exists!xA)$;
- (b) $\Gamma \vdash \text{non}\exists xA \Rightarrow HxA$;
- (c) $\Gamma \vdash \exists!xA \Rightarrow HxA$.

La démonstration de (a), compte tenu de S94 et en appliquant S53, est un simple exercice de logique propositionnelle. C'est encore plus simple pour (c). Pour (b), en choisissant une variable y distincte de x et réversiblement substituable à x dans A , on peut écrire la démonstration

nil		
1°	$\text{non}\exists xA$	hyp
2°	$A \Rightarrow x = y$	1°, S67
3°	$(A \text{ et } (y x)A) \Rightarrow A$	S28
4°	$(A \text{ et } (y x)A) \Rightarrow x = y$	3°, 2°, S20
5°	$\forall x\forall y[(A \text{ et } (y x)A) \Rightarrow x = y]$	4°, S62
6°	HxA	5°, S93
7°	$\text{non}\exists xA \Rightarrow HxA$	6°, $\Rightarrow$ -intro.

Nous rappelons qu'une variable y distincte de x et réversiblement substituable à x dans A ne figure pas librement dans A (§10.4). On peut donc bien généraliser sur y à la ligne 5°, puisque y ne figure pas librement dans l'hypothèse de la ligne 1°.

Changements de variable liée. Le schéma de déduction suivant est semblable au schéma de changement de variable liée S76 pour les quantificateurs existentiels et universels. On peut donc étendre le schéma général de changement de variables liées — deux propositions de même squelette sont équivalentes — en y incluant les quantificateurs d'unicité Hx , $\exists!x$. Il suffit de tenir compte de ces quantificateurs dans la définition générale du squelette d'une proposition. Mais bien entendu, l'équivalence de deux propositions de même squelette contenant des quantificateurs d'unicité est vraie dans la théorie **Lprédég**, étendue avec S93 et S94, alors que si ces propositions ne contiennent que des quantificateurs universels et existentiels, cette équivalence est vraie déjà dans **Lpréd**.

S96.

$$\frac{}{\Gamma \vdash HxA \Leftrightarrow Hy(y|x)A} \quad \frac{}{\Gamma \vdash \exists!xA \Leftrightarrow \exists!y(y|x)A}.$$

Si y est réversiblement substituable à x dans A .

Démonstration. Ces déductions sans prémisses sont triviales si les variables x et y sont identiques, car alors les propositions $Hy(y|x)A$, $\exists!y(y|x)A$ sont identiques respectivement à HxA , $\exists!xA$ et ces déductions sont de la forme S46. Considérons des variables x et y distinctes. Si y est réversiblement substituable à x dans A (§10.4), A est identique à $(x|y)(y|x)A$ et x est réversiblement substituable à y dans $(y|x)A$. On peut donc écrire la démonstration suivante:

nil

$HxA \Leftrightarrow \forall x \forall y [(A \text{ et } (y x)A) \Rightarrow x = y]$	S93
$\Leftrightarrow \forall x \forall y [((y x)A \text{ et } A) \Rightarrow x = y]$	
$\Leftrightarrow \forall x \forall y [((y x)A \text{ et } (x y)(y x)A) \Rightarrow x = y]$	
$\Leftrightarrow \forall y \forall x [((y x)A \text{ et } (x y)(y x)A) \Rightarrow y = x]$	
$\Leftrightarrow Hy(y x)A$	S93.
$\exists! xA \Leftrightarrow (\exists xA \text{ et } HxA)$	S94
$\Leftrightarrow \exists y(y x)A \text{ et } Hy(y x)A$	S76, S96 (premier)
$\Leftrightarrow \exists! y(y x)A$	S94.

Autre expression des quantificateurs d'unicité. Les deux schémas suivants (S97) sont des schémas dérivés. Mais on pourrait les prendre comme schémas élémentaires, c'est-à-dire comme autre définition de HxA et $\exists! xA$ à la place de S93 et S94 et ceux-ci dériveraient alors de cette définition. C'est ce que font certains auteurs.

S97.

1.
$$\frac{\Gamma \vdash HxA \Leftrightarrow \exists y \forall x (A \Rightarrow x = y)}{\text{Si } y \text{ est distincte de } x \text{ et ne figure pas librement dans } A.}$$
2.
$$\frac{\Gamma \vdash HxA \Leftrightarrow \exists y \forall x (A \Rightarrow x = y)}{\Gamma \vdash \exists! xA \Leftrightarrow \exists y \forall x (A \Rightarrow x = y)}$$

La démonstration complète de ces déductions est relativement longue et nous allons la décomposer en démontrant des théorèmes auxiliaires. Premièrement, une démonstration de S97.1 est donnée dans la figure 11.8, mais dans le cas particulier où y est distincte de x et *réversiblement substituable* à x dans A . Cette condition est plus forte que celle du schéma S97. Cela veut dire (§10.4) non seulement que y ne figure pas librement dans A , mais encore que y est librement substituable à x dans A . Cela permet de faire le changement de variable liée de la ligne 3° et d'appliquer S93 aux lignes 7° et 26°. Cela permet aussi la particularisation de la ligne 21°. Le lecteur vérifiera que les conditions des schémas S62 et S64 sont bien satisfaites aux lignes 9°, 14°, 25°, respectivement 11°, 24°.

Le cas général de S97.1, où l'on suppose seulement que y ne figure pas librement dans A , se déduit du cas particulier démontré dans la figure 11.8 par changement de variable liée. Ayant démontré $\Gamma \vdash HxA \Leftrightarrow \exists y \forall x (A \Rightarrow x = y)$ avec une variable y qui ne figure pas librement dans A et qui satisfait à une condition supplémentaire, si y' est une autre variable qui ne figure pas librement dans A , mais qui ne satisfait pas nécessairement à cette condition supplémentaire, on peut déduire $\Gamma \vdash HxA \Leftrightarrow \exists y' \forall x (A \Rightarrow x = y')$ du théorème démontré avec y grâce au fait que la proposition $\exists y' \forall x (A \Rightarrow x = y')$ est identique à un changement de variable liée près, donc équivalente suivant S76, à la proposition $\exists y \forall x (A \Rightarrow x = y)$.

nil		
1°	HxA	hyp
2°	$\exists xA$	hyp
3°	$\exists y(y x)A$	2°, S76
4°	$(y x)A$	hyp
5°	A	hyp
6°	A et $(y x)A$	4°, 5°
7°	$x = y$	1°, 6°, S93, S65
8°	$A \Rightarrow x = y$	7°, $\Rightarrow$ -intro
9°	$\forall x(A \Rightarrow x = y)$	8°, S62
10°	$\exists y \forall x(A \Rightarrow x = y)$	9°, S66
11°	$\exists y \forall x(A \Rightarrow x = y)$	3°, 10°, S64
12°	$\text{non} \exists xA$	hyp
13°	$A \Rightarrow x = y$	12°, S67
14°	$\forall x(A \Rightarrow x = y)$	13°, S62
15°	$\exists y \forall x(A \Rightarrow x = y)$	14°, S66
16°	$\exists y \forall x(A \Rightarrow x = y)$	11°, 15°, S35
17°	$\exists y \forall x(A \Rightarrow x = y)$	hyp
18°	$\exists z \forall x(A \Rightarrow x = z)$	17°, S76
19°	$\forall x(A \Rightarrow x = z)$	hyp
20°	$A \Rightarrow x = z$	19°, S65
21°	$(y x)A \Rightarrow y = z$	19°, S61
22°	$(A \text{ et } (y x)A) \Rightarrow (x = z \text{ et } y = z)$	20°, 21°, S29
23°	$(A \text{ et } (y x)A) \Rightarrow (x = y)$	22°, S85, S48
24°	$(A \text{ et } (y x)A) \Rightarrow (x = y)$	18°, 23°, S64
25°	$\forall x \forall y[(A \text{ et } (y x)A) \Rightarrow (x = y)]$	24°, S62
26°	HxA	25°, S93
27°	$HxA \Leftrightarrow \exists y \forall x(A \Rightarrow x = y)$	16°, 26°, S40.

Figure 11.8 Démonstration de $\vdash HxA \Leftrightarrow \exists y \forall x(A \Rightarrow x = y)$ pour y distincte de x et réversiblement substituable à x dans A . On suppose que z (ligne 18°) est une variable distincte de x et y et ne figure pas librement dans A .

Pour démontrer S97.2, il est commode d'utiliser les schémas de déduction auxiliaires (a) et (b) suivants, qui ne serviront qu'à cette occasion. Ils sont démontrés dans la figure 11.9.

- (a) $\mid \Gamma \vdash (\exists xA \text{ et } \forall x(A \Rightarrow x = T)) \Rightarrow (T|x)A.$
- (b) $\mid \Gamma \vdash (\exists xA \text{ et } \forall x(A \Rightarrow x = T)) \Leftrightarrow \forall x(A \Leftrightarrow x = T).$

Si T est un terme librement substituable à x dans A
et si x ne figure pas librement dans T .

nil		
(a) 1°	$\exists xA$ et $\forall x(A \Rightarrow x = T)$	hyp
2°	$\exists xA$	1°
3°	A	hyp
4°	$x = T$	1°, et-élim, S65, 3°
5°	$(x x)A$	identique à 3°
6°	$(T x)A$	4°, 5°, S83
7°	$(T x)A$	2°, 6°, S64
8°	$(\exists xA \text{ et } \forall x(A \Rightarrow x = T)) \Rightarrow (T x)A$	7°, $\Rightarrow$ -intro.
(b) 1°	$\exists xA$ et $\forall x(A \Rightarrow x = T)$	hyp
2°	$A \Rightarrow x = T$	1°, et-élim, S65
3°	$x = T$	hyp
4°	$(T x)A$	1°, (a), mod ponens
5°	$(x x)A$	3°, 4°, S83
6°	A	identique à 5°
7°	$x = T \Rightarrow A$	6°, $\Rightarrow$ -intro
8°	$A \Leftrightarrow x = T$	2°, 7°, $\Leftrightarrow$ -intro
9°	$\forall x(A \Leftrightarrow x = T)$	8°, S62
10°	$\forall x(A \Leftrightarrow x = T)$	hyp
11°	$A \Leftrightarrow x = T$	10°, S65
12°	$\exists x(x = T)$	S89
13°	$x = T$	hyp
14°	A	11°, 13°, S39
15°	$\exists xA$	14°, S66
16°	$\exists xA$	12°, 15°, S64
17°	$A \Rightarrow x = T$	11°, $\Leftrightarrow$ -élim
18°	$\forall x(A \Rightarrow x = T)$	17°, S62
19°	$\exists xA$ et $\forall x(A \Rightarrow x = T)$	16°, 18°
20°	$(\exists xA \text{ et } \forall x(A \Rightarrow x = T)) \Leftrightarrow \forall x(A \Leftrightarrow x = T)$	9°, 19°, S40.

Figure 11.9 On suppose que x ne figure pas librement dans T et que T est librement substituable à x dans A .

En nous servant de ces schémas auxiliaires, nous pouvons démontrer comme suit le schéma S97.2, *dans le cas particulier* où y satisfait à la condition supplémentaire d'être librement substituable à x dans A .

$\exists! xA \Leftrightarrow (\exists xA \text{ et } HxA)$	S94
$\Leftrightarrow (\exists xA \text{ et } \exists y \forall x(A \Rightarrow x = y))$	S97.1
$\Leftrightarrow \exists y(\exists xA \text{ et } \forall x(A \Rightarrow x = y))$	S81
$\Leftrightarrow \exists y \forall x(A \Leftrightarrow x = y)$	schéma (b)

Le cas général du schéma S97.2, se déduit du cas particulier par changement de variable liée, comme nous l'avons vu pour S97.1. Car si y et y' sont deux

variables distinctes de x et ne figurant pas librement dans A , les propositions $\exists y \forall x (A \Leftrightarrow x = y)$ et $\exists y' \forall x (A \Leftrightarrow x = y')$ sont identiques à un changement de variable liée près.

Cas d'unicité triviale. Lorsqu'une proposition A est équivalente dans un contexte à une égalité $x = T$, où T est un terme dans lequel x ne figure pas librement, il est intuitivement évident qu'il existe un objet x et un seul qui vérifie A , à savoir l'objet désigné par T . C'est ce que nous appelons un cas d'unicité *triviale*. Le schéma suivant décrit les déductions de ce genre.

S98.

$$\frac{\Gamma \vdash A \Rightarrow x = T}{\Gamma \vdash Hx A} \quad \frac{\Gamma \vdash A \Leftrightarrow x = T}{\Gamma \vdash \exists! x A} \quad \text{Si } x \text{ ne figure librement ni dans } \Gamma \text{ ni dans } T.$$

Pour la première de ces déductions on peut écrire la démonstration ci-dessous, où l'on suppose que x satisfait aux conditions du schéma et que y est une variable distincte de x et ne figurant pas librement dans A :

$$\begin{array}{l|ll} \Gamma & & \\ 1^\circ & A \Rightarrow x = T & \text{base} \\ 2^\circ & \forall x (A \Rightarrow x = T) & 1^\circ, \text{S62} \\ 3^\circ & \exists y \forall x (A \Rightarrow x = y) & 2^\circ, \text{S63} \\ 4^\circ & Hx A & 3^\circ, \text{S97.1.} \end{array}$$

La justification détaillée de la troisième ligne est la suivante: premièrement, la proposition de la ligne 2° est identique à $(T|y) \forall x (A \Rightarrow x = y)$ parce que y est distincte de x et ne figure pas librement dans A ; deuxièmement, T est librement substituable à y dans $\forall x (A \Rightarrow x = y)$ parce que x ne figure pas librement dans T . La deuxième déduction de S98 se démontre de manière tout à fait semblable en utilisant S97.2.

Exemples. Les théorèmes de $\mathcal{T}_{\text{ENS}}$

$$\vdash Hx(x \in \emptyset), \quad \vdash \exists! x(x \in \{a\}), \quad \vdash \exists! x \forall a(a \notin x)$$

peuvent se déduire suivant S98 des théorèmes

$$\begin{array}{ll} \vdash x \in \emptyset \Rightarrow x = a & (\text{E8}) \\ \vdash x \in \{a\} \Leftrightarrow x = a & (\text{E60}) \\ \vdash \forall a(a \notin x) \Leftrightarrow x = \emptyset & (\text{E13}). \end{array}$$

Preuve d'égalité par unicité. On peut déduire que deux objets sont égaux s'ils vérifient tous les deux une condition dont on sait qu'elle est satisfaite par au plus un objet. C'est ce qu'exprime le schéma suivant.

S99.

$$\frac{\Gamma \vdash Hx A \quad \Gamma \vdash (T|x) A \quad \Gamma \vdash (U|x) A}{\Gamma \vdash T = U} \quad \text{Si } T \text{ et } U \text{ sont librement substituable à } x \text{ dans } A.$$

Car si T et U sont librement substituables à x dans A , alors en prenant une variable auxiliaire y distincte de x et ne figurant librement ni dans Γ , ni dans A , ni dans T , ni dans U , on peut écrire la démonstration

Γ		
1 ^o	HxA	<i>base</i>
2 ^o	$(T x)A$	<i>base</i>
3 ^o	$(U x)A$	<i>base</i>
4 ^o	$\exists y \forall x (A \Rightarrow x = y)$	1 ^o , S97
5 ^o	$\forall x (A \Rightarrow x = y)$	hyp
6 ^o	$(T x)A \Rightarrow T = y$	5 ^o , S61
7 ^o	$(U x)A \Rightarrow U = y$	5 ^o , S61
8 ^o	$T = y$	2 ^o , 6 ^o
9 ^o	$U = y$	3 ^o , 7 ^o
10 ^o	$T = U$	8 ^o , 9 ^o
11 ^o	$T = U$	4 ^o , 10 ^o , S64.

La ligne 6^o se déduit en effet de la ligne 5^o par S61 parce que la proposition $(T|x)A \Rightarrow T = y$ est identique à $(T|x)(A \Rightarrow x = y)$, du fait que y est distincte de x , et parce que T est librement substituable à x dans $(A \Rightarrow x = y)$ du fait que T est librement substituable à x dans A . La ligne 7^o se justifie de manière semblable. On peut écrire la ligne 11^o parce que y ne figure librement ni dans la proposition $T = U$, ni dans Γ .

Substitutivité de l'équivalence. Les schémas de déduction S100.2–3 ci-dessous, valables si x ne figure pas librement dans Γ , sont les analogues, pour les quantificateurs d'unicité, des schémas S70 pour les quantificateurs existentiels et universels. Ils permettent d'étendre le schéma général de substitutivité de l'équivalence (§10.1) aux propositions dans lesquelles il y a des quantificateurs d'unicité. On rappelle ce schéma général: d'un jugement $\Gamma \vdash C = C'$, on peut déduire le jugement $\Gamma \vdash R \Leftrightarrow R'$ si R' est la proposition obtenue en remplaçant une occurrence de la proposition C dans R par C' et si cette occurrence de C dans R n'est pas dominée par un quantificateur dont la variable figure librement dans Γ . Cette déduction peut se faire aussi dorénavant, dans Lprédég, s'il y a des quantificateurs d'unicité parmi les quantificateurs qui dominent C dans R .

S100.

1.
$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash HxB \Rightarrow HxA}$$
2.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash HxA \Leftrightarrow HxB}$$
3.
$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash \exists! xA \Leftrightarrow \exists! xB}$$

Si x ne figure pas librement dans Γ .

Le schéma S100.1 n'est pas intuitivement évident pour tout le monde, au premier abord. L'inversion du sens de l'implication par rapport à A et B dans

la conclusion peut surprendre. Mais considérons l'exemple de la déduction

$$\frac{\Gamma \vdash x \in a \Rightarrow x \in b}{\Gamma \vdash \text{Hx}(x \in b) \Rightarrow \text{Hx}(x \in a)}. \quad (3)$$

Si x ne figure pas librement dans Γ , on peut déduire $\Gamma \vdash a \subset b$ de la prémisse de (3) d'après E5. Or si $a \subset b$ est vraie, il est vrai que si b possède au plus un élément alors a possède au plus un élément. L'implication de la conclusion va bien dans ce sens.

Pour démontrer une déduction de la forme S100.1, on prend une variable auxiliaire y distincte de x et ne figurant librement ni dans A , ni dans B , ni dans les hypothèses Γ . On peut alors écrire la démonstration

Γ		
1°	$A \Rightarrow B$	<i>base</i>
2°	HxB	<i>hyp</i>
3°	$\exists y \forall x (B \Rightarrow x = y)$	2°, S97
4°	$\forall x (B \Rightarrow x = y)$	<i>hyp</i>
5°	$B \Rightarrow x = y$	4°, S65
6°	$A \Rightarrow x = y$	1°, 5°, S20
7°	$\forall x (A \Rightarrow x = y)$	6°, S62
8°	$\exists y \forall x (A \Rightarrow x = y)$	7°, S66
9°	HxA	8°, S97
10°	HxA	3°, 9°, S64
11°	$\text{HxB} \Rightarrow \text{HxA}$	10°, $\Rightarrow$ -intro.

Le schéma S100.2 se démontre de façon évidente à partir de S100.1 par élimination et introduction de $\Leftrightarrow$. Le schéma S100.3 dérive de S100.2 d'après la démonstration suivante, où l'on admet que x ne figure pas librement dans Γ :

Γ		
1°	$A \Leftrightarrow B$	<i>base</i>
2°	$\text{HxA} \Leftrightarrow \text{HxB}$	1°, S100.2
3°	$\exists x A \Leftrightarrow \exists x B$	1°, S70
4°	$\exists ! x A \Leftrightarrow (\exists x A \text{ et } \text{HxA})$	S94
5°	$\Leftrightarrow (\exists x B \text{ et } \text{HxB})$	2°, 3°
6°	$\Leftrightarrow \exists ! x B$	S94.

11.5 Exercices

Exercice 1. Soit $*$ un symbole fonctionnel binaire. Démontrer le théorème

$$\vdash \text{Ha} \forall x (x * a = x \text{ et } a * x = x)$$

de Lprédég, affirmant l'unicité d'un objet a « neutre » pour l'opération désignée par $*$.

Exercice 2. Démontrer les théorèmes suivants de $\mathcal{T}_{\text{ENS}}$:

$$\begin{aligned} &\vdash \text{Hu}\forall x(x \in u) \\ &\vdash \exists!x\forall a(a \notin x). \end{aligned}$$

Le premier affirme qu'il existe au plus un ensemble u « universel », c'est-à-dire un ensemble dont tout objet x soit élément. Nous verrons qu'en vertu d'un schéma de déduction élémentaire de $\mathcal{T}_{\text{ENS}}$ que nous introduirons au chapitre 14, un tel ensemble u n'existe pas.

Exercice 3. Soit $P(x)$ une proposition exprimant une propriété de x . Démontrer le théorème

$$\vdash \text{Ha}\forall x(P(x) \Leftrightarrow x \in a)$$

de $\mathcal{T}_{\text{ENS}}$, affirmant: il existe au plus un ensemble a ayant pour éléments tous les objets x qui ont la propriété $P(x)$ et seulement ceux-là. Donner un exemple d'une propriété $P(x)$ pour laquelle un tel ensemble a n'existe pas.

Exercice 4. Démontrer le théorème suivant de **Lprédég**:

$$\vdash \exists a\exists b(a \neq b) \Leftrightarrow \forall a\exists b(a \neq b).$$

Ce théorème affirme que s'il y a deux objets différents ($\exists a\exists b(a \neq b)$), dans l'univers du discours, alors pour tout objet a il existe un objet différent de a , et réciproquement. Le théorème $\vdash \forall a\exists b(a \neq b) \Rightarrow \exists a\exists b(a \neq b)$ est immédiat (S68). Pour l'implication inverse, on mettra en forme le raisonnement suivant: soient a et b deux objets différents (on suppose qu'il en existe); alors tout objet x est différent de a ou différent de b (propriétés de l'égalité); donc pour tout objet x , il existe un objet y différent de x . On a démontré ainsi $\forall x\exists y(y \neq x)$ et il n'y a plus qu'à changer de variables liées.

Exercice 5. Soient A, B, C des propositions et x une variable. Démontrer les théorèmes de **Lprédég**:

$$\begin{aligned} &\vdash \text{HxA} \Rightarrow [(\exists x(A \text{ et } B) \text{ et } \exists x(A \text{ et } C)) \Rightarrow \exists x(B \text{ et } C)] \\ &\vdash (\forall xA \text{ et } \exists!xB) \Rightarrow \exists!x(A \text{ et } B). \end{aligned}$$

Pour le premier, l'élimination des quantificateurs de $\exists x(A \text{ et } B)$, $\exists x(A \text{ et } C)$ nécessite un changement de variable liée dans l'une de ces propositions.

Exercice 6. Construire un diagramme de points et de flèches (§10.3, Exercice) vérifiant les propositions $\exists!x(\exists!y(x \rightarrow y))$ et $\text{non}\exists!y(\exists!x(x \rightarrow y))$. Traduire ces propositions en français ou dans une autre langue naturelle sans utiliser de variable.

Extensions définitionnelles

12.1 Axiomes et schémas définitionnels

Le but de ce bref chapitre est de discuter la forme et le rôle des nombreuses *définitions* qui font partie de tout exposé d'une théorie mathématique, et plus généralement de tout exposé scientifique prenant la peine de définir soigneusement son vocabulaire spécifique. Vu le rôle important que jouent les définitions dans ces exposés, il peut paraître surprenant que le concept de définition ne figure pas parmi les concepts logiques fondamentaux présentés au chapitre 5, à savoir: théories, axiomes, déductions, théorèmes, démonstrations.

Mais nous avons déjà signalé à deux reprises que les définitions n'étaient que des axiomes d'une forme particulière. Nous l'avons fait au paragraphe 9.2 en introduisant le langage universel $\mathcal{L}_\Omega$. Chaque définition s'accompagne de l'introduction d'un nouveau symbole, d'une nouvelle forme d'expression de façon générale, et le langage $\mathcal{L}_\Omega$ tient à disposition une infinité de symboles de toutes catégories pour ces extensions. En théorie des ensembles (§11.3), nous avons aussi présenté l'axiome de l'ensemble à un seul élément (E59), $\vdash \{a\} = \{a, a\}$, comme étant une définition. Nous dirons qu'il s'agit d'un axiome *définitionnel*.

Nous avons présenté aussi les schémas de déduction élémentaires S93 et S94 sur les quantificateurs d'unicité $\text{H}x$, $\exists!x$ comme étant des définitions et nous parlerons aussi de schémas de déduction *définitionnels*.

Exemple 1. Nous reprenons l'exemple de l'axiome E59,

$$\vdash \{a\} = \{a, a\},$$

pour mettre en évidence la forme et les propriétés logiques d'une définition.

Pour apprécier la portée ou l'effet d'un axiome, il faut comparer les deux théories que l'on obtient: d'une part en incluant l'axiome en question, d'autre part en l'excluant. De façon précise, dans notre exemple, nous comparerons les deux théories suivantes:

($\mathcal{T}_0$) la théorie ayant pour axiomes tous les axiomes de $\mathcal{T}_{\text{ENS}}$ que nous avons introduits *avant* E59;

($\mathcal{T}_1$) la théorie obtenue en adjoignant l'axiome E59 à $\mathcal{T}_0$.

Les deux théories que nous comparons sont deux stades de développement de $\mathcal{T}_{\text{ENS}}$: celui qui précède immédiatement l'introduction de E59 et celui qui suit immédiatement cette introduction. La théorie $\mathcal{T}_1$ est une extension de $\mathcal{T}_0$ (§5.7) et comme elle résulte de l'adjonction d'un axiome définitionnel, on parle d'une *extension définitionnelle*. Les deux théories admettent les mêmes déductions élémentaires, à savoir toutes celles de la logique des prédicats avec égalité. Nous n'avons pas encore introduit à ce stade les déductions élémentaires spécifiques de $\mathcal{T}_{\text{ENS}}$.

Pour la suite de la discussion, nous utiliserons le symbole ν (lettre grecque nu) comme symbole fonctionnel unaire préfixé, à la place des parenthèses $\{ \}$ du terme $\{a\}$. Donc nous écrivons νa au lieu de $\{a\}$. Nous admettons en somme que ν est le symbole du langage $\mathcal{L}_\Omega$ réservé pour noter les ensembles à un seul élément et qui fait son apparition dans $\mathcal{T}_{\text{ENS}}$ pour la première fois dans l'axiome E59, que nous écrivons:

$$\vdash \nu a = \{a, a\}.$$

La première chose à relever est que le symbole ν est *nouveau* par rapport à la théorie $\mathcal{T}_0$, en ce sens qu'il ne figure dans aucun des axiomes de cette théorie. Une extension définitionnelle introduit toujours un symbole nouveau dans ce sens.

Sur le plan de la déduction logique, une propriété évidente de l'axiome E59 est qu'il permet *d'éliminer le symbole ν* dans n'importe quelle proposition grâce aux propriétés de substitutivité de l'égalité et de l'équivalence. De façon précise, si l'on remplace chaque occurrence de terme de la forme νT dans une proposition R par le terme $\{T, T\}$, on obtient une proposition R' dans laquelle il n'y a pas de symbole ν et telle que

$$\vdash R \Leftrightarrow R'$$

est un théorème de $\mathcal{T}_1$. Cette propriété logique découle, comme nous l'avons dit, de la substitutivité de l'égalité et de l'équivalence. Pour rappeler comment les deux se combinent, nous renvoyons le lecteur à l'exemple de la figure 11.4. Nous pouvons donc énoncer la propriété logique suivante de l'axiome E59:

(a) Pour toute proposition R , on peut construire une proposition R' sans symbole ν telle que

$$\frac{\vdash \nu a = \{a, a\}}{\vdash R \Leftrightarrow R'}$$

soit une déduction de $\mathcal{T}_0$.

Comme on le voit, cette propriété ne tient qu'à la forme de l'axiome E59, à savoir qu'il est une égalité de la forme $\nu a = U$, où U est un terme dans lequel

le nouveau symbole ν ne figure pas — c'est ce qui permet de l'éliminer comme nous venons de le voir. Le fait que le symbole ν soit « nouveau », c'est-à-dire qu'il ne figure pas dans les autres axiomes de $\mathcal{T}_1$ (ceux de $\mathcal{T}_0$), n'intervient pas dans la propriété (a). Mais ce fait intervient bien dans la seconde propriété logique d'un axiome définitionnel, qui est la suivante dans cet exemple :

(b) Si $\Gamma \vdash A$ est un théorème de $\mathcal{T}_1$ et si le symbole ν ne figure ni dans Γ , ni dans A , alors $\Gamma \vdash A$ est un théorème de $\mathcal{T}_0$.

En d'autres termes, si l'on s'est servi de l'axiome E59 ou de théorèmes qui en sont une conséquence pour démontrer un théorème $\Gamma \vdash A$ de $\mathcal{T}_1$ dans lequel il n'y a pas de symbole ν , ce qui veut dire qu'on aura introduit et éliminé des symboles ν dans cette démonstration, alors ce théorème peut être démontré dans $\mathcal{T}_0$, donc sans utiliser E59 ou ses conséquences.

Cette propriété est intuitivement assez évidente, mais elle est nettement plus difficile à prouver que la propriété (a) et une preuve rigoureuse sort du cadre de cet ouvrage. Il faut bien voir que les propriétés (a) et (b) sont des « métathéorèmes » (§2.4). Ce sont des « théorèmes sur les théorèmes de certaines théories » et ces théorèmes sur les théorèmes sont formulés dans le métalangage (le français et les métasymboles) que nous utilisons pour parler du langage formel $\mathcal{L}_\Omega$, des jugements, des théories, etc. La preuve de certains métathéorèmes peut demander des raisonnements assez complexes dans le métalangage. À ce moment-là, on a intérêt à reprendre toute la question différemment, c'est-à-dire à la traiter dans un métalangage *formel* dans lequel on raisonne suivant les schémas de la logique. C'est ce qui se fait dans cette « théorie des théories » qu'on appelle la *Logique mathématique*.

En considérant les deux propriétés (a) et (b), nous voyons ceci : si $\vdash R$ est un théorème de $\mathcal{T}_1$, et si R' est une proposition sans symbole ν équivalente à R au sens de (a), alors $\vdash R'$ est un théorème de $\mathcal{T}_0$. Car R' est un théorème de $\mathcal{T}_1$ (S39) sans symbole ν .

C'est dans ce sens précis qu'une extension définitionnelle, conformément à notre intuition, ne « sert à rien » du point de vue logique et n'est qu'une notation abrégiate. Mais d'un point de vue pratique, les abréviations introduites par les définitions sont indispensables. Actuellement, à peu près toutes les théories mathématiques sont des extensions définitionnelles de la théorie des ensembles. Si l'on voulait éliminer dans certains de leurs théorèmes tous les symboles qui ne figurent pas dans les axiomes non définitionnels de la théorie des ensembles, on obtiendrait des propositions d'une longueur qui rendrait leur écriture pratiquement impossible, pour ne pas parler de leur lecture.

Par exemple, si R est la proposition « F est une fonction », une proposition R' équivalente à R dans $\mathcal{T}_{\text{ENS}}$ et n'utilisant aucun des symboles introduits dans les définitions de $\mathcal{T}_{\text{ENS}}$ de cet ouvrage est la proposition :

$$\begin{aligned} &\forall z(z \in F \Rightarrow \exists x \exists y(z = \{\{x, x\}, \{x, y\}\})) \text{ et} \\ &\forall x \forall y \forall y'((\{\{x, x\}, \{x, y\}\} \in F \text{ et } \{\{x, x\}, \{x, y'\}\} \in F) \Rightarrow y = y'). \end{aligned}$$

On n'imagine pas utiliser pareille expression à la place de chaque occurrence du mot « fonction » dans le langage mathématique. Les définitions sont comme les routines ou sous-programmes en informatique. On peut théoriquement s'en passer, mais non pratiquement.

Extensions conservatives. La propriété (b) s'exprime en disant que la théorie $\mathcal{T}_1$ est une extension *conservative* de la théorie $\mathcal{T}_0$. De façon générale, une théorie $\mathcal{T}'$ est appelée une *extension conservative* d'une théorie $\mathcal{T}$ si $\mathcal{T}'$ est une extension de $\mathcal{T}$ (§5.7) et si elle satisfait à la condition suivante: chaque théorème de $\mathcal{T}'$ dans lequel il n'y pas de symbole « nouveau » par rapport à $\mathcal{T}$ est un théorème de $\mathcal{T}$. Par « symbole nouveau par rapport à $\mathcal{T}$ », on entend un symbole qui ne figure pas dans les axiomes et les schémas de déduction élémentaires de $\mathcal{T}$ et qui n'est pas une variable ou un quantificateur ($\exists x$, $\forall x$) ou un connecteur logique (et, ou, non, $\Rightarrow$, $\Leftrightarrow$). Donc ce peut être une constante individuelle ou un symbole fonctionnel ou un symbole relationnel ou une nouvelle espèce de quantificateur (par exemple les nouveaux quantificateurs Hx , $\exists!x$), ou d'autres « opérateurs » dont nous parlerons au chapitre 13.

Propriété. Si $\mathcal{T}'$ est une extension conservative de $\mathcal{T}$ et si $\mathcal{T}$ n'est pas contradictoire, alors $\mathcal{T}'$ n'est pas contradictoire.

Pour prouver cela, nous montrons la contraposée: si $\mathcal{T}'$ est contradictoire, alors $\mathcal{T}$ est contradictoire. Supposons en effet que $\mathcal{T}'$ soit une extension conservative de $\mathcal{T}$ et soit contradictoire (§5.6). Toute proposition est alors vraie dans $\mathcal{T}'$. En particulier, si A est une proposition qui ne contient aucun symbole nouveau par rapport à $\mathcal{T}$, les jugements $\vdash A$, $\vdash \text{non}A$ sont des théorèmes de $\mathcal{T}'$, donc des théorèmes de $\mathcal{T}$ par définition d'une extension conservative, de sorte que $\mathcal{T}$ est contradictoire.

On voit que l'adjonction d'un axiome définitionnel à une théorie ne peut pas être la cause d'une contradiction. Dans l'exemple 1, toute contradiction qui pourrait être découverte dans $\mathcal{T}_1$ aurait sa cause dans les axiomes de $\mathcal{T}_0$. On peut dire que l'adjonction d'axiomes définitionnels est un mécanisme d'extension *sûr* du point de vue de la non-contradiction.

Exemple 2. La notion d'axiome définitionnel d'une théorie $\mathcal{T}$ n'a de sens que par rapport à un *ordre* dans lequel sont donnés les axiomes de $\mathcal{T}$, puisqu'un axiome définitionnel est comparé à ceux qui le *précèdent*: il contient un symbole qui ne figure pas dans ceux-ci. On peut donc comparer la théorie $\mathcal{T}_0$ dont les axiomes sont les axiomes de $\mathcal{T}$ qui précèdent l'axiome définitionnel en question avec la théorie $\mathcal{T}_1$ dont les axiomes sont ceux de $\mathcal{T}_0$ et l'axiome définitionnel considéré.

Le premier axiome de $\mathcal{T}_{\text{ENS}}$ que nous avons introduit dans cet ouvrage, celui de l'inclusion (E1)

$$\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y),$$

est *déjà* un axiome définitionnel. Dans ce cas, la théorie $\mathcal{T}_0$ est la théorie sans axiomes **Lprédég**al et la théorie $\mathcal{T}_1$ est celle dont le seul axiome est celui-ci.

Le « nouveau symbole » introduit dans cet axiome est le symbole relationnel binaire d'inclusion $\subset$. Tout ce que nous avons dit dans l'exemple 1 peut être répété ici. En vertu de la substitutivité de l'équivalence, l'axiome **E1** permet d'éliminer le symbole $\subset$ dans toute proposition, en ce sens que pour toute proposition **R** on peut construire une proposition **R'** sans symbole d'inclusion telle que

$$\frac{\vdash x \subset y \Leftrightarrow \forall a(a \in x \Rightarrow a \in y)}{\vdash R \Leftrightarrow R'}$$

soit une déduction de $\mathcal{T}_0$. En outre, si **R** est un théorème de $\mathcal{T}_1$, **R'** est un théorème de $\mathcal{T}_0$ (extension conservative).

Le schéma de déduction **E3** indique comment éliminer le symbole d'inclusion dans une proposition. Cela consiste à remplacer chaque sous-proposition de la forme $T_1 \subset T_2$ par l'expression $\forall z(z \in T_1 \Rightarrow z \in T_2)$, où z est une variable qui ne figure pas librement dans les termes T_1, T_2 .

Exemple 3. Le schéma de déduction élémentaire **S93**:

$$\frac{}{\Gamma \vdash HxA \Leftrightarrow \forall x \forall y [(A \text{ et } (y|x)A) \Rightarrow x = y]},$$

avec la condition que y soit distincte de x et réversiblement substituable à x dans A , est un schéma de déduction *définitionnel* de **Lprédég**_{al} par rapport aux autres schémas de déduction élémentaires qui le précèdent. Dans ce cas, c'est une *classe* de nouveaux symboles qui est introduite. Ce sont tous les quantificateurs d'unicité Hx (il y en a un pour chaque variable x). Ces symboles ne figurent pas dans les schémas de déduction précédents. La théorie $\mathcal{T}_0$ est ici la théorie **Lprédég**_{al} telle qu'elle est définie au début du chapitre 11 (appelons-la **Lprédég**_{al0}) et la théorie $\mathcal{T}_1$ (**Lprédég**_{al1}) est la théorie obtenue en ajoutant à $\mathcal{T}_0$ le schéma **S93**.

Soit **R** une proposition contenant n quantificateurs d'unicité $Hx_1, \dots, Hx_n$ et soient $Hx_1A_1, \dots, Hx_nA_n$ les n sous-propositions de **R** qui commencent par ces quantificateurs. Désignons par B_i (pour $i = 1, \dots, n$) la proposition

$$\forall x_i \forall y_i [(A_i \text{ et } (y_i|x_i)A_i) \Rightarrow x_i = y_i],$$

où y_i est une variable distincte de x_i et réversiblement substituable à x_i dans A_i . Alors on peut construire une proposition **R'** sans quantificateurs d'unicité telle que

$$\frac{\vdash Hx_1A_1 \Leftrightarrow B_1 \quad \dots \quad \vdash Hx_nA_n \Leftrightarrow B_n}{\vdash R \Leftrightarrow R'} \quad (1)$$

soit une déduction de **Lprédég**_{al0}. Si en outre **R** est un théorème de **Lprédég**_{al1}, **R'** est un théorème de **Lprédég**_{al0} (extension conservative). La démonstration de la déduction (1) procède uniquement par substitutivité de l'équivalence au sens de **Lprédég**_{al0}, c'est-à-dire en utilisant seulement les schémas **S50** et **S70**. On élimine un à un tous les quantificateurs Hx_i en remplaçant Hx_iA_i par B_i ,

mais en faisant cela chaque fois pour une sous-proposition $Hx_i A_i$ qui n'est pas dominée par un autre quantificateur Hx_j pas encore éliminé.

Par rapport à cette théorie $\mathcal{T}_1$, la théorie $\mathcal{T}_2$ obtenue en ajoutant le schéma de déduction S94 :

$$\Gamma \vdash \exists! x A \Leftrightarrow (\exists x A \text{ et } Hx A)$$

est encore une extension définitionnelle, introduisant la classe des nouveaux symboles $\exists! x$.

Forme générale d'un axiome définitionnel. Il y a plusieurs formes d'axiomes définitionnels. Nous commençons par celles des exemples 1 et 2.

Première forme. La première forme d'axiome définitionnel (Exemple 1) est une égalité

$$s x_1 \cdots x_n = T,$$

où s est un nouveau symbole fonctionnel n -aire, $x_1, \dots, x_n$ sont des variables distinctes et T est un terme dans lequel le symbole s ne figure pas et dans lequel il n'y a pas de variable libre distincte de $x_1, \dots, x_n$. Cela ne veut pas dire que les variables $x_1, \dots, x_n$ doivent toutes figurer librement dans T , mais seulement qu'il ne doit pas y en avoir d'autres. On peut exprimer cela en disant que $x_1, \dots, x_n$ sont les seules variables libres de l'axiome. Par « nouveau symbole », nous entendons naturellement un symbole qui ne figure pas dans les axiomes de la théorie qui précèdent l'axiome considéré.

Deuxième forme. La deuxième forme d'axiome définitionnel (Exemple 2) est une équivalence

$$r x_1 \cdots x_n \Leftrightarrow A,$$

où r est un nouveau symbole relationnel n -aire, $x_1, \dots, x_n$ sont des variables distinctes et A est une proposition dans laquelle le symbole r ne figure pas et dans laquelle il n'y a pas de variable libre distincte de $x_1, \dots, x_n$. Ces variables doivent être les seules variables libres de l'axiome.

Troisième et quatrième formes. Ces deux formes d'axiomes définitionnels sont une variante des précédentes que l'on rencontre très souvent. Ce sont des égalités ou des équivalences de la forme ci-dessus mais *conditionnelles*. Nous entendons par là que ce sont des implications de la forme

$$\begin{aligned} C &\Rightarrow (s x_1 \cdots x_n = T), \\ C &\Rightarrow (r x_1 \cdots x_n \Leftrightarrow A), \end{aligned}$$

où C est une proposition dans laquelle le nouveau symbole fonctionnel s , respectivement le nouveau symbole relationnel r , ne figure pas et dans laquelle il n'y a pas non plus de variable libre distincte des variables $x_1, \dots, x_n$. Celles-ci sont les seules variables libres de l'axiome.

Propriété. Les définitions de la troisième ou de la quatrième forme n'ont pas tout à fait les mêmes propriétés que celles de la première et de la deuxième,

que nous avons formulées dans l'exemple 1 sous les lettres (a) et (b). Si $\mathcal{T}_1$ est une extension d'une théorie $\mathcal{T}_0$ résultant par exemple de l'adjonction d'un axiome de la troisième forme dont le symbole fonctionnel **s** ne figure pas dans les axiomes de $\mathcal{T}_0$, alors $\mathcal{T}_1$ est bien une extension conservative de $\mathcal{T}_0$, mais en général l'axiome ne permet pas d'éliminer le symbole **s** dans n'importe quelle proposition **R**, c'est-à-dire qu'on ne peut pas construire pour toute proposition **R** une proposition **R'** sans symbole **s** telle que

$$\frac{\vdash C \Rightarrow (\mathbf{s}x_1 \cdots x_n = \mathbf{T})}{\vdash R \Leftrightarrow R'} \quad (2)$$

soit une déduction de $\mathcal{T}_0$. Mais on peut le faire pour toute proposition **R** telle que $\vdash R$ soit un théorème de $\mathcal{T}_1$.

Exemple 4. Des exemples de définitions de la troisième ou de la quatrième forme viendront plus loin en théorie des ensembles. Le lecteur peut penser ici à certaines définitions qu'il a rencontrées en mathématique, par exemple, en théorie des nombres réels, une définition de $\sqrt{x}$ telle que: si x est un nombre réel positif, on désigne par $\sqrt{x}$ l'unique nombre réel positif a tel que $a^2 = x$. Cette définition peut se formaliser par l'axiome

$$(x \in \mathbb{R} \text{ et } x \geq 0) \Rightarrow \sqrt{x} = \mathbf{S}a(a \in \mathbb{R} \text{ et } a \geq 0 \text{ et } a^2 = x). \quad (3)$$

Cette définition est de la troisième forme avec $n = 1$, **s** étant le symbole fonctionnel unaire $\sqrt{}$ et le terme **T** étant l'expression

$$\mathbf{S}a(a \in \mathbb{R} \text{ et } a \geq 0 \text{ et } a^2 = x). \quad (4)$$

Les opérateurs **Sx**, dits opérateurs de *sélection* ou *sélecteurs*, seront introduits plus loin (Chap. 13). Le terme (4) désigne précisément « l'objet a tel que: » $a \in \mathbb{R}$ et $a \geq 0$ et $a^2 = x$. On peut dire aussi « la solution a » de la proposition $a \in \mathbb{R}$ et $a \geq 0$ et $a^2 = x$. C'est un terme dans lequel la variable a est locale, autrement dit ne figure pas librement. Elle est liée par l'opérateur **Sa**.

Conservation du nom d'une théorie. Le nom d'une théorie, par exemple celui de « théorie des ensembles », est associé à un système d'axiomes et de déductions élémentaires (Chap. 5), si nous admettons que le langage d'une théorie est toujours le même langage universel $\mathcal{L}_\Omega$ (§9.2). Mais pratiquement, le même nom est utilisé aussi pour une quantité de théories qui sont des extensions définitionnelles de la première. Nous allons poser encore dans cet ouvrage un certain nombre de définitions en théorie des ensembles. Chacune d'elles sera un nouvel axiome définitionnel et si l'on prend le mot « théorie » au sens strict du chapitre 5, chaque allongement d'une liste d'axiomes définit une nouvelle théorie. Mais lorsqu'il s'agit d'axiomes définitionnels, on conserve en général le nom de la théorie. Le mot « théorie » recouvre alors une *famille de théories* $\mathcal{T}_0, \mathcal{T}_1, \mathcal{T}_2$, etc. dont chacune est une extension définitionnelle de la précédente.

Il y a cependant des définitions qui sont à l'origine d'embranchements importants des mathématiques, à l'occasion desquelles on introduit un nouveau

nom de théorie. Par exemple, lorsqu'on ajoute à la théorie des ensembles la définition de ce qu'est un espace topologique, on entre dans la *topologie générale*. C'est le nom donné à l'extension de $\mathcal{T}_{\text{ENS}}$ résultant de ce nouvel axiome définitionnel. Ce genre d'embranchement ne se présentera pas dans cet ouvrage. Nous en resterons à des définitions qui ne sortent pas du noyau rudimentaire traditionnel de la théorie des ensembles.

12.2 Couples et graphes

Ce paragraphe est entièrement consacré à la théorie des ensembles. Il contient plusieurs définitions de la forme décrite dans le paragraphe précédent.

Couples. La notion de couple d'objets (a, b) se rencontre partout en mathématiques. Intuitivement, c'est un objet qui se « compose » de deux objets a, b considérés dans un *ordre* déterminé: a est le premier, b est le second. Il n'y a rien d'autre dans l'objet composé (a, b) que ces deux objets dans cet ordre. Les deux objets peuvent d'ailleurs ne pas être distincts. La propriété principale de ce type d'objet composé est que deux couples (a, b) et (a', b') sont égaux si et seulement si $a = a'$ et $b = b'$. Mais cette propriété ne sera pas un axiome. Nous la déduirons d'une *définition* formelle d'un couple (a, b) , à savoir de l'axiome définitionnel suivant.

E78. *Définition de (a, b)*

$$\vdash (a, b) = \{\{a\}, \{a, b\}\}.$$

Le nouveau symbole introduit dans cette définition est constitué par les deux parenthèses et la virgule du terme (a, b) . C'est un symbole fonctionnel binaire que l'on pourrait écrire en préfixe: $(,)ab$. On pourrait aussi utiliser à la place de $(,)$ un symbole tel que $\mathcal{C}$, en écrivant $\mathcal{C}ab$ au lieu de (a, b) . L'arbre du terme (a, b) est donc



E79. *Théorème principal*

$$\vdash (a, b) = (a', b') \Leftrightarrow (a = a' \text{ et } b = b').$$

Ce théorème est dit « principal » pour la raison suivante. La définition E78 ne sera utilisée que pour démontrer ce théorème et ensuite seul celui-ci sera utilisé. Ce théorème représente la propriété caractéristique fondamentale de la notion de couple et de ce point de vue conceptuel il serait naturel de prendre *cette* propriété comme axiome. Cela peut se faire. Mais c'est ici qu'intervient l'avantage logique d'un axiome *définitionnel* par opposition à un axiome d'une autre forme. De par sa forme, l'axiome E78 n'introduit qu'une abréviation et ne peut pas donner lieu à une contradiction avec les axiomes qui le précèdent.

Mais il n'a été imaginé — fort ingénieusement — que dans le but d'obtenir le théorème E79. Une fois celui-ci démontré, nous n'utiliserons plus la définition E78.

Démonstration. La démonstration du théorème E79 est un bon exercice d'application des théorèmes du paragraphe 11.3 sur les ensembles à un et deux éléments. Sa solution est donnée ici :

nil		
1°	$a = a' \text{ et } b = b'$	hyp
2°	$(a, b) = (a', b)$	1°, S88
3°	$(a', b) = (a', b')$	1°, S88
4°	$(a, b) = (a', b')$	2°, 3°, S85
5°	$(a, b) = (a', b')$	hyp
6°	$\{\{a\}, \{a, b\}\} = \{\{a'\}, \{a', b'\}\}$	5°, E78
7°	$\{a\} \in \{\{a\}, \{a, b\}\}$	E52
8°	$\{a\} \in \{\{a'\}, \{a', b'\}\}$	6°, 7°, S83
9°	$\{a\} = \{a'\} \text{ ou } \{a\} = \{a', b'\}$	8°, E51
10°	$\{a\} = \{a'\}$	hyp
11°	$a' = a$	10°, E66
12°	$\{a\} = \{a', b'\}$	hyp
13°	$a' \in \{a', b'\}$	E52
14°	$a' \in \{a\}$	12°, 13°, S83
15°	$a' = a$	14°, E60
16°	$a' = a$	9°, 11°, 15°, disj cas
17°	$\{\{a\}, \{a, b\}\} = \{\{a\}, \{a, b'\}\}$	6°, 16°, S83
18°	$\{a, b\} = \{a, b'\}$	17°, E58
19°	$b = b'$	18°, E58
20°	$a = a' \text{ et } b = b'$	16°, 19°
21°	$(a, b) = (a', b') \Leftrightarrow (a = a' \text{ et } b = b')$	4°, 20°, S40.

E80. $\vdash (a, b) = (b, a) \Rightarrow a = b$.

Ce théorème est un corollaire immédiat de E79. Il est démontré dans l'annexe (Chap. 16). Par contraposition (S31), on en déduit $\vdash a \neq b \Rightarrow (a, b) \neq (b, a)$. C'est en vertu de cette propriété qu'un couple (a, b) mérite le nom de couple *ordonné* qu'on lui donne souvent, par opposition à un ensemble à deux éléments $\{a, b\}$ qui est toujours égal à $\{b, a\}$ (E55).

Le prédicat « est un couple ». Une variable peut être utilisée pour désigner n'importe quelle espèce d'objet — en particulier un objet qui est le couple formé par deux objets a, b . Bien entendu, le terme (a, b) désigne cet objet, mais on peut vouloir le désigner par une simple variable, disons z . On peut poser pour cela $z = (a, b)$. Si l'on ne s'intéresse pas aux objets particuliers a et b qui composent cet objet désigné par z et si l'on veut exprimer uniquement le fait

que celui-ci est un couple, on écrira

$$\exists a \exists b (z = (a, b))$$

et l'on abrégera cette proposition par

z est un couple.

L'affirmation « z est un couple » ne dit pas de quels objets est formé le couple en question, mais seulement qu'il en *existe*. Syntaxiquement, les mots « est un couple » constituent un nouveau symbole relationnel unaire, postfixé, et ce que nous venons de dire est une nouvelle définition, à savoir l'axiome définitionnel suivant :

E81. *Définition*

$$\vdash z \text{ est un couple} \Leftrightarrow \exists a \exists b (z = (a, b)).$$

Les symboles relationnels de cette forme, commençant par les mots « est un » ou « est une » suivis d'un nom commun désignant une certaine espèce ou classe d'objet, sont très nombreux en mathématique, comme nous l'avons déjà signalé au paragraphe 3.5. Le terme de *prédicat* peut être considéré comme un synonyme de *symbole relationnel* (§9.1), mais on l'utilise surtout pour les symboles relationnels unaires de cette forme. Nous introduirons par exemple les prédicats : est un graphe, est une fonction, est une fonction injective. En pensant aux cours de mathématique qu'il a suivis, le lecteur se rappellera bien d'autres exemples tels que : est un nombre complexe, est un espace vectoriel, est une application linéaire, est un groupe commutatif, etc.

Dans les exposés mathématiques usuels, les définitions de ce genre sont présentées d'un point de vue sémantique. L'accent est mis sur les *notions* que recouvrent les noms communs en question : qu'est-ce qu'un couple, qu'est-ce qu'une fonction, qu'est-ce qu'un espace vectoriel, etc. Mais d'un point de vue formel, chacune de ces définitions se résume par une proposition qui est de la même forme que **E81**, à savoir une équivalence dont le membre de gauche est une proposition de la forme « z est un ... » et dont le membre de droite est une proposition qui n'utilise pas ce prédicat. L'axiome dit donc à quelle proposition est équivalente la proposition « z est un ... ». Savoir au sens *formel* ce qu'est un couple, ou un graphe, ou un nombre complexe, etc. c'est savoir remplacer les propositions « z est un couple », « z est un graphe », « z est un nombre complexe », etc. par des propositions équivalentes dans lesquelles ces mots ne figurent pas.

E82. Théorème. $\vdash (x, y)$ est un couple.

Démonstration :

	nil	
1°	$(x, y) = (x, y)$	S82
2°	$\exists b ((x, y) = (x, b))$	1°, S63
3°	$\exists a \exists b ((x, y) = (a, b))$	2°, S63

4°	$\exists a \exists b((x, y) = (a, b)) \Leftrightarrow (x, y) \text{ est un couple}$	E81
5°	$(x, y) \text{ est un couple}$	3°, 4°, S39.

Projections d'un couple. Les projections d'un couple z sont les deux objets qui composent ce couple et qui sont uniques d'après E79. On les désigne par $\mathbf{pr}_1 z$ et $\mathbf{pr}_2 z$. Intuitivement, il s'agit encore d'une définition. Mais l'axiome suivant n'est pas de l'une des formes dont nous avons parlé au paragraphe 12.1. Nous le déclarons donc simplement comme axiome. Nous verrons plus loin (Chap. 13) comment le remplacer par de véritables définitions $\mathbf{pr}_1 z = \dots$, $\mathbf{pr}_2 z = \dots$, grâce à une extension de la logique des prédicats avec égalité.

E83. *Axiome des projections d'un couple*

$$\vdash z \text{ est un couple} \Rightarrow z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z).$$

Grâce à cet axiome, lorsque la proposition « z est un couple» est vraie dans un contexte, on peut utiliser les composantes de ce couple sans avoir, pour les nommer, à procéder par élimination des quantificateurs existentiels de $\exists a \exists b(z = (a, b))$. Comparer les deux plans de démonstration suivants, où B représente une proposition dans laquelle a et b ne figurent pas librement.

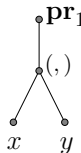
1°	$z \text{ est un couple}$	<i>base</i>	1°	$z \text{ est un couple}$	<i>base</i>
2°	$\exists a \exists b(z = (a, b))$	1°, E81	2°	$z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	1°, E83
3°	$\exists b(z = (a, b))$	hyp		$\vdots$	
4°	$z = (a, b)$	hyp		$B.$	
	$\vdots$				
	B				
	B	S64			
	B	S64.			

Les théorèmes suivants (démonstration dans l'annexe) sont intuitivement évidents.

E84. $\vdash \mathbf{pr}_1(x, y) = x \quad \vdash \mathbf{pr}_2(x, y) = y.$

E85. $\vdash z \text{ est un couple} \Leftrightarrow z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z).$

Les symboles $\mathbf{pr}_1$, $\mathbf{pr}_2$ sont des symboles fonctionnels *unaires*. Dans les expressions $\mathbf{pr}_1(x, y)$, $\mathbf{pr}_2(x, y)$, ces symboles sont combinés avec *un* terme, le terme (x, y) , lequel se compose du symbole fonctionnel binaire « $(,)$ » et des termes x , y . L'arbre du terme $\mathbf{pr}_1(x, y)$ est donc le suivant :



Triplets, quadruplets, etc. On pose les définitions suivantes:

$$\begin{aligned} \vdash (a, b, c) &= (a, (b, c)); \\ \vdash (a, b, c, d) &= (a, (b, c, d)); \\ &\text{etc.} \end{aligned} \tag{1}$$

Le théorème $\vdash (a, b, c) = (a', b', c') \Leftrightarrow (a = a' \text{ et } b = b' \text{ et } c = c')$ se démontre facilement d'après (1) et E79. L'objet désigné par le terme (a, b, c) est appelé un *triplet*. Formellement, on pose $\vdash z$ est un triplet $\Leftrightarrow \exists a \exists b \exists c (z = (a, b, c))$ comme définition du prédicat «est un triplet». D'après (1) on a donc le théorème

$$\vdash z \text{ est un triplet} \Leftrightarrow (z \text{ est un couple et } \mathbf{pr}_2 z \text{ est un couple}).$$

De manière analogue, la proposition « z est un quadruplet» sera équivalente à « z est un couple et $\mathbf{pr}_2 z$ est un triplet», et ainsi de suite.

La manière dont se poursuit la série de définitions (1) devrait être claire. Ces définitions introduisent les symboles fonctionnels « $(,)$ », « $(,,)$ », « $(,,,)$ », etc. dont l'arité est égale au nombre de virgules plus 1. L'objet représenté par un terme $(x_1, \dots, x_n)$ à $n - 1$ virgules ($n \geq 2$) est appelé un *n-uplet* ou *n-uple*. Une égalité de *n-uples*

$$(x_1, \dots, x_n) = (x_1', \dots, x_n')$$

est équivalente à la conjonction des n égalités $x_i = x_i'$ ($i = 1, \dots, n$).

Graphes. Un ensemble dont chaque élément est un couple, autrement dit un ensemble de couples, est appelé un *graphe*. Formellement, nous introduisons le prédicat *est un graphe* (symbole relationnel unaire postfixé) et nous posons l'axiome définitionnel suivant:

E86. *Définition*

$$\vdash G \text{ est un graphe} \Leftrightarrow \forall z (z \in G \Rightarrow z \text{ est un couple}).$$

Le terme de *graphe* provient de l'expression «graphe d'une fonction», désignant la courbe représentative de cette fonction (Fig. 12.1(a)), en tant qu'ensemble G de points du plan $\mathbb{R}^2$, les points du plan étant des couples de nombres réels (coordonnées). N'importe quel ensemble G de points du plan $\mathbb{R}^2$ (Fig. 12.1(b)) est d'ailleurs un graphe au sens de E86, et pas seulement l'ensemble des points de la courbe d'une fonction.

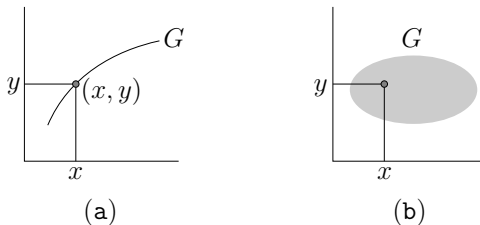


Figure 12.1 Graphes dans le plan $\mathbb{R}^2$.

Nous identifierons plus loin le concept de fonction avec celui du graphe d'une fonction en tant qu'ensemble de couples. Une fonction en général, pas seulement dans le domaine des nombres réels, sera définie comme un ensemble de couples, donc un graphe, mais un graphe qui satisfait à une condition particulière.

Représentation d'un graphe fini. Un graphe fini est un ensemble de couples fini, et un ensemble fini est un ensemble dont le nombre d'éléments est fini. La définition formelle de cette notion fait intervenir l'ensemble $\mathbb{N}$ des entiers naturels, puisqu'il est question de compter les éléments d'un ensemble. Mais nous n'avons pas besoin ici d'une définition formelle. L'ensemble de couples suivant est un graphe fini :

$$G = \{(1, 2), (2, 3), (3, 2), (3, 5), (4, 4)\}.$$

On représente ce graphe par le diagramme de la figure 12.2, dont les « points » ou *sommets* représentent les objets 1, 2, 3, 4, 5 et les flèches ou *arcs* représentent les couples appartenant à G . Un couple (a, b) est représenté graphiquement par la flèche $a \rightarrow b$. La plupart des théorèmes sur les graphes que nous verrons dans la suite sont évidents lorsqu'on se représente mentalement les graphes par de tels diagrammes. Mais ces théorèmes ne font pas l'hypothèse de graphes finis.

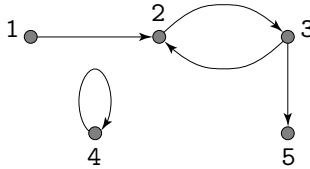


Figure 12.2 Représentation d'un graphe fini.

Théorèmes. Parmi les théorèmes qui suivent sur les graphes, le plus important est E91, en raison de la méthode de démonstration d'inclusion ou d'égalité de graphes qui en découle. La plupart de ces théorèmes sont démontrés dans l'annexe (Chap. 16).

E87. $\vdash \emptyset$ est un graphe.

1°	$\forall x(x \in \emptyset \Rightarrow x \text{ est un couple})$	E8
2°	$\forall x(x \in \emptyset \Rightarrow x \text{ est un couple}) \Leftrightarrow \emptyset \text{ est un graphe}$	E86, (exemplif.)
3°	$\emptyset \text{ est un graphe}$	1°, 2°.

Lorsqu'un prédicat, comme le prédicat « est un graphe », signifie que tous les éléments d'un ensemble ont une certaine propriété — dans le cas d'un graphe cette propriété des éléments est celle d'être un couple — alors ce prédicat s'applique *ipso facto* à l'ensemble vide en vertu de E8. Puisque cet ensemble n'a pas d'élément, il est trivialement vrai que tout élément de cet ensemble satisfait à telle ou telle condition. Voir notre discussion de E8 (§10.7). Pour

cette raison, l'ensemble vide appartient à beaucoup d'espèces d'ensembles. Nous verrons plus loin que c'est non seulement un graphe, mais encore que c'est une fonction (une fonction est un cas particulier de graphe) et même une fonction injective.

E88. $\vdash$ Si G est un graphe, alors $z \in G \Leftrightarrow (z \text{ est un couple et } z \in G)$.

Ce théorème est immédiat (S60.3). Nous utiliserons assez fréquemment dans la suite la tournure « si A , alors B » au lieu de $A \Rightarrow B$. Cela permet souvent d'éviter des parenthèses et améliore un peu la lisibilité. L'application principale de ce théorème consiste à remplacer une proposition $z \in G$ par $(z \in G \text{ et } z \text{ est un couple})$ ou vice-versa dans une chaîne d'équivalences, donc à introduire ou à éliminer « et z est un couple », lorsque la proposition (G est un graphe) est vraie dans le contexte. Un exemple se présente plus bas dans la démonstration de **E91**.

Les deux théorèmes suivants sont aussi évidents. La phrase « G et G' sont des graphes » est une abréviation de la conjonction : G est un graphe et G' est un graphe.

E89. $\vdash$ Si G est un graphe et $F \subset G$, alors F est un graphe.

E90. $\vdash$ Si G et G' sont des graphes, alors $G \cup G'$ et $G \cap G'$ sont des graphes.

L'énoncé suivant (**E91**) contient en fait deux théorèmes : 1. $\vdash$ si G et G' sont des graphes, alors $G \subset G' \Leftrightarrow \dots$; 2. $\vdash$ si G et G' sont des graphes, alors $G = G' \Leftrightarrow \dots$ Nous utiliserons plus d'une fois cette manière de résumer deux ou plusieurs théorèmes qui ont la même condition « si $\dots$ ». Pour **E91**, deux démonstrations du premier théorème sont données dans l'annexe (Chap. 16), l'une utilisant **E85**. La démonstration du deuxième, compte tenu du premier et de **E10**, est une simple question de logique des prédicats.

E91. $\vdash$ Si G et G' sont des graphes, alors

$$G \subset G' \Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G');$$

$$G = G' \Leftrightarrow \forall x \forall y ((x, y) \in G \Leftrightarrow (x, y) \in G').$$

Méthode de démonstration. Les théorèmes **E91** fournissent une importante méthode pour démontrer une inclusion $G \subset G'$ ou une égalité $G = G'$ dans un contexte où les propositions $G \text{ est un graphe}$, $G' \text{ est un graphe}$ sont vraies. Si les variables x, y ne figurent pas librement dans les hypothèses Γ du contexte, il suffit de démontrer l'implication $(x, y) \in G \Rightarrow (x, y) \in G'$, respectivement l'équivalence $(x, y) \in G \Leftrightarrow (x, y) \in G'$. Dans le second cas, on procède si possible par une chaîne d'équivalences. Ces plans de démonstration sont représentés dans la figure 12.4. On omet en général les lignes marquées d'un astérisque.

Bien entendu, pour démontrer une inclusion ou une égalité de graphes G, G' on peut procéder comme pour deux ensembles quelconques, suivant les méthodes introduites au paragraphe 10.7, découlant de **E5** et **E17**. Au lieu de l'hypothèse $(x, y) \in G$ on fait alors l'hypothèse $z \in G$ et l'on démontre $z \in G'$. Pour l'égalité $G = G'$, on démontre l'équivalence $z \in G \Leftrightarrow z \in G'$. Mais dans la

Γ		Γ	
	G et G' sont des graphes	G et G' sont des graphes	
	$(x, y) \in G$ hyp	$(x, y) \in G \Leftrightarrow \dots$	
	$\vdots$	$\Leftrightarrow \dots$	
	$(x, y) \in G'$	$\Leftrightarrow (x, y) \in G'$	
*	$(x, y) \in G \Rightarrow (x, y) \in G'$	*	$(x, y) \in G \Leftrightarrow (x, y) \in G'$
*	$\forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G')$	*	$\forall x \forall y ((x, y) \in G \Leftrightarrow (x, y) \in G')$
	$G \subset G'$ E91.		$G = G'$ E91.

Figure 12.4 Plans de démonstration d'une inclusion et d'une égalité de graphes suivant E91. On suppose que x et y ne figurent pas librement dans Γ . Les lignes marquées d'un astérisque sont omises en général.

plupart des cas il faut utiliser dans la démonstration le fait que G et G' sont des graphes, donc que z est un couple, et introduire les composantes (projections) de ce couple dans le raisonnement. Lorsqu'on suit le plan de la figure 12.4, utilisant E91, on raisonne dès le départ avec un couple dont les composantes sont nommées explicitement. Cela raccourcit en général la démonstration.

E92. $\vdash$ Si G est un graphe, alors

$$G = \emptyset \Leftrightarrow \forall x \forall y ((x, y) \notin G);$$

$$G \neq \emptyset \Leftrightarrow \exists x \exists y ((x, y) \in G).$$

Démonstration:

nil	
1°	G est un graphe hyp
2°	$(x, y) \notin \emptyset$ E2
	$G = \emptyset \Leftrightarrow G \subset \emptyset$ E12
	$\Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in \emptyset)$ E91
	$\Leftrightarrow \forall x \forall y ((x, y) \notin G)$ 2°, S60.6
	$G \text{ est un graphe} \Rightarrow (G = \emptyset \Leftrightarrow \forall x \forall y ((x, y) \notin G)).$

Le deuxième théorème E92 se déduit du premier par S71.

Pour démontrer $G = \emptyset$ dans un contexte où la proposition « G est un graphe » est vraie, on peut toujours procéder suivant E15, comme pour un ensemble quelconque, c'est-à-dire supposer que $z \in G$ et aboutir à $\perp$. Mais si l'on doit utiliser dans la démonstration le fait que G est un graphe, et introduire les projections de z , il est plus court de partir de l'hypothèse $(x, y) \in G$ pour aboutir à $\perp$. Car on en déduit $(x, y) \notin G$ et si les variables x, y ne figurent pas librement dans les hypothèses du contexte, il vient $\forall x \forall y ((x, y) \notin G)$ et l'on peut conclure suivant E92.

13

Langages du premier ordre à opérateurs généraux

13.1 Opérateurs généraux

Les langages du premier ordre tels que nous les avons définis au chapitre 4 sont très simples. Ils ne permettent pas de formaliser toutes les expressions du langage mathématique courant, ni d'ailleurs de n'importe quel langage plus ou moins précis utilisé dans un domaine technique ou scientifique. Le but de ce chapitre est d'élargir la définition des langages du premier ordre de manière à pouvoir y inclure la plupart des expressions mathématiques usuelles.

Bien entendu, de même que les grammairiens ne peuvent pas prévoir toutes les tournures de langage dont sont capables les écrivains — pour ne pas parler de celles de l'homme de la rue — de même il n'est pas possible de prévoir toutes les formes d'expression qui peuvent être utilisées dans la pratique des sciences mathématiques et techniques. Cependant, il est possible de définir un langage formel dont le répertoire d'expressions est suffisamment riche pour que l'on puisse facilement traduire dans ce langage toutes les phrases de la langue mathématique courante. Cet exercice de traduction met d'ailleurs suffisamment en évidence les imperfections de cette langue pour que l'on finisse par ne considérer celle-ci que comme du langage formel parlé avec passablement de liberté.

Exemple 1. La première lacune des langages formels du premier ordre que nous avons considérés jusqu'ici (Chap. 4) est l'absence de termes contenant des variables *locales*. Nous rappelons que suivant les règles de syntaxe du paragraphe 4.1 un terme est soit une variable individuelle, soit une constante individuelle, soit une expression de la forme $sT_1 \cdots T_n$ où s est un symbole fonctionnel n -aire et $T_1, \dots, T_n$ sont des termes construits suivant ces mêmes règles. Il n'y a pas de variables locales dans un tel terme, parce qu'il n'y a pas d'opérateurs liant les variables à la manière des quantificateurs.

Un terme mathématique usuel tel que

$$\sum_{n=1}^p n^2 \quad (1)$$

entre dans la catégorie des termes à variable liée. La variable n de ce terme est évidemment *locale* (§4.4). Celui-ci peut être vu comme étant composé de trois parties :

$$\sum_n [1 \dots p] \quad n^2.$$

La première est un opérateur de « sommation sur n », qui contient la variable n (comme le quantificateur $\forall x$ contient la variable x) et que l'on pourrait écrire Σn comme un quantificateur. La deuxième partie est un *terme*, représentant le domaine de variation de n , à savoir l'ensemble des entiers naturels $1, 2, \dots, p$ ou intervalle $[1 \dots p]$ de $\mathbb{N}$. La troisième composante de l'expression est aussi un *terme*, le terme n^2 . L'opérateur de sommation sur n *lie* la variable n comme le ferait un quantificateur $\forall n$. Nous pouvons récrire le terme (1) de la manière suivante, en traçant un lien explicite :

$$\Sigma n \overline{[1 \dots p]} (n^2) . \quad (2)$$

L'opérateur Σx associé à une variable x (Σn , Σk , etc.) est un opérateur *binaire*. Il se combine avec deux *termes* et l'expression résultante est un *terme*. Nous dirons que cet opérateur est du *type*

$$T \leftarrow (T, T).$$

La lettre T signifie « terme ». Nous verrons plus bas un exemple d'opérateur *ternaire* du type

$$T \leftarrow (P, T, T),$$

où P signifie « proposition ». Un opérateur de ce type se combine avec trois expressions. La première doit être une proposition et les deux autres des termes et l'expression résultante est un terme. De façon générale, un opérateur n -aire se combine avec n expressions dont chacune doit être un terme ou une proposition (cela dépend de l'opérateur) et le résultat de la composition est un terme ou une proposition (cela dépend aussi de l'opérateur).

Un quantificateur $\forall x$ est un opérateur du type $P \leftarrow P$. Il se combine avec une proposition, disons A , et l'expression résultante, $\forall x A$, est une proposition. Un connecteur logique binaire, par exemple $\Leftrightarrow$, est un opérateur du type $P \leftarrow (P, P)$. Un symbole relationnel binaire, le symbole d'égalité par exemple, est un opérateur du type $P \leftarrow (T, T)$. Il se combine avec deux termes et le résultat est une proposition. Un symbole fonctionnel binaire, par exemple le symbole de réunion $\cup$, est un opérateur du type $T \leftarrow (T, T)$. Son type est donc le même que celui de l'opérateur Σn (2), mais contrairement à Σn , l'opérateur $\cup$ ne lie pas de variable.

Le type d'un opérateur inclut donc: son arité (le nombre d'expressions avec lesquelles il se combine); le type (terme ou proposition) de chacune de ces expressions; le type (terme ou proposition) de l'expression résultant de la combinaison.

Un opérateur est caractérisé non seulement par son type, mais encore par les variables qu'il lie éventuellement. Ici intervient une certaine « subtilité » syntaxique, à savoir qu'un opérateur peut lier une variable x seulement dans *certaines* des expressions avec lesquelles il se combine.

L'opérateur de sommation Σn peut être défini de cette manière. Dans l'exemple (2), le premier des deux termes auxquels il est appliqué, le terme $[1 \dots p]$, ne contient pas la variable n . Que dire, par exemple, du terme

$$\Sigma n[1 \dots n](n^2) \quad \text{en notation usuelle : } \sum_{n=1}^n n^2.$$

La première réponse — ou du moins la réponse normale d'un pédagogue soucieux de la clarté de ses notations — est que l'on *n'écrit pas* de termes de ce genre, parce qu'ils n'ont pas de sens: l'intervalle que parcourt la variable de sommation n doit être fixe, il ne doit pas varier lui-même avec n . Mais il y a une autre réponse, plus subtile. C'est de considérer la variable n du terme $[1 \dots n]$ comme globale, et de ne considérer comme locale que la variable n du terme n^2 . Autrement dit, c'est d'admettre que l'opérateur Σn ne lie la variable n que dans le deuxième des termes auxquels il est appliqué:

$$\Sigma n \overline{[1 \dots n]} (n^2) . \quad (3)$$

Le terme (3) a donc la même signification que le terme suivant, où l'on a fait un changement de variable liée:

$$\Sigma x \overline{[1 \dots n]} (x^2) \quad \text{en notation usuelle : } \sum_{x=1}^n x^2. \quad (4)$$

Bien entendu, même si l'on adopte cette propriété de l'opérateur de sommation, on ne peut que recommander à l'utilisateur d'écrire (4) plutôt que (3). Ainsi il ne courra pas le risque d'être mal compris des lecteurs qui ne sont pas versés dans ces subtilités syntaxiques.

Exemple 2. Comme deuxième exemple de terme à variable liée du langage mathématique usuel, nous considérons l'expression ensembliste

$$\{x \mid x \in \mathbb{R} \text{ et } a < x \text{ et } x < b\}, \quad (5)$$

qui se lit: l'ensemble des x tels que $x \in \mathbb{R}$ et $a < x < b$. La forme générale des termes de ce genre est

$$\{x \mid \mathbf{R}\}, \quad (6)$$

où x est une variable et $\mathbf{R}$ une proposition. Certains auteurs écrivent $\{x : \mathbf{R}\}$. L'interprétation d'un terme de cette forme est toujours la même: il désigne

l'ensemble des objets qui ont une certaine propriété, à savoir la propriété de x exprimée par la proposition R . Cette notion est fondamentale. La théorie $\mathcal{T}_{\text{ENS}}$ possède un schéma de déduction élémentaire spécifique à ce sujet (§14.2).

Il est clair que la variable x est locale dans le terme (5). Notamment le terme $\{y \mid y \in \mathbb{R} \text{ et } a \leq y \text{ et } y \leq b\}$ a la même signification. Syntaxiquement, un terme de la forme (6) est composé de deux parties: premièrement l'opérateur « l'ensemble des x tels que », et deuxièmement la proposition R . Si nous abrégeons cet opérateur par $\mathbf{E}x$, le terme (5) s'écrit, avec les liens:

$$\mathbf{E}x \overbrace{(x \in \mathbb{R} \text{ et } a < x \text{ et } x < b)}.$$

Dans la suite, nous utiliserons généralement la notation usuelle (6), au lieu de $\mathbf{E}xR$. Nous n'utilisons cette dernière que pour les discussions syntaxiques telles que la présente. Mais nous donnerons néanmoins un nom à l'opérateur $\mathbf{E}x$ qui est « dissimulé » dans la notation usuelle (6). Nous l'appellerons l'opérateur de *collection sur* x . En l'appliquant à une proposition R , on « collectionne » les objets qui possèdent une certaine propriété.

Les opérateurs de collection $\mathbf{E}x$ ($\mathbf{E}x$, $\mathbf{E}y$, $\mathbf{E}a$, etc.) sont de type $T \leftarrow P$. Ils se combinent toujours avec une proposition R et l'expression qui en résulte, $\mathbf{E}xR$, est un *terme*. Les occurrences libres de x dans R sont liées à cet opérateur dans $\mathbf{E}xR$, comme s'il s'agissait d'un quantificateur.

Exemple 3. Considérons l'intégrale définie

$$\int_a^{b+1} F(x)dx, \quad (7)$$

en supposant que a et b désignent des nombres réels et F une application de $\mathbb{R}$ dans $\mathbb{R}$. Cette expression est encore un exemple de terme à variable liée, car la variable x est locale. Notamment l'intégrale

$$\int_a^{b+1} F(z)dz,$$

a la même valeur que (7). Le terme (7) peut être analysé syntaxiquement comme étant composé de quatre parties:

$$\int dx \quad a \quad b+1 \quad F(x).$$

La première est un opérateur *d'intégration sur* x que nous écrirons Ix pour la discussion. Cet opérateur est combiné avec trois *termes*: les termes a , $b+1$ représentant les limites de l'intégration, et le terme $F(x)$ représentant la valeur de la fonction F correspondant au nombre x . La variable du terme $F(x)$ est liée à l'opérateur d'intégration. Nous pouvons écrire le terme (7):

$$\overbrace{Ix \ a \ (b+1) \ (F(x))}.$$

L'opérateur d'intégration Ix est de type $T \leftarrow (T, T, T)$. Comme dans le cas de l'opérateur de sommation (Exemple 1), nous devons prendre une décision

sur les liens qu'il y a dans une intégrale lorsque la variable d'intégration figure elle-même dans les limites d'intégration, par exemple dans l'intégrale

$$\int_x^{x+1} F(x) dx. \quad (8)$$

C'est un genre d'écriture que la plupart des auteurs actuels évitent, mais pour accommoder ceux qui ne le font pas, il convient de considérer la variable x des limites de l'intégration comme une variable globale. Cela veut dire que l'opérateur d'intégration ne lie la variable d'intégration que dans le troisième terme :

$$\overline{\text{Ix } x(x+1)(F(x))}.$$

Si l'on craint que le terme (8) soit mal interprété, on peut toujours changer de variable liée (variable d'intégration) et écrire

$$\overline{\text{Iu } x(x+1)(F(u))} \quad \text{en notation usuelle : } \int_x^{x+1} F(u) du.$$

Exemple 4. Considérons la définition usuelle de la valeur absolue d'un nombre réel x :

$$|x| = \begin{cases} x & \text{si } x \geq 0; \\ -x & \text{si } x < 0. \end{cases}$$

Sachant que dans ce contexte la proposition $x < 0$ équivaut à la négation de $x \geq 0$, nous pouvons écrire cette définition sur une ligne :

$$|x| = (\text{Si } x \geq 0 \text{ alors } x, \text{ sinon } -x). \quad (9)$$

Le terme de droite de cette égalité est un *terme conditionnel*, de la forme générale

$$\text{Si } A \text{ alors } T_1 \text{ sinon } T_2. \quad (10)$$

Il se compose : d'un opérateur ternaire de *test de condition*, représenté par les mots « si alors sinon » ; d'une proposition A qui est la « condition testée » ; de deux termes T_1 , T_2 qui représentent la valeur du terme conditionnel lorsque la condition est vraie, respectivement fausse. Nous traiterons des propriétés logiques des termes conditionnels au paragraphe 13.3. Seule leur syntaxe nous intéresse ici.

Si nous abrégeons l'opérateur de test par le seul mot *Si*, un terme conditionnel (10) s'écrit $Si AT_1T_2$. L'opérateur *Si* est de type $T \leftarrow (P, T, T)$. Il ne lie aucune variable.

Les opérateurs du langage $\mathcal{L}_\Omega$. Nous admettons que le langage universel $\mathcal{L}_\Omega$ comporte non seulement une infinité de constantes individuelles, de symboles fonctionnels et de symboles relationnels de chaque arité, mais plus généralement qu'il comporte une infinité d'opérateurs de chaque *espèce*. Nous devons préciser ce que nous entendons ici par une *espèce d'opérateur*.

Chaque opérateur Ξ (lettre grecque Xi majuscule) possède une arité n et un type $t_0 \leftarrow (t_1, \dots, t_n)$ où chaque t_i est soit T soit P. L'opérateur se combine

toujours avec n expressions $A_1, \dots, A_n$ sous la forme

$$\Xi A_1 \cdots A_n. \quad (11)$$

L'expression A_i doit être de type t_i , c'est-à-dire un terme si t_i est T, une proposition si t_i est P. L'expression (11) est alors de type t_0 (T ou P). On sait en outre quelles variables sont liées par l'opérateur dans chacune des expressions $A_1, \dots, A_n$. Autrement dit on peut écrire une liste de variables (éventuellement vide) $x_1^1, \dots, x_{p_1}^1$ qui sont les variables liées par l'opérateur Ξ dans le premier de ses arguments (A_1); une liste de variables (éventuellement vide et éventuellement identique à la première) $x_1^2, \dots, x_{p_2}^2$ qui sont les variables liées par Ξ dans A_2 , et ainsi de suite. La i -ème de ces listes de variables contient les variables liées par Ξ dans son argument de rang i (A_i).

La donnée d'une arité n (0, 1, 2, etc.), d'un type $t_0 \leftarrow (t_1, \dots, t_n)$ et d'une liste de n listes de variables caractérise une *espèce* d'opérateur. Par exemple, le connecteur logique « $\Rightarrow$ » appartient à l'espèce d'opérateur suivante:

arité: 2
 type: $P \leftarrow (P, P)$
 variables liées par Ξ : (nil, nil).

L'opérateur d'intégration sur x (Exemple 3) est de l'espèce:

arité: 3
 type: $T \leftarrow (T, T, T)$
 variables liées par Ξ : (nil, nil, x).

Tous les symboles de $\mathcal{L}_\Omega$ autres que les variables individuelles sont des opérateurs. Une constante individuelle peut être vue en effet comme un opérateur d'arité 0, de type $T \leftarrow \text{nil}$ et dont la liste de listes de variables liées est nil.

Cette description des opérateurs est relativement compliquée. Par contre, la définition des expressions de $\mathcal{L}_\Omega$ est très simple. Les règles de syntaxe du paragraphe 4.1 sont remplacées par les deux règles suivantes:

(a) Les variables individuelles sont des termes.

(b) Si Ξ est un opérateur de $\mathcal{L}_\Omega$ d'arité n et de type $t_0 \leftarrow (t_1, \dots, t_n)$ et si $A_1, \dots, A_n$ sont des expressions de $\mathcal{L}_\Omega$ de types respectifs $t_1, \dots, t_n$, l'expression $\Xi A_1 \cdots A_n$ est une expression de $\mathcal{L}_\Omega$ de type t_0 .

Nous n'irons pas plus loin dans cette description générale du langage $\mathcal{L}_\Omega$. Un petit nombre de règles suffisent à définir: les variables libres des expressions de $\mathcal{L}_\Omega$ (§4.4); le sens de l'opération de substitution $(T|x)$ d'un terme T à une variable x dans les expressions de $\mathcal{L}_\Omega$ (§4.5); les expressions dans lesquelles T est librement substituable à x (§4.6).

Ces règles sont évidentes d'après ce qui précède. Une variable x figure librement dans l'expression $\Xi A_1 \cdots A_n$ si et seulement si elle figure librement dans l'une des expressions A_i et n'appartient pas à la i -ème liste de variables liées par l'opérateur Ξ . L'expression $(T|x)(\Xi A_1 \cdots A_n)$ est identique par définition

à l'expression $\Xi A_1' \cdots A_i'$ dans laquelle l'expression A_i' (pour $i = 1, \dots, p$) est identique à A_i si x appartient à la i -ème liste de variables liées par Ξ , sinon A_i' est identique à $(T|x)A_i$. Le terme T n'est pas librement substituable à x dans $\Xi A_1 \cdots A_n$ si et seulement si x figure librement dans l'une des expressions A_i et en même temps la i -ème liste de variables liées par Ξ ne contient pas x mais contient une variable qui figure librement dans T .

13.2 Sélecteurs

Exemple 1. Si a et y sont deux nombres réels strictement positifs, on définit généralement le logarithme de base a de y comme étant

$$\text{le nombre réel } x \text{ tel que } a^x = y \quad (1)$$

et l'on convient de noter ce nombre $\log_a(y)$. Le symbole « log » est un symbole fonctionnel binaire. On pourrait écrire simplement $\log ax$. Les expressions de la forme (1) sont très fréquentes dans le langage mathématique — et aussi d'une certaine manière dans le langage courant. Un objet est très souvent désigné comme étant « l'objet tel que ... », c'est-à-dire l'objet qui vérifie une certaine proposition. Observons que la variable x dans (1) est locale. Elle peut être changée sans que cela change le sens de l'expression: « le nombre réel z tel que $a^z = y$ » désigne aussi le logarithme de y . L'*opérateur de sélection* ou *sélecteur* Sx est utilisé pour abréger l'expression (1):

$$Sx \overbrace{(x \in \mathbb{R} \text{ et } a^{\exp} x = y)}. \quad (2)$$

Le « S » de l'opérateur Sx peut faire penser aussi au mot « Solution ». L'objet représenté par le terme (2) peut-être décrit aussi comme la « solution en x » de la proposition $(x \in \mathbb{R} \text{ et } a^x = y)$. Si l'univers du discours est restreint aux nombres réels, on peut omettre $x \in \mathbb{R}$ dans la formalisation du terme (1). La définition du logarithme de y s'exprime alors formellement par l'axiome

$$(a > 0 \text{ et } y > 0) \Rightarrow \log_a(y) = Sx(a^x = y).$$

Cet axiome est de la forme

$$C \Rightarrow sx_1 \cdots x_n = T$$

notée au paragraphe 12.1 comme forme typique d'axiome définitionnel. Dans cet exemple $n = 2$, s est le symbole fonctionnel binaire « log », T est le terme $Sx(a^x = y)$, x_1 et x_2 sont les variables a et y , le nouveau symbole s (log) ne figure ni dans le terme T , ni dans la condition C et les seules variables libres de l'axiome sont x_1 et x_2 (a et y). La variable x ne figure pas librement dans l'axiome.

Termes descriptifs. Nous admettons que pour toute variable individuelle x , Sx est un opérateur du langage $\mathcal{L}_\Omega$ de type $T \leftarrow P$. Si A est une proposition et si x est une variable, l'expression SxA est un terme dans lequel la variable x

ne figure pas librement. Les occurrences libres de x dans A deviennent liées au sélecteur Sx dans le terme SxA . Dans la définition des variables libres et des variables liées d'un terme ou d'une proposition, un sélecteur est traité comme un quantificateur. Il en va de même dans la définition des termes librement substituables à une variable dans un terme ou dans une proposition. Par exemple, un terme qui contient des occurrences libres de x n'est pas librement substituable à y dans le terme $Sx(a^x = y)$.

Nous appelons *termes descriptifs* les termes de la forme SxA . D'autres auteurs les appellent des « descriptions ». Nous procédons plus bas à une nouvelle extension de la théorie **Lprédégal** (logique des prédicats avec égalité) en admettant dans cette théorie les trois nouvelles formes de déductions décrites par les schémas de **SEL1**–**SEL3**. Auparavant, nous voulons discuter l'interprétation des termes descriptifs qui se traduit formellement par ces schémas.

L'interprétation d'un terme descriptif SxA est clairement évidente dans un contexte où la proposition $\exists!x A$ est vraie. Dans ce cas on peut dire que la proposition A possède une solution x et une seule et que SxA désigne cette solution. Par exemple, si a et y sont deux nombres réels strictement positifs, on sait qu'il existe un nombre x et un seul tel que $a^x = y$. Le terme $Sx(a^x = y)$ désigne ce nombre.

L'interprétation de SxA est moins évidente dans les deux autres cas, c'est-à-dire lorsque la proposition A n'a pas de solution en x ou lorsqu'elle en a plusieurs. Dans le premier de ces deux cas, la proposition $\text{non}\exists x A$ est vraie et le terme SxA est *sans interprétation*. Un terme sans interprétation dans un contexte n'est pas une rareté dans le langage mathématique, ni dans le langage courant. Le terme « le roi de France » avait une interprétation bien claire à une certaine époque, il n'en a plus de nos jours quand bien même il est syntaxiquement correct. En théorie des nombres réels, le terme $\log_a(-1)$ n'a pas d'interprétation. Le terme $\sqrt{a}$ n'a d'interprétation que si a est positif. Le terme $\lim_{n \rightarrow \infty} x_n$ n'a d'interprétation que si la suite (x_n) est convergente, etc. Il n'y a pas d'inconvénient à introduire une notation qui n'a pas d'interprétation dans tous les contextes.

En mathématique, un terme est utilisé normalement dans les contextes où il possède une interprétation. Mais les schémas de déduction de la logique permettent de formuler des théorèmes indépendants de toute interprétation. Par exemple, pour le terme sans interprétation $\log_a(-1)$, on a d'après **S82** et **S89** les théorèmes

$$\vdash \log_a(-1) = \log_a(-1), \quad \vdash \exists x(x = \log_a(-1)). \quad (3)$$

Celui qui n'a jamais étudié la logique peut trouver ces théorèmes contraires à l'intuition. C'est un fait qu'on ne rencontre guère de théorèmes de ce genre dans un exposé usuel de mathématiques. Un lecteur non averti pourrait même penser que le deuxième de ces théorèmes est « faux », car on entend parfois dire que « le logarithme de -1 n'existe pas ». Mais cet usage du mot « existe »

n'a rien à voir avec le quantificateur existentiel de (3). Nous en parlerons plus loin dans ce paragraphe. Les théorèmes (3) ne sont nullement « contraires » à notre interprétation de « $\log_a(x)$ » en général, puisque nous n'avons pas d'interprétation de $\log_a(-1)$. Ce sont des vérités formelles indépendantes de toute interprétation.

Le dernier cas à discuter est celui d'un contexte dans lequel la proposition A admet plusieurs solutions x distinctes, voire une infinité. Certains auteurs¹ considèrent que le terme SxA n'a pas non plus d'interprétation dans ce cas. Mais il est possible d'admettre une interprétation « partielle » en disant que SxA désigne dans ce cas *une* solution de A , mais une solution *indéterminée*. Par exemple, le terme

$$Sx(x \in \{0, 1\})$$

désigne l'un des objets x qui vérifient la proposition $x \in \{0, 1\}$, autrement dit on a

$$\vdash Sx(x \in \{0, 1\}) = 0 \text{ ou } Sx(x \in \{0, 1\}) = 1.$$

Mais on ne peut rien dire de plus de cet objet. C'est en raison de cette interprétation que nous parlons d'opérateurs de *sélection* ou *sélecteurs*. Appliqué à la proposition $x \in \{0, 1\}$, l'opérateur Sx « *sélectionne* » un objet x qui vérifie cette proposition — mais sans que nous sachions lequel.

De manière générale, étant donné une proposition A et une variable x , nous pouvons dire ceci de l'interprétation du terme SxA : s'il existe un objet possédant la propriété exprimée par A pour x , alors SxA désigne un tel objet et c'est tout ce qu'on en sait.

Termes descriptifs dans le langage naturel. Le langage naturel n'utilise pas de variables donc n'a pas de termes descriptifs au sens syntaxique exact que nous venons de définir. Mais certains termes se traduisent naturellement par des termes descriptifs si l'on introduit des variables. Par exemple

l'homme qui a découvert l'Amérique

est un terme qui peut se traduire par

$$Sx(x \text{ est un homme et } x \text{ a découvert l'Amérique}).$$

Lorsqu'on parle de cet homme on pense traditionnellement à Christophe Colomb. C'est un exemple où la proposition à laquelle le sélecteur est appliqué est perçue comme ayant une solution et une seule, encore que cela puisse être discuté.

Il est plus difficile de donner des exemples de termes dont la formalisation soit de toute évidence un terme descriptif SxA dont la proposition A admet plusieurs solutions en x . À notre avis, un terme tel que

l'utilisateur

dans le « Manuel de l'utilisateur » d'un appareil, ou les termes

¹ Par exemple Willard V. O. Quine, *Méthodes de logique*, Paris, Armand Colin, 1972.

le lecteur de cet ouvrage,
l'Américain moyen,
l'homme de la rue

appartiennent à cette catégorie. Lorsqu'on utilise ce genre de terme, on sélectionne par la pensée un représentant indéterminé d'une certaine classe d'individus. On pourrait s'exprimer au pluriel, en parlant *des* utilisateurs, *des* lecteurs de cet ouvrage, *des* Américains moyens et *des* hommes de la rue. Mais on préfère penser à *un* représentant de chacune de ces classes et ce sont précisément de tels représentants que désignent les termes descriptifs

$Sx(x \text{ utilise l'appareil}),$
 $Sx(x \text{ lit cet ouvrage}),$
 $Sx(x \text{ est un Américain moyen}),$
 $Sx(x \text{ est un homme de la rue}).$

Schémas de déduction élémentaires. Les trois schémas de déduction suivants sont les schémas de déduction élémentaires de la logique des prédicats avec égalité relatifs aux sélecteurs. Notons que ce ne sont pas des schémas définitionnels, contrairement aux schémas élémentaires relatifs aux quantificateurs d'unicité (S93, S94). L'extension de Lprédégál définie par l'adjonction de ces trois schémas est néanmoins conservative².

SEL1.

$$\frac{\Gamma \vdash \exists xA \quad \Gamma \vdash T = SxA}{\Gamma \vdash (T|x)A} \quad \begin{array}{l} \text{Si } T \text{ est librement substituable} \\ \text{à } x \text{ dans } A. \end{array}$$

SEL2.

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash SxA = SxB} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array}$$

SEL3.

$$\frac{}{\Gamma \vdash SxA = Sy((y|x)A)} \quad \begin{array}{l} \text{Si } y \text{ est réversiblement} \\ \text{substituable à } x \text{ dans } A. \end{array}$$

Le schéma SEL1 correspond à ce que nous avons dit de l'interprétation du terme SxA . L'égalité $T = SxA$ signifie que les termes T et SxA désignent le même objet. Si T est librement substituable à x dans A , la proposition $(T|x)A$ exprime que l'objet désigné par T possède la propriété qui est exprimée par A pour x . Donc le schéma exprime que s'il existe un objet x ayant cette propriété, le terme SxA en désigne un.

Le schéma SEL2 est le schéma de substitutivité de l'équivalence dans un terme descriptif. L'objet désigné par SxA est invariant si l'on remplace A par une proposition équivalente B , à condition qu'il n'y ait pas d'hypothèse faite sur x dans le contexte. Par exemple on peut faire la déduction

$$\frac{\vdash x \in a \cup b \Leftrightarrow (x \in a \text{ ou } x \in b)}{\vdash Sx(x \in a \cup b) = Sx(x \in a \text{ ou } x \in b)}.$$

² A. C. Leisenring, *Mathematical Logic and Hilbert's ε -Symbol*, London, MacDonald, 1969.

Le schéma **SEL3** est le schéma de changement de variable liée dans les termes descriptifs. Par exemple, suivant ce schéma, on a dans **Lprédég** le théorème

$$\vdash Sx(a + x = b) = Sy(a + y = b).$$

Exemple 2. Au paragraphe 12.2 nous avons posé comme axiome de $\mathcal{T}_{\text{ENS}}$ sur les projections d'un couple:

E83. $\vdash z$ est un couple $\Rightarrow z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$.

Nous n'avons pas appelé cet axiome une définition car il n'est pas de l'une des formes définitionnelles notées au paragraphe 12.1. En utilisant des sélecteurs, on peut remplacer cet axiome par les deux *définitions* suivantes:

(a) $\vdash z$ est un couple $\Rightarrow \mathbf{pr}_1 z = Sa\exists b(z = (a, b))$.

(b) $\vdash z$ est un couple $\Rightarrow \mathbf{pr}_2 z = Sb\exists a(z = (a, b))$.

Ces axiomes sont de la troisième forme définitionnelle (§12.1). Leur signification est claire. La première projection d'un couple z est définie comme la solution a de la proposition $\exists b(z = (a, b))$ et la deuxième projection de z est définie symétriquement. Il faut bien voir la nécessité des quantificateurs $\exists a$, $\exists b$ dans ces définitions. Le terme $\mathbf{pr}_1 z$ doit être défini comme égal à un terme dans lequel le symbole $\mathbf{pr}_1$ ne figure pas et dans lequel il n'y a pas d'autre variable libre que z .

Le schéma **SEL1** permet de *démontrer* E83 à partir des définitions (a), (b):

	nil	
1°	z est un couple	hyp
2°	$\exists a\exists b(z = (a, b))$	1°, E81
3°	$\mathbf{pr}_1 z = Sa\exists b(z = (a, b))$	1°, (a)
4°	$\exists b(z = (\mathbf{pr}_1 z, b))$	2°, 3°, SEL1
5°	$\exists b\exists a(z = (a, b))$	2°, S72
6°	$\mathbf{pr}_2 z = Sb\exists a(z = (a, b))$	1°, (b)
7°	$\exists a(z = (a, \mathbf{pr}_2 z))$	5°, 6°, SEL1
8°	$z = (a, \mathbf{pr}_2 z)$	hyp
9°	$\exists b(z = (\mathbf{pr}_1 z, b))$	4°, a fortiori
10°	$z = (\mathbf{pr}_1 z, b)$	hyp
11°	$b = \mathbf{pr}_2 z$	8°, 10°, E79
12°	$z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	10°, 11°
13°	$z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	9°, 12°, S64
14°	$z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	7°, 13°, S64
15°	z est un couple $\Rightarrow z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	14°, $\Rightarrow$ -intro.

La déduction de la ligne 4° peut être mise en rapport comme suit avec le schéma **SEL1**, en désignant par Γ la liste d'hypothèses comprenant la seule

hypothèse « z est un couple » :

$$\frac{(2^\circ) \Gamma \vdash \exists a \overbrace{\exists b(z = (a, b))}^A \quad (3^\circ) \Gamma \vdash \overbrace{\mathbf{pr}_1 z = Sa}^T \exists b \overbrace{(z = (a, b))}^A}{(4^\circ) \Gamma \vdash \underbrace{\exists b(z = (\mathbf{pr}_1 z, b))}_{(T|a)A}}.$$

Symboles primordiaux. Nous appelons *symboles primordiaux* du langage $\mathcal{L}_\Omega$ les symboles suivants: les variables individuelles, les connecteurs logiques (et, ou, non, $\Rightarrow$, $\Leftrightarrow$), les quantificateurs ($\forall x$, $\exists x$), et les sélecteurs (Sx).

Ce qui vaut à ces symboles le qualificatif de « primordial » est le fait remarquable que l'on peut construire une version de la théorie des ensembles dans laquelle tous les axiomes et schémas de déduction qui ne sont pas des définitions n'utilisent que ces symboles. Comme toutes les théories mathématiques actuelles sont des extensions définitionnelles de la théorie des ensembles, cela signifie que l'on peut en principe éliminer dans tous les théorèmes de toutes les mathématiques tous les symboles qui n'appartiennent pas à ce petit répertoire de symboles primordiaux.

On peut aller plus loin et éliminer encore: les connecteurs logiques $\Rightarrow$, $\Leftrightarrow$ au moyen de **S53** et **S59.1**; le connecteur « et » au moyen de **S52.4**; les quantificateurs existentiels au moyen de **S71.4**. En principe, tous les théorèmes mathématiques peuvent ainsi se traduire dans un langage formel dont les seuls symboles autres que les variables individuelles sont $\neg$, $\vee$, $\forall$, S , $=$, $\in$. Ce langage n'a ni constantes individuelles, ni symboles fonctionnels. Nous verrons plus bas qu'on peut encore se passer du symbole $\forall$.

Pour pouvoir éliminer par exemple la constante individuelle $\emptyset$, il suffit de remplacer notre axiome de l'ensemble vide $\vdash \forall x(x \notin \emptyset)$ (**E2**) par les deux axiomes suivants, le deuxième étant une définition:

- (a) $\vdash \exists y \forall x(x \notin y)$.
- (b) $\vdash \emptyset = Sy \forall x(x \notin y)$.

Le premier ne contient que des symboles primordiaux. Il affirme seulement l'existence d'un ensemble sans élément et le deuxième définit l'ensemble $\emptyset$ comme la solution en y de la proposition $\forall x(x \notin y)$. De ces deux axiomes on déduit $\vdash \forall x(x \notin \emptyset)$ suivant **SEL1**:

$$\frac{\vdash \exists y \overbrace{\forall x(x \notin y)}^A \quad \vdash \emptyset = Sy \overbrace{\forall x(x \notin y)}^A}{\vdash \underbrace{\forall x(x \notin \emptyset)}_{(\emptyset|y)A}}.$$

Pour pouvoir éliminer le symbole fonctionnel binaire $\cup$, il suffit de remplacer notre axiome $\vdash x \in a \cup b \Leftrightarrow (x \in a \text{ ou } x \in b)$ (**E18**) par les deux axiomes suivants:

- (c) $\vdash \exists y \forall x (x \in y \Leftrightarrow (x \in a \text{ ou } x \in b))$.
 (d) $\vdash a \cup b = \mathbf{S}y \forall x (x \in y \Leftrightarrow (x \in a \text{ ou } x \in b))$.

Le premier n'utilise que des symboles primordiaux. Il affirme seulement l'existence d'un ensemble ayant la propriété caractéristique de la réunion des ensembles a et b . Le deuxième est une définition de $a \cup b$. Notre axiome E18 se *déduit* de (c) et (d) suivant SEL1 et S65.

Nous avons déjà remarqué que notre axiome de l'inclusion (E1) est en fait une définition. Tous les autres axiomes non définitionnels de $\mathcal{T}_{\text{ENS}}$ que nous avons introduits jusqu'ici — intersection de deux ensembles (E28), ensemble complémentaire (E40), ensembles à deux éléments (E51), ensemble des parties d'un ensemble (E71) — peuvent être remplacés, comme nous venons de le faire pour celui de la réunion de deux ensembles, par un axiome d'existence n'utilisant que des symboles primordiaux et une définition au moyen d'un sélecteur. Dans le cas de l'axiome E71 il faut commencer par éliminer le symbole non primordial $\subset$ au moyen de E1.

Sélecteurs et quantificateurs. Les deux schémas dérivés suivants présentent une remarquable symétrie. Ils permettent d'éliminer tous les quantificateurs d'une proposition, en utilisant à leur place des sélecteurs.

SEL4.

$$\frac{}{\Gamma \vdash \exists x A \Leftrightarrow (\mathbf{S}x A | x) A} \quad \frac{}{\Gamma \vdash \forall x A \Leftrightarrow (\mathbf{S}x \text{non} A | x) A}$$

Si $\mathbf{S}x A$ est librement substituable à x dans A .

Démonstration. Si le terme $\mathbf{S}x A$ est librement substituable à x dans A , alors de l'hypothèse $\exists x A$ et de l'égalité $\mathbf{S}x A = \mathbf{S}x A$ on déduit $(\mathbf{S}x A | x) A$ suivant SEL1. Inversement, de l'hypothèse $(\mathbf{S}x A | x) A$ on déduit $\exists x A$ suivant S63.

Si le terme $\mathbf{S}x A$ est librement substituable à x dans A , il en est de même du terme $\mathbf{S}x \text{non} A$ car les deux termes possèdent les mêmes variables libres. D'autre part, si ce terme T est librement substituable à x dans A , il est aussi librement substituable à x dans $\text{non} A$ et la proposition $(T | x) \text{non} A$ est identique à la proposition $\text{non}(T | x) A$ (§4.5, §4.6). Dans ces conditions, on peut écrire la démonstration suivante sous une liste d'hypothèses vide:

$$\begin{array}{ll} \forall x A \Leftrightarrow \text{non}(\exists x \text{non} A) & \text{S71} \\ \Leftrightarrow \text{non}((\mathbf{S}x \text{non} A | x) \text{non} A) & \text{SEL4 (premier)} \\ \Leftrightarrow \text{nonnon}((\mathbf{S}x \text{non} A | x) A) & \text{S46} \\ \Leftrightarrow (\mathbf{S}x \text{non} A | x) A & \text{S51} \end{array}$$

La troisième ligne est justifiée par S46, puisque les propositions

$$\text{non}((\mathbf{S}x \text{non} A | x) \text{non} A), \quad \text{nonnon}((\mathbf{S}x \text{non} A | x) A)$$

sont identiques.

Les schémas SEL4 permettent d'éliminer le quantificateur initial d'une proposition $\exists x A$, resp. $\forall x A$, malgré la condition « $\mathbf{S}x A$ est librement substi-

	nil		
(SEL6)	1°	$\exists xA$	hyp
	2°	$x = SxA$	hyp
	3°	$(x x)A$	1°, 2°, SEL1
	4°	$x = SxA \Rightarrow A$	3°, $\Rightarrow$ -intro
	5°	$\forall x(x = SxA \Rightarrow A)$	4°, S62
	6°	$\forall x(x = SxA \Rightarrow A)$	hyp
	7°	$x = SxA \Rightarrow A$	6°, S65
	8°	$\exists x(x = SxA) \Rightarrow \exists xA$	7°, S69
	9°	$\exists x(x = SxA)$	S89
	10°	$\exists xA$	8°, 9°
	11°	$\exists xA \Leftrightarrow \forall x(x = SxA \Rightarrow A)$	5°, 10°, S40.
(SEL7.1)	1°	HxA	hyp
	2°	A	hyp
	3°	$y = SxA$	hyp
		où y est une variable distincte de x et réversiblement substituable à x dans A .	
	4°	$\exists xA$	2°, S66
	5°	$(y x)A$	3°, 4°, SEL1
	6°	A et $(y x)A$	2°, 5°
	7°	$x = y$	1°, 6°, S93
	8°	$x = SxA$	7°, 3°
	9°	$x = SxA$	8°, S90
	10°	$\forall x(A \Rightarrow x = SxA)$	9°, $\Rightarrow$ -intro, S62
	11°	$\forall x(A \Rightarrow x = SxA)$	hyp
	12°	HxA	11°, S65, S98
	13°	$HxA \Leftrightarrow \forall x(A \Rightarrow x = SxA)$	10°, 12°, S40.
(SEL7.2)		$\exists! xA \Leftrightarrow (\exists xA \text{ et } HxA)$	S94
		$\Leftrightarrow \forall x(x = SxA \Rightarrow A) \text{ et } \forall x(A \Rightarrow x = SxA)$	SEL6, SEL7.1
		$\Leftrightarrow \forall x(A \Leftrightarrow x = SxA)$	S74, S41.

Figure 13.1 Démonstrations. Le fait que x ne figure pas librement dans SxA intervient à la ligne 9° de la première démonstration (S89) et aux lignes 9° (S90) et 12° (S98) de la deuxième.

tuable à x dans A », car on peut toujours effectuer dans A des changements de variables liées de manière à ce que cette condition soit satisfaite.

SEL5.

$$\frac{}{\Gamma \vdash (Sx(x = T)) = T}$$

Si x ne figure pas librement dans T .

Ce schéma est intuitivement bien évident : la solution en x d'une équation de la forme $x = T$ dans laquelle x ne figure librement qu'à gauche du signe

d'égalité est égale à T . Démonstration: Si x ne figure pas librement dans le terme T , la proposition $\exists x(x = T)$ est vraie (S89) et elle est équivalente à $(Sx(x = T)|x)(x = T)$ d'après SEL4. Or cette dernière proposition n'est autre que $(Sx(x = T)) = T$.

Sélecteurs et quantificateurs d'unicité. Les quantificateurs d'unicité Hx , $\exists!x$ peuvent s'exprimer très simplement au moyen de quantificateurs universels et de sélecteurs. On commence par exprimer de cette manière les quantificateurs existentiels, au moyen du schéma dérivé suivant.

SEL6.

$$\Gamma \vdash \exists xA \Leftrightarrow \forall x(x = SxA \Rightarrow A).$$

Nous rappelons que la variable x ne figure pas librement dans le terme SxA . Le x qui est à gauche du signe $=$ est lié au quantificateur $\forall x$ et n'a pas de rapport avec les occurrences de x dans le terme SxA .

Démonstration. Si le terme SxA est librement substituable à x dans A , la démonstration de SEL6 est immédiate suivant S92 et SEL4. La démonstration donnée dans la figure 13.1 ne demande pas que cette condition soit satisfaite. Les deux schémas suivants sont également démontrés dans la figure 13.1.

SEL7.

$$\begin{array}{l} 1. \frac{}{\Gamma \vdash HxA \Leftrightarrow \forall x(A \Rightarrow x = SxA)}. \\ 2. \frac{}{\Gamma \vdash \exists!xA \Leftrightarrow \forall x(A \Leftrightarrow x = SxA)}. \end{array}$$

Schémas de résolution d'une proposition univoque. Si x est une variable, on dit qu'une proposition A est *univoque en x* dans un contexte si la proposition HxA (il existe au plus un x tel que A) est vraie dans ce contexte. Si en outre la proposition $\exists xA$ est vraie, donc si $\exists!xA$ est vraie dans le contexte, alors la proposition A peut se « résoudre par rapport à x », en ce sens qu'elle est équivalente à l'égalité $x = SxA$ dans laquelle x ne figure librement qu'à gauche du signe $=$. C'est ce qu'exprime le schéma SEL8.2 ci-dessous.

Si l'on a démontré seulement HxA dans un contexte, la proposition A n'est pas équivalente à $x = SxA$, mais à la conjonction de $\exists xA$ et $x = SxA$ (SEL8.1).

Les démonstrations de SEL8.1–2 sont données dans la figure 13.2.

SEL8.

$$\begin{array}{ll} 1. \frac{\Gamma \vdash HxA}{\Gamma \vdash A \Leftrightarrow (\exists xA \text{ et } x = SxA)}. & 2. \frac{\Gamma \vdash \exists!xA}{\Gamma \vdash A \Leftrightarrow x = SxA}. \end{array}$$

Le pseudo-prédicat « existe ». Le langage mathématique courant n'utilise pas formellement d'opérateurs de sélection Sx , mais ceux-ci sont présents sous la forme d'expressions dans lesquelles il est question de l'objet x qui possède une certaine propriété ou qui satisfait à une certaine condition. Ces expressions se trouvent généralement dans les définitions. Nous avons vu la définition de $\log_a(y)$, qui est *le nombre x tel que $a^x = y$* .

Γ			Γ		
1°	HxA	<i>base</i>	1°	$\exists!xA$	<i>base</i>
2°	A	hyp	2°	$\exists xA$	1°, S94
3°	$\exists xA$	2°, S66	3°	HxA	1°, S94
4°	$A \Rightarrow x = SxA$	1°, SEL7	$A \Leftrightarrow (\exists xA \text{ et } x = SxA)$		3°, SEL8.1
5°	$\exists xA \text{ et } x = SxA$	2°–4°	$\Leftrightarrow x = SxA$		2°, S60.1.
6°	$\exists xA \text{ et } x = SxA$	hyp			
7°	$x = SxA \Rightarrow A$	6°, SEL6			
8°	A	6°, 7°			
9°	$A \Leftrightarrow (\exists xA \text{ et } x = SxA).$				

Figure 13.2 Démonstration de SEL8.1–2.

Nous voulons parler ici d'une tournure de langage fréquemment employée à la suite d'une définition de ce genre. Après une définition de la forme

$$\vdash T = SxA \quad (4)$$

ou de la forme conditionnelle $\vdash C \Rightarrow T = SxA$, dans laquelle le terme T contient un nouvel opérateur, on écrit souvent « T existe» à la place de $\exists xA$.

En général, une définition de la forme (4) est posée pour une proposition A et une variable x telles que l'on ait le théorème $\vdash HxA$. Suivant SEL8.1, on a alors le théorème $\vdash A \Leftrightarrow (\exists xA \text{ et } x = T)$ et on l'écrit souvent

$$\vdash A \Leftrightarrow (T \text{ existe et } x = T).$$

De même, une définition conditionnelle de la forme $\vdash C \Rightarrow T = SxA$ est posée en général lorsque $\vdash C \Rightarrow HxA$ est un théorème. On a alors d'après SEL8.1 le théorème $C \vdash A \Leftrightarrow (\exists xA \text{ et } x = T)$ que l'on écrit

$$C \vdash A \Leftrightarrow (T \text{ existe et } x = T).$$

En outre, on commet assez souvent l'*abus* de supprimer « T existe» dans la conjonction $(T \text{ existe et } x = T)$, en convenant que « $x = T$ » signifie « $x = T$ et T existe». Un exemple typique est donné ci-dessous (Exemple 3) et le lecteur peut vouloir l'examiner avant de lire ce qui suit à propos du pseudo-prédicat «*existe*».

Le mot «*existe*» dans l'expression « T existe» n'est pas un quantificateur, mais une sorte de prédicat (symbole relationnel unaire) post-fixé. Il ne faut surtout pas confondre une assertion de la forme « T existe», qui représente la proposition $\exists xA$ lorsqu'on a posé $T = SxA$, avec la proposition $\exists x(x = T)$ qui est toujours vraie (S89) lorsque x ne figure pas librement dans T —et c'est toujours le cas lorsqu'on pose une définition de la forme $T = SxA$.

Nous parlons d'un pseudo-prédicat, parce qu'il ne s'emploie qu'avec certains termes particuliers, dans le genre de situation que nous venons de décrire. Nous l'emploierons au chapitre 14, mais chaque emploi fera l'objet d'une définition

particulière. Il n'y a pas de définition générale de la forme $\vdash x$ existe $\Leftrightarrow R(x)$ ou de la forme conditionnelle $\vdash C \Rightarrow (x \text{ existe} \Leftrightarrow R(x))$. On ne peut pas adopter non plus un nouveau schéma de déduction élémentaire de **Lprédég** tel que

$$\frac{}{\Gamma \vdash Sx A \text{ existe} \Leftrightarrow \exists x A} \quad (?)$$

car cela rendrait cette théorie contradictoire. En effet on démontrerait alors pour n'importe quel terme T le théorème $\vdash T$ existe, en prenant pour A la proposition $x = T$ où x est une variable ne figurant pas librement dans T et en appliquant **SEL5** et **S89**.

Exemple 3. Nous prenons comme exemple la définition de la limite d'une suite (x_n) de nombres réels. Ce nombre, noté $\lim_{n \rightarrow \infty} x_n$ est défini brièvement comme *le nombre a tel que la suite (x_n) converge vers a , s'il existe*. Si nous désignons par C la proposition « (x_n) est une suite de nombres réels» et par $Conv(x, a)$ la proposition «la suite (x_n) converge vers a », cette définition conditionnelle est de la forme

$$\vdash C \Rightarrow \lim_{n \rightarrow \infty} x_n = Sa Conv(x, a). \quad (5)$$

Nous ne rappelons pas la définition de $Conv(x, a)$, mais nous rappelons qu'une suite de nombres réels converge vers *au plus* un nombre a , autrement dit qu'on a le théorème d'unicité:

$$\vdash C \Rightarrow Ha Conv(x, a). \quad (6)$$

De (5) et (6), suivant **SEL8.1**, on déduit le théorème

$$C \vdash Conv(x, a) \Leftrightarrow (\exists a Conv(x, a) \text{ et } a = \lim_{n \rightarrow \infty} x_n). \quad (7)$$

Sous la plume de maint auteur ce théorème s'écrit

$$C \vdash Conv(x, a) \Leftrightarrow (\lim_{n \rightarrow \infty} x_n \text{ existe et } a = \lim_{n \rightarrow \infty} x_n) \quad (8)$$

et même, par *abus de langage*:

$$C \vdash Conv(x, a) \Leftrightarrow (a = \lim_{n \rightarrow \infty} x_n). \quad (9)$$

Le jugement (9) *n'est pas un théorème* de la théorie considérée. On le voit immédiatement par le fait qu'on peut en déduire

$$C \vdash \exists a Conv(x, a) \Leftrightarrow \exists a (a = \lim_{n \rightarrow \infty} x_n) \quad (10)$$

suivant **S70**. Or

$$C \vdash \exists a (a = \lim_{n \rightarrow \infty} x_n) \quad (11)$$

est un théorème de **Lprédég** (**S89**). Donc si l'on admet (10), n'importe quelle suite de nombres réels est convergente. Ni le jugement (9), ni le jugement (10) ne sont des théorèmes, mais bien le jugement (8), si l'on convient que « $\lim_{n \rightarrow \infty} x_n$ existe» représente la proposition $\exists a Conv(x, a)$.

C'est pourtant une habitude bien ancrée dans le langage mathématique courant que de « faire comme si » (9) et (10) étaient des théorèmes. Le lecteur non averti peut voir ici une divergence entre la logique et les mathématiques usuelles à propos des propriétés de l'égalité. Mais ce n'est pas le cas. Lorsqu'on omet dans le membre de droite de (8) la sous-proposition $\lim_{n \rightarrow \infty} x_n$ existe, on donne au signe $=$ un sens spécial qu'il n'a pas normalement et qui en fait un « autre » signe d'égalité. Mentalement, on lit (9) avec un signe $='$ pour lequel on convient que

$$a = ' \lim_{n \rightarrow \infty} x_n \Leftrightarrow (\text{la suite } (x_n) \text{ converge et } a = \lim_{n \rightarrow \infty} x_n).$$

Exemple 4. Un exemple important en théorie des ensembles est celui d'un terme de la forme $\{x \mid P(x)\}$ où $P(x)$ est une proposition exprimant une propriété de x et le terme désigne l'ensemble des objets qui possèdent cette propriété. Cet ensemble est défini comme *l'ensemble y tel que $\forall x(x \in y \Leftrightarrow P(x))$* , autrement dit par

$$\{x \mid P(x)\} = \text{Sy} \forall x(x \in y \Leftrightarrow P(x)).$$

On suppose que la variable y ne figure pas librement dans la proposition P . Dire que « $\{x \mid P(x)\}$ existe » est une manière de dire: $\exists y \forall x(x \in y \Leftrightarrow P(x))$. On peut affirmer par exemple que l'ensemble $\{x \mid x \notin x\}$ n'existe pas, puisque $\vdash \text{non} \exists y \forall x(x \in y \Leftrightarrow x \notin x)$ est un théorème de **Lpréd** (§10.3, Théorème du barbier).

13.3 Termes conditionnels

Nous posons dans ce paragraphe la définition générale (schéma définitionnel) des termes conditionnels. Nous en avons donné un exemple précédemment (§13.1, Exemple 4), à savoir le terme

$$\text{Si } x \geq 0 \text{ alors } x \text{ sinon } -x, \quad (1)$$

utilisé dans la définition usuelle de la valeur absolue d'un nombre réel x . Ce terme est composé de: l'opérateur ternaire *Si alors sinon* que nous abrègerons par *Si*; la proposition $x \geq 0$; les deux termes x , $-x$. Le nombre (1) sera défini comme *le nombre y tel que*

$$(x \geq 0 \Rightarrow y = x) \text{ et } (\text{non}(x \geq 0) \Rightarrow y = -x),$$

autrement dit le nombre $\text{Sy}[(x \geq 0 \Rightarrow y = x) \text{ et } (\text{non}(x \geq 0) \Rightarrow y = -x)]$.

Ce paragraphe contient quelques schémas de déduction sur les termes conditionnels. Ils sont numérotés TC1–TC7. Le schéma TC1 est un schéma définitionnel. Il est pris comme schéma de déduction élémentaire de **Lprédég**, et ceci est la dernière extension de cette théorie à laquelle il sera procédé dans cet ouvrage. Nous rappelons que les précédentes ont été l'adjonction des schémas définitionnels S93, S94 sur les quantificateurs d'unicité et des schémas

SEL1, SEL2, SEL3 sur les sélecteurs. Le schéma TC0 ci-dessous est un schéma dérivé préalable de Lprédégal.

Existence des objets conditionnels. Le schéma dérivé suivant de Lprédégal garantit que *l'objet* x *tel que* $(A \Rightarrow x = T_1)$ et $(\text{non}A \Rightarrow x = T_2)$ existe, sous les conditions indiquées.

TC0.

$$\frac{}{\Gamma \vdash \exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]}.$$

Si x ne figure librement ni dans A , ni dans T_1 , ni dans T_2 .

Démonstration:

nil

1°	A	hyp
2°	$x = T_1$	hyp
3°	$A \Rightarrow x = T_1$	2°, S21
4°	$\text{non}A \Rightarrow x = T_2$	1°, S22
5°	$\exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$	2°, 3°, et-intro, S66
6°	$\exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$	5°, S90
7°	$\text{non}A$	hyp
8°	$x = T_2$	hyp
9°	$A \Rightarrow x = T_1$	7°, S22
10°	$\text{non}A \Rightarrow x = T_2$	8°, S21
11°	$\exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$	9°, 10°, et-intro, S66
12°	$\exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$	11°, S90
13°	$\exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$	6°, 12°, S35.

On peut d'ailleurs renforcer le schéma TC0 en remplaçant le quantificateur $\exists x$ par $\exists!x$. Sous les conditions du schéma, on démontre en effet facilement le théorème $\vdash Hx[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$, par disjonction des cas A , $\text{non}A$. Car si R désigne la proposition $[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]$, alors $R \Rightarrow x = T_1$ est vraie lorsque A est vraie; $R \Rightarrow x = T_2$ est vraie lorsque $\text{non}A$ est vraie. Dans les deux cas HxR est vraie d'après S98.

Définition. Si A est une proposition et si T_1 et T_2 sont des termes, le terme $SiAT_1T_2$ est défini par le schéma de déduction élémentaire suivant:

TC1.

$$\frac{}{\Gamma \vdash SiAT_1T_2 = Sx[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)]}$$

Si x ne figure librement ni dans A , ni dans T_1 , ni dans T_2 .

Le terme $SiAT_1T_2$ s'écrit aussi $Si(A, T_1, T_2)$, ou encore $Si A \text{ alors } T_1 \text{ sinon } T_2$, ou encore verticalement

$$\left\{ \begin{array}{l} T_1 \quad si \ A; \\ T_2 \quad sinon \end{array} \right\}.$$

Le schéma **TC1** fait désormais partie des schémas de déduction élémentaires de **Lprédég**, de même que les schémas **SEL1**, **SEL2**, **SEL3** sur les sélecteurs.

Substitutivité de l'équivalence. Le schéma dérivé suivant est le schéma de substitutivité de l'équivalence dans les termes conditionnels.

TC2.

$$\frac{\Gamma \vdash A \Leftrightarrow B}{\Gamma \vdash SiAT_1T_2 = SiBT_1T_2.}$$

Pour démontrer une déduction de cette forme, on prend une variable x qui ne figure librement ni dans A , ni dans B , ni dans T_1 , ni dans T_2 , ni dans Γ et l'on écrit la séquence de jugements suivante. La deuxième ligne se déduit de la première suivant **S50** (en plusieurs pas). La troisième se déduit de la deuxième suivant **SEL2**.

$$\begin{array}{l|l} \Gamma & \\ \hline 1^o & A \Leftrightarrow B \quad \text{base} \\ 2^o & [(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)] \Leftrightarrow [(B \Rightarrow x = T_1) \text{ et } (\text{non}B \Rightarrow x = T_2)] \\ 3^o & Sx[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)] = Sx[(B \Rightarrow x = T_1) \text{ et } (\text{non}B \Rightarrow x = T_2)] \\ 4^o & SiAT_1T_2 = SiBT_1T_2 \quad 3^o, \text{ TC1.} \end{array}$$

Schémas de déduction principaux. Les schémas de déduction suivants sont ceux que l'on utilise le plus souvent lorsqu'on raisonne sur des termes conditionnels.

TC3.

$$\frac{\Gamma \vdash A}{\Gamma \vdash SiAT_1T_2 = T_1} \quad \frac{\Gamma \vdash \text{non}A}{\Gamma \vdash SiAT_1T_2 = T_2.}$$

Pour démontrer ces deux déductions, dans lesquelles A est une proposition, T_1 et T_2 sont des termes et Γ est une liste de propositions quelconque, il suffit de démontrer le théorème:

$$\Gamma \vdash (A \Rightarrow SiAT_1T_2 = T_1) \text{ et } (\text{non}A \Rightarrow SiAT_1T_2 = T_2).$$

Or celui-ci est une conséquence immédiate de **TC0** et **TC1** suivant **SEL1**. En effet, en prenant une variable x qui ne figure librement ni dans A , ni dans T_1 , ni dans T_2 , on peut écrire la démonstration:

$$\begin{array}{l|l} \Gamma & \\ \hline 1^o & \exists x[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)] \quad \text{TC0} \\ 2^o & SiAT_1T_2 = Sx[(A \Rightarrow x = T_1) \text{ et } (\text{non}A \Rightarrow x = T_2)] \quad \text{TC1} \\ 3^o & (A \Rightarrow SiAT_1T_2 = T_1) \text{ et } (\text{non}A \Rightarrow SiAT_1T_2 = T_2) \quad 1^o, 2^o, \text{ SEL1.} \end{array}$$

Autres schémas dérivés. Les autres schémas de déduction dérivés sur les termes conditionnels se démontrent facilement par disjonction des cas suivant les schémas **TC3**. Cette démonstration est laissée au lecteur. On raisonne par disjonction des cas A , $\text{non}A$ pour **TC4** et **TC6**; par disjonction des cas B , $\text{non}B$ pour **TC7**.

TC4.

$$\frac{}{\Gamma \vdash SiAT\ T = T.}$$

TC5.

$$\frac{\Gamma \vdash A \Rightarrow T_1 = T'_1}{\Gamma \vdash SiAT_1T_2 = SiAT'_1T_2} \quad \frac{\Gamma \vdash \text{non}A \Rightarrow T_2 = T'_2}{\Gamma \vdash SiAT_1T_2 = SiAT_1T'_2.}$$

TC6. *Substitutivité de l'égalité*

$$\frac{\Gamma \vdash T_1 = T'_1}{\Gamma \vdash SiAT_1T_2 = SiAT'_1T_2} \quad \frac{\Gamma \vdash T_2 = T'_2}{\Gamma \vdash SiAT_1T_2 = SiAT_1T'_2.}$$

TC7.

$$\frac{\Gamma \vdash A \Rightarrow B}{\Gamma \vdash SiAT_1(SiBT_2T_3) = SiB(SiAT_1T_2)T_3.}$$

Exemple. En théorie des nombres réels, on définit la valeur absolue d'un nombre x en posant

$$\vdash |x| = Si(x \geq 0, x, -x).$$

On définit aussi parfois la *partie positive* x_+ et la *partie négative* x_- de x en posant

$$\begin{aligned} x_+ &= Si(x \geq 0, x, 0); \\ x_- &= Si(x \geq 0, 0, -x). \end{aligned} \tag{2}$$

Alors on peut démontrer comme suit le théorème $\vdash |x| = x_+ + x_-$.

nil		
1°	$x \geq 0 \Rightarrow x_+ = x$	(2), TC3
2°	$x \geq 0 \Rightarrow x_- = 0$	(2), TC3
3°	$\text{non}(x \geq 0) \Rightarrow x_+ = 0$	(2), TC3
4°	$\text{non}(x \geq 0) \Rightarrow x_- = -x$	(2), TC3
	$x_+ + x_- = Si(x \geq 0, x_+ + x_-, x_+ + x_-)$	TC4
	$= Si(x \geq 0, x + 0, 0 + (-x))$	1°, 2°, 3°, 4°, TC5
	$= Si(x \geq 0, x, -x)$	TC6
	$= x .$	

14

Opérateurs de réunion et de collection de la théorie des ensembles

14.1 Opérateurs de réunion

Introduction. Le présent chapitre et le suivant sont entièrement consacrés à la théorie des ensembles. Dans l'interprétation de cette théorie, un terme T quelconque représente toujours un ensemble (§8.2). Désignons par $T(x)$ un terme « dépendant de la variable x ». Cela ne signifie rien de spécial quant au terme lui-même, mais indique seulement une manière de le considérer. Il n'est pas nécessaire que x figure librement dans un terme T pour que nous puissions considérer T comme « dépendant de x ». L'indépendance est le cas limite de la dépendance. Nous allons définir l'ensemble

$$\bigcup_{x \in A} T(x), \quad (1)$$

pour un ensemble A quelconque. Pour chaque élément x de l'ensemble A , le terme $T(x)$ représente un ensemble que nous considérons comme associé à x (Fig. 14.1). Au cas où x ne figure pas librement dans T , c'est le même ensemble T qui est associé à tous les éléments de A . L'ensemble désigné par l'expression (1) est la réunion de tous les ensembles $T(x)$ associés aux divers éléments x de A . Il est symbolisé par la zone grise de la figure 14.1.

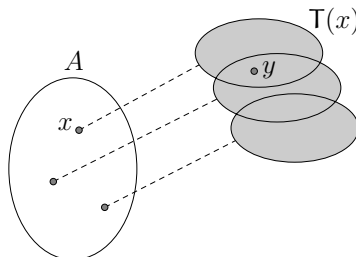


Figure 14.1 Réunion d'une famille d'ensembles.

Un objet y est élément de cette réunion d'ensembles si et seulement s'il appartient à l'un au moins des ensembles de cette famille, autrement dit, si et seulement s'il existe un x , élément de A , tel que $y \in T(x)$:

$$\vdash y \in \bigcup_{x \in A} T(x) \Leftrightarrow \exists x(x \in A \text{ et } y \in T(x)). \quad (2)$$

Par exemple, pour un ensemble A à deux éléments, $A = \{a, b\}$, nous aurons le théorème

$$\vdash \bigcup_{x \in \{a, b\}} T(x) = T(a) \cup T(b),$$

où $T(a)$, $T(b)$ désignent les termes $(a|x)T$, $(b|x)T$ et l'on suppose que les variables a et b sont librement substituables à x dans T .

Le théorème (2) peut se déduire d'une *définition* de l'ensemble (1) comme étant l'ensemble z tel que $\forall y(y \in z \Leftrightarrow \exists x(x \in A \text{ et } y \in T(x)))$, autrement dit de la définition

$$\vdash \bigcup_{x \in A} T(x) = Sz[\forall y(y \in z \Leftrightarrow \exists x(x \in A \text{ et } y \in T(x)))]. \quad (3)$$

Il suffit de postuler l'*existence* d'un tel ensemble z , à savoir de prendre comme axiome:

$$\vdash \exists z[\forall y(y \in z \Leftrightarrow \exists x(x \in A \text{ et } y \in T(x)))]. \quad (4)$$

De (4) et (3) on déduit (2) suivant SEL1.

Opérateurs de réunion. Dans le terme (1) la variable x est locale. Si x' est réversiblement substituable à x dans T , les termes

$$\bigcup_{x \in A} T(x), \quad \bigcup_{x' \in A} T(x') \quad (5)$$

désignent le même ensemble. Le premier de ces termes se compose de trois parties:

$$\bigcup_x A \quad T(x).$$

Le symbole $\in$ de (5) n'est que du « sucre syntaxique » sans importance logique. La première des trois parties est l'opérateur de *réunion sur x* que l'on peut noter Ux comme un quantificateur. La deuxième partie du terme est dans cet exemple la variable A . Dans le cas général, cette partie peut être un terme E quelconque.

De façon générale, à chaque variable individuelle x est associé un opérateur de réunion Ux de type $T \leftarrow (T, T)$. Si E et T sont deux termes quelconques, l'expression $UxET$ est un terme que l'on écrit généralement

$$\bigcup_{x \in E} T. \quad (6)$$

L'opérateur Ux lie la variable x dans son deuxième argument seulement (T). Les occurrences libres de x dans le terme E restent libres dans le terme $UxET$ (§13.1, Exemples 1 et 3).

Le schéma de réunion. Le schéma de déduction suivant est le premier schéma de déduction élémentaire spécifique de la théorie $\mathcal{T}_{\text{ENS}}$. Pour une variable x et des termes E, T , ce dernier étant considéré comme dépendant de x et désigné par $T(x)$, le schéma affirme l'existence d'un ensemble qui est la réunion des ensembles $T(x)$ correspondant aux éléments x de l'ensemble E .

E93.

$$\frac{}{\Gamma \vdash \exists u [\forall y (y \in u \Leftrightarrow \exists x (x \in E \text{ et } y \in T(x)))]}$$

Si x, u, y sont des variables distinctes et si:
 x ne figure pas librement dans E ;
 u et y ne figurent librement ni dans E , ni dans T .

Définition. Pour une variable x et des termes E, T quelconques, l'ensemble désigné par le terme (6) est défini par le schéma de déduction définitionnel suivant de $\mathcal{T}_{\text{ENS}}$:

E94.

$$\frac{}{\Gamma \vdash \bigcup_{x \in E} T(x) = \text{Su} [\forall y (y \in u \Leftrightarrow \exists x' (x' \in E \text{ et } y \in T(x')))]}$$

Si x', u, y sont des variables distinctes et si: x' est réversiblement substituable à x dans T et ne figure pas librement dans E ; $T(x')$ est le terme $(x'|x)T$; u et y ne figurent librement ni dans E , ni dans T .

Si x ne figure pas librement dans E , x' peut être la variable x . Sinon x' doit être distincte de x et ce changement de variable a pour but de ne pas lier les occurrences libres de x dans E , puisqu'elles ne sont pas liées par l'opérateur Ux à gauche de l'égalité. Les conditions sur les variables x', u, y font que les deux membres de l'égalité ont les mêmes variables libres, à savoir toutes les variables qui figurent librement dans E et toutes les variables autres que x qui figurent librement dans T .

Schéma dérivé principal. Le schéma qui découle directement des deux précédents suivant SEL1 est le schéma E95 ci-dessous. C'est celui qui est utilisé pratiquement dans la démonstration de tous les théorèmes dans lesquels interviennent des opérateurs de réunion.

E95.

$$\frac{}{\Gamma \vdash y \in \bigcup_{x \in E} T(x) \Leftrightarrow \exists x' (x' \in E \text{ et } y \in T(x'))}$$

Si x' et y sont des variables distinctes et si: x' est réversiblement substituable à x dans T et ne figure pas librement dans E ; $T(x')$ est le terme $(x'|x)T$; y ne figure librement ni dans E , ni dans T .

Démonstration. Si ces conditions sont satisfaites et si u est une variable distincte de x' et de y et ne figurant librement ni dans E , ni dans T , on peut écrire:

Γ		
1°	$\exists u [\forall y (y \in u \Leftrightarrow \exists x' (x' \in E \text{ et } y \in T(x')))]$	E93
2°	$\bigcup_{x \in E} T(x) = Su [\forall y (y \in u \Leftrightarrow \exists x' (x' \in E \text{ et } y \in T(x')))]$	E94
3°	$\forall y (y \in \bigcup_{x \in E} T(x) \Leftrightarrow \exists x (x \in E \text{ et } y \in T(x)))$	1°, 2°, SEL1
4°	$y \in \bigcup_{x \in E} T(x) \Leftrightarrow \exists x (x \in E \text{ et } y \in T(x))$	3°, S65.

Substitutivité de l'égalité et changement de variable liée. Les schémas généraux de substitutivité de l'égalité et de changement de variable liée sont applicables aux expressions qui contiennent des opérateurs de réunion. La définition E94 et le schéma SEL2 permettent de démontrer en effet les schémas de déduction suivants dans lesquels E , E' , $T(x)$, $T'(x)$ sont des termes et x , z sont des variables. Les démonstrations sont données dans l'annexe (Chap. 16).

E96.

$$\frac{\Gamma \vdash E = E'}{\Gamma \vdash \bigcup_{x \in E} T(x) = \bigcup_{x \in E'} T(x)}$$

$$\frac{\Gamma \vdash T(x) = T'(x)}{\Gamma \vdash \bigcup_{x \in E} T(x) = \bigcup_{x \in E} T'(x)} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas librement} \\ \text{dans } \Gamma. \end{array}$$

E97.

$$\frac{}{\Gamma \vdash \bigcup_{x \in E} T(x) = \bigcup_{z \in E} T(z)} \quad \begin{array}{l} \text{Si } z \text{ est réversiblement sub-} \\ \text{stituable à } x \text{ dans } T. \end{array}$$

Théorèmes de réunion. Les théorèmes de $\mathcal{T}_{\text{ENS}}$ qui suivent sont démontrés dans l'annexe. La plupart sont en fait des *schémas* de théorèmes puisqu'il font intervenir un terme $T(x)$ indéterminé, considéré comme dépendant de la variable x . On désigne par $T(a)$, $T(b)$ les termes $(a|x)T$, $(b|x)T$ et l'on suppose que a et b sont librement substituables à x dans T .

$$\text{E98.} \quad \vdash \bigcup_{x \in \{a,b\}} T(x) = T(a) \cup T(b).$$

$$\vdash \bigcup_{x \in \{a\}} T(x) = T(a).$$

$$\text{E99.} \quad \vdash \bigcup_{x \in \{a,b\}} x = a \cup b.$$

$$\vdash \bigcup_{x \in \{a\}} x = a.$$

$$\text{E100. } \vdash \bigcup_{x \in \emptyset} T(x) = \emptyset.$$

$$\text{E101. } \vdash A \subset B \Rightarrow \bigcup_{x \in A} T(x) \subset \bigcup_{x \in B} T(x).$$

$$\text{E102. } \vdash \bigcup_{x \in A \cup B} T(x) = \bigcup_{x \in A} T(x) \cup \bigcup_{x \in B} T(x).$$

$$\text{E103. } \vdash x \in A \Rightarrow T(x) \subset \bigcup_{x \in A} T(x).$$

$$\text{E104. } \vdash x \in A \Rightarrow T(x) \in \bigcup_{x \in A} \{T(x)\}.$$

$$\text{E105. } \vdash \bigcup_{x \in A} T(x) \subset Y \Leftrightarrow \forall x (x \in A \Rightarrow T(x) \subset Y).$$

Produit cartésien de deux ensembles. Les opérateurs de réunion permettent de définir d'autres opérations fréquentes de construction d'ensembles. L'une d'elles est le produit cartésien $X \times Y$ de deux ensembles. Le symbole $\times$ est un nouveau symbole fonctionnel binaire. Il ne figure dans aucun de nos axiomes précédents de $\mathcal{T}_{\text{ENS}}$, ni dans les schémas de déduction élémentaires de $\mathcal{T}_{\text{ENS}}$ introduits jusqu'ici — qui sont tous ceux de **Lprédég**, y compris les schémas définitionnels sur les quantificateurs d'unicité, les schémas sur les sélecteurs (**SEL1**–**SEL3**) et le schéma définitionnel **TC1** sur les termes conditionnels.

E106. Définition du produit cartésien

$$\vdash X \times Y = \bigcup_{x \in X} \left(\bigcup_{y \in Y} \{(x, y)\} \right).$$

Théorèmes. Les théorèmes suivants sont démontrés ci-dessous (**E107**) ou dans l'annexe (Chap. 16). Le théorème **E109** est dit « principal » car c'est le plus utilisé. En vertu de **E108**, on peut appliquer aux produits cartésiens les méthodes de démonstration d'inclusion ou d'égalité de graphes reposant sur **E91** et **E92** (§12.2, Fig. 12.4). Par exemple, pour démontrer **E111**, on fait l'hypothèse $X \subset Y$ et $Y \subset Y'$ et sous cette hypothèse on démontre que $(x, y) \in X \times Y \Rightarrow (x, y) \in X' \times Y'$.

$$\text{E107. } \vdash z \in X \times Y \Leftrightarrow (z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y).$$

La démonstration de ce théorème est la chaîne d'équivalences suivante, dans un cadre sans hypothèses :

$$z \in X \times Y \Leftrightarrow z \in \bigcup_{x \in X} \left(\bigcup_{y \in Y} \{(x, y)\} \right) \quad \text{E106, S86}$$

$$\Leftrightarrow \exists x (x \in X \text{ et } z \in \bigcup_{y \in Y} \{(x, y)\}) \quad \text{E95}$$

$$\Leftrightarrow \exists x (x \in X \text{ et } \exists y (y \in Y \text{ et } z \in \{(x, y)\})) \quad \text{E95}$$

$$\Leftrightarrow \exists x \exists y (x \in X \text{ et } y \in Y \text{ et } z = (x, y)) \quad \text{S81, E60}$$

- $\Leftrightarrow \exists x \exists y (x \in X \text{ et } y \in Y \text{ et } z = (x, y) \text{ et } x = \mathbf{pr}_1(x, y) \text{ et } y = \mathbf{pr}_2(x, y))$ E84, S60.1
 $\Leftrightarrow \exists x \exists y (x \in X \text{ et } y \in Y \text{ et } z = (x, y) \text{ et } x = \mathbf{pr}_1 z \text{ et } y = \mathbf{pr}_2 z)$ S87.1
 $\Leftrightarrow \exists x (x \in X \text{ et } \mathbf{pr}_2 z \in Y \text{ et } z = (x, \mathbf{pr}_2 z) \text{ et } x = \mathbf{pr}_1 z)$ S92
 $\Leftrightarrow \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y \text{ et } z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$ S92
 $\Leftrightarrow z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y$ E85.

E108. $\vdash X \times Y$ est un graphe.

E109. *Théorème principal*

$$\vdash (x, y) \in X \times Y \Leftrightarrow (x \in X \text{ et } y \in Y).$$

E110. $\vdash X \times Y = \emptyset \Leftrightarrow (X = \emptyset \text{ ou } Y = \emptyset).$

E111. $\vdash (X \subset X' \text{ et } Y \subset Y') \Rightarrow X \times Y \subset X' \times Y'.$

E112. $\vdash$ Si $X \neq \emptyset$ et $Y \neq \emptyset$, alors

$$(X \subset X' \text{ et } Y \subset Y') \Leftrightarrow X \times Y \subset X' \times Y'.$$

E113. $\vdash (X \times Y) \cap (X' \times Y') = (X \cap X') \times (Y \cap Y').$

E114. $\vdash X \times (A \cup B) = (X \times A) \cup (X \times B).$

$$\vdash (A \cup B) \times X = (A \times X) \cup (B \times X).$$

E115. $\vdash \{a\} \times \{b\} = \{(a, b)\}.$

$$\vdash \{a\} \times \{b, c\} = \{(a, b), (a, c)\}.$$

$$\vdash \{a, b\} \times \{c\} = \{(a, c), (b, c)\}.$$

$$\vdash \{a, b\} \times \{c, d\} = \{(a, c), (a, d), (b, c), (b, d)\}.$$

Projections d'un graphe. On appelle *première projection* d'un graphe G ou $\mathbf{Pr}_1 G$ l'ensemble dont les éléments sont les premières projections des couples appartenant à G . Symétriquement, on appelle *deuxième projection* de G ou $\mathbf{Pr}_2 G$ l'ensemble dont les éléments sont les deuxièmes projections des couples appartenant à G . Ces notions sont illustrées par les figures 14.2 et 14.3. Dans la deuxième, les projections de G sont les ensembles de nombres réels représentés par les segments épais des axes de coordonnées.

Formellement, les projections d'un graphe sont définies au moyen d'opérateurs de réunion.

E116. *Définition des projections d'un graphe*

$$\vdash G \text{ est un graphe} \Rightarrow \left(\mathbf{Pr}_1 G = \bigcup_{z \in G} \{\mathbf{pr}_1 z\} \text{ et } \mathbf{Pr}_2 G = \bigcup_{z \in G} \{\mathbf{pr}_2 z\} \right).$$

E117. *Théorèmes principaux*

$$G \text{ est un graphe} \vdash x \in \mathbf{Pr}_1 G \Leftrightarrow \exists y ((x, y) \in G).$$

$$G \text{ est un graphe} \vdash y \in \mathbf{Pr}_2 G \Leftrightarrow \exists x ((x, y) \in G).$$

E118. $\vdash G \text{ est un graphe} \Rightarrow G \subset \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G.$

E119. $\vdash$ Si G et G' sont des graphes, alors

$$G \subset G' \Rightarrow (\mathbf{Pr}_1 G \subset \mathbf{Pr}_1 G' \text{ et } \mathbf{Pr}_2 G \subset \mathbf{Pr}_2 G').$$

E120. $\vdash G \subset X \times Y \Rightarrow (\mathbf{Pr}_1 G \subset X \text{ et } \mathbf{Pr}_2 G \subset Y).$

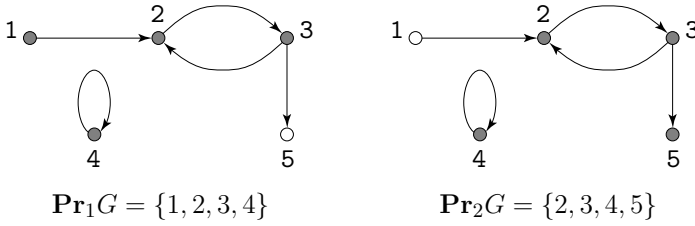
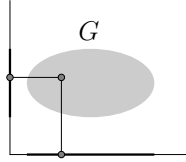


Figure 14.2 Projections d'un graphe fini.

Figure 14.3 Projections d'un ensemble G de points du plan en tant qu'ensemble de couples de coordonnées.E121. $\vdash Y \neq \emptyset \Rightarrow \text{Pr}_1(X \times Y) = X$. $\vdash X \neq \emptyset \Rightarrow \text{Pr}_2(X \times Y) = Y$.E122. G est un graphe $\vdash G = \emptyset \Leftrightarrow \text{Pr}_1 G = \emptyset$. G est un graphe $\vdash G = \emptyset \Leftrightarrow \text{Pr}_2 G = \emptyset$.

Grappe réciproque d'un graphe. On appelle *grappe réciproque* d'un graphe G et l'on note G^{-1} le graphe qui a pour éléments les couples *opposés* des couples appartenant à G . Par « couple opposé » d'un couple (x, y) nous désignons le couple (y, x) . Cette notion est illustrée par la figure 14.4. Le diagramme du graphe G^{-1} s'obtient en inversant le sens des flèches de celui de G .

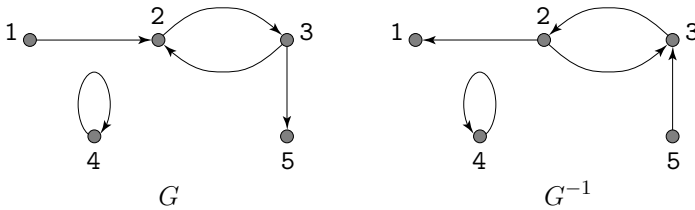


Figure 14.4 Grappe réciproque d'un graphe fini.

Syntaxiquement, le terme G^{-1} se compose d'un symbole fonctionnel *unaire* noté « $^{-1}$ » et du terme G . Il ne s'agit pas d'un cas particulier d'une opération binaire d'exponentiation G^n . Certains auteurs préfèrent parler du graphe *opposé* de G et le noter G^{op} .

Formellement, le graphe G^{-1} est défini au moyen d'un opérateur de réunion.

E123. *Définition*

$$\vdash G \text{ est un graphe} \Rightarrow G^{-1} = \bigcup_{z \in G} \{(\text{pr}_2 z, \text{pr}_1 z)\}.$$

Les théorèmes suivants sont démontrés dans l'annexe. Le théorème E125 est le principal, car c'est celui qui le plus utilisé dans les démonstrations.

E124. $\vdash G$ est un graphe $\Rightarrow G^{-1}$ est un graphe.

E125. *Théorème principal*

$\vdash$ Si G est un graphe, alors $(x, y) \in G^{-1} \Leftrightarrow (y, x) \in G$.

E126. G est un graphe $\vdash (G^{-1})^{-1} = G$.

E127. G est un graphe $\vdash \mathbf{Pr}_1(G^{-1}) = \mathbf{Pr}_2 G$ et $\mathbf{Pr}_2(G^{-1}) = \mathbf{Pr}_1 G$.

E128. G et G' sont des graphes $\vdash (G \cup G')^{-1} = G^{-1} \cup G'^{-1}$.

G et G' sont des graphes $\vdash (G \cap G')^{-1} = G^{-1} \cap G'^{-1}$.

E129. $\vdash \emptyset^{-1} = \emptyset$.

Graphe identique d'un ensemble. Le graphe identique d'un ensemble X , ou Id_X , est l'ensemble des couples de la forme (x, x) tels que $x \in X$. Il est défini formellement ci-dessous (E130) au moyen d'un opérateur de réunion. Cette notion est illustrée dans la figure 14.5 par le diagramme du graphe identique d'un ensemble à quatre éléments et par la « courbe » du graphe identique d'un intervalle X de $\mathbb{R}$. Dans cet exemple, Id_X est l'ensemble des points (couples de coordonnées) situés sur la diagonale du carré qui est l'ensemble de points $X \times X$. En raison de cet exemple, l'ensemble Id_X est appelé parfois le *sous-ensemble diagonal* de l'ensemble $X \times X$ même lorsque X n'est pas un intervalle de la droite réelle.

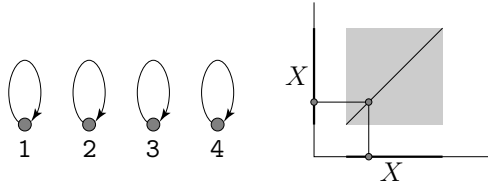


Figure 14.5 Graphe identique de l'ensemble $\{1, 2, 3, 4\}$ et graphe identique d'un intervalle X de $\mathbb{R}$.

E130. *Définition*

$$\vdash \text{Id}_X = \bigcup_{a \in X} \{(a, a)\}.$$

Théorèmes

E131. $\vdash \text{Id}_X \subset X \times X$.

E132. $\vdash \text{Id}_X$ est un graphe.

E133. *Théorème principal*

$\vdash (x, y) \in \text{Id}_X \Leftrightarrow (x \in X \text{ et } y = x)$.

E134. $\vdash \mathbf{Pr}_1(\text{Id}_X) = X$ et $\mathbf{Pr}_2(\text{Id}_X) = X$.

E135. $\vdash \text{Id}_\emptyset = \emptyset$.

E136. $\vdash \text{Id}_{(X \cup Y)} = \text{Id}_X \cup \text{Id}_Y$.

$\vdash \text{Id}_{(X \cap Y)} = \text{Id}_X \cap \text{Id}_Y$.

E137. $\vdash (\text{Id}_X)^{-1} = \text{Id}_X$.

14.2 Opérateurs de collection simples

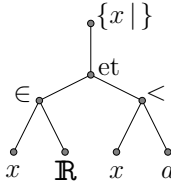
Schéma de définition. Nous avons présenté au paragraphe 13.1 (Exemple 2) les opérateurs de *collection* avec lesquels sont construits les termes tels que

$$\{x \mid x \in \mathbb{R} \text{ et } x < a\}. \quad (1)$$

Ce terme se lit «l'ensemble des x tels que: $x \in \mathbb{R}$ et $x < a$ ». Le préfixe «l'ensemble des x tels que» est un opérateur de type $T \leftarrow P$ qui lie la variable x dans la proposition avec laquelle il se combine — dans cet exemple, la proposition $(x \in \mathbb{R} \text{ et } x < a)$. Dans la notation traditionnelle (1), l'opérateur de collection se compose des symboles $\{x \mid \}$. On pourrait utiliser à la place un opérateur $\mathfrak{E}x$ de forme semblable à celle d'un quantificateur et écrire

$$\mathfrak{E}x(x \in \mathbb{R} \text{ et } x < a).$$

Nous utiliserons la notation traditionnelle, mais il faut bien voir que l'arbre du terme (1) est



L'interprétation du terme (1) est qu'il s'agit de *l'ensemble y tel que*

$$\forall x(x \in y \Leftrightarrow (x \in \mathbb{R} \text{ et } x < a)). \quad (2)$$

C'est dire que l'ensemble (1) est défini au moyen d'un opérateur de sélection. Il est la solution y de la proposition (2), c'est-à-dire que l'on a:

$$\{x \mid x \in \mathbb{R} \text{ et } x < a\} = Sy \forall x(x \in y \Leftrightarrow (x \in \mathbb{R} \text{ et } x < a)).$$

La définition générale de l'ensemble $\{x \mid R\}$, pour une proposition R et une variable x quelconques, est donnée par le schéma de déduction définitionnel suivant.

E138. Définition

$$\frac{}{\Gamma \vdash \{x \mid R\} = Sy \forall x(x \in y \Leftrightarrow R)}$$

Si y est distincte de x et ne figure pas librement dans R .

Un terme de la forme $\{x \mid R\}$ est appelé une *collection*. L'opérateur $\{x \mid \}$ est appelé un opérateur de collection *simple*. Nous introduirons plus loin des opérateurs de collection plus généraux.

Existence d'une collection. Le sens de l'expression « $\{x \mid R\}$ existe» est celui que nous avons défini au paragraphe 13.2 (Le pseudo-prédicat «existe»).

E139. Définition

$$\frac{}{\Gamma \vdash \{x \mid R\} \text{ existe} \Leftrightarrow \exists y \forall x(x \in y \Leftrightarrow R)}$$

Si y est distincte de x et ne figure pas librement dans R .

Par exemple, on sait que $\vdash \text{non} \exists y \forall x (x \in y \Leftrightarrow x \notin x)$ est un théorème de Lpréd (§10.3, Théorème du barbier). Suivant E139, on a donc dans $\mathcal{T}_{\text{ENS}}$ le théorème

$$\vdash \{x \mid x \notin x\} \text{ n'existe pas.}$$

Unicité d'une collection. Si l'ensemble $\{x \mid R\}$ existe au sens de E139, pour une proposition R et une variable x données, il est unique d'après le schéma de déduction suivant.

E140.

$$\frac{\Gamma \vdash Hy \forall x (x \in y \Leftrightarrow R)}{\text{Si } y \text{ est distincte de } x \text{ et ne figure pas librement dans } R.}$$

On démontre en effet facilement le théorème

$$\vdash \forall y \forall y' [(\underbrace{\forall x (x \in y \Leftrightarrow R)}_A \text{ et } \underbrace{\forall x (x \in y' \Leftrightarrow R)}_{(y'|y)A}) \Rightarrow y = y'],$$

pour deux variables y et y' distinctes l'une de l'autre, distinctes de x et ne figurant pas librement dans R . On en tire $\vdash Hy \forall x (x \in y \Leftrightarrow R)$ suivant S93, car y' est réversiblement substituable à y dans la proposition désignée par A .

nil		
1°	$\forall x (x \in y \Leftrightarrow R) \text{ et } \forall x (x \in y' \Leftrightarrow R)$	hyp
2°	$x \in y \Leftrightarrow x \in y'$	1°, et-élim, S65, S48
3°	$y = y'$	2°, E17
4°	$(\forall x (x \in y \Leftrightarrow R) \text{ et } \forall x (x \in y' \Leftrightarrow R)) \Rightarrow y = y'$	3°, $\Rightarrow$ -intro.

Utilisation principale des collections. Les collections sont utilisées, lorsqu'elles existent bien entendu, principalement dans les définitions de la forme

$$\vdash T = \{x \mid R\}$$

où T est un terme dans lequel x ne figure pas librement, ou dans les définitions conditionnelles de la forme $\vdash C \Rightarrow T = \{x \mid R\}$. Le schéma de déduction suivant s'applique alors. C'est aussi le schéma principal sur les collections.

E141.

$$\frac{\Gamma \vdash \forall x (x \in T \Leftrightarrow R) \Leftrightarrow (\{x \mid R\} \text{ existe et } T = \{x \mid R\})}{\text{Si } x \text{ ne figure pas librement dans } T.}$$

Démonstration. Si x ne figure pas librement dans T et si y est une variable distincte de x et ne figurant pas librement dans R , on a

nil		
1°	$Hy \forall x (x \in y \Leftrightarrow R)$	E140
2°	$\forall x (x \in y \Leftrightarrow R) \Leftrightarrow (\exists y \forall x (x \in y \Leftrightarrow R) \text{ et } y = Sy \forall x (x \in y \Leftrightarrow R))$	1°, SEL8
3°	$\forall x (x \in y \Leftrightarrow R) \Leftrightarrow (\{x \mid R\} \text{ existe et } y = \{x \mid R\})$	2°, E139, E138
4°	$\forall x (x \in T \Leftrightarrow R) \Leftrightarrow (\{x \mid R\} \text{ existe et } T = \{x \mid R\})$	3°, S79.

La condition « x ne figure pas librement dans T » est nécessaire pour la déduction de la ligne 4°, c'est-à-dire pour que T soit librement substituable à y dans la proposition de la ligne 3°.

Exemples. Du théorème $\vdash \forall x(x \in a \Leftrightarrow x \in a)$, on déduit suivant **E141** le théorème

$$\vdash \{x \mid x \in a\} \text{ existe et } a = \{x \mid x \in a\}.$$

Des déductions analogues peuvent se faire à partir des théorèmes suivants (cf. **E18**, **E51**, **E40**, **E71**, **E107**):

$$\begin{aligned} &\vdash \forall x(x \in a \cup b \Leftrightarrow (x \in a \text{ ou } x \in b)) \\ &\vdash \forall x(x \in \{a, b\} \Leftrightarrow (x = a \text{ ou } x = b)) \\ &\vdash \forall x(x \in \mathbb{C}_a(b) \Leftrightarrow (x \in a \text{ et } x \notin b)). \\ &\vdash \forall x(x \in \mathcal{P}(a) \Leftrightarrow x \subset a). \\ &\vdash \forall z(z \in X \times Y \Leftrightarrow (z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y)). \end{aligned}$$

On en déduit d'après **E141**:

$$\begin{aligned} &\vdash \{x \mid x \in a \text{ ou } x \in b\} \text{ existe et } a \cup b = \{x \mid x \in a \text{ ou } x \in b\}. \\ &\vdash \{x \mid x = a \text{ ou } x = b\} \text{ existe et } \{a, b\} = \{x \mid x = a \text{ ou } x = b\}. \\ &\vdash \{x \mid x \in a \text{ et } x \notin b\} \text{ existe et } \mathbb{C}_a(b) = \{x \mid x \in a \text{ et } x \notin b\}. \\ &\vdash \{x \mid x \subset a\} \text{ existe et } \mathcal{P}(a) = \{x \mid x \subset a\}. \\ &\vdash \{z \mid z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y\} \text{ existe et } \\ &\quad X \times Y = \{z \mid z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y\}. \end{aligned}$$

Enfin, du théorème $\vdash G \text{ est un graphe} \Rightarrow \forall x(x \in \mathbf{Pr}_1 G \Leftrightarrow \exists y((x, y) \in G))$ (cf. **E117**) on peut tirer: $\vdash$ Si G est un graphe, alors

$$\{x \mid \exists y((x, y) \in G)\} \text{ existe et } \mathbf{Pr}_1 G = \{x \mid \exists y((x, y) \in G)\}.$$

Exercice. Montrer que le jugement $\vdash \mathbb{R}_+ = \mathbb{R}_+^* \cup \{0\}$ se déduit dans $\mathcal{T}_{\text{ENS}}$ des jugements suivants:

$$\begin{aligned} &\vdash \mathbb{R}_+ = \{x \mid x \in \mathbb{R} \text{ et } x \geq 0\} \text{ et } \{x \mid x \in \mathbb{R} \text{ et } x \geq 0\} \text{ existe.} \\ &\vdash \mathbb{R}_+^* = \{x \mid x \in \mathbb{R} \text{ et } x > 0\} \text{ et } \{x \mid x \in \mathbb{R} \text{ et } x > 0\} \text{ existe.} \\ &\vdash 0 \in \mathbb{R}. \\ &\vdash x \in \mathbb{R} \Rightarrow (x \geq 0 \Leftrightarrow (x > 0 \text{ ou } x = 0)). \end{aligned}$$

Corollaires. Nous mentionnons ici deux schémas de déduction qui sont des corollaires immédiats de **E141**.

E142.

$$\Gamma \vdash \{x \mid R\} \text{ existe} \Leftrightarrow \forall x(x \in \{x \mid R\} \Leftrightarrow R).$$

Il suffit d'appliquer **E141** en prenant pour T le terme $\{x \mid R\}$ lui-même, dans lequel x ne figure pas librement:

$$\begin{aligned} \forall x(x \in \{x \mid R\} \Leftrightarrow R) &\Leftrightarrow (\{x \mid R\} \text{ existe et } \{x \mid R\} = \{x \mid R\}) && \text{E141} \\ &\Leftrightarrow \{x \mid R\} \text{ existe} && \text{S82.} \end{aligned}$$

E143.

$$\frac{\Gamma \vdash \{x \mid R\} \text{ existe et } T = \{x \mid R\}}{\Gamma \vdash U \in T \Leftrightarrow (U \mid x)R}$$

Si x ne figure pas librement dans T et si
 U est librement substituable à x dans R .

De la prémisse de cette déduction on déduit en effet suivant E141 le jugement $\Gamma \vdash \forall x(x \in T \Leftrightarrow R)$, dont on déduit $\Gamma \vdash U \in T \Leftrightarrow (U \mid x)R$ par particularisation.

Substitutivité de l'équivalence

E144.

$$\frac{\Gamma \vdash R \Leftrightarrow R'}{\Gamma \vdash \{x \mid R\} \text{ existe} \Leftrightarrow \{x \mid R'\} \text{ existe}} \quad \frac{\Gamma \vdash R \Leftrightarrow R'}{\Gamma \vdash \{x \mid R\} = \{x \mid R'\}}$$

Si x ne figure pas librement dans Γ .

En effet, si x ne figure pas librement dans Γ , alors en prenant une variable y distincte de x et ne figurant librement ni dans R , ni dans Γ , on peut écrire la démonstration suivante de ces deux déductions. On peut d'ailleurs déduire les lignes 3^o et 4^o directement de la ligne 1^o en invoquant la règle générale de substitutivité de l'équivalence de la logique des prédicats (§10.1) puisque x et y ne figurent pas librement dans Γ .

Γ	
1 ^o	$R \Leftrightarrow R'$ <i>base</i>
2 ^o	$(x \in y \Leftrightarrow R) \Leftrightarrow (x \in y \Leftrightarrow R')$ 1^o, S50.9
3 ^o	$\forall x(x \in y \Leftrightarrow R) \Leftrightarrow \forall x(x \in y \Leftrightarrow R')$ 2^o, S70
4 ^o	$\exists y \forall x(x \in y \Leftrightarrow R) \Leftrightarrow \exists y \forall x(x \in y \Leftrightarrow R')$ 3^o, S70
5 ^o	$\{x \mid R\} \text{ existe} \Leftrightarrow \{x \mid R'\} \text{ existe}$ 4^o, E139
6 ^o	$Sy \forall x(x \in y \Leftrightarrow R) = Sy \forall x(x \in y \Leftrightarrow R')$ 3^o, SEL2
7 ^o	$\{x \mid R\} = \{x \mid R'\}$ 7^o, E138.

Changements de variable liée

E145.

$$\frac{}{\Gamma \vdash \{x \mid R\} \text{ existe} \Leftrightarrow \{z \mid (z \mid x)R\} \text{ existe}} \quad \frac{}{\Gamma \vdash \{x \mid R\} = \{z \mid (z \mid x)R\}}$$

Si z est réversiblement substituable à x dans R .

En effet, si z est réversiblement substituable à x dans R , alors en choisissant une variable y distincte de x et de z et ne figurant librement ni dans R , ni dans Γ , on peut écrire la démonstration suivante. La première ligne se justifie par le fait que z est réversiblement substituable à x dans la proposition $x \in y \Leftrightarrow R$.

Γ	
1 ^o	$\forall x(x \in y \Leftrightarrow R) \Leftrightarrow \forall z((z \mid x)(x \in y \Leftrightarrow R))$ S76
2 ^o	$\forall x(x \in y \Leftrightarrow R) \Leftrightarrow \forall z(z \in y \Leftrightarrow (z \mid x)R)$ identique à 1^o
3 ^o	$\exists y \forall x(x \in y \Leftrightarrow R) \Leftrightarrow \exists y \forall z(z \in y \Leftrightarrow (z \mid x)R)$ 2^o, S70

4°	$\{x \mid R\} \text{ existe} \Leftrightarrow \{z \mid (z \mid x)R\} \text{ existe}$	3°, E139
5°	$Sy \forall x (x \in y \Leftrightarrow R) = Sy \forall z (z \in y \Leftrightarrow R)$	2°, SEL2
6°	$\{x \mid R\} = \{z \mid (z \mid x)R\}$	5°, E138.

Le schéma de séparation de $\mathcal{T}_{\text{ENS}}$. Le schéma de déduction suivant est le *deuxième schéma de déduction élémentaire spécifique* non définitionnel de la théorie $\mathcal{T}_{\text{ENS}}$. Nous rappelons que le premier est le schéma de réunion (E93). Tous les autres schémas de déduction de cette théorie sont des schémas dérivés ou des définitions (schémas définitionnels).

E146.

$$\Gamma \vdash \{x \mid x \in E \text{ et } R\} \text{ existe}$$

Si R est une proposition, E est un terme et x est une variable qui ne figure pas librement dans E .

Si R , E et x satisfont à ces conditions, on peut affirmer encore que l'ensemble $\{x \mid x \in E \text{ et } R\}$ est un sous-ensemble de E , à savoir qu'on a le théorème

$$\vdash \{x \mid x \in E \text{ et } R\} \subset E.$$

Démonstration:

	nil	
1°	$x \in \{x \mid x \in E \text{ et } R\} \Leftrightarrow (x \in E \text{ et } R)$	E146, E142
2°	$x \in \{x \mid x \in E \text{ et } R\} \Rightarrow x \in E$	1°, S28, S49
3°	$\{x \mid x \in E \text{ et } R\} \subset E$	2°, E5.

La déduction de la ligne 3° (E5) est possible parce que x ne figure pas librement dans le terme $\{x \mid x \in E \text{ et } R\}$, ni dans le terme E .

Le nom de *schéma de séparation* donné à E146 correspond à l'interprétation suivante de ce schéma. Si x ne figure pas librement dans le terme E et si R est une proposition, l'ensemble $\{x \mid x \in E \text{ et } R\}$ est l'ensemble des éléments de E qui satisfont à la condition exprimée par R pour x . Ce schéma permet donc de construire un ensemble en « séparant » dans un ensemble donné les éléments qui satisfont à une condition donnée, des autres éléments de l'ensemble, et en ne gardant que les premiers. On pourrait parler aussi d'un procédé de « filtrage » d'un ensemble suivant une condition. Ce schéma a été proposé en 1908 par Zermelo sous le nom allemand de *Aussonderungsaxiom*, qui traduit aussi l'idée d'un procédé de séparation ou de filtrage.

La proposition de Zermelo marque la fin de la période de gestation de la théorie des ensembles, au cours de laquelle de nombreux « paradoxes » furent découverts. On appelait ainsi des contradictions dont on ne saisissait pas bien l'origine. Elles résultaient de ce que l'on admettait plus ou moins consciemment le schéma trop général: $\vdash \{x \mid R\} \text{ existe}$. Le logicien Frege, principal inventeur de la logique des prédicats, eut le malheur d'être celui qui publia une version formelle de la théorie des ensembles (1884) admettant ce schéma

général et de recevoir en 1902 une lettre de Russel lui signalant le théorème : $\vdash \{x \mid x \notin x\}$ *n'existe pas*, c'est-à-dire $\vdash \text{non} \exists y \forall x (x \in y \Leftrightarrow x \notin x)$ (§10.3).

Variante du schéma de séparation. Lorsqu'on veut démontrer un théorème de la forme « $\Gamma \vdash \{x \mid R\}$ existe », on se sert souvent du schéma suivant qui dérive du schéma de séparation (E146), mais que l'on pourrait prendre comme schéma élémentaire à la place de E146, car chacun de ces deux schémas dérive de l'autre.

E147.

$$\frac{\Gamma \vdash R \Rightarrow x \in E}{\Gamma \vdash \{x \mid R\} \text{ existe}} \quad \begin{array}{l} \text{Si } x \text{ ne figure librement} \\ \text{ni dans } \Gamma, \text{ ni dans } E. \end{array}$$

$$\begin{array}{lcl} \Gamma & & \\ 1^\circ \mid R \Rightarrow x \in E & & \text{base} \\ 2^\circ \mid R \Leftrightarrow (x \in E \text{ et } R) & & 1^\circ, \text{S60.3} \\ 3^\circ \mid \{x \mid R\} \text{ existe} \Leftrightarrow \{x \mid x \in E \text{ et } R\} \text{ existe} & & 2^\circ, \text{E144} \\ 4^\circ \mid \{x \mid R\} \text{ existe} & & 3^\circ, \text{E146.} \end{array}$$

Pour démontrer l'existence d'une collection $\{x \mid R\}$, dans un contexte où l'on n'a fait aucune hypothèse sur x , il suffit de trouver un ensemble ne dépendant pas de x , donc représenté par un terme E dans lequel x ne figure pas librement, tel que la proposition $R \Rightarrow x \in E$ soit vraie dans ce contexte. Un premier exemple se présente ci-dessous avec la collection figurant dans la définition E148.

Transformé d'un ensemble par un graphe. Étant donné un graphe G et un ensemble X , on appelle *transformé de X par G* et l'on note $G\langle X \rangle$ l'ensemble des objets qui correspondent à un élément de l'ensemble X suivant le graphe G . Lorsqu'on dit qu'un objet y *correspond* à un objet x suivant un graphe G , on veut dire simplement que le couple (x, y) appartient à G . Un objet y appartient à l'ensemble $G\langle X \rangle$ si et seulement s'il *existe* un objet x tel que $x \in X$ et tel que y corresponde à x suivant G . Formellement, on pose la définition suivante:

E148. *Définition*

$$\vdash \text{Si } G \text{ est un graphe, alors } G\langle X \rangle = \{y \mid \exists x (x \in X \text{ et } (x, y) \in G)\}.$$

Cette notion est illustrée par la figure 14.6 dans le cas d'un graphe fini représenté par un diagramme de points et de flèches et dans le cas d'un graphe G qui est un ensemble de points du plan en tant que couples de coordonnées. Dans le premier cas les éléments de l'ensemble X sont entourés par un rectangle arrondi, de même que ceux de l'ensemble $G\langle X \rangle$. Dans le deuxième, X est un intervalle de la droite réelle et l'on a marqué un élément y de l'ensemble $G\langle X \rangle$ ainsi qu'un élément de X auquel y correspond suivant G .

Syntaxiquement, le terme $G\langle X \rangle$ se compose d'un symbole fonctionnel *binai-* *re* constitué par les deux parenthèses spéciales $\langle \rangle$. L'arbre de ce terme est



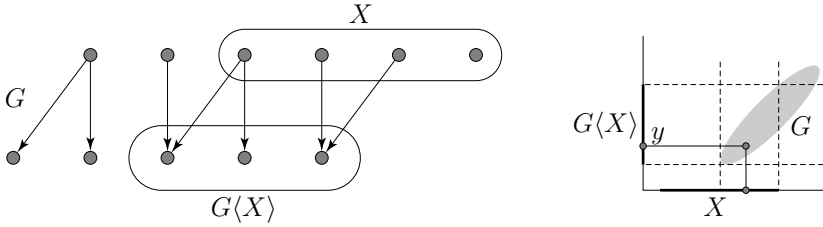


Figure 14.6 Transformé $G\langle X \rangle$ de l'ensemble X par le graphe G .

Lorsqu'on parle du « transformé de X par G », ce sont les mots « le transformé de – par – » qui constituent un symbole fonctionnel binaire. Mais avec celui-ci, les opérands sont notées dans l'ordre X, G .

La définition E148 est de la forme

$$\vdash C \Rightarrow T = \{x \mid R\}, \quad (3)$$

où T est un terme construit avec un nouveau symbole fonctionnel et la variable x ne figure librement ni dans C , ni dans T . Normalement, on ne pose de définition de cette forme que lorsque

$$C \vdash (\{x \mid R\} \text{ existe}) \quad (4)$$

est un théorème. Alors d'après E141 on a le théorème

$$C \vdash x \in T \Leftrightarrow R, \quad (5)$$

qui est le théorème principal sur le nouveau symbole introduit dans T . Dans le cas de la définition E148, ce théorème principal est

$$G \text{ est un graphe } \vdash y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G).$$

C'est celui qui est utilisé dans la plupart des démonstrations où intervient le transformé d'un ensemble par un graphe.

Une définition de la forme (3) est donc toujours accompagnée d'une démonstration du théorème d'existence (4) correspondant, dont découle « automatiquement » (E141) le théorème principal (5). En général, la démonstration de (4) est une application assez évidente du schéma E147, à savoir qu'on trouve assez facilement un terme E dans lequel x ne figure pas librement et tel que l'on ait $C \vdash (R \Rightarrow x \in E)$.

Par ailleurs, dans les ouvrages de mathématiques qui utilisent la théorie des ensembles mais dont le sujet n'est pas cette théorie elle-même, on omet presque toujours la démonstration de (4) et même la mention de ce théorème d'existence — existence considérée comme allant de soi. Le théorème principal (5) va donc aussi de soi et n'est pas mentionné. Il est considéré comme étant compris dans la définition (3).

Dans le cas de la définition E148, les théorèmes (4) et (5) sont les théorèmes E149 et E150, démontrés ci-dessous. Les théorèmes suivants, E151–E163, sont démontrés dans l'annexe (Chap 16).

E149. G est un graphe $\vdash \{y \mid \exists x(x \in X \text{ et } (x, y) \in G)\}$ existe.

E150. *Théorème principal*

G est un graphe $\vdash y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G)$.

Les deux théorèmes se démontrent ensemble comme suit. La ligne essentielle de la démonstration est la ligne 7°, qui permet d'appliquer E147. Le terme $\mathbf{Pr}_2G$ est le terme E qu'il fallait trouver, et qui saute aux yeux lorsqu'on se pose la question de l'existence de $\{y \mid \exists x(x \in X \text{ et } (x, y) \in G)\}$.

nil		
1°	G est un graphe	hyp
2°	$\exists x(x \in X \text{ et } (x, y) \in G)$	hyp
3°	$x \in X \text{ et } (x, y) \in G$	hyp
4°	$(x, y) \in G$	3°
5°	$y \in \mathbf{Pr}_2G$	4°, E118
6°	$y \in \mathbf{Pr}_2G$	2°, 5°, S64
7°	$\exists x(x \in X \text{ et } (x, y) \in G) \Rightarrow y \in \mathbf{Pr}_2G$	6°, $\Rightarrow$ -intro
8°	$\{y \mid \exists x(x \in X \text{ et } (x, y) \in G)\}$ existe	7°, E147
9°	$G\langle X \rangle = \{y \mid \exists x(x \in X \text{ et } (x, y) \in G)\}$	1°, E148
10°	$y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G)$	8°, 9°, E141.

E151. G est un graphe $\vdash G\langle X \rangle \subset \mathbf{Pr}_2G$.

E152. G est un graphe $\vdash G\langle \mathbf{Pr}_1G \rangle = \mathbf{Pr}_2G$.

E153. G est un graphe $\vdash G\langle X \rangle = G\langle X \cap \mathbf{Pr}_1G \rangle$.

E154. G est un graphe $\vdash (x, y) \in G \Leftrightarrow y \in G\langle \{x\} \rangle$.

E155. G est un graphe $\vdash G\langle X \rangle = \bigcup_{x \in X} G\langle \{x\} \rangle$.

E156. G est un graphe $\vdash X \subset Y \Rightarrow G\langle X \rangle \subset G\langle Y \rangle$.

E157. G est un graphe $\vdash G\langle X \cup Y \rangle = G\langle X \rangle \cup G\langle Y \rangle$.

E158. G est un graphe $\vdash G\langle X \cap Y \rangle \subset G\langle X \rangle \cap G\langle Y \rangle$.

E159. G est un graphe $\vdash G\langle X \rangle = \emptyset \Leftrightarrow X \cap \mathbf{Pr}_1G = \emptyset$.

E160. G est un graphe $\vdash x \in \mathbf{Pr}_1G \Leftrightarrow G\langle \{x\} \rangle \neq \emptyset$.

E161. G est un graphe $\vdash G\langle \emptyset \rangle = \emptyset$.

E162. $\vdash \text{Id}_Y\langle X \rangle = X \cap Y$.

E163. $\vdash X \subset Y \Leftrightarrow \text{Id}_Y\langle X \rangle = X$.

Collections vides. Réunion, intersection, inclusion de collections

E164.

$$\Gamma \vdash \text{non} \exists xR \Leftrightarrow (\{x \mid R\} \text{ existe et } \{x \mid R\} = \emptyset).$$

Sous une liste d'hypothèses Γ vide, on peut écrire en effet la démonstration suivante:

$$\begin{aligned} (\{x \mid R\} \text{ existe et } \{x \mid R\} = \emptyset) &\Leftrightarrow \forall x(x \in \emptyset \Leftrightarrow R) && \text{E141} \\ &\Leftrightarrow \forall x(\text{non}R) && \text{E2, S60.8} \\ &\Leftrightarrow \text{non} \exists xR && \text{S71.} \end{aligned}$$

E165.

$$\frac{\Gamma \vdash \{x \mid R\} \text{ existe} \quad \Gamma \vdash \{x \mid S\} \text{ existe}}{\Gamma \vdash \{x \mid R \text{ ou } S\} \text{ existe et } \{x \mid R \text{ ou } S\} = \{x \mid R\} \cup \{x \mid S\}.}$$

$$\frac{\Gamma \vdash \{x \mid R\} \text{ existe} \quad \Gamma \vdash \{x \mid S\} \text{ existe}}{\Gamma \vdash \{x \mid R \text{ et } S\} \text{ existe et } \{x \mid R \text{ et } S\} = \{x \mid R\} \cap \{x \mid S\}.}$$

Comme la variable x ne figure pas dans les assertions des prémisses de ces déductions, il suffit de démontrer celles-ci dans le cas particulier où x ne figure pas librement dans Γ (§9.5). C'est ce que nous supposons dans la démonstration suivante (ligne 6°). La démonstration de la deuxième déduction de E165 est semblable à celle-ci.

Γ		
1°	$\{x \mid R\} \text{ existe}$	<i>base</i>
2°	$\{x \mid S\} \text{ existe}$	<i>base</i>
3°	$x \in \{x \mid R\} \cup \{x \mid S\} \Leftrightarrow (x \in \{x \mid R\} \text{ ou } x \in \{x \mid S\})$	E18
4°	$\Leftrightarrow (R \text{ ou } S)$	1°, 2°, E142
5°	$x \in \{x \mid R\} \cup \{x \mid S\} \Leftrightarrow (R \text{ ou } S)$	3°, 4°
6°	$\forall x(x \in \{x \mid R\} \cup \{x \mid S\} \Leftrightarrow (R \text{ ou } S))$	5°, S62
7°	$\{x \mid R \text{ ou } S\} \text{ existe et } \{x \mid R \text{ ou } S\} = \{x \mid R\} \cup \{x \mid S\}$	6°, E141.

Dans les déductions suivantes (E166) l'assertion « $\{x \mid R\}$ et $\{x \mid S\}$ existent » est évidemment une abréviation de « $\{x \mid R\}$ existe et $\{x \mid S\}$ existe ». On pourrait d'ailleurs remplacer la prémisse de ces déductions par les deux prémisses des déductions E165.

E166.

$$\frac{\Gamma \vdash \{x \mid R\} \text{ et } \{x \mid S\} \text{ existent}}{\Gamma \vdash \forall x(R \Rightarrow S) \Leftrightarrow \{x \mid R\} \subset \{x \mid S\}} \quad \frac{\Gamma \vdash \{x \mid R\} \text{ et } \{x \mid S\} \text{ existent}}{\Gamma \vdash \forall x(R \Leftrightarrow S) \Leftrightarrow \{x \mid R\} = \{x \mid S\}.}$$

Comme x ne figure pas librement dans les assertions de la prémisse de ces déductions, il suffit de les démontrer dans le cas où x ne figure pas librement dans Γ . C'est ce qui est supposé à la ligne 5° de la démonstration suivante de la première déduction.

Γ		
1°	$\{x \mid R\} \text{ et } \{x \mid S\} \text{ existent}$	<i>base</i>
2°	$R \Leftrightarrow x \in \{x \mid R\}$	1°, E142
3°	$S \Leftrightarrow x \in \{x \mid S\}$	1°, E142
4°	$(R \Rightarrow S) \Leftrightarrow (x \in \{x \mid R\} \Rightarrow x \in \{x \mid S\})$	2°, 3°, S50
5°	$\forall x(R \Rightarrow S) \Leftrightarrow \forall x(x \in \{x \mid R\} \Rightarrow x \in \{x \mid S\})$	4°, S70
6°	$\Leftrightarrow \{x \mid R\} \subset \{x \mid S\}$	5°, E3.

La deuxième déduction de E166 peut se démontrer de manière semblable ou au moyen de la première et du théorème

$$\Gamma \vdash \forall x(R \Leftrightarrow S) \Leftrightarrow (\forall x(R \Rightarrow S) \text{ et } \forall x(S \Rightarrow R)).$$

Collections non existantes. Nous traitons finalement de quelques formes de collections qui n'existent pas, au sens de E139. La première est donnée par le théorème de Russel (§10.3, Théorème du barbier)

$$\vdash \text{non} \exists y \forall x (x \in y \Leftrightarrow x \notin x),$$

dont on tire suivant E139:

E167. $\vdash \{x \mid x \notin x\}$ n'existe pas.

Par changement de variables liées dans le théorème de Russel, on peut obtenir le théorème $\vdash \{x \mid x \notin x\}$ pour n'importe quelle variable x . Cela permet de démontrer le schéma de déduction suivant, qui est celui que l'on utilise le plus souvent dans les démonstrations de non-existence de collections.

E168.

$\Gamma \vdash \text{non} \forall x (x \in E)$		Si x ne figure pas librement dans E .
nil		
1°	$\forall x (x \in E)$	hyp
2°	$x \in E$	1°, S65
3°	$x \notin x \Rightarrow x \in E$	2°, S21
4°	$\{x \mid x \notin x\}$ existe	3°, E147
5°	$\perp$	4°, E167
6°	$\text{non} \forall x (x \in E)$	5°, red abs J.

Le schéma E168 peut être utilisé pour démontrer le théorème suivant, qu'on exprime parfois en disant qu'il n'existe pas d'ensemble « universel » dans le sens d'un ensemble dont tous les objets soient éléments:

$$\vdash \text{non} \exists u \forall x (x \in u).$$

La démonstration par réduction à l'absurde (élimination de $\exists u$ et E168) est évidente. Un théorème équivalent, suivant S71, est $\vdash \forall u \exists x (x \notin u)$. Pour tout ensemble u , il existe un objet qui n'est pas élément de u .

Les deux schémas suivants décrivent deux cas simples de collections qui n'existent pas. Comme la variable x ne figure pas librement dans les assertions des prémisses de ces déductions, il suffit de les démontrer dans le cas où x ne figure pas librement dans Γ . C'est ce qui est supposé dans les démonstrations ci-dessous (ligne 4° de la première et ligne 5° de la deuxième).

E169.

$\Gamma \vdash \forall x R$		
$\Gamma \vdash \{x \mid R\}$ n'existe pas.		
Γ		
1°	$\forall x R$	base
2°	$\{x \mid R\}$ existe	hyp
3°	$R \Leftrightarrow x \in \{x \mid R\}$	2°, E142

4°	$\forall xR \Leftrightarrow \forall x(x \in \{x \mid R\})$	3°, S70
5°	$\forall x(x \in \{x \mid R\})$	1°, 4°, S39
6°	$\perp$	5°, E168
7°	$\{x \mid R\}$ n'existe pas	6°, red abs J.

E170.

$$\frac{\Gamma \vdash \{x \mid R\} \text{ existe}}{\Gamma \vdash \{x \mid \text{non}R\} \text{ n'existe pas.}}$$

Il suffit de démontrer cette déduction dans le cas où x ne figure pas librement dans Γ (ligne 5°):

Γ		
1°	$\{x \mid R\}$ existe	<i>base</i>
2°	$\{x \mid \text{non}R\}$ existe	<i>hyp</i>
3°	$\{x \mid R \text{ ou } \text{non}R\}$ existe	1°, 2°, E165
4°	$R \text{ ou } \text{non}R$	S34
5°	$\forall x(R \text{ ou } \text{non}R)$	4°, S62
6°	$\{x \mid R \text{ ou } \text{non}R\}$ n'existe pas	5°, E169
7°	$\perp$	3°, 6°
8°	$\{x \mid \text{non}R\}$ n'existe pas	7°, red abs J.

Le dernier schéma de ce paragraphe fournit une méthode générale pour démontrer qu'une collection $\{x \mid R\}$ n'existe pas et nous en donnerons plus bas quelques exemples.

E171.

$$\frac{\Gamma \vdash \forall y \exists x(R \text{ et } y \in T(x))}{\Gamma \vdash \{x \mid R\} \text{ n'existe pas}}$$

Si y est distincte de x et ne figure librement ni dans R , ni dans T .

Il suffit de démontrer cette déduction dans le cas où x et y ne figurent pas librement dans Γ . C'est ce qu'on suppose dans la démonstration ci-dessous lorsqu'on applique la règle générale de substitutivité de l'équivalence (§10.1) dans les déductions des lignes 4° et 5°. On peut remplacer dans une proposition une sous-proposition par une proposition équivalente, à condition que cette sous-proposition ne soit pas dominée par des quantificateurs dont la variable figure librement dans les hypothèses du contexte. Aux lignes 4° et 5° ces hypothèses sont les propositions de la liste Γ et la proposition de la ligne 2°.

Γ		
1°	$\forall y \exists x(R \text{ et } y \in T(x))$	<i>base</i>
2°	$\{x \mid R\}$ existe	<i>hyp</i>
3°	$R \Leftrightarrow x \in \{x \mid R\}$	2°, E142
4°	$\forall y \exists x(x \in \{x \mid R\} \text{ et } y \in T(x))$	1°, 3°

5°	$\left \begin{array}{c} \forall y(y \in \bigcup_{x \in \{x R\}} T(x)) \\ \perp \end{array} \right.$	4°, E95
6°		5°, E168
10°	$\{x R\}$ n'existe pas	6°, red abs J.

Exemples. Les quatre déductions suivantes sont de la forme E171:

$$\frac{\vdash \forall y \exists G(G \text{ est un graphe et } y \in \mathbf{Pr}_1 G)}{\vdash \{G | G \text{ est un graphe}\} \text{ n'existe pas.}}$$

$$\frac{\vdash \forall y \exists z(z \text{ est un couple et } y \in \{\mathbf{pr}_1 z\})}{\vdash \{z | z \text{ est un couple}\} \text{ n'existe pas.}}$$

$$\frac{\vdash \forall y \exists x(a \subset x \text{ et } y \in \mathcal{P}(x))}{\vdash \{x | a \subset x\} \text{ n'existe pas.}} \quad \frac{\vdash \forall y \exists x(a \cap x = \emptyset \text{ et } y \in \mathcal{P}(a \cup x))}{\vdash \{x | a \cap x = \emptyset\} \text{ n'existe pas.}}$$

Or les prémisses de ces déductions sont des théorèmes faciles à démontrer. Ils se déduisent des théorèmes suivants par introduction de $\exists$ (S63) et généralisation sur y :

$$\begin{aligned} &\vdash \{(y, y)\} \text{ est un graphe et } y \in \mathbf{Pr}_1 \{(y, y)\}. \\ &\vdash (y, y) \text{ est un couple et } y \in \{\mathbf{pr}_1(y, y)\}. \\ &\vdash a \subset a \cup y \text{ et } y \in \mathcal{P}(a \cup y). \\ &\vdash a \cap \mathbb{C}_y(a) = \emptyset \text{ et } y \in \mathcal{P}(a \cup \mathbb{C}_y(a)). \end{aligned}$$

14.3 Opérateurs de collection généraux

Exemple 1. La fameuse conjecture de Goldbach, dont on cherche toujours la démonstration, est que tout entier naturel pair supérieur à 3 est la somme de deux nombres premiers. Nous rappelons que le plus petit nombre premier est 2, et que les suivants sont 3, 5, 7, 11, 13, etc. Personne n'a jamais trouvé de nombre pair supérieur à 3 qui ne soit pas la somme de deux nombres premiers, mais il n'est pas encore démontré qu'il n'en existe pas. La conjecture de Goldbach peut s'exprimer en disant que l'ensemble des entiers naturels pairs plus grands que 2 est contenu dans

$$\begin{aligned} &\text{l'ensemble des entiers naturels de la forme } x + y \\ &\text{tels que } x \text{ et } y \text{ sont des nombres premiers.} \end{aligned} \tag{1}$$

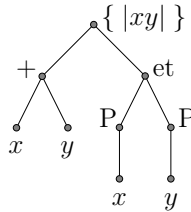
Ce n'est pas la conjecture elle-même qui nous intéresse ici, mais seulement la forme de l'expression (1) qui est utilisée pour décrire un ensemble.

Observons d'abord que les variables x et y sont *locales* dans cette expression. Si l'on remplace ces variables par d'autres, a et b par exemple, on obtient une expression qui désigne le *même* ensemble: l'ensemble des entiers naturels de la forme $a + b$ tels que a et b sont des nombres premiers.

L'expression (1) se compose de trois parties. Elle contient premièrement un opérateur de construction d'ensemble formé par les mots « l'ensemble des – de la forme – tels que »; deuxièmement le terme $x + y$; troisièmement la proposition « x et y sont des nombres premiers ». Les mots utilisés dans l'opérateur n'indiquent pas explicitement que les variables x et y sont locales. Cela est ajouté dans la version formelle de cette expression, qui est la suivante:

$$\{x + y \mid xy \mid \text{Premier}(x) \text{ et Premier}(y)\}. \quad (2)$$

L'opérateur de construction d'ensemble de cette expression est l'assemblage de signes $\{ \mid xy \mid \}$ et l'arbre de l'expression est le suivant, où le prédicat « Premier » est abrégé par P:



Forme générale des opérateurs de collection. L'opérateur $\{ \mid xy \mid \}$ appartient à la classe des *opérateurs de collection généraux*. Ces opérateurs sont différents des opérateurs de collection simples étudiés précédemment (§14.2). Les opérateurs de collection généraux sont de la forme $\{ \mid x_1 \cdots x_n \mid \}$, où $x_1, \dots, x_n$ est une liste de variables *distinctes*. Ils se combinent toujours avec un terme E et une proposition R et l'expression résultante $\{E \mid x_1 \cdots x_n \mid R\}$ est un terme. Le tableau suivant compare ces opérateurs avec les opérateurs de collection simples $\{ \mid x \mid \}$.

Opérateur:	$\{ \mid x \mid \}$	$\{ \mid x_1 \cdots x_n \mid \}$
Arité:	1	2
Type:	$T \leftarrow P$	$T \leftarrow (T, P)$
Expression typique:	$\{ \mid x \mid R \}$	$\{E \mid x_1 \cdots x_n \mid R\}$
Variables liées:	(x)	$(x_1 \cdots x_n, x_1 \cdots x_n)$

Notre terminologie — opérateurs de collection simples, opérateurs de collection généraux — suggère que les premiers sont un cas particulier des seconds. Le tableau ci-dessus montre qu'il n'en est rien, puisque ces opérateurs ne sont pas de même arité. Mais une *expression* $\{ \mid x \mid R \}$ construite avec un opérateur de collection simple peut être assimilée à un cas particulier d'expression $\{E \mid x_1 \cdots x_n \mid R\}$, à savoir à l'expression $\{ \mid x \mid x \mid R \}$. Cette assimilation se fonde sur le théorème $\vdash \{ \mid x \mid R \} = \{ \mid x \mid x \mid R \}$ que nous démontrerons plus loin.

Définition d'une collection $\{E \mid x_1 \cdots x_n \mid R\}$. Considérons encore l'expression (1) ou l'expression formelle correspondante (2). Dire qu'un nombre z appartient à cet ensemble, ou dire simplement que z est la somme de deux nombres premiers, signifie qu'il *existe* deux nombres premiers x et y tels que

$z = x + y$. C'est cette *quantification existentielle* que l'on exprime par les mots « de la forme ». Ce sont aussi ces quantificateurs $\exists x, \exists y$, sous-entendus dans l'expression (1), qui font que les variables x et y de celles-ci sont liées. Nous voyons que l'interprétation naturelle de l'expression (1), ou de l'expression formelle (2) qui lui correspond, est que cette expression désigne l'ensemble

$$\{z \mid \exists x \exists y (z = x + y \text{ et Premier}(x) \text{ et Premier}(y))\}. \quad (3)$$

L'égalité de (1) et (3) est un cas particulier du schéma général suivant, qui constitue la définition de la notation $\{E \mid x_1 \cdots x_n \mid R\}$.

E172. Définition

$$\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} = \{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\}$$

Si E est un terme et R est une proposition et si $x_1, \dots, x_n, z$ sont des variables distinctes et z ne figure librement ni dans E ni dans R .

L'ensemble $\{E \mid x_1 \cdots x_n \mid R\}$ est donc défini au moyen d'un opérateur de collection simple $\{z \mid \}$. Le sens de la proposition

$$\{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} \text{ existe}$$

est défini par **E139**. En accord avec **E172**, il est normal de définir le sens de « $\{E \mid x_1 \cdots x_n \mid R\}$ existe » comme suit :

E173. Définition

$$\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe} \Leftrightarrow \{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} \text{ existe}$$

Si $E, R, x_1, \dots, x_n, z$ satisfont aux conditions de **E172**.

Exemple 2. Soient a et b des entiers naturels, $a \in \mathbb{N}$ et $b \in \mathbb{N}$. L'ensemble

$$\{ax + 1 \mid x \in \mathbb{N} \text{ et } x \leq b\} \quad (4)$$

peut être décrit comme

$$\begin{aligned} &\text{l'ensemble des nombres de la forme } ax + 1, \\ &\text{où } x \text{ est un entier naturel inférieur à } b. \end{aligned} \quad (5)$$

Cependant, à défaut de conventions générales sur ce genre de description, celle-ci est ambiguë — contrairement à l'expression formelle (4). Dans cette dernière, on voit que les variables a et b sont libres, alors que la variable x est liée par l'opérateur $\{ \mid x \mid \}$. Rien n'indique cette différence de statut entre la variable x et les variables a et b dans l'expression (5). Cette différence est certes perçue d'après le contexte. Ayant fait les hypothèses $a \in \mathbb{N}, b \in \mathbb{N}$, on « sent » bien que dans l'expression (5) les variables a et b représentent deux nombres « fixes ». Ce sont des variables globales de l'expression, alors que x est une variable locale.

D'après **E172**, l'ensemble (4) est égal à

$$\{z \mid \exists x (z = ax + 1 \text{ et } x \in \mathbb{N} \text{ et } x \leq b)\}. \quad (6)$$

Convention de notation. La notation formelle $\{E \mid x_1 \cdots x_n \mid R\}$ n'est guère utilisée dans les mathématiques courantes. La liste de variables $x_1 \cdots x_n$ est simplement omise. On écrit $\{E \mid R\}$. Ainsi les termes (2) et (4) sont abrégés

$$\{x + y \mid \text{Premier}(x) \text{ et } \text{Premier}(y)\} \quad (7)$$

$$\{ax + 1 \mid x \in \mathbb{N} \text{ et } x \leq b\}. \quad (8)$$

La liste de variables liées $x_1 \cdots x_n$ omise dans cette notation abrégée $\{E \mid R\}$ est fixée par la *convention* suivante: cette liste se compose de toutes les variables qui figurent librement à la fois dans le terme E et dans la proposition R , autrement dit toutes les variables libres *communes* à E et R .

Cette convention permet de reconstruire exactement les termes (2) et (4) abrégés par (7) et (8).

Nous utiliserons cette convention dans les applications, car elle est commode et naturelle. Mais dans les schémas de déduction que nous allons donner sur les opérateurs de collection généraux, nous utilisons la notation explicite $\{E \mid x_1 \cdots x_n \mid R\}$.

D'autre part il faut savoir que cette convention n'est pas toujours respectée dans les mathématiques usuelles. On rencontre des expressions de la forme $\{E \mid R\}$ dans lesquelles seul le contexte et la compréhension du sujet permettent de savoir quelles sont les variables $x_1 \cdots x_n$ omises dans l'expression. En général, cela ne présente pas de difficulté.

Par contre, cette liberté de notation n'est pas possible dans un langage de spécification formelle de logiciel tel que le langage Z^1 . Celui-ci utilise une notation analogue à $\{E \mid x_1 \cdots x_n \mid R\}$.

Utilisation principale des collections. Tout comme les opérateurs de collection simples (§14.2, Utilisation principale), les opérateurs de collection généraux sont utilisés principalement dans les définitions de la forme $\vdash T = \{E \mid x_1 \cdots x_n \mid R\}$. Le schéma de déduction principal qui s'applique alors est le suivant. Il correspond au schéma E141 sur les collections simples.

E174.

$$\frac{\Gamma \vdash \forall z(z \in T \Leftrightarrow \exists x_1 \cdots \exists x_n(z = E \text{ et } R))}{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe et } T = \{E \mid x_1 \cdots x_n \mid R\}}$$

$$\frac{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe et } T = \{E \mid x_1 \cdots x_n \mid R\}}{\Gamma \vdash \forall z(z \in T \Leftrightarrow \exists x_1 \cdots \exists x_n(z = E \text{ et } R))}$$

Si $x_1, \dots, x_n, z$ sont des variables distinctes et si z ne figure librement ni dans E , ni dans T , ni dans R .

La démonstration de ces deux déductions est une simple application du schéma E141, suivant lequel la proposition

$$\forall z(z \in T \Leftrightarrow \exists x_1 \cdots \exists x_n(z = E \text{ et } R))$$

¹ Voir Bibliographie.

est *équivalente* à

$$\{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} \text{ existe et } T = \{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\}.$$

On conclut suivant E172 et E173.

Substitutivité de l'égalité et de l'équivalence

E175.

$$\frac{\Gamma \vdash E = E' \quad \Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} = \{E' \mid x_1 \cdots x_n \mid R\}}{\Gamma \vdash R \Leftrightarrow R'} \quad \Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} = \{E \mid x_1 \cdots x_n \mid R'\}$$

Si $x_1, \dots, x_n$ sont des variables distinctes qui ne figurent pas librement dans Γ (condition pour les deux schémas).

Démonstration. Si cette condition est satisfaite et si z est une variable distincte de $x_1, \dots, x_n$ et ne figurant librement ni dans E , ni dans R , ni dans Γ , on peut écrire :

Γ		
1°	$E = E'$	<i>base</i>
2°	$(z = E \text{ et } R) \Leftrightarrow (z = E' \text{ et } R)$	1°, S86
3°	$\exists x_1 \cdots \exists x_n (z = E \text{ et } R) \Leftrightarrow \exists x_1 \cdots \exists x_n (z = E' \text{ et } R)$	2°, S70
4°	$\{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} = \{z \mid \exists x_1 \cdots \exists x_n (z = E' \text{ et } R)\}$	3°, E144
5°	$\{E \mid x_1 \cdots x_n \mid R\} = \{E' \mid x_1 \cdots x_n \mid R\}$	4°, E172.

Le deuxième schéma se démontre de manière semblable, en utilisant S50 au lieu de S86.

Changement de variables liées. Considérons l'exemple suivant. Si P est un symbole relationnel unaire, disons le prédicat « est un nombre premier », on a le théorème de changement de variables liées :

$$\vdash \{x + y \mid xy \mid P(x) \text{ et } P(y)\} = \{a + b \mid ab \mid P(a) \text{ et } P(b)\}$$

Cette classe de théorèmes est décrite par le schéma de déduction suivant.

E176.

$$\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} = \{E' \mid y_1 \cdots y_n \mid R'\}$$

Si $x_1, \dots, x_n$ sont des variables distinctes; $y_1, \dots, y_n$ sont des variables distinctes qui ne figurent pas librement dans $\{E \mid x_1 \cdots x_n \mid R\}$; E' est le terme $(y_1 \mid x_1, \dots, y_n \mid x_n)E$; R' est la proposition $(y_1 \mid x_1, \dots, y_n \mid x_n)R$; y_i est librement substituable à x_i dans E et dans R (pour $i = 1, \dots, n$).

Pour démontrer commodément cette égalité, il est indiqué de généraliser d'abord le schéma de déduction S76 de la logique des prédicats sous la forme du schéma à n variables suivant dans lequel chaque Q_i peut être l'un quelconque

des symboles $\forall, \exists$.

$$\frac{}{\Gamma \vdash Q_1 x_1 \cdots Q_n x_n A \Leftrightarrow Q_1 y_1 \cdots Q_n y_n (y_1 | x_1, \dots, y_n | x_n) A} \quad (9)$$

Si $x_1, \dots, x_n$ sont des variables distinctes; $y_1, \dots, y_n$ sont des variables distinctes qui ne figurent pas librement dans $Q_1 x_1 \cdots Q_n x_n A$; y_i est librement substituable à x_i dans A (pour $i = 1, \dots, n$).

Nous admettrons (9) sans démonstration. Il est assez évident que les propositions de gauche et de droite dans cette équivalence ont le même squelette.

En particulier, si A est une proposition de la forme $(z = E \text{ et } R)$, et si $z, x_1, \dots, x_n, y_1, \dots, y_n, E'$ et R' satisfont à aux conditions de E176, on peut écrire la démonstration suivante:

$$\begin{array}{ll} \text{nil} & \\ 1^\circ & \left| \begin{array}{l} \exists x_1 \cdots \exists x_n (z = E \text{ et } R) \Leftrightarrow \exists y_1 \cdots \exists y_n (z = E' \text{ et } R') \end{array} \right. \quad (9) \\ 2^\circ & \left| \begin{array}{l} \{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} = \{z \mid \exists y_1 \cdots \exists y_n (z = E' \text{ et } R')\} \end{array} \right. \quad 1^\circ, \text{E144} \\ 3^\circ & \left| \begin{array}{l} \{E \mid x_1 \cdots x_n \mid R\} = \{E' \mid y_1 \cdots y_n \mid R'\} \end{array} \right. \quad 2^\circ, \text{E172.} \end{array}$$

Schéma d'inclusion d'une collection dans un ensemble. Considérons par exemple l'ensemble

$$T = \{x + y \mid P(x) \text{ et } P(y) \text{ et } x \neq 2 \text{ et } y \neq 2\} \quad (10)$$

où P est le prédicat «est un nombre premier». Une propriété évidente de cet ensemble est qu'il est contenu dans l'ensemble des nombres (entiers naturels) pairs, car tout nombre premier différent de 2 est impair et la somme de deux nombres impairs est un nombre pair. Dans ce bref raisonnement, on admet que pour démontrer

$$\forall z (z \in T \Rightarrow z \text{ est pair}), \quad (11)$$

il suffit de démontrer

$$\forall x \forall y ((P(x) \text{ et } P(y) \text{ et } x \neq 2 \text{ et } y \neq 2) \Rightarrow x + y \text{ est pair}), \quad (12)$$

et l'on déduit tacitement (11) de (12). En fait les propositions (11) et (12) sont *équivalentes*. Les équivalences de cette forme, à propos de collections $\{E \mid x_1 \cdots x_n \mid R\}$ qui *existent*, sont très fréquentes. Elles sont décrites dans le schéma suivant.

E177.

$$\frac{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe} \quad \Gamma \vdash T = \{E \mid x_1 \cdots x_n \mid R\}}{\Gamma \vdash \forall z (z \in T \Rightarrow A) \Leftrightarrow \forall x_1 \cdots \forall x_n (R \Rightarrow (E | z) A)}$$

Si $x_1, \dots, x_n, z$ sont des variables distinctes et si: $x_1, \dots, x_n$ ne figurent librement ni dans T , ni dans A ; z ne figure librement ni dans E , ni dans T , ni dans R ; E est librement substituable à z dans A .

Dans l'exemple qui précède, A est la proposition « z est pair»; E est le terme $x + y$; R est la proposition $(P(x) \text{ et } P(y) \text{ et } x \neq 2 \text{ et } y \neq 2)$.

Comme les variables $x_1, \dots, x_n, z$ ne figurent pas librement dans les assertions des prémisses de la déduction E177, il suffit de démontrer celle-ci dans le cas particulier où ces variables ne figurent pas librement dans Γ . C'est ce qui est supposé dans la démonstration suivante où l'on remplace des propositions dominées par les quantificateurs $\forall x_1, \dots, \forall x_n, \forall z$ par des propositions équivalentes.

Γ		
1°	$\{E \mid x_1 \cdots x_n \mid R\}$ existe	<i>base</i>
2°	$T = \{E \mid x_1 \cdots x_n \mid R\}$	<i>base</i>
	$\forall z(z \in T \Rightarrow A)$	
	$\Leftrightarrow \forall z(\exists x_1 \cdots \exists x_n(z = E \text{ et } R) \Rightarrow A)$	1°, 2°, E174
	$\Leftrightarrow \forall z(\text{non} \exists x_1 \cdots \exists x_n(z = E \text{ et } R) \text{ ou } A)$	S53
	$\Leftrightarrow \forall z(\forall x_1 \cdots \forall x_n \text{non}(z = E \text{ et } R) \text{ ou } A)$	S71
	$\Leftrightarrow \forall x_1 \cdots \forall x_n \forall z(\text{non}(z = E \text{ et } R) \text{ ou } A)$	S81, S72
	$\Leftrightarrow \forall x_1 \cdots \forall x_n \forall z(z \neq E \text{ ou } (\text{non} R \text{ ou } A))$	S52
	$\Leftrightarrow \forall x_1 \cdots \forall x_n \forall z(z = E \Rightarrow (R \Rightarrow A))$	S53
	$\Leftrightarrow \forall x_1 \cdots \forall x_n((E z)(R \Rightarrow A))$	S92
	$\Leftrightarrow \forall x_1 \cdots \forall x_n(R \Rightarrow (E z)A)$	z ne fig. pas librement dans R .

E178.

$$\frac{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe} \quad \Gamma \vdash T = \{E \mid x_1 \cdots x_n \mid R\}}{\Gamma \vdash T \subset U \Leftrightarrow \forall x_1 \cdots \forall x_n(R \Rightarrow E \in U)}$$

Si $x_1, \dots, x_n$ sont des variables distinctes et ne figurent librement ni dans T , ni dans U .

Il suffit de voir que l'on peut remplacer $T \subset U$ par $\forall z(z \in T \Rightarrow z \in U)$ en prenant une variable z qui ne figure librement ni dans T ni dans U . Si en outre z ne figure librement ni dans R , ni dans Γ , on peut appliquer E177 comme suit :

Γ		
1°	$\{E \mid x_1 \cdots x_n \mid R\}$ existe	<i>base</i>
2°	$T = \{E \mid x_1 \cdots x_n \mid R\}$	<i>base</i>
	$T \subset U \Leftrightarrow \forall z(z \in T \Rightarrow z \in U)$	E3
	$\Leftrightarrow \forall x_1 \cdots \forall x_n(R \Rightarrow E \in U)$	1°, 2°, E177.

Exemple. Pour l'ensemble $T = \{x + y \mid P(x) \text{ et } P(y)\}$ et un ensemble U quelconque, la proposition $T \subset U$ est équivalente à

$$\forall x \forall y((P(x) \text{ et } P(y)) \Rightarrow x + y \in U).$$

En particulier, la proposition

$$\forall x \forall y((P(x) \text{ et } P(y)) \Rightarrow x + y \in T)$$

est vraie pour cet ensemble T , puisqu'elle équivaut à $T \subset T$. C'est un exemple du schéma de déduction suivant.

E179.

$$\frac{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe} \quad \Gamma \vdash T = \{E \mid x_1 \cdots x_n \mid R\}}{\Gamma \vdash \forall x_1 \cdots \forall x_n (R \Rightarrow E \in T)}$$

Si $x_1, \dots, x_n$ sont des variables distinctes ne figurant pas librement dans T .

Cas particulier des collections de couples. Le symbole d'implication dans la conclusion du schéma E179 peut être remplacé par un symbole d'équivalence dans le cas particulier où E est le terme $(x_1, \dots, x_n)$ (§12.2, n -uples). Ceci est formulé pour $n = 2$ par le schéma suivant.

E180.

$$\frac{\Gamma \vdash \{(x, y) \mid xy \mid R\} \text{ existe} \quad \Gamma \vdash T = \{(x, y) \mid xy \mid R\}}{\Gamma \vdash T \text{ est un graphe et } \forall x \forall y ((x, y) \in T \Leftrightarrow R)}$$

Si x et y sont des variables distinctes qui ne figurent pas librement dans T .

Il suffit de démontrer cette déduction dans le cas particulier où x et y ne figurent pas librement dans Γ . La démonstration nécessite un changement de variables liées. On choisit pour cela deux variables distinctes x' et y' , distinctes de x et y et réversiblement substituables à x et y dans R , en ce sens que les propositions

$$R \quad (x \mid x', y \mid y')((x' \mid x, y' \mid y)R)$$

soient identiques. Nous désignons par R' la proposition $(x' \mid x, y' \mid y)R$. Celle-ci est identique d'ailleurs à $(x' \mid x)((y' \mid y)R)$ puisque x' et y' sont distinctes (§10.5). Nous supposons enfin que x' et y' (ainsi que x et y) ne figurent pas librement dans T et que z est une variable distincte de x et y et ne figure librement ni dans T , ni dans R , ni dans Γ .

Γ		
1°	$\{(x, y) \mid xy \mid R\} \text{ existe.}$	<i>base</i>
2°	$T = \{(x, y) \mid xy \mid R\}$	<i>base</i>
3°	$z \in T \Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } R)$	1°, 2°, E174
4°	$z \in T \Rightarrow z \text{ est un couple}$	3°, E81
5°	$\forall z (z \in T \Rightarrow z \text{ est un couple})$	4°, S62
6°	$T \text{ est un graphe}$	5°, E86
7°	$(x', y') \in T \Leftrightarrow \exists x \exists y ((x', y') = (x, y) \text{ et } R)$	3°, S79
8°	$\Leftrightarrow \exists x \exists y (x = x' \text{ et } y = y' \text{ et } R)$	E79
9°	$\Leftrightarrow \exists x (x = x' \text{ et } \exists y (y = y' \text{ et } R))$	S81
10°	$\Leftrightarrow (x' \mid x)((y' \mid y)R)$	S92
11°	$(x', y') \in T \Leftrightarrow R'$	7°–10°
12°	$(x \mid x', y \mid y')((x', y') \in T \Leftrightarrow R')$	11°, S79
13°	$(x, y) \in T \Leftrightarrow R$	identique à 12°
14°	$\forall x \forall y ((x, y) \in T \Leftrightarrow R).$	13°, S62.

Condition d'existence d'une collection. Le schéma suivant est celui que l'on applique généralement pour montrer qu'une collection $\{E \mid x_1 \cdots x_n \mid R\}$ existe.

E181.

$$\frac{\Gamma \vdash R \Rightarrow E \in U}{\Gamma \vdash \{E \mid x_1 \cdots x_n \mid R\} \text{ existe}}$$

Si $x_1, \dots, x_n$ sont des variables distinctes et ne figurent librement ni dans Γ , ni dans U .

En prenant une variable z distincte de $x_1, \dots, x_n$ et ne figurant ni dans E , ni dans R , ni dans U , ni dans Γ , on peut écrire la démonstration suivante. La condition « $x_1, \dots, x_n$ ne figurent pas librement dans Γ » intervient lorsqu'on remplace une proposition dominée par les quantificateurs $\exists x_1, \dots, \exists x_n$ par une proposition équivalente.

Γ		
1°	$R \Rightarrow E \in U$	<i>base</i>
2°	$\exists x_1 \cdots \exists x_n (z = E \text{ et } R) \Leftrightarrow \exists x_1 \cdots \exists x_n (z = E \text{ et } R \text{ et } E \in U)$	1°, S60.3
3°	$\Leftrightarrow \exists x_1 \cdots \exists x_n (z = E \text{ et } R \text{ et } z \in U)$	S87.1
4°	$\Leftrightarrow \exists x_1 \cdots \exists x_n (z = E \text{ et } R) \text{ et } z \in U$	S81
5°	$\Rightarrow z \in U$	S28
6°	$\exists x_1 \cdots \exists x_n (z = E \text{ et } R) \Rightarrow z \in U$	2°–5°
7°	$\{z \mid \exists x_1 \cdots \exists x_n (z = E \text{ et } R)\} \text{ existe}$	6°, E147
8°	$\{E \mid x_1 \cdots x_n \mid R\} \text{ existe}$	E173.

On a souvent affaire au cas particulier d'une collection de la forme $\{E \mid x \mid R\}$ pour laquelle on connaît un terme V dans lequel x ne figure pas librement et tel que la proposition $R \Rightarrow x \in V$ est vraie. Le schéma suivant s'applique alors.

E182.

$$\frac{\Gamma \vdash R \Rightarrow x \in V}{\Gamma \vdash \{E \mid x \mid R\} \text{ existe}} \quad \begin{array}{l} \text{Si } x \text{ ne figure librement} \\ \text{ni dans } V, \text{ ni dans } \Gamma. \end{array}$$

Car si x ne figure pas librement dans V , on a le théorème (E104)

$$\vdash x \in V \Rightarrow E(x) \in \bigcup_{x \in V} \{E(x)\}$$

où $E(x)$ est une désignation du terme E suggérant la dépendance de x . De la prémisses de E182 on déduit donc $\Gamma \vdash R \Rightarrow E \in \bigcup_{x \in V} \{E\}$ et l'on peut conclure sui-

vant E181, puisque la variable x ne figure pas librement dans le terme $\bigcup_{x \in V} \{E\}$.

Les collections simples comme cas particulier. Lorsqu'on abrège la notation $\{E \mid x_1 \cdots x_n \mid R\}$ par $\{E \mid R\}$ en appliquant la convention qui précède (à savoir que $x_1, \dots, x_n$ sont les variables libres communes à E et R), alors une collection simple $\{x \mid R\}$ où x figure librement dans R est aussi l'abréviation de

la collection $\{x \mid x \mid R\}$. Cette ambiguïté n'est pas gênante, en vertu des schémas suivants :

$$\Gamma \vdash \{x \mid x \mid R\} \text{ existe} \Leftrightarrow \{x \mid R\} \text{ existe} \qquad \Gamma \vdash \{x \mid x \mid R\} = \{x \mid R\}.$$

Ces déductions se démontrent en choisissant une variable z distincte de x et réversiblement substituable à x dans R , donc librement substituable à x dans R et ne figurant pas librement dans R . On peut écrire alors la démonstration :

nil		
1°	$\exists x(z = x \text{ et } R) \Leftrightarrow (z \mid x)R$	S92
	$\{x \mid x \mid R\} \text{ existe} \Leftrightarrow \{z \mid \exists x(z = x \text{ et } R)\} \text{ existe}$	E173
	$\Leftrightarrow \{z \mid (z \mid x)R\} \text{ existe}$	1°, E144
	$\Leftrightarrow \{x \mid R\} \text{ existe}$	E145
	$\{x \mid x \mid R\} = \{z \mid \exists x(z = x \text{ et } R)\}$	E172
	$= \{z \mid (z \mid x)R\}$	1°, E144
	$= \{x \mid R\}$	E145.

Composé de deux graphes. L'opération de composition de deux graphes G, G' est illustrée par la figure 14.7. Le graphe composé $G' \circ G$ a pour éléments les couples (a, b) qui vérifient la condition suivante : il existe un x tel que $(a, x) \in G$ et $(x, b) \in G'$. Au vu de cette définition, il semblerait naturel de noter ce graphe $G \circ G'$ plutôt que $G' \circ G$. Certains auteurs le font, mais la notation $G' \circ G$ est plus traditionnelle. La tradition est de noter $F' \circ F$ la composée de deux fonctions F, F' , avec F' à gauche, car il s'agit de la fonction qui retourne pour un argument x la valeur $F'(F(x))$, et F' vient à gauche de F dans cette expression. Les fonctions étant des graphes particuliers, il faut adopter pour tous les graphes la notation « à gauche » de la composition, si l'on ne veut pas rompre avec la tradition des notations fonctionnelles.

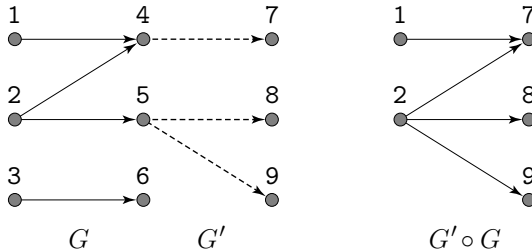


Figure 14.7 Composé $G' \circ G$ de deux graphes.

La définition formelle de l'opération de composition de graphes est la suivante. L'ensemble $G' \circ G$ est défini comme une collection de couples. Cela nous fournit l'occasion d'appliquer les schémas de déduction qui précèdent.

E183. *Définition*

$\vdash$ Si G et G' sont des graphes,
 $G' \circ G = \{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}.$

La démonstration détaillée des théorèmes suivants est donnée dans l'annexe (Chap 16). Nous esquissons ici celle des deux premiers, qui est une application standard des schémas de déduction sur les collections. On ne pose de définition de la forme

$$\vdash C \Rightarrow T = \{E \mid x_1 \cdots x_n \mid R\}$$

que lorsque la proposition $C \Rightarrow (\{E \mid x_1 \cdots x_n \mid R\} \text{ existe})$ est vraie. Dans le cas présent :

C est la proposition « G et G' sont des graphes »,
 T est le terme $G' \circ G$,
 E est le terme (a, b) ,
 R est la proposition $\exists x((a, x) \in G \text{ et } (x, b) \in G')$.

La liste de variables liées $x_1 \cdots x_n$ est omise dans le terme $\{E \mid x_1 \cdots x_n \mid R\}$. Suivant la convention faite plus haut, cette liste se compose de toutes les variables figurant librement dans E et dans R , donc des variables a et b .

Sous l'hypothèse « G et G' sont des graphes », on démontre d'abord que $R \Rightarrow (a, b) \in \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$, ce qui se voit d'ailleurs immédiatement. Suivant E181, on en déduit que $\{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}$ existe. On a donc le théorème

E184. $\vdash$ Si G et G' sont des graphes, alors
 $\{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}$ existe.

Le théorème principal qui suit découle de E183 et E184 suivant E180 :

E185. *Théorème principal*

$\vdash$ Si G et G' sont des graphes, alors $G' \circ G$ est un graphe et
 $(a, b) \in G' \circ G \Leftrightarrow \exists x((a, x) \in G \text{ et } (x, b) \in G').$

E186. $\vdash G$ et G' sont des graphes $\Rightarrow G' \circ G \subset \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$.

E187. $\vdash G, G', G''$ sont des graphes $\Rightarrow G'' \circ (G' \circ G) = (G'' \circ G') \circ G$.

E188. $\vdash G$ et G' sont des graphes $\Rightarrow (G' \circ G)^{-1} = G^{-1} \circ G'^{-1}$.

E189. $\vdash$ Si G et G' sont des graphes, alors

$$\mathbf{Pr}_1(G' \circ G) = G^{-1} \langle \mathbf{Pr}_1 G' \rangle;$$

$$\mathbf{Pr}_2(G' \circ G) = G' \langle \mathbf{Pr}_2 G \rangle;$$

$$\mathbf{Pr}_1(G' \circ G) \subset \mathbf{Pr}_1 G \text{ et } \mathbf{Pr}_2(G' \circ G) \subset \mathbf{Pr}_2 G'.$$

E190. $\vdash G$ est un graphe $\Rightarrow (G \circ \emptyset = \emptyset \text{ et } \emptyset \circ G = \emptyset)$.

E191. $\vdash (G \text{ est un graphe et } \mathbf{Pr}_1 G \subset X) \Rightarrow G \circ \text{Id}_X = G$.

$\vdash (G \text{ est un graphe et } \mathbf{Pr}_2 G \subset Y) \Rightarrow \text{Id}_Y \circ G = G$.

E192. $\vdash G$ est un graphe $\Rightarrow \text{Id}_{\mathbf{Pr}_1 G} \subset G^{-1} \circ G$.

$\vdash G$ est un graphe $\Rightarrow \text{Id}_{\mathbf{Pr}_2 G} \subset G \circ G^{-1}$.

E193. $\vdash G$ est un graphe $\Rightarrow G \langle X \rangle = \mathbf{Pr}_2(G \circ \text{Id}_X)$.

E194. $\vdash G$ et G' sont des graphes $\Rightarrow (G' \circ G) \langle X \rangle = G' \langle G \langle X \rangle \rangle$.

15

Fonctions

15.1 Fonctions et applications

Remarque préliminaire. Le sens du mot « fonction » oscille dans la littérature entre deux définitions apparentées mais différentes. Nous avons fait ici le choix de la plus simple, suivant laquelle une fonction n'est qu'un cas particulier de graphe. Dans l'autre définition, une fonction est un objet composé dont l'une des composantes est un graphe, et l'on doit distinguer entre une fonction et le graphe de cette fonction. Mais dans de nombreuses circonstances on ne s'intéresse qu'au graphe d'une fonction et les autres composantes de la définition ne jouent aucun rôle. Un exposé formel de ces notions est donc simplifié lorsqu'on identifie le concept de fonction à celui de graphe d'une fonction. Nous utiliserons le terme d'*application* pour parler d'une fonction dans le sens composé. Il ne sera donc pas synonyme de « fonction ».

Les définitions, théorèmes et schémas de déduction de ce paragraphe sont numérotés **F1**, **F2**, etc. Les démonstrations omises se trouvent dans l'annexe (Chap. 16).

Définition d'une fonction. Un graphe F est appelé une fonction s'il satisfait à la condition suivante: pour tout objet x , il existe au plus un objet y tel que $(x, y) \in F$. Un graphe G n'est pas une fonction s'il vérifie la négation de cette condition, à savoir s'il existe un objet x et deux objets *distincts* y, y' tels que $(x, y) \in G$ et $(x, y') \in G$. La figure 15.1 représente un graphe F qui est une fonction et un graphe G qui n'en est pas une.

F1. *Définition*

$\vdash F$ est une fonction $\Leftrightarrow [F \text{ est un graphe et } \forall x \text{Hy}((x, y) \in F)]$.

Le théorème suivant est une autre manière de caractériser une fonction. Il se démontre facilement d'après **F1** et **E117**, suivant **S94** et **S95**.

F2. $\vdash F$ est une fonction $\Leftrightarrow$

$[F \text{ est un graphe et } \forall x[x \in \mathbf{Pr}_1 F \Rightarrow \exists! y((x, y) \in F)]]$.

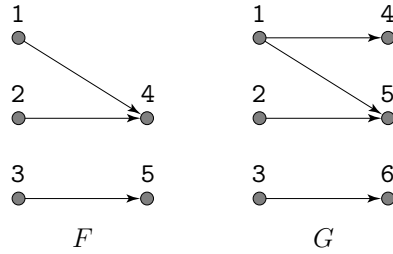


Figure 15.1 Le graphe F est une fonction; le graphe G n'est pas une fonction.

Enfin, une troisième manière de caractériser une fonction est donnée par le théorème suivant, que l'on comparera avec E192.

F3. $\vdash F$ est une fonction $\Leftrightarrow (F \text{ est un graphe et } F \circ F^{-1} = \text{Id}_{\mathbf{Pr}_2 F})$.

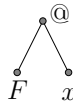
Domaine et portée d'une fonction. Si F est une fonction, l'ensemble $\mathbf{Pr}_1 F$ est appelé traditionnellement le *domaine* de F . On parle aussi du *domaine de définition* de F . Nous utiliserons donc la notation $\text{Dom} F$ au lieu de $\mathbf{Pr}_1 F$. Il n'y a pas de terme largement répandu en français pour l'ensemble $\mathbf{Pr}_2 F$. En anglais, on l'appelle «range of F » et l'on écrit $\text{Ran} F$. Nous parlerons de la *portée* de F , ce qui est une traduction littérale de «range», et nous écrirons $\text{Prt} F$. Nous posons donc les définitions suivantes.

F4. *Définition*

$\vdash F$ est une fonction $\Rightarrow (\text{Dom} F = \mathbf{Pr}_1 F \text{ et } \text{Prt} F = \mathbf{Pr}_2 F)$.

Image d'un objet par une fonction. Pour un objet x qui appartient au domaine d'une fonction F , l'unique objet y tel que $(x, y) \in F$ est noté traditionnellement $F(x)$ et appelé de diverses manières: l'image de x par F ; l'objet associé à x par F ; l'objet qui correspond à x suivant F ; la valeur de F pour x ou correspondant à x ; la valeur retournée par F (langage des informaticiens) pour l'argument x ; le résultat de l'application de F à l'argument x .

La notation $F(x)$ n'est pas très heureuse, et nous utiliserons à sa place la notation $F @ x$ (lire: F at x) dans laquelle $@$ est un symbole fonctionnel binaire. Ce symbole représente l'opération d'application d'une fonction à un argument. L'arbre du terme $F @ x$ est



L'arbre du terme $F(x)$ est le même, mais ce sont les deux parenthèses $()$ qui jouent le rôle du symbole $@$. Dans les deux notations, F et x sont des *variables* et se trouvent sur les feuilles de l'arbre.

Ce qui n'est pas très bon dans la notation $F(x)$ c'est qu'elle utilise des parenthèses ordinaires dans un rôle spécial qui n'a rien à voir avec leur rôle usuel de groupement ou d'association. Lorsqu'on a affaire à des fonctions dont

le domaine et la portée sont des ensembles de fonctions, autrement dit des fonctions qui prennent des fonctions comme arguments et retournent des fonctions, on peut avoir des expressions telles que

$$(h @ ((f @ g) @ u)) @ x.$$

L'expression traditionnelle correspondante est $(h((f(g))(u)))(x)$ et sa structure hiérarchique est beaucoup moins claire à cause des parenthèses qui remplacent le symbole $@$ et qu'on ne distingue pas immédiatement des parenthèses de groupement.

Certains auteurs utilisent la notation multiplicative Fx au lieu de $F(x)$. Elle convient parfaitement lorsqu'elle ne se mélange pas avec d'autres sortes de produits. Cette notation multiplicative Fx , avec un espace entre F et x , est aussi celle de certains langages de programmation fonctionnels [30].

Il nous reste à poser la définition formelle de $F @ x$. Comme nous l'avons dit, il s'agit de *l'objet y tel que $(x, y) \in F$* . L'opérateur de sélection Sy est donc utilisé dans cette définition.

F5. Définition

$$\vdash F \text{ est une fonction} \Rightarrow F @ x = Sy((x, y) \in F).$$

Le théorème suivant découle de cette définition et de la définition d'une fonction (F1) d'après SEL8. C'est le théorème principal sur la notation $F @ x$.

F6. $\vdash$ Si F est une fonction, alors

$$(x, y) \in F \Leftrightarrow (x \in \text{Dom } F \text{ et } y = F @ x).$$

nil		
1°	$F \text{ est une fonction}$	hyp
2°	$Hy((x, y) \in F)$	1°, F1
3°	$(x, y) \in F \Leftrightarrow [\exists y((x, y) \in F) \text{ et } y = Sy((x, y) \in F)]$ $\Leftrightarrow x \in \text{Dom } F \text{ et } y = F @ x$	2°, SEL8 F4, E117, F5.

En utilisant le pseudo-prédicat « existe » (§13.2), l'équivalence de F6 peut s'écrire

$$(x, y) \in F \Leftrightarrow (F @ x \text{ existe et } y = F @ x),$$

puisque'on a posé la définition $F @ x = Sy((x, y) \in F)$ et puisque la proposition $x \in \text{Dom } F$ équivaut à $\exists y((x, y) \in F)$. Dans ce cas on dit plutôt « est défini » que « existe »:

$$(x, y) \in F \Leftrightarrow (F @ x \text{ est défini et } y = F @ x).$$

Mais l'omission de « $F @ x$ existe » ou de « $F @ x$ est défini » dans ces équivalences est un *abus*.

F7. $\vdash$ Si F est une fonction, alors

$$x \in \text{Dom } F \Leftrightarrow (x, F @ x) \in F;$$

$$\text{Prt } F = \{F @ x \mid x \in \text{Dom } F\}.$$

En effet, sous l'hypothèse « F est une fonction », on a les équivalences

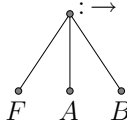
$$\begin{aligned}
 (x, F @ x) \in F &\Leftrightarrow (x \in \text{Dom} F \text{ et } F @ x = F @ x) && \text{F6} \\
 &\Leftrightarrow x \in \text{Dom} F && \text{S82, S60.1.}
 \end{aligned}$$

Le deuxième théorème de F7 est démontré dans l'annexe (Chap. 16).

Applications. On dit que F est une *application* de A dans B si F est une fonction et si $\text{Dom} F = A$ et $\text{Pr} F \subset B$. Formellement, cette définition introduit un nouveau prédicat (symbole relationnel) ternaire formé par les mots « est une application de ... dans ... ». La proposition « F est une application de A dans B » se note aussi

$$F : A \rightarrow B. \quad (1)$$

Le symbole relationnel ternaire est constitué ici par les signes « : » et « $\rightarrow$ ». L'arbre de la proposition (1) est le suivant :



F8. *Définition*

$$\vdash (F : A \rightarrow B) \Leftrightarrow (F \text{ est une fonction et } \text{Dom} F = A \text{ et } \text{Pr} F \subset B).$$

Théorèmes :

$$\text{F9. } \vdash (F : A \rightarrow B) \Leftrightarrow \left(F \text{ est une fonction et } \text{Dom} F = A \text{ et } \left(\forall x (x \in A \Rightarrow F @ x \in B) \right) \right).$$

$$\text{F10. } \vdash F \text{ est une fonction} \Leftrightarrow (F : \text{Dom} F \rightarrow \text{Pr} F).$$

Autre terminologie. De nombreux auteurs ne définissent pas séparément les notions de fonction et d'application. Ils n'introduisent qu'une seule notion, appelée indifféremment fonction ou application. Cette notion est celle d'un triplet (F, A, B) dans lequel F est un graphe tel que $\text{pr}_1 F = A$ et $\text{Pr}_2 F \subset B$ et tel que $\forall x \text{Hy}((x, y) \in F)$. Autrement dit, une « fonction » au sens de ces auteurs est un triplet (F, A, B) tel que l'on ait $F : A \rightarrow B$ au sens de *nos* définitions. En principe ces auteurs devraient toujours distinguer une *fonction* (F, A, B) au sens de leur définition, du *graphe* de cette fonction, le graphe F , qui n'est qu'une des trois composantes du triplet (F, A, B) . Mais dans la pratique, cette distinction n'est pas toujours respectée et le sens du mot « fonction » varie : on l'utilise tantôt pour désigner (F, A, B) , tantôt pour désigner seulement le graphe F .

15.2 Opérateurs d'abstraction

Caractérisation d'une fonction par des termes. Une fonction particulière F est souvent définie par la donnée de deux *termes* E, T au moyen d'équations de la forme

$$\begin{aligned}
 \text{Dom} F &= E; \\
 F @ x &= T \quad (\text{pour tout } x \in \text{Dom} F).
 \end{aligned} \quad (1)$$

Il est admis sans démonstration, dans presque tous les ouvrages de mathématiques, que des équations de cette forme admettent une solution F et une seule qui est une fonction, autrement dit que ces équations peuvent être prises comme définition d'une fonction particulière F bien déterminée.

Par exemple, il est admis sans autre que si a, b, c sont trois nombres réels, il existe une fonction f et une seule ayant pour domaine l'intervalle $[a, b]$ et telle que $f @ x = x^2 + c$ pour tout élément x de ce domaine. Le terme E de cet exemple est le terme $[a, b]$; le terme T est $x^2 + c$. La fonction f est déterminée par la donnée de ces deux termes et par l'indication que le second est considéré comme dépendant de x (et non de c).

En admettant cela, on effectue une déduction dont nous allons donner ici la justification générale, sous la forme de quelques schémas de déduction dérivés de $\mathcal{T}_{\text{ENS}}$. Dans ces schémas, le terme désigné par T est toujours considéré comme dépendant d'une variable x , même si celle-ci ne figure pas librement dans ce terme, et nous le désignons par $T(x)$ pour indiquer cette manière de le considérer. Si x' est une autre variable, nous désignons par $T(x')$ le terme $(x'|x)T$. Le schéma de base est le suivant.

F11.

$$\Gamma \vdash \left[\begin{array}{l} f \text{ est une fonction et} \\ \text{Dom } f = E \text{ et} \\ \forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x)) \end{array} \right] \Leftrightarrow f = \{(x, T(x)) \mid x \mid x \in E\}$$

Si E et $T(x)$ sont des termes; f et x sont des variables distinctes;
 x ne figure pas librement dans E .

La démonstration de ce schéma est relativement longue. Elle est donnée dans la figure 15.2. Le lecteur est invité à la vérifier minutieusement, car ce schéma joue un rôle clef dans ce chapitre. Tous ceux qui le suivent en dérivent très simplement. Il s'agit surtout de vérifier qu'en vertu des hypothèses faites sur les variables utilisées (légende de la figure 15.2) les conditions d'application des schémas de déduction utilisés dans la démonstration sont toujours satisfaites. L'assertion de la ligne 1^o se justifie par E182 puisque la proposition $x \in E \Rightarrow x \in E$ est vraie et l'on suppose que x ne figure pas librement dans E . Quelques lignes évidentes sont omises entre 13^o et 14^o.

Comme première application de ce schéma, nous avons le théorème suivant :

F12. $\vdash F$ est une fonction $\Leftrightarrow F = \{(x, F @ x) \mid x \mid x \in \text{Dom } F\}$.

Démonstration. Il suffit d'appliquer le schéma F11 en prenant: pour f la variable F ; pour E le terme $\text{Dom } F$; pour $T(x)$ le terme $F @ x$. Les propositions vraies $\text{Dom } F = \text{Dom } F$ et $\forall x (x \in \text{Dom } F \Rightarrow F @ x = F @ x)$ que contient alors le membre de gauche de l'équivalence F11 peuvent être supprimées suivant S60.1.

Le schéma suivant fournit la justification de la méthode de définition d'une fonction au moyen de deux termes par les équations (1).

nil		
1°	$\{(x, T(x)) \mid x \mid x \in E\}$ existe	E182
2°	f est une fonction et $\text{Dom } f = E$ et $\forall x(x \in \text{Dom } f \Rightarrow f @ x = T(x))$	hyp
3°	$z \in f \Leftrightarrow (z \text{ est un couple et } z \in f)$	E88
4°	$z \in f \Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } z \in f)$	E81, S81
5°	$z \in f \Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } (x, y) \in f)$	S87.1
6°	$\Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } x \in \text{Dom } f \text{ et } y = f @ x)$	2°, F6
7°	$\Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } x \in \text{Dom } f \text{ et } y = T(x))$	2°, S87.3
8°	$\Leftrightarrow \exists x \exists y (z = (x, y) \text{ et } x \in E \text{ et } y = T(x))$	2°, S86
9°	$\Leftrightarrow \exists x (z = (x, T(x)) \text{ et } x \in E)$	2°, S92
10°	$\Leftrightarrow z \in \{(x, T(x)) \mid x \mid x \in E\}$	1°, E174
11°	$f = \{(x, T(x)) \mid x \mid x \in E\}$	5°–10°, E17
12°	$f = \{(x, T(x)) \mid x \mid x \in E\}$	hyp
13°	$z \in f \Leftrightarrow \exists x (z = (x, T(x)) \text{ et } x \in E)$	1°, 12°, E174
14°	$\forall z (z \in f \Rightarrow z \text{ est un couple})$	13°
15°	f est un graphe	14°, E86
16°	$(x', y) \in f \Leftrightarrow \exists x ((x', y) = (x, T(x)) \text{ et } x \in E)$	13°, S79
17°	$\Leftrightarrow \exists x (x = x' \text{ et } y = T(x) \text{ et } x \in E)$	E79
18°	$(x', y) \in f \Leftrightarrow (y = T(x') \text{ et } x' \in E)$	S92
19°	$(x, y) \in f \Leftrightarrow (y = T(x) \text{ et } x \in E)$	18°, S79
20°	$(x, y) \in f \Rightarrow y = T(x)$	19°
21°	$\text{Hy}((x, y) \in f)$	20°, S98
22°	f est une fonction	15°, 21°, F1
23°	$x \in \text{Dom } f \Leftrightarrow \exists y ((x, y) \in f)$	E117
24°	$\Leftrightarrow \exists y (y = T(x) \text{ et } x \in E)$	19°
25°	$\Leftrightarrow x \in E$	S92
26°	$\text{Dom } f = E$	23°–25°, E17
27°	$x \in \text{Dom } f$	hyp
28°	$(x, f @ x) \in f$	26°, F7
29°	$(x, f @ x) \in f \Rightarrow f @ x = T(x)$	20°, S79
30°	$f @ x = T(x)$	28°, 29°
31°	$\forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x))$	$\Rightarrow$ -intro, S62
32°	f est une fonction et $\text{Dom } f = E$ et $\forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x))$	22°, 26°, 31°
33°	$(f \text{ est une fonction et } \dots) \Leftrightarrow f = \{(x, T(x)) \mid x \mid x \in E\}$	11°, 32°, S40.

Figure 15.2 Démonstration de F11. On suppose que f, x, x', y, z sont des variables distinctes et que

- x et x' ne figurent pas librement dans E ;
- x' est réversiblement substituable à x dans T ;
- y, z ne figurent librement ni dans E , ni dans T .

F13.

$$\Gamma \vdash \exists! f (f \text{ est une fonction et } \text{Dom } f = E \text{ et } \forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x)))$$

Si f et x sont des variables distinctes et si
 f ne figure librement ni dans E , ni dans T .

Démonstration. Désignons par A la proposition

$$f \text{ est une fonction et } \text{Dom } f = E \text{ et } \forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x)).$$

Par changement de variable liée dans la troisième partie de cette proposition, on peut toujours se ramener au cas où x ne figure pas librement dans E . Il nous suffit donc de démontrer le schéma dans ce cas particulier. Or dans ce cas particulier on a le théorème

$$\vdash A \Leftrightarrow f = \{(x, T(x)) \mid x \mid x \in E\}$$

d'après F11. On en déduit $\vdash \exists! f A$ suivant S98, puisque f ne figure pas librement dans le terme $\{(x, T(x)) \mid x \mid x \in E\}$.

Opérateurs d'abstraction. Étant donné un terme E , une variable x et un terme $T(x)$, on désigne par

$$(\lambda x \in E).T(x) \tag{2}$$

la fonction f ayant pour domaine l'ensemble E et telle que $f @ x = T(x)$ pour tout $x \in E$. Formellement, on pose le schéma définitionnel suivant:

F14. *Définition*

$$\Gamma \vdash (\lambda x \in E).T(x) = Sf \left[\begin{array}{l} f \text{ est une fonction et} \\ \text{Dom } f = E \text{ et} \\ \forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x)) \end{array} \right]$$

Si f et x sont des variables distinctes et si
 f ne figure librement ni dans E , ni dans T .

Le corollaire immédiat de cette définition, compte de tenu de F14 et suivant SEL8, est le schéma:

F15.

$$\Gamma \vdash f = (\lambda x \in E).T(x) \Leftrightarrow \left[\begin{array}{l} f \text{ est une fonction et} \\ \text{Dom } f = E \text{ et} \\ \forall x (x \in \text{Dom } f \Rightarrow f @ x = T(x)) \end{array} \right]$$

Si f et x sont des variables distinctes et si
 f ne figure librement ni dans E , ni dans T .

Les trois schémas suivants (F16) sont des corollaires tout aussi immédiats de F15. On les démontre pour Γ vide en remplaçant f par $(\lambda x \in E).T(x)$ dans l'équivalence de F15, dont le premier membre devient alors l'égalité vraie

$(\lambda x \in E).T(x) = (\lambda x \in E).T(x)$. Le troisième schéma de F16 s'obtient en particulierisant avec U la proposition $\forall x(x \in \text{Dom } f \Rightarrow f @ x = T(x))$, dans laquelle on a remplacé f par $(\lambda x \in E).T(x)$.

F16.

$$\frac{}{\Gamma \vdash (\lambda x \in E).T(x) \text{ est une fonction}} \quad \frac{}{\Gamma \vdash \text{Dom}((\lambda x \in E).T(x)) = E}$$

$$\frac{}{\Gamma \vdash U \in \text{Dom}((\lambda x \in E).T(x)) \Rightarrow ((\lambda x \in E).T(x)) @ U = T(U)}$$

Si E , $T(x)$, U sont des termes; $T(U)$ est le terme $(U|x)T$ et U est librement substituable à x dans $T(x)$;

Syntaxiquement, le terme (2) se compose essentiellement de l'opérateur λx (λ est la lettre grecque lambda) et des deux termes E et T . Le reste n'est que de la décoration. On peut écrire simplement $\lambda x ET$ au lieu de $(\lambda x \in E).T(x)$. L'opérateur λx est appelé l'opérateur d'*abstraction fonctionnelle* sur x et le terme $\lambda x ET$ est appelé une *abstraction*.

L'opérateur λx est un opérateur d'arité 2 et de type $T \leftarrow (T, T)$ qui lie la variable x mais seulement dans son deuxième argument $T(x)$. Les occurrences libres de x dans E , s'il y en a, restent libres dans le terme composé $\lambda x ET(x)$. La raison de cette convention est que ces occurrences de x sont libres dans le membre de droite de l'égalité de F14, car le terme E qui figure dans ce membre de droite n'est pas dominé par le quantificateur $\forall x$, ni par un quelconque opérateur liant la variable x .

Le schéma de déduction suivant est utile dans la démonstration de propositions de la forme $(x, y) \in U \Leftrightarrow (x, y) \in (\lambda x \in E).T(x)$ au moyen desquelles on montre qu'un graphe U est égal à une fonction donnée par une abstraction.

F17.

$$\frac{}{\Gamma \vdash (x, y) \in (\lambda x \in E).T(x) \Leftrightarrow (x \in E \text{ et } y = T(x))}.$$

Pour la démonstration, on prend une variable f distincte de x et y et ne figurant librement ni dans E , ni dans T . On peut alors écrire:

nil		
1°	$f = (\lambda x \in E).T(x)$	hyp
2°	f est une fonction	1°, F15
3°	$\text{Dom } f = E$	1°, F15
4°	$x \in \text{Dom } f \Rightarrow f @ x = T(x)$	1°, F15
5°	$(x, y) \in f \Leftrightarrow (x \in \text{Dom } f \text{ et } y = f @ x)$	2°, F6
6°	$\Leftrightarrow (x \in \text{Dom } f \text{ et } y = T(x))$	4°, S87.3
7°	$\Leftrightarrow (x \in E \text{ et } y = T(x))$	3°
8°	$(x, y) \in f \Leftrightarrow (x \in E \text{ et } y = T(x))$	5°–7°
9°	$(x, y) \in (\lambda x \in E).T(x) \Leftrightarrow (x \in E \text{ et } y = T(x))$	8°, 1°
10°	$(x, y) \in (\lambda x \in E).T(x) \Leftrightarrow (x \in E \text{ et } y = T(x))$	9°, S90.

Substitutivité de l'égalité et changement de variable liée. À chaque introduction d'un nouvel opérateur dans un schéma définitionnel, il faut s'assurer que les règles générales de substitutivité de l'équivalence et de l'égalité et de changement de variable liée s'appliquent aux expressions construites avec le nouvel opérateur. Nous l'avons fait au chapitre 14 pour les opérateurs de réunion (E96, E97), les opérateurs de collection simples (E144, E145) et les opérateurs de collection généraux (E175, E176). Nous le faisons ici pour les opérateurs d'abstraction.

F18.

$$\frac{\Gamma \vdash E = E'}{\Gamma \vdash (\lambda x \in E).T(x) = (\lambda x \in E').T(x)}$$

$$\frac{\Gamma \vdash T(x) = T'(x)}{\Gamma \vdash (\lambda x \in E).T(x) = (\lambda x \in E).T'(x)} \quad \begin{array}{l} \text{Si } x \text{ ne figure pas} \\ \text{librement dans } \Gamma. \end{array}$$

F19.

$$\frac{}{\Gamma \vdash (\lambda x \in E).T(x) = (\lambda x' \in E).T(x')} \quad \begin{array}{l} \text{Si } x' \text{ est réversiblement} \\ \text{substituable à } x \text{ dans } T. \end{array}$$

Démonstration. Soit f une variable distincte de x et ne figurant librement ni dans Γ , ni dans E , ni dans T . Désignons par A la proposition

$$f \text{ est une fonction et } \text{Dom} f = E \text{ et } \forall x (x \in \text{Dom} f \Rightarrow f @ x = T(x))$$

et par A' la proposition obtenue en remplaçant le terme E par un terme E' dans cette proposition. De $\Gamma \vdash E = E'$ on déduit $\Gamma \vdash A \Leftrightarrow A'$ suivant S86, puis $\Gamma \vdash \text{Sf} A = \text{Sf} A'$ suivant SEL2, donc $\Gamma \vdash (\lambda x \in E).T(x) = (\lambda x \in E').T(x)$ suivant F14.

On raisonne semblablement si A' est la proposition obtenue en remplaçant le terme $T(x)$ dans A par un terme $T'(x)$. De $\Gamma \vdash T(x) = T'(x)$ on peut déduire $\Gamma \vdash A \Leftrightarrow A'$ suivant S86, S70, S50, à condition que x ne figure pas librement dans Γ , puisque $T(x)$ est dominé dans A par le quantificateur $\forall x$.

Enfin, si x' est une variable réversiblement substituable à x dans $T(x)$ et si A' est la proposition

$$f \text{ est une fonction et } \text{Dom} f = E \text{ et } \forall x' (x' \in \text{Dom} f \Rightarrow f @ x' = T(x')),$$

on a $\Gamma \vdash A \Leftrightarrow A'$ suivant S76, en supposant que l'on ait choisi f distincte de x et de x' , et l'on conclut suivant SEL2 et F14.

Cas usuel. Par changement de variable liée, on peut toujours transformer une abstraction en une abstraction égale (F19) de la forme $(\lambda x \in E).T(x)$ où la variable x ne figure pas librement dans E . Pratiquement c'est toujours à ce cas particulier que l'on a affaire. Le schéma suivant permet d'exprimer dans ce cas une abstraction par une collection dont la signification est particulièrement claire.

F20.

$$\Gamma \vdash (\lambda x \in E).T(x) = \{(x, T(x)) \mid x \mid x \in E\}$$

Si x ne figure pas librement dans E .

En effet, si cette condition est vérifiée, et si f est une variable distincte de x et ne figurant librement ni dans E , ni dans T , on a d'après F11 et F15 le théorème

$$\vdash f = \{(x, T(x)) \mid x \mid x \in E\} \Leftrightarrow f = (\lambda x \in E).T(x)$$

En y remplaçant f par $(\lambda x \in E).T(x)$ on voit l'égalité de F20 est équivalente à l'égalité vraie $(\lambda x \in E).T(x) = (\lambda x \in E).T(x)$.

D'après F20, on déduit de F12 le théorème:

F21. $\vdash F$ est une fonction $\Leftrightarrow F = (\lambda x \in \text{Dom} F).(F @ x)$.

Autres notations. Les mathématiques usuelles n'utilisent pas la notation λx . Les opérateurs d'abstraction fonctionnelle y sont notés de diverses manières plus ou moins précises. L'expression la plus répandue pour $(\lambda x \in E).T(x)$, lorsque x ne figure pas librement dans E , est

$$\text{la fonction } x \mapsto T(x) \quad (x \in E),$$

voire simplement

$$\text{la fonction } T(x) \quad (x \in E).$$

On va même jusqu'à omettre $(x \in E)$ lorsqu'il n'y a pas de confusion possible sur la variable x dont le terme T est considéré comme dépendant, ni sur le domaine E considéré.

Par exemple, si a, b, c sont trois nombres réels on parlera de

$$\text{la fonction } x \mapsto (x^2 + c) \quad (x \in [a, b]) \quad (3)$$

au lieu de $(\lambda x \in [a, b]).(x^2 + c)$.

La notation $(\lambda x \in [a, b]).(x^2 + c)$ présente des avantages par rapport à (3). Premièrement, il est clair que la variable x n'y figure pas librement, liée qu'elle est par l'opérateur λx . Les variables libres du terme $(\lambda x \in [a, b]).(x^2 + c)$ sont a, b, c . Cela est moins clair dans l'expression (3).

Deuxièmement, le terme $(\lambda x \in [a, b]).(x^2 + c)$ peut être utilisé à l'intérieur d'autres expressions, comme n'importe quel terme. L'usage le plus simple de ce genre est de poser

$$f = (\lambda x \in [a, b]).(x^2 + c). \quad (4)$$

Avec la notation (3), on s'exprime plutôt avec les mots: «soit f la fonction $x \mapsto (x^2 + c)$ ($x \in [a, b]$)». Mais lorsqu'on a posé (4) on peut aussi remplacer f par le terme $(\lambda x \in [a, b]).(x^2 + 1)$ dans toute autre expression dans laquelle il est librement substituable à f . Par exemple, au lieu de

$$y \in [a, b] \Rightarrow f @ y = y^2 + c \quad (\text{F16})$$

on peut écrire $y \in [a, b] \Rightarrow ((\lambda x \in [a, b]).(x^2 + c)) @ y = y^2 + c$.

Les opérateurs d'abstraction sont parmi les opérateurs fondamentaux des langages de programmation dits *fonctionnels*.

15.3 Propriétés de fonctions

Égalité de deux fonctions. L'égalité de deux fonctions F, G est celle de deux ensembles de couples, ou graphes (E91). En tenant compte de la propriété particulière de ces graphes, on a le théorème intuitivement évident suivant :

F22. $\vdash$ Si F et G sont des fonctions,

$$F = G \Leftrightarrow [\text{Dom} F = \text{Dom} G \text{ et } \forall x (x \in \text{Dom} F \Rightarrow F @ x = G @ x)].$$

La première moitié de la démonstration qui suit n'est qu'une affaire de logique des prédicats avec égalité. La théorie des ensembles n'y joue aucun rôle. La deuxième moitié utilise les schémas de déduction du paragraphe précédent.

nil		
1°	F et G sont des fonctions	hyp
2°	$F = G$	hyp
3°	$\text{Dom} F = \text{Dom} G$	2°, S88
4°	$F @ x = G @ x$	2°, S88
5°	$x \in \text{Dom} F \Rightarrow F @ x = G @ x$	4°, S21
6°	$\text{Dom} F = \text{Dom} G \text{ et } \forall x (x \in \text{Dom} F \Rightarrow F @ x = G @ x)$	3°, 5°, S62
7°	$\text{Dom} F = \text{Dom} G \text{ et } \forall x (x \in \text{Dom} F \Rightarrow F @ x = G @ x)$	hyp
8°	F est une fonction	1°
	$F = (\lambda x \in \text{Dom} G).(G @ x)$	8°, 7°, F15
	$= G$	1°, F21

Fonctions identiques; la fonction vide. Le graphe identique d'un ensemble X (§14.1) est une fonction. Plus précisément, on a les théorèmes suivants :

F23. $\vdash \text{Id}_X = (\lambda x \in X).x$
 $\vdash \text{Id}_X$ est une fonction.
 $\vdash \text{Dom} \text{Id}_X = X$.
 $\vdash x \in X \Rightarrow \text{Id}_X @ x = x$.

Démonstration. Les trois derniers de ces théorèmes se déduisent évidemment du premier d'après F15. Pour le premier, comme Id_X et $(\lambda x \in X).x$ sont des graphes (E132, F15), il suffit de montrer que les propositions $(x, y) \in \text{Id}_X$, $(x, y) \in (\lambda x \in X).x$ sont équivalentes :

nil	
$(x, y) \in (\lambda x \in X).x \Leftrightarrow (x \in x \text{ et } y = x)$	F17
$\Leftrightarrow (x, y) \in \text{Id}_X$	E133.

En particulier, le graphe $\text{Id}_\emptyset$ est une fonction et d'après E135 et E122 on a

- F24. $\vdash F = \emptyset \Leftrightarrow (F \text{ est une fonction et } \text{Dom}F = \emptyset).$
 $\vdash F = \emptyset \Leftrightarrow (F \text{ est une fonction et } \text{Pr}tF = \emptyset).$
 $\vdash \emptyset \text{ est une fonction et } \text{Dom}\emptyset = \emptyset \text{ et } \text{Pr}t\emptyset = \emptyset.$

D'après F8, on démontre encore:

- F25. $\vdash (F : \emptyset \rightarrow B) \Leftrightarrow F = \emptyset.$
 $\vdash (F : A \rightarrow \emptyset) \Leftrightarrow (A = \emptyset \text{ et } F = \emptyset).$
 $\vdash \exists! F(F : \emptyset \rightarrow B).$
 $\vdash A \neq \emptyset \Rightarrow \text{non}\exists F(F : A \rightarrow \emptyset).$

Composition de fonctions. Il est intuitivement assez évident que le graphe composé $G \circ F$ de deux *fonctions* F et G est une fonction, comme cela est illustré par la figure 15.3. Le théorème F26 donne plus d'informations sur cette fonction. La figure 15.3 illustre notamment la deuxième et la troisième assertions de F26.

- F26. $\vdash$ Si F et G sont des fonctions, alors
 $G \circ F$ est une fonction;
 $\text{Dom}(G \circ F) = \{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\};$
 $\forall x(x \in \text{Dom}(G \circ F) \Rightarrow (G \circ F) @ x = G @(F @ x)).$

Démonstration. D'après F15, la conjonction de ces trois propositions est équivalente à

$$G \circ F = (\lambda x \in \{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\}).(G @(F @ x)).$$

C'est cette proposition que nous démontrons, sous l'hypothèse que F et G sont des fonctions. Les deux membres de cette équivalence sont des graphes. Il suffit donc de montrer qu'un couple (x, y) appartient au premier membre si et seulement s'il appartient au second.

nil		
1°	$\{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\}$ existe	E146
2°	F et G sont des fonctions	hyp
3°	$A = \{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\}$	hyp
4°	$(x, y) \in G \circ F \Leftrightarrow \exists a((x, a) \in F \text{ et } (a, y) \in G)$	E185
5°	$\Leftrightarrow \exists a(x \in \text{Dom}F \text{ et } a = F @ x \text{ et } a \in \text{Dom}G \text{ et } y = G @ a)$	2°, F6
6°	$\Leftrightarrow (x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G \text{ et } y = G @(F @ x))$	S92
7°	$\Leftrightarrow (x \in A \text{ et } y = G @(F @ x))$	1°, 3°, E141
8°	$\Leftrightarrow (x, y) \in (\lambda x \in A).(G @(F @ x))$	F17
9°	$G \circ F = (\lambda x \in A).(G @(F @ x))$	4°–8°
10°	$G \circ F = (\lambda x \in \{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\}).(G @(F @ x))$	3°
	On conclut suivant S90 et $\Rightarrow$ -intro.	

Le théorème suivant décrit un cas particulier important. La plupart des fonctions composées auxquelles on a affaire en mathématique sont dans ce cas.

- F27. $\vdash$ Si F et G sont des fonctions et si $\text{Prt} F \subset \text{Dom} G$, alors
 $\text{Dom}(G \circ F) = \text{Dom} F$.
- F28. $\vdash (F : A \rightarrow B \text{ et } G : B \rightarrow C) \Rightarrow (G \circ F : A \rightarrow C)$.
- F29. $\vdash (F : A \rightarrow B) \Rightarrow (\text{Id}_B \circ F = F \text{ et } F \circ \text{Id}_A = F)$.

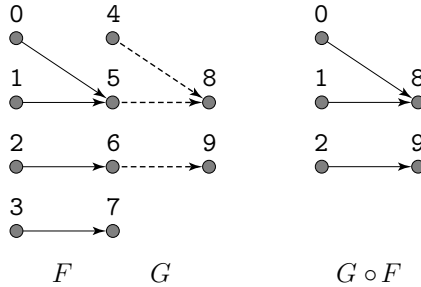


Figure 15.3 Composition de deux fonctions.

L'axiome de choix. Le théorème ci-dessous (F30) est l'une des nombreuses versions de ce qu'on appelle en théorie des ensembles l'*axiome de choix*. Dans notre construction logique ce n'est pas un axiome mais un théorème qui se démontre. Ce qui nous permet de le faire, ce sont nos schémas de déduction sur les sélecteurs (SEL1–SEL3). Les mathématiques usuelles n'utilisent pas formellement de sélecteurs et doivent admettre cet énoncé comme axiome.

- F30. $\vdash$ Si G est un graphe, alors
 $\exists F (F \text{ est une fonction et } F \subset G \text{ et } \text{Dom} F = \mathbf{Pr}_1 G)$.

L'idée de la démonstration est illustrée par la figure 15.4. Celle-ci explique aussi le nom « axiome de choix » donné à cette proposition. Étant donné un graphe G , on peut construire une fonction F telle que $F \subset G$ et $\text{Dom} F = \mathbf{Pr}_1 G$ en *choisissant*, pour chaque élément x de l'ensemble $\mathbf{Pr}_1 G$ un objet y unique tel que $(x, y) \in G$. Dans cet exemple on a choisi $y = 5$ pour $x = 1$; $y = 5$ pour $x = 2$; $y = 6$ (pas d'autre choix possible) pour $x = 3$.

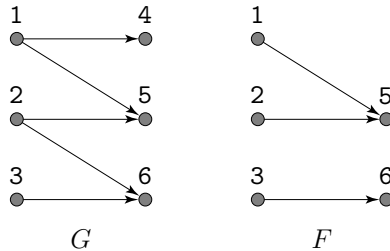


Figure 15.4 Une fonction F telle que $F \subset G$ et $\text{Dom} F = \mathbf{Pr}_1 G$.

Démonstration. Si x appartient à l'ensemble $\mathbf{Pr}_1 G$, il existe (E117) un objet y tel que $(x, y) \in G$. Or dans ce cas, le terme $\mathbf{S}y((x, y) \in G)$ représente

précisément un tel objet. La démonstration ci-dessous consiste donc à montrer que la fonction

$$F = (\lambda x \in \mathbf{Pr}_1 G). \text{Sy}((x, y) \in G)$$

possède les propriétés $F \subset G$ et $\text{Dom} F = \mathbf{Pr}_1 G$.

nil		
1°	G est un graphe	hyp
2°	$F = (\lambda x \in \mathbf{Pr}_1 G). \text{Sy}((x, y) \in G)$	hyp
3°	$\text{Dom} F = \mathbf{Pr}_1 G$	1°, F15
4°	$(x, y) \in F$	hyp
5°	$x \in \mathbf{Pr}_1 G$ et $y = \text{Sy}((x, y) \in G)$	4°, 2°, F17
6°	$x \in \mathbf{Pr}_1 G$	5°
7°	$\exists y((x, y) \in G)$	6°, E117
8°	$y = \text{Sy}((x, y) \in G)$	5°
9°	$(x, y) \in G$	7°, 8°, SEL1
10°	$F \subset G$	9°, E91
11°	$F \subset G$ et $\text{Dom} F = \mathbf{Pr}_1 G$	3°, 10°
12°	$\exists F(F \subset G \text{ et } \text{Dom} F = \mathbf{Pr}_1 G)$	S91.

Comme corollaire de F30, on a le théorème suivant.

F31. $\vdash$ Si G est une fonction, alors il existe une fonction F telle que $G \circ F = \text{Id}_{\text{Pr}G}$.

Il se démontre en prenant une fonction F telle que $F \subset G^{-1}$ et $\text{Dom} F = \mathbf{Pr}_1(G^{-1})$, autrement dit par élimination du quantificateur $\exists F$ de la proposition $\exists F(F \text{ est une fonction et } F \subset G^{-1} \text{ et } \text{Dom} F = \mathbf{Pr}_1(G^{-1}))$, qui est vraie d'après F30 (et E124) si G est un graphe.

Fonctions injectives. Une fonction F est dite *injective* si le graphe F^{-1} est une fonction. Cette notion est illustrée par la figure 15.5. Dire que F est une fonction injective équivaut à dire que F est une fonction et que pour deux éléments x, x' quelconques du domaine de F , on a

$$F @ x = F @ x' \Rightarrow x = x'.$$

Ceci est la manière la plus courante de caractériser une fonction injective.

F32. *Définition*

$$\vdash F \text{ est une fonction injective} \Leftrightarrow \left(\begin{array}{l} F \text{ est une fonction et} \\ F^{-1} \text{ est une fonction} \end{array} \right).$$

Théorèmes

F33. $\vdash F$ est une fonction injective $\Leftrightarrow [F \text{ est une fonction et } \forall x \forall x' ((x \in \text{Dom} F \text{ et } x' \in \text{Dom} F \text{ et } F @ x = F @ x') \Rightarrow x = x')]$.

F34. $\vdash \text{Id}_X$ est une fonction injective.

$\vdash \emptyset$ est une fonction injective.

- F35. $\vdash (F \text{ est une fonction injective}) \Rightarrow (F^{-1} \text{ est une fonction injective}).$
 F36. $\vdash$ Si F et G sont des fonctions injectives, alors
 $G \circ F$ est une fonction injective.
 F37. $\vdash F$ est une fonction injective $\Leftrightarrow$
 $(F \text{ est une fonction et } F^{-1} \circ F = \text{Id}_{\text{Dom}F}).$

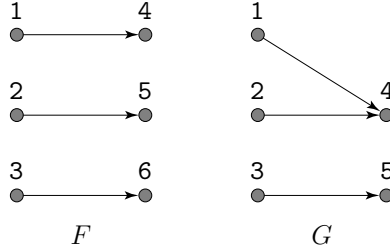


Figure 15.5 La fonction F est injective; le fonction G ne l'est pas.

Applications injectives, surjectives, bijectives. On dit que F est une *application injective* de A dans B si F est une application de A dans B (F8) et si F est en outre une fonction injective. Nous abrègerons par

$$F : A \xrightarrow{\text{inj}} B$$

la proposition « F est une application injective de A dans B ». On dit que F est une *application surjective* de A dans B , et nous écrivons $F : A \xrightarrow{\text{sur}} B$, si F est une application de A dans B et si $\text{Prt} F = B$. On dit enfin que F est une *application bijective* de A dans B , et nous écrivons $F : A \xrightarrow{\text{bij}} B$, si F est à la fois une application injective et une application surjective de A dans B .

F38. *Définitions*

- $\vdash (F : A \xrightarrow{\text{inj}} B) \Leftrightarrow (F : A \rightarrow B \text{ et } F \text{ est une fonction injective}).$
 $\vdash (F : A \xrightarrow{\text{sur}} B) \Leftrightarrow (F : A \rightarrow B \text{ et } \text{Prt} F = B).$
 $\vdash (F : A \xrightarrow{\text{bij}} B) \Leftrightarrow (F : A \xrightarrow{\text{inj}} B \text{ et } F : A \xrightarrow{\text{sur}} B).$

La « parenté logique » de ces trois notions apparaît dans le théorème suivant dans lequel elles sont exprimées au moyen de quantificateurs.

F39. $\vdash$ Si $F : A \rightarrow B$, alors

- $(F : A \xrightarrow{\text{inj}} B) \Leftrightarrow \forall y(y \in B \Rightarrow \text{H}x(x \in A \text{ et } y = F @ x));$
 $(F : A \xrightarrow{\text{sur}} B) \Leftrightarrow \forall y(y \in B \Rightarrow \exists x(x \in A \text{ et } y = F @ x));$
 $(F : A \xrightarrow{\text{bij}} B) \Leftrightarrow \forall y(y \in B \Rightarrow \exists! x(x \in A \text{ et } y = F @ x)).$

- F40. $\vdash (F : A \xrightarrow{\text{inj}} B \text{ et } G : B \xrightarrow{\text{inj}} C) \Rightarrow (G \circ F : A \xrightarrow{\text{inj}} C).$
 $\vdash (F : A \xrightarrow{\text{sur}} B \text{ et } G : B \xrightarrow{\text{sur}} C) \Rightarrow (G \circ F : A \xrightarrow{\text{sur}} C).$
 $\vdash (F : A \xrightarrow{\text{bij}} B \text{ et } G : B \xrightarrow{\text{bij}} C) \Rightarrow (G \circ F : A \xrightarrow{\text{bij}} C).$

- F41. $\vdash (F : A \xrightarrow{\text{bij}} B) \Rightarrow (F^{-1} : B \xrightarrow{\text{bij}} A \text{ et } F^{-1} \circ F = \text{Id}_A).$

16

Annexe

Démonstrations de théorie des ensembles

E19	nil		
	$x \in a \cup a \Leftrightarrow (x \in a \text{ ou } x \in a)$	E18	
	$\Leftrightarrow x \in a$	S52	
	$a \cup a = a$	E17.2	
E20			
	$x \in a \cup b \Leftrightarrow (x \in a \text{ ou } x \in b)$	E18	
	$\Leftrightarrow (x \in b \text{ ou } x \in a)$	S52	
	$\Leftrightarrow x \in b \cup a$	E18.	
E21			
	$x \in a \cup (b \cup c) \Leftrightarrow (x \in a \text{ ou } x \in (b \cup c))$	E18	
	$\Leftrightarrow (x \in a \text{ ou } (x \in b \text{ ou } x \in c))$	E18	
	$\Leftrightarrow ((x \in a \text{ ou } x \in b) \text{ ou } x \in c)$	S52	
	$\Leftrightarrow (x \in a \cup b \text{ ou } x \in c)$	E18	
	$\Leftrightarrow x \in (a \cup b) \cup c$	E18.	
E22			
	$x \in a \cup \emptyset \Leftrightarrow (x \in a \text{ ou } x \in \emptyset)$	E18	
	$\Leftrightarrow x \in a$	E2, S60.2.	
E23			
1°	$x \in a$	hyp	
2°	$x \in a \text{ ou } x \in b$	1°, ou-intro	
3°	$x \in a \cup b$	2°, E18	
4°	$a \subset a \cup b$	3°, E5.2	
E24			
	$(a \subset c \text{ et } b \subset c) \Leftrightarrow (\forall x(x \in a \Rightarrow x \in c) \text{ et } \forall x(x \in b \Rightarrow x \in c))$	E3	
	$\Leftrightarrow \forall x((x \in a \Rightarrow x \in c) \text{ et } (x \in b \Rightarrow x \in c))$	S74	
	$\Leftrightarrow \forall x((x \in a \text{ ou } x \in b) \Rightarrow x \in c)$	S58	
	$\Leftrightarrow \forall x(x \in a \cup b \Rightarrow x \in c)$	E18	
	$\Leftrightarrow a \cup b \subset c$	E3.	
E25			
	$a \subset b \Leftrightarrow (a \subset b \text{ et } b \subset b)$	E6, S60.1	
	$\Leftrightarrow (a \cup b \subset b)$	E24 (exempl)	

	$\begin{array}{ l} \Leftrightarrow (a \cup b \subset b \text{ et } b \subset a \cup b) \\ \Leftrightarrow (a \cup b = b) \end{array}$	E23, S60.1 E10.
E26	$\begin{array}{ l} 1^{\circ} \quad a \subset a' \text{ et } b \subset b' \\ 2^{\circ} \quad a \subset a' \\ 3^{\circ} \quad a' \subset a' \cup b' \\ 4^{\circ} \quad a \subset a' \cup b' \\ 5^{\circ} \quad b \subset a' \cup b' \\ 6^{\circ} \quad a \cup b \subset a' \cup b' \\ 7^{\circ} \quad (a \subset a' \text{ et } b \subset b') \Rightarrow a \cup b \subset a' \cup b' \end{array}$	hyp 1 ^o , et-élim E23 2 ^o , 3 ^o , E9 même manière 4 ^o , 5 ^o , E24 $\Rightarrow$ -intro.
E27	$\begin{array}{ l} a \cup b = \emptyset \Leftrightarrow a \cup b \subset \emptyset \\ \Leftrightarrow (a \subset \emptyset \text{ et } b \subset \emptyset) \\ \Leftrightarrow (a = \emptyset \text{ et } b = \emptyset) \end{array}$	E12 E24 E12.
E29	$\begin{array}{ l} x \in a \cap a \Leftrightarrow (x \in a \text{ et } x \in a) \\ \Leftrightarrow x \in a \\ a \cap a = a \end{array}$	E28 S52 E17.2
E30	$\begin{array}{ l} x \in a \cap b \Leftrightarrow (x \in a \text{ et } x \in b) \\ \Leftrightarrow (x \in b \text{ et } x \in a) \\ \Leftrightarrow x \in b \cap a \end{array}$	E28 S52 E28.
E31	$\begin{array}{ l} x \in a \cap (b \cap c) \Leftrightarrow (x \in a \text{ et } x \in (b \cap c)) \\ \Leftrightarrow (x \in a \text{ et } (x \in b \text{ et } x \in c)) \\ \Leftrightarrow ((x \in a \text{ et } x \in b) \text{ et } x \in c) \\ \Leftrightarrow (x \in a \cap b \text{ et } x \in c) \\ \Leftrightarrow x \in (a \cap b) \cap c \end{array}$	E28 E28 S52 E28 E28.
E32	$\begin{array}{ l} 1^{\circ} \quad x \in a \cap \emptyset \\ 2^{\circ} \quad x \in a \text{ et } x \in \emptyset \\ 3^{\circ} \quad x \in \emptyset \\ 4^{\circ} \quad \perp \\ 5^{\circ} \quad a \cap \emptyset = \emptyset \end{array}$	hyp 1 ^o , E28 2 ^o , et-élim 3 ^o , E2 4 ^o , E15.
E33	$\begin{array}{ l} 1^{\circ} \quad x \in a \cap b \\ 2^{\circ} \quad x \in a \text{ et } x \in b \\ 3^{\circ} \quad x \in a \\ 4^{\circ} \quad a \cap b \subset a \end{array}$	hyp 1 ^o , E28 2 ^o , et-élim 3 ^o , E5.2
E34	$\begin{array}{ l} (c \subset a \text{ et } c \subset b) \Leftrightarrow (\forall x(x \in c \Rightarrow x \in a) \text{ et } \forall x(x \in c \Rightarrow x \in b)) \\ \Leftrightarrow \forall x((x \in c \Rightarrow x \in a) \text{ et } (x \in c \Rightarrow x \in b)) \\ \Leftrightarrow \forall x(x \in c \Rightarrow (x \in a \text{ et } x \in b)) \\ \Leftrightarrow \forall x(x \in c \Rightarrow x \in a \cap b) \\ \Leftrightarrow c \subset a \cap b \end{array}$	E3 S74 S58 E28 E3.
E35	$a \subset b \Leftrightarrow (a \subset a \text{ et } a \subset b)$	E6, S60.1

	$\Leftrightarrow (a \subset a \cap b)$	E34 (exempl)
	$\Leftrightarrow (a \subset a \cap b \text{ et } a \cap b \subset a)$	E33, S60.1
	$\Leftrightarrow (a \cap b = a)$	E10.
E36		
1°	$a \subset a' \text{ et } b \subset b'$	hyp
2°	$a \cap b \subset a$	E33
3°	$a \subset a'$	1°, et-élim
4°	$a \cap b \subset a'$	2°, 3°, E9
5°	$a \cap b \subset b'$	même manière
6°	$a \cap b \subset a' \cap b'$	4°, 5°, E34
7°	$(a \subset a' \text{ et } b \subset b') \Rightarrow a \cap b \subset a' \cap b'$	6°, $\Rightarrow$ -intro.
E37		
	$x \in a \cap (b \cup c) \Leftrightarrow x \in a \text{ et } x \in b \cup c$	E28
	$\Leftrightarrow x \in a \text{ et } (x \in b \text{ ou } x \in c)$	E18
	$\Leftrightarrow (x \in a \text{ et } x \in b) \text{ ou } (x \in a \text{ et } x \in c)$	S52.11
	$\Leftrightarrow x \in a \cap b \text{ ou } x \in a \cap c$	E28
	$\Leftrightarrow x \in (a \cup b) \cup (a \cap c)$	E18.
E38		
	$x \in a \cup (b \cap c) \Leftrightarrow x \in a \text{ ou } x \in b \cap c$	E18
	$\Leftrightarrow x \in a \text{ ou } (x \in b \text{ et } x \in c)$	E28
	$\Leftrightarrow (x \in a \text{ ou } x \in b) \text{ et } (x \in a \text{ ou } x \in c)$	S52.12
	$\Leftrightarrow x \in a \cup b \text{ et } x \in a \cup c$	E18
	$\Leftrightarrow x \in (a \cup b) \cap (a \cup c)$	E28.
E41		
1°	$x \in \mathbb{C}_a(b) \Leftrightarrow (x \in a \text{ et } x \notin b)$	E40
2°	$\Rightarrow x \in a$	S28.
3°	$\mathbb{C}_a(b) \subset a$	1°-2°, S49, E5.
E42		
	$\mathbb{C}_a(b) = a \Leftrightarrow (\mathbb{C}_a(b) \subset a \text{ et } a \subset \mathbb{C}_a(b))$	E10
	$\Leftrightarrow a \subset \mathbb{C}_a(b)$	E41, S60.1
	$\Leftrightarrow \forall x(x \in a \Rightarrow x \in \mathbb{C}_a(b))$	E3
	$\Leftrightarrow \forall x(x \in a \Rightarrow (x \in a \text{ et } x \notin b))$	E40
	$\Leftrightarrow \forall x((x \in a \Rightarrow x \in a) \text{ et } (x \in a \Rightarrow x \notin b))$	S58
	$\Leftrightarrow \forall x(x \in a \Rightarrow x \notin b)$	S19, S60.1
	$\Leftrightarrow a \cap b = \emptyset$	E39.
E43		
	$\mathbb{C}_a(b) = \emptyset \Leftrightarrow \forall x(x \notin \mathbb{C}_a(b))$	E13
	$\Leftrightarrow \forall x \text{ non}(x \in a \text{ et } x \notin b)$	E40
	$\Leftrightarrow \forall x(x \notin a \text{ ou } x \in b)$	De Morgan
	$\Leftrightarrow \forall x(x \in a \Rightarrow x \in b)$	S53
	$\Leftrightarrow a \subset b$	E3.
E44		
	$x \in \mathbb{C}_a(\emptyset) \Leftrightarrow (x \in a \text{ et } x \notin \emptyset)$	E40
	$\Leftrightarrow x \in a$	E2, S60.1
	$\mathbb{C}_a(\emptyset) = a$	E5.
1°	$x \in \mathbb{C}_a(a)$	hyp
2°	$x \in a \text{ et } x \notin a$	1°, E40

3°		$\text{non}(x \in a \text{ et } x \notin a)$	S17
4°		$\perp$	2°, 3°, $\perp$ -intro
5°		$\mathbb{C}_a(a) = \emptyset$	4°, E15.
E45			
1°		$x \in \mathbb{C}_{\emptyset}(b)$	hyp
2°		$x \in \emptyset \text{ et } x \notin b$	1°, E40
3°		$x \in \emptyset$	2°, et-élim
4°		$\perp$	3°, E2, $\perp$ -intro
5°		$\mathbb{C}_{\emptyset}(b) = \emptyset$	4°, E15.
E46			
		$x \in \mathbb{C}_X(a \cup b) \Leftrightarrow (x \in X \text{ et } x \notin a \cup b)$	E40
		$\Leftrightarrow x \in X \text{ et } (x \notin a \text{ et } x \notin b)$	E18, De Morgan
		$\Leftrightarrow (x \in X \text{ et } x \notin a) \text{ et } (x \in X \text{ et } x \notin b)$	S52.6, S52.8, S52.10
		$\Leftrightarrow x \in \mathbb{C}_X(a) \text{ et } x \in \mathbb{C}_X(b)$	E40
		$\Leftrightarrow x \in \mathbb{C}_X(a) \cap \mathbb{C}_X(b)$	E28.
E47			
		$x \in \mathbb{C}_X(a \cap b) \Leftrightarrow (x \in X \text{ et } x \notin a \cap b)$	E40
		$\Leftrightarrow x \in X \text{ et } (x \notin a \text{ ou } x \notin b)$	E28, De Morgan
		$\Leftrightarrow (x \in X \text{ et } x \notin a) \text{ ou } (x \in X \text{ et } x \notin b)$	S52.11
		$\Leftrightarrow x \in \mathbb{C}_X(a) \text{ ou } x \in \mathbb{C}_X(b)$	E40
		$\Leftrightarrow x \in \mathbb{C}_X(a) \cup \mathbb{C}_X(b)$	E18.
E48.1			
1°		$a \subset b$	hyp
2°		$x \in \mathbb{C}_X(b)$	hyp
3°		$x \in X \text{ et } x \notin b$	2°, E40
4°		$x \notin b$	3°, et-élim
5°		$x \notin a$	1°, 4°, E4.3
6°		$x \in X \text{ et } x \notin a$	3°, 5°
7°		$x \in \mathbb{C}_X(a)$	6°, E40
8°		$\mathbb{C}_X(b) \subset \mathbb{C}_X(a)$	7°, E5
9°		$a \subset b \Rightarrow (\mathbb{C}_X(b) \subset \mathbb{C}_X(a))$	8°, $\Rightarrow$ -intro.
E48.2			
1°		$a \subset X$	hyp
2°		$b \subset X$	hyp
3°		$\mathbb{C}_X(b) \subset \mathbb{C}_X(a)$	hyp
4°		$x \in a$	hyp
5°		$x \notin b$	hyp
6°		$x \in X$	4°, 1°, E4
7°		$x \in \mathbb{C}_X(b)$	5°, 6°, E40
8°		$x \in \mathbb{C}_X(a)$	7°, 3°, E4
9°		$x \notin a$	8°, E40, et-élim
10°		$\perp$	4°, 9°
11°		$x \in b$	10°, red abs K
12°		$a \subset b$	11°, E5
13°		$\mathbb{C}_X(b) \subset \mathbb{C}_X(a) \Rightarrow a \subset b$	12°, $\Rightarrow$ -intro
14°		$a \subset b \Rightarrow \mathbb{C}_X(b) \subset \mathbb{C}_X(a)$	E48 (premier)
15°		$a \subset b \Leftrightarrow \mathbb{C}_X(b) \subset \mathbb{C}_X(a)$	13°, 14°, $\Leftrightarrow$ -intro.

E49	$ \begin{aligned} & x \in \mathbb{C}_X(\mathbb{C}_X(a)) \Leftrightarrow x \in X \text{ et } x \notin \mathbb{C}_X(a) \\ & \Leftrightarrow (x \in X \text{ et } (x \notin X \text{ ou } x \in a)) \\ & \Leftrightarrow ((x \in X \text{ et } x \notin X) \text{ ou } (x \in X \text{ et } x \in a)) \\ & \Leftrightarrow (x \in X \text{ et } x \in a) \\ & \Leftrightarrow x \in X \cap a \end{aligned} $	E40 E40, De Morgan S52.11 S52.15 E28.
E50	$ \begin{aligned} & \mathbb{C}_X(\mathbb{C}_X(a)) = a \Leftrightarrow \forall x(x \in \mathbb{C}_X(\mathbb{C}_X(a)) \Leftrightarrow x \in a) \\ & \Leftrightarrow \forall x(x \in X \cap a \Leftrightarrow x \in a) \\ & \Leftrightarrow \forall x(x \in a \cap X \Leftrightarrow x \in a) \\ & \Leftrightarrow a \cap X = a \\ & \Leftrightarrow a \subset X \end{aligned} $	E16 E49 E30 E16 E35.
E52	$ \begin{aligned} 1^\circ & a = a \\ 2^\circ & a = a \text{ ou } a = b \\ 3^\circ & (a = a \text{ ou } a = b) \Leftrightarrow a \in \{a, b\} \\ 4^\circ & a \in \{a, b\} \end{aligned} $	S82 1°, ou-intro E51 (exemplifié) 2°, 3°, S39.
E53	$ \begin{aligned} 1^\circ & a \in \{a, b\} \\ 2^\circ & \exists x(x \in \{a, b\}) \\ 2^\circ & \{a, b\} \neq \emptyset \end{aligned} $	E52 1°, S63 2°, E13.
E54	$ \begin{aligned} 1^\circ & a \in \{a, b\} \\ 2^\circ & \exists x(a \in x) \\ 3^\circ & \forall a \exists x(a \in x) \end{aligned} $	E52 1°, S63 2°, S62.
E55	$ \begin{aligned} & x \in \{a, b\} \Leftrightarrow (x = a \text{ ou } x = b) \\ & \Leftrightarrow (x = b \text{ ou } x = a) \\ & \Leftrightarrow x \in \{b, a\} \end{aligned} $	E51 S52 E51.
E56.1	$ \begin{aligned} & \exists x(x \in \{a, b\} \text{ et } R) \\ & \Leftrightarrow \exists x((x = a \text{ ou } x = b) \text{ et } R) \\ & \Leftrightarrow \exists x((x = a \text{ et } R) \text{ ou } (x = b \text{ et } R)) \\ & \Leftrightarrow \exists x(x = a \text{ et } R) \text{ ou } \exists x(x = b \text{ et } R) \\ & \Leftrightarrow ((a x)R \text{ ou } (b x)R) \end{aligned} $	E51 S52.11 S74 S92.
E56.2	$ \begin{aligned} & \forall x(x \in \{a, b\} \Rightarrow R) \\ & \Leftrightarrow \forall x((x = a \text{ ou } x = b) \Rightarrow R) \\ & \Leftrightarrow \forall x((x = a \Rightarrow R) \text{ et } (x = b \Rightarrow R)) \\ & \Leftrightarrow \forall x(x = a \Rightarrow R) \text{ et } \forall x(x = b \Rightarrow R) \\ & \Leftrightarrow ((a x)R \text{ et } (b x)R) \end{aligned} $	E51 S58 S74 S92.
E57	$ \begin{aligned} & \{a, b\} \subset X \Leftrightarrow \forall x(x \in \{a, b\} \Rightarrow x \in X) \\ & \Leftrightarrow (a \in X \text{ et } b \in X) \end{aligned} $	E3 E56.
E58	$ \begin{aligned} 1^\circ & b = c \\ 2^\circ & \{a, b\} = \{a, c\} \end{aligned} $	hyp 1°, S88

3°	$\{a, b\} = \{a, c\}$	hyp
4°	$b \neq c$	hyp
5°	$b \in \{a, b\}$	E52
6°	$b \in \{a, c\}$	5°, 3°, S83
7°	$b = a$ ou $b = c$	6°, E51
8°	$b = a$	4°, 7°, S18
9°	$c \in \{a, c\}$	E52
10°	$c \in \{a, b\}$	9°, 3°, S83
11°	$c = a$ ou $c = b$	10°, E51
12°	$c = a$	4°, 11°, S18
13°	$b = c$	8°, 12°, S85, S84
14°	$\perp$	$\perp$ -intro, 4°, 13°
15°	$b = c$	14°, red abs K
16°	$\{a, b\} = \{a, c\} \Leftrightarrow b = c$	2°, 15°, S40.
E60	$x \in \{a\} \Leftrightarrow x \in \{a, a\}$ $\Leftrightarrow x = a$ ou $x = a$ $\Leftrightarrow x = a$	E59, S86 E51 S52.
E61	1° $a = a$ 2° $a = a \Leftrightarrow a \in \{a\}$ 3° $a \in \{a\}$	S82 E60 (exemplifié) 1°, 2°, S39.
E62	1° $\{a\} = \{a, a\}$ 2° $\{a, a\} \neq \emptyset$ 3° $\{a\} \neq \emptyset$	E59 E53 1°, 2°, S83.
E63	$\{a\} \subset X \Leftrightarrow \{a, a\} \subset X$ $\Leftrightarrow (a \in X \text{ et } a \in X)$ $\Leftrightarrow a \in X$	E59 E57 S52.
E64	$\{a\} \cap X \neq \emptyset \Leftrightarrow \exists x(x \in \{a\} \cap X)$ $\Leftrightarrow \exists x(x \in \{a\} \text{ et } x \in X)$ $\Leftrightarrow \exists x(x = a \text{ et } x \in X)$ $\Leftrightarrow a \in X$	E13 E28 E60 S92.
E65	$\{a\} \subset \{b\} \Leftrightarrow a \in \{b\}$ $\Leftrightarrow a = b$	E63 E60.
E66	$\{a\} = \{b\} \Leftrightarrow (\{a\} \subset \{b\} \text{ et } \{b\} \subset \{a\})$ $\Leftrightarrow (a = b \text{ et } b = a)$ $\Leftrightarrow a = b$	E10 E65 S84, S52.
E67	$\{a\} = \{a, b\} \Leftrightarrow \{a, a\} = \{a, b\}$ $\Leftrightarrow a = b$	E59 E58, exempl.

E68

$x \in \{a, b\} \Leftrightarrow (x = a \text{ ou } x = b)$	E51
$\Leftrightarrow (x \in \{a\} \text{ ou } x \in \{b\})$	E60
$\Leftrightarrow x \in \{a\} \cup \{b\}$	E18.

E69

$\{a\} \cap \{b\} \neq \emptyset \Leftrightarrow a \in \{b\}$	E64
$\Leftrightarrow a = b$	E60.

E70

1°	$X = \emptyset \text{ ou } X = \{a\}$	hyp
2°	$X = \emptyset$	hyp
3°	$\emptyset \subset \{a\}$	E7
4°	$X \subset \{a\}$	2°, 3°, S83
5°	$X = \{a\}$	hyp
6°	$X \subset \{a\}$	5°, E10
7°	$X \subset \{a\}$	1°, 4°, 6°, disj cas
8°	$X \subset \{a\}$	hyp
9°	$X \neq \emptyset$	hyp
10°	$\exists x(x \in X)$	9°, E13
11°	$x \in X$	hyp
12°	$x \in \{a\}$	8°, 11°, E4
13°	$x = a$	12°, E60
14°	$a \in X$	11°, 13°, S83
15°	$a \in X$	10°, 14°, S64
16°	$\{a\} \subset X$	15°, E63
17°	$X = \{a\}$	8°, 16°, E10
18°	$X \neq \emptyset \Rightarrow X = \{a\}$	17°, $\Rightarrow$ -intro
19°	$X = \emptyset \text{ ou } X = \{a\}$	18°, S54
20°	$X \subset \{a\} \Leftrightarrow (X = \emptyset \text{ ou } X = \{a\})$	7°, 19°, S40.

E72

1°	$a \subset a$	E6
2°	$a \subset a \Leftrightarrow a \in \mathcal{P}(a)$	E71
3°	$a \in \mathcal{P}(a)$	1°, 2°, S39.
1°	$\emptyset \subset a$	E7
2°	$\emptyset \in \mathcal{P}(a)$	1°, E71.

E73

1°	$a \in \mathcal{P}(a)$	E72
2°	$\exists x(x \in \mathcal{P}(a))$	1°, S63
3°	$\mathcal{P}(a) \neq \emptyset$	2°, E13.

E74

$x \in \mathcal{P}(\emptyset) \Leftrightarrow x \subset \emptyset$	E71
$\Leftrightarrow x = \emptyset$	E12
$\Leftrightarrow x \in \{\emptyset\}$	E60
$\mathcal{P}(\emptyset) = \{\emptyset\}$	E5.2.

E75

$a \in \mathcal{P}(\{x\}) \Leftrightarrow a \subset \{x\}$	E71
$\Leftrightarrow (a = \emptyset \text{ ou } a = \{x\})$	E70
$\Leftrightarrow a \in \{\emptyset, \{x\}\}$	E51.

E77

1°	$a \subset b$	hyp
2°	$x \in \mathcal{P}(a)$	hyp
3°	$x \subset a$	2°, E71
4°	$x \subset b$	1°, 3°, E9
5°	$x \in \mathcal{P}(b)$	4°, E71
6°	$\mathcal{P}(a) \subset \mathcal{P}(b)$	5°, E5.2
7°	$\mathcal{P}(a) \subset \mathcal{P}(b)$	hyp
8°	$a \in \mathcal{P}(a)$	E72
9°	$a \in \mathcal{P}(b)$	7°, 8°, E4
10°	$a \subset b$	9°, E71
11°	$a \subset b \Leftrightarrow \mathcal{P}(a) \subset \mathcal{P}(b)$	6°, 10°, S40.

E80

1°	$(a, b) = (b, a)$	hyp
2°	$a = b$ et $b = a$	1°, E79
3°	$a = b$	2°
4°	$(a, b) = (b, a) \Rightarrow a = b$	3°, $\Rightarrow$ -intro.

E84

1°	(x, y) est un couple	E82
2°	$(x, y) = (\mathbf{pr}_1(x, y), \mathbf{pr}_2(x, y))$	1°, E83
3°	$x = \mathbf{pr}_1(x, y)$ et $y = \mathbf{pr}_2(x, y)$	2°, E79.

E85

z est un couple $\Leftrightarrow \exists a \exists b (z = (a, b))$	E81
$\Leftrightarrow \exists a \exists b (z = (a, b) \text{ et } a = \mathbf{pr}_1(a, b) \text{ et } b = \mathbf{pr}_2(a, b))$	E84, S60.1
$\Leftrightarrow \exists a \exists b (z = (a, b) \text{ et } a = \mathbf{pr}_1 z \text{ et } b = \mathbf{pr}_2 z)$	S87.1
$\Leftrightarrow \exists a (z = (a, \mathbf{pr}_2 z) \text{ et } a = \mathbf{pr}_1 z)$	S92
$\Leftrightarrow z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	S92.

E88

1°	G est un graphe	hyp
2°	$\forall z (z \in G \Rightarrow z \text{ est un couple})$	1°, E86
3°	$z \in G \Rightarrow z \text{ est un couple}$	2°, S65
4°	$z \in G \Leftrightarrow (z \in G \text{ et } z \text{ est un couple})$	3°, S60.3
5°	$G \text{ est un graphe} \Rightarrow (z \in G \Leftrightarrow (z \text{ est un couple et } z \in G))$	4°, $\Rightarrow$ -intro.

E91.1

G et G' sont des graphes	hyp
$G \subset G' \Leftrightarrow \forall z (z \in G \Rightarrow z \in G')$	E3
$\Leftrightarrow \forall z ((z \text{ est un couple et } z \in G) \Rightarrow z \in G')$	E88
$\Leftrightarrow \forall z ((\exists x \exists y (z = (x, y)) \text{ et } z \in G) \Rightarrow z \in G')$	E81
$\Leftrightarrow \forall z (\exists x \exists y (z = (x, y) \text{ et } z \in G) \Rightarrow z \in G')$	S81
$\Leftrightarrow \forall z (\text{non} \exists x \exists y (z = (x, y) \text{ et } z \in G) \text{ ou } z \in G')$	S53
$\Leftrightarrow \forall z (\forall x \forall y (z \neq (x, y) \text{ ou } z \notin G) \text{ ou } z \in G')$	S71, Morgan
$\Leftrightarrow \forall z \forall x \forall y (z \neq (x, y) \text{ ou } z \notin G \text{ ou } z \in G')$	S81
$\Leftrightarrow \forall x \forall y \forall z (z = (x, y) \Rightarrow (z \in G \Rightarrow z \in G'))$	S72, S53
$\Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G')$	S92
$G \text{ et } G' \text{ sont des graphes} \Rightarrow (G \subset G' \Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G'))$	

E91.1 Autre démonstration

1°	G et G' sont des graphes	hyp
2°	$G \subset G'$	hyp
3°	$(x, y) \in G \Rightarrow (x, y) \in G'$	2°, E4
4°	$\forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G')$	3°, S62
5°	$\forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G')$	hyp
6°	$z \in G$	hyp
7°	z est un couple	6°, 1°, E86
8°	$z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	7°, E83
9°	$(\mathbf{pr}_1 z, \mathbf{pr}_2 z) \in G$	6°, 8°, S83
10°	$(\mathbf{pr}_1 z, \mathbf{pr}_2 z) \in G \Rightarrow (\mathbf{pr}_1 z, \mathbf{pr}_2 z) \in G'$	5°, S77
11°	$(\mathbf{pr}_1 z, \mathbf{pr}_2 z) \in G'$	9°, 10°
12°	$z \in G'$	11°, 8°
13°	$G \subset G'$	12°, E5.2
14°	$G \subset G' \Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G')$	4°, 13°, S40
15°	G et G' sont des graphes $\Rightarrow$ $(G \subset G' \Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G'))$.	$\Rightarrow$ -intro.

E91.2

	G et G' sont des graphes	hyp
	$G = G' \Leftrightarrow (G \subset G' \text{ et } G' \subset G)$	E10
	$\Leftrightarrow \forall x \forall y ((x, y) \in G \Rightarrow (x, y) \in G') \text{ et } \forall x \forall y ((x, y) \in G' \Rightarrow (x, y) \in G)$	E91.1
	$\Leftrightarrow \forall x \forall y [((x, y) \in G \Rightarrow (x, y) \in G') \text{ et } ((x, y) \in G' \Rightarrow (x, y) \in G)]$	S74 (2 fois)
	$\Leftrightarrow \forall x \forall y ((x, y) \in G \Leftrightarrow (x, y) \in G')$	S41.

E96

Soient $E, E', T(x)$ des termes, x une variable. Soient u, x', y des variables satisfaisant aux conditions de E95 et ne figurant librement ni dans E' , ni dans Γ . On désigne par

- R la proposition $\exists x' (x' \in E \text{ et } y \in T(x'))$
R' la proposition $\exists x' (x' \in E' \text{ et } y \in T(x'))$

Γ

1°	$E = E'$	base
2°	$\forall y (y \in u \Leftrightarrow R) \Leftrightarrow \forall y (y \in u \Leftrightarrow R')$	1°, S86
3°	$Su \forall y (y \in u \Leftrightarrow R) = Su \forall y (y \in u \Leftrightarrow R')$	2°, SEL2
4°	$\bigcup_{x \in E} T(x) = \bigcup_{x \in E'} T(x)$	3°, E94

Le deuxième schéma E96 se démontre de manière semblable.

E97

Soient $E, T(x)$ des termes, x une variable. Soient u, x', y des variables satisfaisant aux conditions de E95. Si z est une variable réversiblement substituable à x dans $T(x)$, alors $T(x')$ est identique à $(x'|x)T(x)$ et aussi à $(x'|z)T(z)$, de sorte que l'on peut écrire:

1°	$\bigcup_{x \in E} T(x) = Su \forall y (y \in u \Leftrightarrow \exists x' (x' \in E \text{ et } y \in T(x')))$	E94
	$= \bigcup_{z \in E} T(z)$	E94.

E98

Dans les démonstrations de E98 à E105, on suppose que la variable y ne figure pas librement dans $T(x)$.

nil	
$y \in \bigcup_{x \in \{a,b\}} T(x)$	
$\Leftrightarrow \exists x(x \in \{a,b\} \text{ et } y \in T(x))$	E95
$\Leftrightarrow y \in T(a) \text{ ou } y \in T(b)$	E56
$\Leftrightarrow y \in T(a) \cup T(b)$	E18
$\bigcup_{x \in \{a,b\}} T(x) = T(a) \cup T(b)$	E17.
$\bigcup_{x \in \{a\}} T(x) = \bigcup_{x \in \{a,a\}} T(x)$	E59, S88
$= T(a) \cup T(a)$	E98 (premier)
$= T(a)$	E19.

E99 Il suffit d'appliquer E98 en prenant pour $T(x)$ le terme x .

E100

1°	$y \in \bigcup_{x \in \emptyset} T(x)$	hyp
2°	$\exists x(x \in \emptyset \text{ et } y \in T(x))$	1°, E95
3°	$\exists x(x \in \emptyset)$	2°, S75
4°	$\text{non} \exists x(x \in \emptyset)$	E2
5°	$\perp$	3°, 4°, $\perp$ -intro
6°	$\bigcup_{x \in \emptyset} T(x) = \emptyset$	5°, E15.

E101

1°	$A \subset B$	hyp
2°	$y \in \bigcup_{x \in A} T(x)$	hyp
3°	$\exists x(x \in A \text{ et } y \in T(x))$	2°, E95
4°	$x \in A \text{ et } y \in T(x)$	hyp
5°	$x \in B \text{ et } y \in T(x)$	1°, 4° ($x \in A$), E4
6°	$\exists x(x \in B \text{ et } y \in T(x))$	5°, S66
7°	$\exists x(x \in B \text{ et } y \in T(x))$	3°, 6°, S64
8°	$y \in \bigcup_{x \in B} T(x)$	7°, E95
9°	$\bigcup_{x \in A} T(x) \subset \bigcup_{x \in B} T(x)$	8°, E5
10°	$A \subset B \Rightarrow \bigcup_{x \in A} T(x) \subset \bigcup_{x \in B} T(x)$	$\Rightarrow$ -intro.

E102

$y \in \bigcup_{x \in A \cup B} T(x) \Leftrightarrow \exists x(x \in A \cup B \text{ et } y \in T(x))$	E95
--	-----

$\Leftrightarrow \exists x((x \in A \text{ ou } x \in B) \text{ et } y \in T(x))$	E18
$\Leftrightarrow \exists x((x \in A \text{ et } y \in T(x)) \text{ ou } (x \in B \text{ et } y \in T(x)))$	S52.11
$\Leftrightarrow \exists x(x \in A \text{ et } y \in T(x)) \text{ ou } \exists x(x \in B \text{ et } y \in T(x))$	S74
$\Leftrightarrow y \in \bigcup_{x \in A} T(x) \text{ ou } y \in \bigcup_{x \in B} T(x)$	E95
$\Leftrightarrow y \in \bigcup_{x \in A} T(x) \cup \bigcup_{x \in B} T(x)$	E18.

E103

1°	$x \in A$	hyp
2°	$y \in T(x)$	hyp
3°	$x \in A \text{ et } y \in T(x)$	1°, 2°, et-intro
4°	$\exists x(x \in A \text{ et } y \in T(x))$	3°, S66
5°	$y \in \bigcup_{x \in A} T(x)$	4°, E95
6°	$T(x) \subset \bigcup_{x \in A} T(x)$	5°, E5
7°	$x \in A \Rightarrow T(x) \subset \bigcup_{x \in A} T(x)$	6°, $\Rightarrow$ -intro.

E104

1°	$x \in A \Rightarrow \{T(x)\} \subset \bigcup_{x \in A} \{T(x)\}$	E103
	$\Leftrightarrow T(x) \in \bigcup_{x \in A} \{T(x)\}$	E63

E105

$\bigcup_{x \in A} T(x) \subset Y \Leftrightarrow \forall y(y \in \bigcup_{x \in A} T(x) \Rightarrow y \in Y)$	E3
$\Leftrightarrow \forall y(\exists x(x \in A \text{ et } y \in T(x)) \Rightarrow y \in Y)$	E95
$\Leftrightarrow \forall y(\text{non} \exists x(x \in A \text{ et } y \in T(x)) \text{ ou } y \in Y)$	S53
$\Leftrightarrow \forall y \forall x(\text{non}(x \in A \text{ et } y \in T(x)) \text{ ou } y \in Y)$	S71, S81
$\Leftrightarrow \forall x \forall y(x \notin A \text{ ou } y \notin T(x) \text{ ou } y \in Y)$	S72, De Morgan
$\Leftrightarrow \forall x(x \notin A \text{ ou } \forall y(y \in T(x) \Rightarrow y \in Y))$	S81, S53
$\Leftrightarrow \forall x(x \in A \Rightarrow T(x) \subset Y)$	E3, S53.

E105 Autre démonstration

1°	$\bigcup_{x \in A} T(x) \subset Y$	hyp
2°	$x \in A \Rightarrow T(x) \subset \bigcup_{x \in A} T(x)$	E103
3°	$x \in A \Rightarrow T(x) \subset Y$	1°, 2°, E9, logique prop
4°	$\forall x(x \in A \Rightarrow T(x) \subset Y)$	3°, S62
5°	$\forall x(x \in A \Rightarrow T(x) \subset Y)$	hyp
6°	$y \in \bigcup_{x \in A} T(x)$	hyp
7°	$\exists x(x \in A \text{ et } y \in T(x))$	6°, E95
8°	$x \in A \text{ et } y \in T(x)$	hyp
9°	$T(x) \subset Y$	8° ($x \in A$), 5°
10°	$y \in Y$	8°, 9°, E4

11°	$y \in Y$	7°, 10°, S64
12°	$\bigcup_{x \in A} \mathsf{T}(x) \subset Y$	11°, E5
13°	$\bigcup_{x \in A} \mathsf{T}(x) \subset Y \Leftrightarrow \forall x(x \in A \Rightarrow \mathsf{T}(x) \subset Y)$	4°, 12°, S40.

E107

$z \in X \times Y \Leftrightarrow z \in \bigcup_{x \in X} \left(\bigcup_{y \in Y} \{(x, y)\} \right)$	E106, S86
$\Leftrightarrow \exists x(x \in X \text{ et } z \in \bigcup_{y \in Y} \{(x, y)\})$	E95
$\Leftrightarrow \exists x(x \in X \text{ et } \exists y(y \in Y \text{ et } z \in \{(x, y)\}))$	E95
$\Leftrightarrow \exists x \exists y(x \in X \text{ et } y \in Y \text{ et } z = (x, y))$	S81, E60
$\Leftrightarrow \exists x \exists y(x \in X \text{ et } y \in Y \text{ et } z = (x, y) \text{ et } x = \mathbf{pr}_1(x, y) \text{ et } y = \mathbf{pr}_2(x, y))$	E84, S60.1
$\Leftrightarrow \exists x \exists y(x \in X \text{ et } y \in Y \text{ et } z = (x, y) \text{ et } x = \mathbf{pr}_1 z \text{ et } y = \mathbf{pr}_2 z)$	S87.1
$\Leftrightarrow \exists x(x \in X \text{ et } \mathbf{pr}_2 z \in Y \text{ et } z = (x, \mathbf{pr}_2 z) \text{ et } x = \mathbf{pr}_1 z)$	S92
$\Leftrightarrow \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y \text{ et } z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	S92
$\Leftrightarrow z \text{ est un couple et } \mathbf{pr}_1 z \in X \text{ et } \mathbf{pr}_2 z \in Y$	E85.

E108

1°	$\forall z(z \in X \times Y \Rightarrow z \text{ est un couple})$	E107
2°	$X \times Y \text{ est un graphe}$	1°, E86.

E109

$(x, y) \in X \times Y$	
$\Leftrightarrow ((x, y) \text{ est un couple et } \mathbf{pr}_1(x, y) \in X \text{ et } \mathbf{pr}_2(x, y) \in Y)$	E107
$\Leftrightarrow \mathbf{pr}_1(x, y) \in X \text{ et } \mathbf{pr}_2(x, y) \in Y$	E82, S60.1
$\Leftrightarrow x \in X \text{ et } y \in Y$	E84, S86.

E110

$X \times Y = \emptyset \Leftrightarrow \forall x \forall y((x, y) \notin X \times Y)$	E108, E92
$\Leftrightarrow \forall x \forall y(x \notin X \text{ ou } y \notin Y)$	E109, De Morgan
$\Leftrightarrow \forall x(x \notin X) \text{ ou } \forall y(y \notin Y)$	S81 (2 fois)
$\Leftrightarrow X = \emptyset \text{ ou } Y = \emptyset$	E13.

E111

1°	$X \subset X' \text{ et } Y \subset Y'$	hyp
2°	$(x, y) \in X \times Y$	hyp
3°	$x \in X \text{ et } y \in Y$	2°, E109
4°	$x \in X' \text{ et } y \in Y'$	1°, 3°, E4
5°	$(x, y) \in X' \times Y'$	4°, E109
6°	$(x, y) \in X \times Y \Rightarrow (x, y) \in X' \times Y'$	5°, $\Rightarrow$ -intro
7°	$\forall x \forall y((x, y) \in X \times Y \Rightarrow (x, y) \in X' \times Y')$	6°, S62
lignes 6° et 7° omises en général		
8°	$X \times Y \subset X' \times Y'$	7°, E108, E91
9°	$(X \subset X' \text{ et } Y \subset Y') \Rightarrow X \times Y \subset X' \times Y'$	8°, $\Rightarrow$ -intro.

E112

1°	$X \neq \emptyset \text{ et } Y \neq \emptyset$	hyp
2°	$X \times Y \subset X' \times Y'$	hyp

3°			$x \in X$	hyp
4°			$\exists y(y \in Y)$	1°, E13
5°			$y \in Y$	hyp
6°			$(x, y) \in X \times Y$	3°, 5°, E109
7°			$(x, y) \in X' \times Y'$	6°, 2°, E4
8°			$x \in X'$	7°, E109
9°			$x \in X'$	4°, 8°, S64
10°			$X \subset X'$	9°, E5
			$\vdots$	lignes analogues à 3°–9°
18°			$Y \subset Y'$	
19°			$X \subset X'$ et $Y \subset Y'$	10°, 18°
20°			$(X \times Y \subset X' \times Y') \Rightarrow (X \subset X' \text{ et } Y \subset Y')$	19°, $\Rightarrow$ -intro
21°			$(X \subset X' \text{ et } Y \subset Y') \Leftrightarrow X \times Y \subset X' \times Y'$	20°, E111, $\Leftrightarrow$ -intro
22°			...	20°, $\Rightarrow$ -intro.

E113

$(x, y) \in (X \times Y) \cap (X' \times Y')$	
$\Leftrightarrow (x, y) \in X \times Y \text{ et } (x, y) \in X' \times Y'$	E28
$\Leftrightarrow x \in X \text{ et } y \in Y \text{ et } x \in X' \text{ et } y \in Y'$	E109
$\Leftrightarrow x \in X \text{ et } x \in X' \text{ et } y \in Y \text{ et } y \in Y'$	S52
$\Leftrightarrow x \in X \cap X' \text{ et } y \in Y \cap Y'$	E28
$\Leftrightarrow (x, y) \in (X \cap X') \times (Y \cap Y')$	E109
$(X \times Y) \cap (X' \times Y') = (X \cap X') \times (Y \cap Y')$	E108, E91.

E114

$(x, y) \in X \times (A \cup B)$	
$\Leftrightarrow x \in X \text{ et } (y \in A \text{ ou } y \in B)$	E109, E18
$\Leftrightarrow (x \in X \text{ et } y \in A) \text{ ou } (x \in X \text{ et } y \in B)$	S52.11
$\Leftrightarrow (x, y) \in X \times A \text{ ou } (x, y) \in X \times B$	E109
$\Leftrightarrow (x, y) \in (X \times A) \cup (X \times B)$	E18.

E115.1

$(x, y) \in \{a\} \times \{b\} \Leftrightarrow x \in \{a\} \text{ et } y \in \{b\}$	E109
$\Leftrightarrow x = a \text{ et } y = b$	E60
$\Leftrightarrow (x, y) = (a, b)$	E79
$\Leftrightarrow (x, y) \in \{(a, b)\}$	E60.

E115.2

$(x, y) \in \{a\} \times \{b, c\} \Leftrightarrow x \in \{a\} \text{ et } y \in \{b, c\}$	E109
$\Leftrightarrow x = a \text{ et } (y = b \text{ ou } y = c)$	E60, E51
$\Leftrightarrow (x = a \text{ et } y = b) \text{ ou } (x = a \text{ et } y = c)$	S52.11
$\Leftrightarrow (x, y) = (a, b) \text{ ou } (x, y) = (a, c)$	E79
$\Leftrightarrow (x, y) \in \{(a, b), (a, c)\}$	E51.

E115.4

$z \in \{(a, c), (a, d), (b, c), (b, d)\}$	
$\Leftrightarrow z \in \{(a, c)\} \cup \{(a, d)\} \cup \{(b, c)\} \cup \{(b, d)\}$	§11.3
$\Leftrightarrow z \in \{(a, c), (a, d)\} \cup \{(b, c), (b, d)\}$	E68
$\Leftrightarrow z \in (\{a\} \times \{c, d\}) \cup (\{b\} \times \{c, d\})$	E115 (deuxième)
$\Leftrightarrow z \in (\{a\} \cup \{b\}) \times \{c, d\}$	E114
$\Leftrightarrow z \in \{a, b\} \times \{c, d\}$	E68.

E117

$x \in \mathbf{Pr}_1 G \Leftrightarrow x \in \bigcup_{z \in G} \{\mathbf{pr}_1 z\}$	E116
$\Leftrightarrow \exists z(z \in G \text{ et } x \in \{\mathbf{pr}_1 z\})$	E95
$\Leftrightarrow \exists z(z \text{ est un couple et } z \in G \text{ et } x = \mathbf{pr}_1 z)$	E88, E60
$\Leftrightarrow \exists z(\exists a \exists y(z = (a, y)) \text{ et } z \in G \text{ et } x = \mathbf{pr}_1 z)$	E81
$\Leftrightarrow \exists y \exists a \exists z(z = (a, y) \text{ et } z \in G \text{ et } x = \mathbf{pr}_1 z)$	S81, S72
$\Leftrightarrow \exists y \exists a((a, y) \in G \text{ et } x = \mathbf{pr}_1(a, y))$	S92
$\Leftrightarrow \exists y \exists a((a, y) \in G \text{ et } x = a)$	E84, S86
$\Leftrightarrow \exists y((x, y) \in G)$	S92.

E118

1°	$G \text{ est un graphe}$	hyp
2°	$(x, y) \in G$	hyp
3°	$\exists y((x, y) \in G)$	2°, S66
4°	$x \in \mathbf{Pr}_1 G$	3°, E117
5°	$\exists x((x, y) \in G)$	2°, S66
6°	$y \in \mathbf{Pr}_2 G$	5°, E117
7°	$(x, y) \in \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G$	4°, 6°, E109
8°	$G \subset \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G$	7°, $\Rightarrow$ -intro, E91
9°	...	8°, $\Rightarrow$ -intro.

E119

1°	$G \text{ et } G' \text{ sont des graphes}$	hyp
2°	$G \subset G'$	hyp
3°	$\bigcup_{z \in G} \{\mathbf{pr}_1 z\} \subset \bigcup_{z \in G'} \{\mathbf{pr}_1 z\}$	2°, E101
4°	$\mathbf{Pr}_1 G \subset \mathbf{Pr}_1 G'$	3°, E116
5°	$G \subset G' \Rightarrow \mathbf{Pr}_1 G \subset \mathbf{Pr}_1 G'$	4°, $\Rightarrow$ -intro.

Autre démonstration

1°	$G \text{ et } G' \text{ sont des graphes}$	hyp
2°	$G \subset G'$	hyp
3°	$x \in \mathbf{Pr}_1 G$	hyp
4°	$\exists y((x, y) \in G)$	3°, E117
5°	$(x, y) \in G$	hyp
6°	$(x, y) \in G'$	2°, 5°, E4
7°	$\exists y((x, y) \in G')$	6°, S66
8°	$x \in \mathbf{Pr}_1 G'$	7°, E117
9°	$x \in \mathbf{Pr}_1 G'$	4°, 8°, S64
10°	$\mathbf{Pr}_1 G \subset \mathbf{Pr}_1 G'$	9°, E5
11°	$G \subset G' \Rightarrow \mathbf{Pr}_1 G \subset \mathbf{Pr}_1 G'$	10°, $\Rightarrow$ -intro.

E120

1°	$G \subset X \times Y$	hyp
2°	$x \in \mathbf{Pr}_1 G$	hyp
3°	$G \text{ est un graphe}$	1°, E108, E89
4°	$\exists y((x, y) \in G)$	2°, 3°, E117
5°	$(x, y) \in G$	hyp
6°	$(x, y) \in X \times Y$	5°, 1°, E4
7°	$x \in X$	6°, E109
8°	$x \in X$	7°, 4°, S64

9°	$\mathbf{Pr}_1 G \subset X$	8°, E5
10°	$\mathbf{Pr}_2 G \subset Y$	démonstr. analogue à 2°–9°
	...	

E121

1°	$Y \neq \emptyset$	hyp
2°	$\exists y(y \in Y)$	1°, E13
3°	$x \in \mathbf{Pr}_1(X \times Y) \Leftrightarrow \exists y((x, y) \in X \times Y)$	E117, E108
4°	$\Leftrightarrow \exists y(x \in X \text{ et } y \in Y)$	E109
5°	$\Leftrightarrow (x \in X \text{ et } \exists y(y \in Y))$	S81
6°	$\Leftrightarrow x \in X$	2°, S60.1
7°	$\mathbf{Pr}_1(X \times Y) = X$	E17
8°	$Y \neq \emptyset \Rightarrow \mathbf{Pr}_1(X \times Y) = X$	7°, $\Rightarrow$ -intro.

E122

	G est un graphe	hyp
	$G = \emptyset \Leftrightarrow \forall x \forall y((x, y) \notin G)$	E92
	$\Leftrightarrow \forall x \text{ non } \exists y((x, y) \in G)$	S71
	$\Leftrightarrow \forall x(x \notin \mathbf{Pr}_1 G)$	E117
	$\Leftrightarrow \mathbf{Pr}_1 G = \emptyset$	E13
	$G = \emptyset \Leftrightarrow \forall y \forall x((x, y) \notin G)$	E92, S72
	$\Leftrightarrow \forall y \text{ non } \exists x((x, y) \in G)$	S71
	$\Leftrightarrow \forall y(y \notin \mathbf{Pr}_2 G)$	E117
	$\Leftrightarrow \mathbf{Pr}_2 G = \emptyset$	E13.

E124

1°	G est un graphe	hyp
2°	$y \in G^{-1}$	hyp
3°	$y \in \bigcup_{z \in G} \{(\mathbf{pr}_1 z, \mathbf{pr}_2 z)\}$	2°, E123
4°	$\exists z(z \in G \text{ et } y \in \{(\mathbf{pr}_1 z, \mathbf{pr}_2 z)\})$	3°, E95
5°	$z \in G \text{ et } y = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	hyp
6°	$y = (\mathbf{pr}_1 z, \mathbf{pr}_2 z)$	5°, 1°
7°	$(\mathbf{pr}_1 z, \mathbf{pr}_2 z)$ est un couple	E82
8°	y est un couple	6°, 7°
9°	y est un couple	8°, 4°, S64
10°	$\forall y(y \in G^{-1} \Rightarrow y \text{ est un couple})$	9°, $\Rightarrow$ -intro, S62
12°	G^{-1} est un graphe	10°, E86

E125

	G est un graphe	hyp
	$(x, y) \in G^{-1} \Leftrightarrow (x, y) \in \bigcup_{z \in G} \{(\mathbf{pr}_2 z, \mathbf{pr}_1 z)\}$	S86
	$\Leftrightarrow \exists z(z \in G \text{ et } (x, y) \in \{(\mathbf{pr}_2 z, \mathbf{pr}_1 z)\})$	E95
	$\Leftrightarrow \exists z(z \text{ est un couple et } z \in G \text{ et } (x, y) = (\mathbf{pr}_2 z, \mathbf{pr}_1 z))$	E88, E60
	$\Leftrightarrow \exists z(z = (\mathbf{pr}_1 z, \mathbf{pr}_2 z) \text{ et } z \in G \text{ et } x = \mathbf{pr}_2 z \text{ et } y = \mathbf{pr}_1 z)$	E85, E79
	$\Leftrightarrow \exists z(z = (y, x) \text{ et } z \in G \text{ et } y = \mathbf{pr}_1 z \text{ et } x = \mathbf{pr}_2 z)$	S87.1
	$\Leftrightarrow (y, x) \in G \text{ et } y = \mathbf{pr}_1(y, x) \text{ et } x = \mathbf{pr}_2(y, x)$	S92
	$\Leftrightarrow (y, x) \in G$	E84, S60.1

E126	$ \begin{array}{l} G \text{ est un graphe} \\ (x, y) \in (G^{-1})^{-1} \Leftrightarrow (y, x) \in G^{-1} \\ \Leftrightarrow (x, y) \in G \end{array} $	hyp E125 E125.
E127	$ \begin{array}{l} G \text{ est un graphe} \\ x \in \mathbf{Pr}_1(G^{-1}) \Leftrightarrow \exists y((x, y) \in G^{-1}) \\ \Leftrightarrow \exists y((y, x) \in G) \\ \Leftrightarrow x \in \mathbf{Pr}_2 G \end{array} $	hyp E117 E125 E117.
E128	$ \begin{array}{l} G \text{ et } G' \text{ sont des graphes} \\ (x, y) \in (G \cup G')^{-1} \Leftrightarrow (y, x) \in G \cup G' \\ \Leftrightarrow (y, x) \in G \text{ ou } (y, x) \in G' \\ \Leftrightarrow (x, y) \in G^{-1} \text{ ou } (x, y) \in G'^{-1} \\ \Leftrightarrow (x, y) \in G^{-1} \cup G'^{-1} \\ (x, y) \in (G \cap G')^{-1} \Leftrightarrow (y, x) \in G \cap G' \\ \Leftrightarrow (y, x) \in G \text{ et } (y, x) \in G' \\ \Leftrightarrow (x, y) \in G^{-1} \text{ et } (x, y) \in G'^{-1} \\ \Leftrightarrow (x, y) \in G^{-1} \cap G'^{-1} \end{array} $	hyp E90, E125 E18 E125 E18 E90, E125 E28 E125 E28.
E129	$ \begin{array}{l} \emptyset^{-1} = \bigcup_{z \in \emptyset} \{(\mathbf{pr}_2 z, \mathbf{pr}_1 z)\} \\ = \emptyset \end{array} $	E87, E123 E100.
E131	$ \begin{array}{l} 1^\circ \quad z \in \text{Id}_X \\ 2^\circ \quad z \in \bigcup_{a \in X} \{(a, a)\} \\ 3^\circ \quad \exists a(a \in X \text{ et } z \in \{(a, a)\}) \\ 4^\circ \quad \exists a(a \in X \text{ et } z = (a, a)) \\ 5^\circ \quad \left \begin{array}{l} a \in X \text{ et } z = (a, a) \\ (a, a) \in X \times X \end{array} \right. \\ 6^\circ \quad \left \begin{array}{l} a \in X \text{ et } z = (a, a) \\ (a, a) \in X \times X \end{array} \right. \\ 7^\circ \quad \left \begin{array}{l} a \in X \text{ et } z = (a, a) \\ z \in X \times X \end{array} \right. \\ 8^\circ \quad z \in X \times X \\ 9^\circ \quad \text{Id}_X \subset X \times X \end{array} $	hyp 1°, E130 2°, E95 3°, E60 hyp 5°, E109 5°, 6° 4°, 7°, S64 8°, E5.
E132	Découle de E131, E108, E89.	
E133	$ \begin{array}{l} (x, y) \in \text{Id}_X \Leftrightarrow \exists a(a \in X \text{ et } (x, y) \in \{(a, a)\}) \\ \Leftrightarrow \exists a(a \in X \text{ et } (x, y) = (a, a)) \\ \Leftrightarrow \exists a(a \in X \text{ et } x = a \text{ et } y = a) \\ \Leftrightarrow \exists a(a = x \text{ et } a \in X \text{ et } y = a) \\ \Leftrightarrow x \in X \text{ et } y = x \end{array} $	E130, E95 E60 E79 S92.
E134	$ \begin{array}{l} x \in \mathbf{Pr}_1(\text{Id}_X) \Leftrightarrow \exists y((x, y) \in \text{Id}_X) \\ \Leftrightarrow \exists y(x \in X \text{ et } y = x) \\ \Leftrightarrow \exists y(y = x \text{ et } x \in X) \\ \Leftrightarrow x \in X \end{array} $	E117 E133 S92

	$y \in \mathbf{Pr}_2(\text{Id}_X) \Leftrightarrow \exists x((x, y) \in \text{Id}_X)$ $\Leftrightarrow \exists x(x \in X \text{ et } y = x)$ $\Leftrightarrow \exists x(x = y \text{ et } x \in X)$ $\Leftrightarrow y \in X$	E117 E133 S92.
E135	$\text{Id}_\emptyset = \bigcup_{a \in \emptyset} \{(a, a)\}$ $= \emptyset$	E130 E100.
E136	$(x, y) \in \text{Id}_{X \cap Y} \Leftrightarrow x \in X \cap Y \text{ et } y = x$ $\Leftrightarrow x \in X \text{ et } x \in Y \text{ et } y = x$ $\Leftrightarrow x \in X \text{ et } y = x \text{ et } x \in Y \text{ et } y = x$ $\Leftrightarrow (x, y) \in \text{Id}_X \text{ et } (x, y) \in \text{Id}_Y$ $\Leftrightarrow (x, y) \in \text{Id}_X \cap \text{Id}_Y$	E133 E28 S52 E133 E28.
E137	$(x, y) \in \text{Id}_X^{-1} \Leftrightarrow (y, x) \in \text{Id}_X$ $\Leftrightarrow y \in X \text{ et } x = y$ $\Leftrightarrow x \in X \text{ et } x = y$ $\Leftrightarrow x \in X \text{ et } y = x$ $\Leftrightarrow (x, y) \in \text{Id}_X$	E132, E125 E133 S87.1 S84 E133.
E151	$1^\circ \quad G \text{ est un graphe}$ $y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G)$ $\Rightarrow \exists x((x, y) \in G)$ $\Leftrightarrow y \in \mathbf{Pr}_2 G$ $G\langle X \rangle \subset \mathbf{Pr}_2 G$	hyp E150 S75 E117 E5.
E152	$1^\circ \quad G \text{ est un graphe}$ $2^\circ \quad (x, y) \in G \Rightarrow x \in \mathbf{Pr}_1 G$ $y \in G\langle \mathbf{Pr}_1 G \rangle \Leftrightarrow \exists x((x, y) \in G \text{ et } x \in \mathbf{Pr}_1 G)$ $\Leftrightarrow \exists x((x, y) \in G)$ $\Leftrightarrow y \in \mathbf{Pr}_2 G$	hyp 1° , E118 1° , E150 2° , S60.3 E117.
E153	$1^\circ \quad G \text{ est un graphe}$ $2^\circ \quad (x, y) \in G \Rightarrow x \in \mathbf{Pr}_1 G$ $y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G)$ $\Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G \text{ et } x \in \mathbf{Pr}_1 G)$ $\Leftrightarrow \exists x(x \in X \cap \mathbf{Pr}_1 G \text{ et } (x, y) \in G)$ $\Leftrightarrow y \in G\langle X \cap \mathbf{Pr}_1 G \rangle$	hyp 1° , E118 1° , E150 2° , S60.3 E28 E150.
E154	$1^\circ \quad G \text{ est un graphe}$ $y \in G\langle \{x\} \rangle \Leftrightarrow \exists a(a \in \{x\} \text{ et } (a, y) \in G)$ $\Leftrightarrow \exists a(a = x \text{ et } (a, y) \in G)$ $\Leftrightarrow (x, y) \in G$	hyp 1° , E150 E60 S92.

E155

1°	G est un graphe	hyp
	$y \in G\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in G)$	1°, E150
	$\Leftrightarrow \exists x(x \in X \text{ et } y \in G\langle \{x\} \rangle)$	E154
	$\Leftrightarrow y \in \bigcup_{x \in X} G\langle \{x\} \rangle$	E95.

E156 D'après E155 et E101.

E157 D'après E155 et E102.

E158

1°	G est un graphe	hyp
2°	$X \cap Y \subset X$	E33
3°	$G\langle X \cap Y \rangle \subset G\langle X \rangle$	2°, E156
4°	$X \cap Y \subset Y$	E33
5°	$G\langle X \cap Y \rangle \subset G\langle Y \rangle$	4°, E156
6°	$G\langle X \cap Y \rangle \subset G\langle X \rangle \cap G\langle Y \rangle$	3°, 5°, E34.

E159

1°	G est un graphe	hyp
	$G\langle X \rangle = \emptyset \Leftrightarrow \text{non} \exists y(y \in G\langle X \rangle)$	E13
	$\Leftrightarrow \text{non} \exists y \exists x(x \in X \text{ et } (x, y) \in G)$	E150
	$\Leftrightarrow \text{non} \exists x(x \in X \text{ et } \exists y((x, y) \in G))$	S72, S81
	$\Leftrightarrow \text{non} \exists x(x \in X \text{ et } x \in \mathbf{Pr}_1 G)$	E117
	$\Leftrightarrow \text{non} \exists x(x \in X \cap \mathbf{Pr}_1 G)$	E28
	$\Leftrightarrow X \cap \mathbf{Pr}_1 G = \emptyset$	E13.

E160

1°	G est un graphe	hyp
	$G\langle \{x\} \rangle \neq \emptyset \Leftrightarrow \{x\} \cap \mathbf{Pr}_1 G \neq \emptyset$	E159, S50.1
	$\Leftrightarrow x \in \mathbf{pr}_1 G$	E64.

E161

D'après E155 et E100.

E162

	Id_Y est un graphe	E132
	$y \in \text{Id}_Y\langle X \rangle \Leftrightarrow \exists x(x \in X \text{ et } (x, y) \in \text{Id}_Y)$	E150
	$\Leftrightarrow \exists x(x \in X \text{ et } x \in Y \text{ et } y = x)$	E133
	$\Leftrightarrow \exists x(x = y \text{ et } x \in X \cap Y)$	E28
	$\Leftrightarrow y \in X \cap Y$	S92.

E163

	$X \subset Y \Leftrightarrow X \cap Y = X$	E35
	$\Leftrightarrow \text{Id}_Y\langle X \rangle = X$	E162.

E184

1°	G et G' sont des graphes	hyp
2°	$\exists x((a, x) \in G \text{ et } (x, b) \in G')$	hyp
3°	$\exists x((a, x) \in G) \text{ et } \exists x((x, b) \in G')$	2°, S75
4°	$a \in \mathbf{Pr}_1 G \text{ et } b \in \mathbf{Pr}_2 G'$	3°, E117
5°	$(a, b) \in \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$	4°, E109
6°	$\exists x((a, x) \in G \text{ et } (x, b) \in G') \Rightarrow (a, b) \in \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$	5°, $\Rightarrow$ -intro
7°	$\{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}$ existe	6°, E181.

E185

1°	G et G' sont des graphes	hyp
2°	$G' \circ G = \{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}$	E183
3°	$\{(a, b) \mid \exists x((a, x) \in G \text{ et } (x, b) \in G')\}$ existe	E184
4°	$z \in G' \circ G$	hyp
5°	$\exists a \exists b (z = (a, b) \text{ et } \exists x((a, x) \in G \text{ et } (x, b) \in G'))$	2°, 3°, E174
6°	$\exists a \exists b (z = (a, b))$	élimination et intro de $\exists a \exists b$
7°	z est un couple	6°, E81
8°	$z \in G' \circ G \Rightarrow z$ est un couple	7°, $\Rightarrow$ -intro
9°	$G' \circ G$ est un graphe	8°, S62, E86
10°	$(a, b) \in G' \circ G \Leftrightarrow \exists x((a, x) \in G \text{ et } (x, b) \in G')$	2°, 3°, E180.

E186

1°	G et G' sont des graphes	hyp
2°	$(a, b) \in G' \circ G$	hyp
3°	$\exists x((a, x) \in G \text{ et } (x, b) \in G')$	1°, 2°, E185
4°	$\exists x((a, x) \in G) \text{ et } \exists x((x, b) \in G')$	3°, S75
5°	$a \in \mathbf{Pr}_1 G \text{ et } b \in \mathbf{Pr}_2 G'$	4°, E117
6°	$(a, b) \in \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$	5°, E109
7°	$G' \circ G \subset \mathbf{Pr}_1 G \times \mathbf{Pr}_2 G'$	6°, E91.

E187

	G, G', G'' sont des graphes	hyp
	$(a, b) \in G'' \circ (G' \circ G) \Leftrightarrow \exists y((a, y) \in G' \circ G \text{ et } (y, b) \in G'')$	E185
	$\Leftrightarrow \exists y(\exists x((a, x) \in G \text{ et } (x, y) \in G') \text{ et } (y, b) \in G'')$	E185
	$\Leftrightarrow \exists y \exists x((a, x) \in G \text{ et } ((x, y) \in G' \text{ et } (y, b) \in G''))$	S81, S52.10
	$\Leftrightarrow \exists x((a, x) \in G \text{ et } \exists y((x, y) \in G' \text{ et } (y, b) \in G''))$	S72, S81
	$\Leftrightarrow \exists x((a, x) \in G \text{ et } (x, b) \in G'' \circ G')$	E185
	$\Leftrightarrow (a, b) \in (G'' \circ G') \circ G$	E185.

E188

	G et G' sont des graphes	hyp
	$(a, b) \in (G' \circ G)^{-1} \Leftrightarrow (b, a) \in G' \circ G$	E125
	$\Leftrightarrow \exists x((b, x) \in G \text{ et } (x, a) \in G')$	E185
	$\Leftrightarrow \exists x((a, x) \in G'^{-1} \text{ et } (x, b) \in G^{-1})$	E125, S52.8
	$\Leftrightarrow (a, b) \in G^{-1} \circ G'^{-1}$	E185.

E189

1°	G et G' sont des graphes	hyp
2°	$x \in \mathbf{Pr}_1(G' \circ G) \Leftrightarrow \exists z((x, z) \in G' \circ G)$	E117
3°	$\Leftrightarrow \exists z \exists y((x, y) \in G \text{ et } (y, z) \in G')$	E185
4°	$\Leftrightarrow \exists y(\exists z((y, z) \in G') \text{ et } (x, y) \in G)$	S72, S81
5°	$\Leftrightarrow \exists y(y \in \mathbf{Pr}_1 G' \text{ et } (x, y) \in G)$	E117
6°	$\Leftrightarrow \exists y(y \in \mathbf{Pr}_1 G' \text{ et } (y, x) \in G^{-1})$	E125
7°	$\Leftrightarrow x \in G^{-1} \langle \mathbf{Pr}_1 G' \rangle$	E148
8°	$\mathbf{Pr}_1(G' \circ G) = G^{-1} \langle \mathbf{Pr}_1 G' \rangle$	E17
9°	$\mathbf{Pr}_2(G' \circ G) = \mathbf{Pr}_1((G' \circ G)^{-1})$	E127
10°	$= \mathbf{Pr}_1(G^{-1} \circ G'^{-1})$	E188
11°	$= (G'^{-1})^{-1} \langle \mathbf{Pr}_1 G^{-1} \rangle$	E189 premier
12°	$= G' \langle \mathbf{Pr}_2 G \rangle$	E126, E127
13°	$\mathbf{Pr}_1(G' \circ G) \subset \mathbf{Pr}_2(G^{-1})$	8°, E151

14°	$\mathbf{Pr}_1(G' \circ G) \subset \mathbf{Pr}_1 G$	13°, E127
15°	$\mathbf{Pr}_2(G' \circ G) \subset \mathbf{Pr}_2 G'$	12°, E151.
E190		
1°	G est un graphe	hyp
2°	$\emptyset$ est un graphe	E87
3°	$\mathbf{Pr}_1(G \circ \emptyset) \subset \mathbf{Pr}_1 \emptyset$	E189
4°	$\mathbf{Pr}_1 \emptyset = \emptyset$	E122
5°	$\mathbf{Pr}_1(G \circ \emptyset) = \emptyset$	3°, 4°, E12
6°	$G \circ \emptyset = \emptyset$	5°, E122.
	$\emptyset \circ G = ((\emptyset \circ G)^{-1})^{-1}$	E126
	$= (G^{-1} \circ \emptyset^{-1})^{-1}$	E188
	$= (G^{-1} \circ \emptyset)^{-1}$	E129
	$= \emptyset^{-1}$	E190 premier
	$= \emptyset$	E129.
E191		
1°	G est un graphe et $\mathbf{Pr}_1 G \subset X$	hyp
2°	$(a, b) \in G \Rightarrow a \in X$	1°, E118
	$(a, b) \in G \circ \text{Id}_X \Leftrightarrow \exists x((a, x) \in \text{Id}_X \text{ et } (x, b) \in G)$	E185
	$\Leftrightarrow \exists x(x = a \text{ et } a \in X \text{ et } (x, b) \in G)$	E133
	$\Leftrightarrow a \in X \text{ et } (a, b) \in G$	S92
	$\Leftrightarrow (a, b) \in G$	2°, S60.3
	$G \circ \text{Id}_X = G$	E91.
1°	G est un graphe et $\mathbf{Pr}_2 G \subset Y$	hyp
2°	G^{-1} est un graphe et $\mathbf{Pr}_1 G^{-1} \subset Y$	1°, E124, E127
3°	$G^{-1} \circ \text{Id}_Y = G^{-1}$	E191 premier
4°	$(G^{-1} \circ \text{Id}_Y)^{-1} = G$	E126
5°	$\text{Id}_Y \circ G = G$	E188, E137.
E192		
1°	G est un graphe	hyp
2°	$\text{Id}_{\mathbf{Pr}_1 G}$ et G^{-1} et $G^{-1} \circ G$ sont des graphes	E132, E124, E185
3°	$(x, x') \in \text{Id}_{\mathbf{Pr}_1 G}$	hyp
4°	$x \in \mathbf{Pr}_1 G$ et $x' = x$	3°, E133
5°	$\exists y((x, y) \in G)$	4°, et-élim, E117
6°	$\exists y((x, y) \in G \text{ et } (x, y) \in G)$	5°, S52
7°	$\exists y((x, y) \in G \text{ et } (y, x) \in G^{-1})$	6°, E125
8°	$(x, x) \in G^{-1} \circ G$	7°, E185
9°	$(x, x') \in G^{-1} \circ G$	8°, 4°
10°	$\text{Id}_{\mathbf{Pr}_1 G} \subset G^{-1} \circ G$	E91.
11°	G est un graphe $\Rightarrow \text{Id}_{\mathbf{Pr}_1 G} \subset G^{-1} \circ G$	$\Rightarrow$ -intro
1°	G est un graphe	hyp
2°	G^{-1} est un graphe	1°, E124
3°	$\text{Id}_{\mathbf{Pr}_1 G^{-1}} \subset (G^{-1})^{-1} \circ G^{-1}$	E192 (premier)
4°	$\text{Id}_{\mathbf{Pr}_2 G} \subset G \circ G^{-1}$	E127, E126.
E193		
	G est un graphe	hyp
	$y \in \mathbf{Pr}_2(G \circ \text{Id}_X) \Leftrightarrow \exists x((x, y) \in G \circ \text{Id}_X)$	E117
	$\Leftrightarrow \exists x \exists a((x, a) \in \text{Id}_X \text{ et } (a, y) \in G)$	E185

	$\Leftrightarrow \exists x \exists a (x \in X \text{ et } a = x \text{ et } (a, y) \in G)$	E133
	$\Leftrightarrow \exists x (x \in X \text{ et } (x, y) \in G)$	S92
	$\Leftrightarrow y \in G\langle X \rangle$	E150.
E194	$G \text{ et } G' \text{ sont des graphes}$	hyp
	$z \in (G' \circ G)\langle X \rangle \Leftrightarrow \exists x (x \in X \text{ et } (x, z) \in G' \circ G)$	E150
	$\Leftrightarrow \exists x (x \in X \text{ et } \exists y ((x, y) \in G \text{ et } (y, z) \in G'))$	E185
	$\Leftrightarrow \exists y \exists x (x \in X \text{ et } (x, y) \in G \text{ et } (y, z) \in G')$	S81, S72
	$\Leftrightarrow \exists y (\exists x (x \in X \text{ et } (x, y) \in G) \text{ et } (y, z) \in G')$	S81
	$\Leftrightarrow \exists y (y \in G\langle X \rangle \text{ et } (y, z) \in G')$	E150
	$\Leftrightarrow z \in G'\langle G\langle X \rangle \rangle$	E150.
F2	$F \text{ est une fonction} \Leftrightarrow [F \text{ est un graphe et } \forall x \text{Hy}((x, y) \in F)]$	F1
	$\Leftrightarrow [F \text{ est un graphe et } \forall x (\exists y ((x, y) \in F) \Rightarrow \exists ! y ((x, y) \in F))]$	S95
	$\Leftrightarrow [F \text{ est un graphe et } \forall x (x \in \mathbf{Pr}_1 F \Rightarrow \exists ! y ((x, y) \in F))]$	E117.
F3	Se déduit selon la logique propositionnelle de F1 et du théorème:	
	$\vdash F \text{ est un graphe} \Rightarrow [\forall x \text{Hy}((x, y) \in F) \Leftrightarrow F \circ F^{-1} = \text{Id}_{\mathbf{Pr}_2 F}]$	(*)
Démonstration de (*)		
1°	$F \text{ est un graphe}$	hyp
2°	$\forall x \text{Hy}((x, y) \in F)$	hyp
3°	$((x, y) \in F \text{ et } (x, y') \in F) \Rightarrow y = y'$	2°, S93, S65
4°	$(y, y') \in F \circ F^{-1} \Leftrightarrow \exists x ((y, x) \in F^{-1} \text{ et } (x, y') \in F)$	E185
5°	$\Leftrightarrow \exists x ((x, y) \in F \text{ et } (x, y') \in F)$	E125
6°	$\Leftrightarrow \exists x ((x, y) \in F \text{ et } (x, y') \in F \text{ et } y = y')$	3°, S60.3
7°	$\Leftrightarrow \exists x ((x, y) \in F \text{ et } (x, y) \in F \text{ et } y = y')$	S87
8°	$\Leftrightarrow \exists x ((x, y) \in F) \text{ et } y = y'$	S52.6, S81
9°	$\Leftrightarrow y \in \mathbf{Pr}_2 F \text{ et } y = y'$	E117
10°	$\Leftrightarrow (y, y') \in \text{Id}_{\mathbf{Pr}_2 F}$	E133
11°	$F \circ F^{-1} = \text{Id}_{\mathbf{Pr}_2 F}$	4°–10°, E91
12°	$F \circ F^{-1} = \text{Id}_{\mathbf{Pr}_2 F}$	hyp
13°	$(x, y) \in F \text{ et } (x, y') \in F$	hyp
14°	$(y, x) \in F^{-1} \text{ et } (x, y') \in F$	13°, E125
15°	$\exists x ((y, x) \in F^{-1} \text{ et } (x, y') \in F)$	14°, S66
16°	$(y, y') \in F \circ F^{-1}$	15°, E185
17°	$(y, y') \in \text{Id}_{\mathbf{Pr}_2 F}$	12°, 16°
18°	$y = y'$	17°, E133
19°	$\forall x \forall y \forall y' ((x, y) \in F \text{ et } (x, y') \in F \Rightarrow y = y')$	18°, $\Rightarrow$ -intro, S62
20°	$\forall x \text{Hy}((x, y) \in F)$	19°, S93
21°	$\forall x \text{Hy}((x, y) \in F) \Leftrightarrow F \circ F^{-1} = \text{Id}_{\mathbf{Pr}_2 F}$	11°, 20°, S40.
F7.2	$F \text{ est une fonction}$	hyp
1°	$y \in \text{Prt } F \Leftrightarrow \exists x ((x, y) \in F)$	F4, E117
2°	$\Leftrightarrow \exists x (y = F @ x \text{ et } x \in \text{Dom } F)$	F6
3°	$\forall y (y \in \text{Prt } F \Leftrightarrow \exists x (y = F @ x \text{ et } x \in \text{Dom } F))$	2°–3°, S62
4°	$\text{Prt } F = \{F @ x \mid x \in \text{Dom } F\}$	E174

F9 Se déduit de **F8** et du théorème suivant démontré ici :

$$\vdash F \text{ est une fonction} \Rightarrow (\text{Prt}F \subset B \Leftrightarrow \forall x(x \in \text{Dom}F \Rightarrow F @ x \in B)).$$

1°	$F \text{ est une fonction}$	hyp
2°	$\text{Prt}F = \{F @ x \mid x \in \text{Dom}F\}$	F7
3°	$\{F @ x \mid x \in \text{Dom}F\} \text{ existe}$	E182
4°	$\text{Prt}F \subset B \Leftrightarrow \forall x(x \in \text{Dom}F \Rightarrow F @ x \in B)$	2°, 3°, E178.

F24

1°	$F = \emptyset$	hyp
2°	$F = \text{Id}_{\emptyset}$	1°, E135
3°	$F \text{ est une fonction et } \text{Dom}F = \emptyset$	2°, F23
4°	$F \text{ est une fonction et } \text{Dom}F = \emptyset$	hyp
5°	$\text{Dom}F = \emptyset \Leftrightarrow F = \emptyset$	4°, E122
6°	$F = \emptyset$	4°, 5°.

F25

1°	$F : \emptyset \rightarrow B$	hyp
2°	$F \text{ est une fonction et } \text{Dom}F = \emptyset$	1°, F8
3°	$F = \emptyset$	2°, F24
4°	$F = \emptyset$	hyp
5°	$F \text{ est une fonction et } \text{Dom}F = \emptyset$	4°, F24
6°	$\text{Prt}F = \emptyset$	4°, F24
7°	$\text{Prt}F \subset B$	6°, E7
8°	$F : \emptyset \rightarrow B$	5°, 7°, F8.
9°	$(F : \emptyset \rightarrow B) \Leftrightarrow F = \emptyset$	3°, 8°, S40
10°	$\exists! F(F : \emptyset \rightarrow B)$	9°, S98
11°	$(F : A \rightarrow \emptyset) \Leftrightarrow (F \text{ est une fonction et } \text{Dom}F = A \text{ et } \text{Pr}tF \subset \emptyset)$	F8
12°	$\Leftrightarrow (F \text{ est une fonction et } \text{Pr}tF = \emptyset \text{ et } \text{Dom}F = A)$	E12
13°	$\Leftrightarrow (F = \emptyset \text{ et } \text{Dom}F = A)$	F24
14°	$\Leftrightarrow (F \text{ est une fonction et } \text{Dom}F = \emptyset \text{ et } \text{Dom}F = A)$	F24
15°	$\Leftrightarrow (F \text{ est une fonction et } \text{Dom}F = \emptyset \text{ et } \emptyset = A)$	S87
16°	$\Leftrightarrow (F = \emptyset \text{ et } A = \emptyset)$	F24
17°	$(F : A \rightarrow \emptyset) \Leftrightarrow (A = \emptyset \text{ et } F = \emptyset)$	11°–16°
18°	$(F : A \rightarrow \emptyset) \Rightarrow A = \emptyset$	17°
19°	$A \neq \emptyset$	hyp
20°	$\text{non}(F : A \rightarrow \emptyset)$	18°, 19°, S32
21°	$\forall F \text{non}(F : A \rightarrow \emptyset)$	20°, S62
22°	$\text{non}\exists F(F : A \rightarrow \emptyset)$	21°, S71
23°	$A \neq \emptyset \Rightarrow \text{non}\exists F(F : A \rightarrow \emptyset)$	16°, $\Rightarrow$ -intro.

F27

1°	$F \text{ et } G \text{ sont des fonctions et } \text{Pr}tF \subset \text{Dom}G$	hyp
2°	$\text{Pr}tF = \{F @ x \mid x \in \text{Dom}F\}$	F7
3°	$\{F @ x \mid x \in \text{Dom}F\} \text{ existe}$	E182
4°	$\text{Pr}tF \subset \text{Dom}G$	1°
5°	$x \in \text{Dom}F \Rightarrow F @ x \in \text{Dom}G$	2°, 3°, 4°, E178
	$x \in \text{Dom}(G \circ F) \Leftrightarrow x \in \{x \mid x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G\}$	F26
	$\Leftrightarrow x \in \text{Dom}F \text{ et } F @ x \in \text{Dom}G$	E146, E142
	$\Leftrightarrow x \in \text{Dom}F$	5°, S60.3.

F28

1°	$F : A \rightarrow B$ et $G : B \rightarrow C$	hyp
2°	F et G sont des fonctions	1°, F8
3°	$G \circ F$ est une fonction	2°, F26
4°	$\text{Prt} F \subset B$ et $B = \text{Dom} G$	1°, F8
5°	$\text{Prt} F \subset \text{Dom} G$	4°
6°	$\text{Dom}(G \circ F) = \text{Dom} F$	5°, F27
7°	$\text{Dom}(G \circ F) = A$	6°, 1°, F8
8°	$\text{Prt}(G \circ F) \subset \text{Prt} G$	E189 (3ème)
9°	$\text{Prt}(G \circ F) \subset C$	8°, 1°, F8
10°	$G \circ F : A \rightarrow C$	3°, 7°, 9°, F8.

F29

1°	$F : A \rightarrow B$	hyp
2°	F est une fonction et $\text{Dom} F = A$ et $\text{Prt} F \subset B$	1°, F8
3°	F est un graphe	2°, F1
4°	$\mathbf{Pr}_2 F \subset B$	2°, F4
5°	$\text{Id}_B \circ F = F$	3°, 4°, E191
6°	$\mathbf{Pr}_1 F = A$	2°, F4
7°	$\mathbf{Pr}_1 F \subset A$	6°, E10
8°	$F \circ \text{Id}_A = F$	3°, 7°, E191
9°	$\text{Id}_B \circ F = F$ et $F \circ \text{Id}_A = F$	5°, 8°.

F31

1°	G est une fonction	hyp
2°	G^{-1} est un graphe	1°, F1, E124
3°	$\exists F (F \text{ est une fonction et } F \subset G^{-1} \text{ et } \text{Dom} F = \mathbf{Pr}_1(G^{-1}))$	2°, F30
4°	F est une fonction et $F \subset G^{-1}$ et $\text{Dom} F = \mathbf{Pr}_1(G^{-1})$	hyp
5°	$G \circ F$ est une fonction	1°, 4°, F26
6°	$\text{Dom} F = \mathbf{Pr}_2 G$	4°, E127
7°	$F \subset G^{-1}$	4°
8°	$\text{Prt} F \subset \mathbf{Pr}_2 G^{-1}$	7°, E119
9°	$\text{Prt} F \subset \text{Dom} G$	8°, E127
10°	$\text{Dom}(G \circ F) = \text{Dom} F$	9°, F27
11°	$\text{Dom}(G \circ F) = \text{Prt} G$	10°, 6°
12°	$\text{Dom}(G \circ F) = \text{Dom}(\text{Id}_{\text{Prt} G})$	11°, F23
13°	$x \in \text{Dom}(G \circ F)$	hyp
14°	$x \in \text{Dom} F$	13°, 10°
15°	$(x, F @ x) \in F$	14°, F7
16°	$(F @ x, x) \in G$	15°, 7°, E125
17°	$x = G @ (F @ x)$	16°, F6
18°	$x \in \text{Prt} G$	13°, 11°
19°	$x = \text{Id}_{\text{Prt} G} @ x$	18°, F23
20°	$G @ (F @ x) = \text{Id}_{\text{Prt} G} @ x$	17°, 19°
21°	$\forall x (x \in \text{Dom}(G \circ F) \Rightarrow G @ (F @ x) = \text{Id}_{\text{Prt} G} @ x)$	20°, $\Rightarrow$ -intro, S62
22°	$G \circ F = \text{Id}_{\text{Prt} G}$	5°, 12°, 21°, F22
23°	F est une fonction et $G \circ F = \text{Id}_{\text{Prt} G}$	4°, 22°
24°	$\exists F (F \text{ est une fonction et } G \circ F = \text{Id}_{\text{Prt} G})$	23°, S66
25°	$\exists F (F \text{ est une fonction et } G \circ F = \text{Id}_{\text{Prt} G})$	3°, 24°, S64.

F33 Se déduit par logique propositionnelle de F32 et du théorème suivant :

$\vdash$ Si F est une fonction, alors $[F^{-1}$ est une fonction $\Leftrightarrow \forall x \forall x' ((x \in \text{Dom} F \text{ et } x' \in \text{Dom} F \text{ et } F @ x = F @ x') \Rightarrow x = x')]$. (*)

Démonstration de (*)

1°	F est une fonction	hyp
2°	F^{-1} est un graphe	F1, E124
	F^{-1} est une fonction $\Leftrightarrow (F^{-1}$ est un graphe et $\forall y \text{Hx}((y, x) \in F^{-1}))$	F1
	$\Leftrightarrow \forall y \text{Hx}((y, x) \in F^{-1}))$	2°, S60.1
	$\Leftrightarrow \forall y \forall x \forall x' (((y, x) \in F^{-1} \text{ et } (y, x') \in F^{-1}) \Rightarrow x = x')$	S93
	$\Leftrightarrow \forall y \forall x \forall x' (((x, y) \in F \text{ et } (x', y) \in F) \Rightarrow x = x')$	E125
	$\Leftrightarrow \forall x \forall x' \forall y (\text{non}((x, y) \in F \text{ et } (x', y) \in F) \text{ ou } x = x')$	S72, S53
	$\Leftrightarrow \forall x \forall x' (\exists y((x, y) \in F \text{ et } (x', y) \in F) \Rightarrow x = x')$	S81, S71
	$\Leftrightarrow \forall x \forall x' (\exists y(x \in \text{Dom} F \text{ et } y = F @ x \text{ et } x' \in \text{Dom} F \text{ et } y = F @ x') \Rightarrow x = x')$	F6
	$\Leftrightarrow \forall x \forall x' ((x \in \text{Dom} F \text{ et } x' \in \text{Dom} F \text{ et } F @ x = F @ x') \Rightarrow x = x')$	S92.

F34 Immédiat d'après F32, F23 et E137.

F35 Immédiat d'après F32 et E126.

F36

1°	F et G sont des fonctions injectives	hyp
2°	$G \circ F$ est une fonction	1°, F32, F26
3°	$(G \circ F)^{-1} = F^{-1} \circ G^{-1}$	E188
4°	F^{-1} et G^{-1} sont des fonctions	1°, F32
5°	$F^{-1} \circ G^{-1}$ est une fonction	4°, F26
6°	$(G \circ F)^{-1}$ est une fonction	3°, 5°
7°	$(G \circ F)^{-1}$ est une fonction injective	2°, 6°, F32.

F37

1°	F est une fonction	hyp
2°	F^{-1} est un graphe	F1, E124
	F^{-1} est une fonction $\Leftrightarrow F^{-1} \circ (F^{-1})^{-1} = \text{Id}_{\text{Pr} F^{-1}}$	F3, 2°, S60.1
	$\Leftrightarrow F^{-1} \circ F = \text{Id}_{\text{Dom} F}$	E126, E127.

F39.1

1°	$F : A \rightarrow B$	hyp
2°	$F : A \xrightarrow{\text{inj}} B$	hyp
3°	F est une fonction injective	2°, F38
4°	F^{-1} est une fonction	3°, F32
5°	$\forall y \text{Hx}((y, x) \in F^{-1})$	4°, F1
6°	$\text{Hx}((x, y) \in F)$	5°, S65, E125
7°	$y \in B \Rightarrow \text{Hx}((x, y) \in F)$	6°, S21
8°	$\forall y(y \in B \Rightarrow \text{Hx}((x, y) \in F))$	7°, S62
9°	$\forall y(y \in B \Rightarrow \text{Hx}((x, y) \in F))$	hyp
10°	$y \in B$	hyp
11°	$\text{Hx}((x, y) \in F)$	9°, 10°
12°	$y \notin B$	hyp
13°	$\text{Pr} F \subset B$	1°, F8
14°	$y \notin \text{Pr} F$	12°, 13°, E4
15°	$\text{non} \exists x((x, y) \in F)$	14°, E117

16°	$\text{non} \exists x((x, y) \in F) \text{ ou } \exists ! x((x, y) \in F)$	15°, ou-intro
17°	$\text{Hx}((x, y) \in F)$	16°, S95
18°	$\text{Hx}((x, y) \in F)$	11°, 17°, disj cas
19°	$\forall y \text{Hx}((y, x) \in F^{-1})$	18°, S62, E125
20°	F est une fonction	1°
21°	F^{-1} est un graphe	20°, F1
22°	F^{-1} est une fonction	19°, 21°, F1
23°	F est une fonction injective	20°, 22°, F32
24°	$F : A \xrightarrow{\text{inj}} B$	1°, 23°, F38
25°	$F : A \xrightarrow{\text{inj}} B \Leftrightarrow \forall y(y \in B \Rightarrow \text{Hx}((x, y) \in F))$	8°, 24°, S40.

F39.2

1°	$F : A \rightarrow B$	hyp
2°	$F : A \xrightarrow{\text{sur}} B$	hyp
3°	$\text{Prt} F = B$	2°, F38
4°	$B \subset \text{Prt} F$	3°
5°	$\forall y(y \in B \Rightarrow y \in \text{Prt} F)$	4°, E3
6°	$\forall y(y \in B \Rightarrow \exists x((x, y) \in F))$	5°, E117
7°	$\forall y(y \in B \Rightarrow \exists x(x \in \text{Dom} F \text{ et } y = F @ x))$	6°, F6
8°	$\forall y(y \in B \Rightarrow \exists x(x \in A \text{ et } y = F @ x))$	7°, 1°
9°	$\forall y(y \in B \Rightarrow \exists x(x \in A \text{ et } y = F @ x))$	hyp
10°	$B \subset \text{Prt} F$	déductions inverses de 4°–8°
11°	$\text{Prt} F \subset B$	1°, F8
12°	$\text{Prt} F = B$	10°, 11°
13°	$F : A \xrightarrow{\text{sur}} B$	1°, 12°, F38
14°	$(F : A \xrightarrow{\text{sur}} B) \Leftrightarrow \forall y(y \in B \Rightarrow \exists x(x \in A \text{ et } y = F @ x))$	8°, 13°, S40.

F39.3

On désigne par R la proposition $(x, y) \in F$

1°	$F : A \rightarrow B$	hyp
	$(F : A \xrightarrow{\text{bij}} B) \Leftrightarrow (F : A \xrightarrow{\text{inj}} B \text{ et } F : A \xrightarrow{\text{sur}} B)$	F38
	$\Leftrightarrow (\forall y(y \in A \Rightarrow \text{Hx} R) \text{ et } \forall y(y \in A \Rightarrow \exists x R))$	F39 (1er et 2ème)
	$\Leftrightarrow \forall y((y \in A \Rightarrow \text{Hx} R) \text{ et } (y \in A \Rightarrow \exists x R))$	S74
	$\Leftrightarrow \forall y(y \in A \Rightarrow (\text{Hx} R \text{ et } \exists x R))$	S58
	$\Leftrightarrow \forall y(y \in A \Rightarrow \exists ! x R)$	S94.

F40.1

1°	$F : A \xrightarrow{\text{inj}} B \text{ et } G : B \xrightarrow{\text{inj}} C$	hyp
2°	$F : A \rightarrow B \text{ et } G : B \rightarrow C$	1°, F38
3°	$G \circ F : A \rightarrow C$	2°, F28
4°	F et G sont des fonctions injectives	1°, F38
5°	$G \circ F$ est une fonction injective	4°, F36
6°	$G \circ F : A \xrightarrow{\text{inj}} C$	3°, 5°, F38.

F40.2

1°	$F : A \xrightarrow{\text{sur}} B \text{ et } G : B \xrightarrow{\text{sur}} C$	hyp
2°	$F : A \rightarrow B \text{ et } G : B \rightarrow C$	1°, F38
3°	$G \circ F : A \rightarrow C$	2°, F28
4°	$\text{Prt} F = B \text{ et } \text{Prt} G = C$	1°, F38
5°	$\text{Prt}(G \circ F) = G(\text{Prt} F)$	E189

6°	$= G\langle B \rangle$	4°
7°	$= G\langle \text{Dom} G \rangle$	2°
8°	$= \text{Prt} G$	E152
9°	$\text{Prt}(G \circ F) = C$	8°, 4°
10°	$G \circ F : A \xrightarrow{\text{sur}} C$	3°, 9°, F38.

F41

1°	$F : A \xrightarrow{\text{bij}} B$	hyp
2°	F est une fonction injective	1°, F38
3°	$\text{Dom} F = A$ et $\text{Prt} F = B$	1°, F38
4°	F^{-1} est une fonction injective	2°, F35
5°	$\text{Dom}(F^{-1}) = B$ et $\text{Prt}(F^{-1}) = A$	3°, E127
6°	$F^{-1} : B \xrightarrow{\text{bij}} A$	4°, 5°, F38
7°	$F^{-1} \circ F = \text{Id}_A$	2°, 3°, F37.

BIBLIOGRAPHIE

Cette bibliographie sommaire n'a pour but que d'indiquer au lecteur intéressé quelques points de départ, soit pour une étude plus approfondie de la logique mathématique, soit pour une exploration des applications de la logique en informatique. La littérature sur ces applications—à part celle qui concerne le domaine de la spécification et de la vérification formelle—suppose une certaine formation en logique mathématique et notre recommandation est de commencer dans ce domaine par un ouvrage classique tel que ceux qui figurent ci-dessous.

Ouvrages généraux. Nous mentionnons ici quelques ouvrages généraux sur la logique. Le premier [1] est récent et comporte des applications à l'informatique (logique des programmes). Les références [2–7] sont des introductions classiques à la logique mathématique. Notre manière de présenter les déductions logiques est influencée par [7]. Le recueil [8] est une encyclopédie des domaines actuels de la logique, notamment des extensions de la logique classique.

- [1] P. Gochet, P. Gribomont, *Logique*, Vol. 1–2, Hermès, Paris, 1991 et 1994.
- [2] S. C. Kleene, *Logique mathématique*, Armand Colin, Paris, 1971.
- [3] W. V. O. Quine, *Méthodes de logique*, Armand Colin, Paris, 1972.
- [4] A. Church, *Introduction to Mathematical Logic*, Princeton University Press, Princeton, 1956.
- [5] H. B. Enderton, *A Mathematical Introduction to Logic*, Academic Press, New York, 1970.
- [6] J. R. Shoenfield, *Mathematical Logic*, Addison-Wesley, Reading, 1967.
- [7] H. Hermes, *Einführung in die Mathematische Logik*, B. G. Teubner, Stuttgart, 1972.
- [8] D. Gabbay, F. Guenther (Editors), *Handbook of Philosophical Logic*, Vol. 1–4, D. Reidel Publishing Company, Dordrecht, 1983–1986.

Sujets particuliers. Gentzen [9] a introduit en logique formelle le système de la déduction naturelle. La notation emboîtée des démonstrations naturelles est utilisée dans [10]. Leisenring [11] est un ouvrage de logique mathématique consacré aux opérateurs de sélection.

- [9] G. Gentzen, *Untersuchungen über das logische Schließen*, Mathematische Zeitschrift, Vol. 39 (1934–35), pp. 176–210, 405–431.

- [10] J.-B. Grize, *Logique moderne*, Fascicule I, Mouton et Gauthier-Villars, Paris, 1972.
- [11] A. C. Leisenring, *Mathematical Logic And Hilbert's ε -Symbol*, MacDonal, London, 1969.

Spécification et vérification formelle des programmes. Les livres suivants présentent des méthodes formelles de développement de programmes basées sur la logique et la théorie des ensembles.

- [12] J.-R. Abrial, *The B-Book*, Cambridge University Press, Cambridge, 1996.
- [13] J. M. Spivey, *The Z Notation: A Reference Manual*, Second Edition, Prentice Hall, New York, 1992.
- [14] J. Woodcock, J. Davies, *Using Z: Specification, Refinement, and Proof*, Prentice Hall, London, 1996.
- [15] A. Diller, *Z: An Introduction to Formal Methods*, John Wiley & Sons, Chichester, 1994.
- [16] C. B. Jones, *Systematic Software Development using VDM*, Second Edition, Prentice Hall, Englewood Cliffs NJ, 1990.
- [17] R. C. Backhouse, *Construction et vérification des programmes*, Masson, Paris, 1989.
- [18] Z. Manna, R. Waldinger, *The Logical Basis for Computer Programming*, Vol. 1–2, Addison-Wesley, Reading, 1985 et 1990.
- [19] E. W. Dijkstra, C. S. Scholten, *Predicate Calculus and Program Semantics*, Springer-Verlag, New York, 1990.

Logique et mathématique discrète. De nombreux livres de « mathématique discrète » contiennent une introduction à la logique. Nous en mentionnons trois. Le premier [20] s'adresse spécialement aux informaticiens. Le livre [21] traite différents sujets de mathématique discrète de manière très formelle, mais basée sur une présentation de la logique différente de la nôtre.

- [20] M. Marchand, *Mathématique discrète*, De Boeck-Wesmael, Bruxelles, 1989.
- [21] D. Gries, F. B. Schneider, *A Logical Approach to Discrete Math*, Springer-Verlag, New York, 1993.
- [22] S. S. Epp, *Discrete Mathematics with Applications*, Wadsworth, Belmont CA, 1990.

Logique et intelligence artificielle. Un domaine de recherche en informatique dans lequel la logique joue un rôle important est celui de l'intelligence artificielle. Les quatre volumes de [23] contiennent un survol de ces applications. Une introduction au domaine particulier de la démonstration automatique est fournie par [24], alors que [25] est une référence classique dans ce domaine. La programmation logique (PROLOG) est étroitement liée à la démonstration

automatique. Nous mentionnons deux ouvrages [26], [27] consacrés aux fondements logiques de cette forme de programmation. Le recueil [28] fournit une vue d'ensemble des applications de la logique en intelligence artificielle.

- [23] A. Thayse et co-auteurs, *Approche logique de l'intelligence artificielle*, Vol. 1–4, Dunod, Paris, 1988–1991.
- [24] L. Wos, R. Overbeek, E. Lusk, J. Boyle, *Automated Reasoning : Introduction and Applications*, Second Edition, McGraw-Hill Inc., New York, 1992.
- [25] C.-L. Chang, R. Lee, *Symbolic Logic and Mechanical Theorem Proving*, Academic Press, New York, 1973.
- [26] J. W. Lloyd, *Fondements de la programmation logique*, Eyrolles, Paris, 1988.
- [27] A. Nerode, R. A. Shore, *Logic for Applications*, Springer-Verlag, New York, 1993.
- [28] D. Gabbay, C. J. Hogger, J. A. Robinson (Editors), *Handbook of Logic in Artificial Intelligence and Logic Programming*, Vol. 1–4, Clarendon Press, Oxford, 1993–1995.

Programmation fonctionnelle. La notation « λ » des fonctions introduite par Church [29] est à la base des langages de programmation fonctionnels, sur lesquels nous donnons ici trois références.

- [29] A. Church, *The Calculi of Lambda-conversion*, Princeton University Press, 1941.
- [30] R. Bird, P. Wadler, *Introduction to Functional Programming*, Prentice Hall, Englewood Cliffs NJ, 1988.
- [31] G. Cousineau, M. Mauny, *Approche fonctionnelle de la programmation*, Ediscience international, Paris, 1995.
- [32] S. L. Peyton Jones, *The Implementation of Functional Programming Languages*, Prentice Hall, Englewood Cliffs NJ, 1987.

INDEX

A fortiori, 104, 106, 121, 123, 133

Application, 384

- bijective, 395
- injective, 395
- surjective, 395

Arbre, 59

- d’une démonstration, 113
- d’une expression, 58

Arité

- d’un symbole fonctionnel, 40
- d’un symbole relationnel, 41

Arithmétique, 51, 90, 198

Assertion, 86

Association, 132, 165

Associativité, 162

Axiome, 88

- de choix, 393
- définitionnel, 313, 316, 318
- de l’ensemble vide, 190, 263
- de l’inclusion, 188, 263
- d’extensionnalité, 267
- en tant que théorème, 98

Barbier, 241

Barre

- de déduction, 81
- de négation, 43

Base d’une démonstration, 94, 110

Bloc d’une démonstration, 117

Chaîne

- d’équivalences, 153
- d’implications et équivalences, 154

Changement de variable liée, 245, 249, 253, 304, 339, 354, 362, 374, 389

Collection 351, 359, 371, 373

- de couples, 377
- non existante, 368
- simple, 359, 360, 366

Commutativité, 162

Composition

- de fonctions, 392
- de graphes, 379

Conclusion, 81

Condition

- de réversibilité d’une substitution, 247
- nécessaire, 177, 180
- suffisante, 177, 180

Conjonction, 45

Connecteur logique, 29, 42

- booléen, 158

Constante individuelle, 29, 30, 33, 37

Construction génératrice, 54

- de propositions, 57
- de termes, 56

Contexte, 23, 79, 98

- contradictoire, 99

Contradiction, 140

Contraposée, 123

Contraposition, 146, 176

Corps d’une démonstration, 94, 110

Couple, 320, 322

Déduction, 79, 96

- dérivée, 93
- élémentaire, 89
- identique, 97, 133
- structurelle, 131, 208
- triviale, 96

Définition, 188, 198, 313

Démonstration, 93, 94, 125

- d’une déduction, 98
- d’une proposition, 98
- d’un théorème, 98
- naturelle, 116, 117

De Morgan A., 159, 161, 167, 232, 234

Diagramme syntaxique, 19

- des termes et propositions, 53

- Disjonction, 44
 - de cas contraires, 147
 - des cas, 104, 107, 121
- Distribution
 - de quantificateurs, 237
 - d’une implication, 176
- Distributivité, 162
- Domaine
 - de l’interprétation, 29
 - du discours, 29
 - d’une fonction, 382
- Dualité, 159, 161, 166, 169, 232, 236
- Élimination
 - de *et*, 104
 - de quantificateur avec une égalité, 294
 - de $\Rightarrow$, 131
 - de $\Leftrightarrow$, 104
 - de $\exists$, 192, 195, 220
 - de $\forall$, 203
- Ensemble
 - à deux éléments, 295
 - à trois éléments et plus, 297
 - à un seul élément, 296
 - complémentaire, 273
 - des parties d’un ensemble, 298
 - vide, 190, 263
- Équivalence, 47, 176
- $\exists!x$, 301
- Exemplification, 259, 262
- Ex falso quodlibet*, 104, 109
- Existence, 343
 - d’une collection, 359, 363, 364, 372, 378
- Expressions, 15
- Extension, 100
 - conservative, 316
 - définitionnelle, 313
- Fonction, 381
 - identique, 391
 - injective, 394
 - vide, 391
- Forme prénex, 274, 276
- Généralisation, 192, 194, 204
- Gödel K., 100
- Graphe, 320, 324
 - identique, 358
 - réciproque, 357
- Hilbert D., 13, 131
- Hx, 301
- Hypothèse, 84, 104, 106, 135
 - de désignation, 292
- Idempotence, 162
- Implication, 45, 170
- Intersection, 271
- Introduction,
 - de *et*, 104
 - de *ou*, 104
 - de $\Rightarrow$, 104, 108, 121
 - de $\Leftrightarrow$, 104, 150
 - de $\perp$, 104, 108
 - de $\exists$, 192, 193, 212, 218
- Jugement, 84, 86
 - sans hypothèses, 133
- Jugements équivalents, 97
- Lambda (λ), 387
- Langage, 15
 - déclaratif, 23
 - de spécification, 2, 185
 - du premier ordre, 27, 53
 - d’ordre supérieur, 34
 - formel, 7, 15, 17
 - universel, 198, 199
- Lien, 62, 63
- Logique
 - des prédicats, 191
 - des prédicats avec égalité, 281
 - élémentaire, 10
 - intuitionniste, 126, 128
 - mathématique, 10, 315
 - propositionnelle, 103
- Lpréd, 191
- Lprédégal, 200, 281
- Lprop, 103
- Mathématique
 - constructive, 43, 130
 - formelle, 8
- Métalangage, 24, 46, 57
- Métavariable, 42, 48, 107
- Méthode formelle, 1
- Modus ponens*, 104, 150
- Négation, 42
- Notation
 - emboîtée, 119

- infixée, 38
 - postfixée, 38
 - polonaise, 37
 - préfixée, 37, 44
- Occurrence, 16
- libre d'une variable, 63
 - liée d'une variable, 63
- Opérateur, 329
- d'abstraction fonctionnelle, 384, 387
 - de collection général, 370, 371
 - de collection simple, 332, 359
 - de réunion, 351, 352
 - de sommation, 330
 - de test de condition, 333
 - d'intégration, 332
- Opération d'application d'une fonction à un argument, 38, 382
- Particularisation, 192, 193, 200
- Particularisations simultanées, 253, 256
- Permutation
- de jugements, 97
 - de lignes d'une démonstration, 112
 - de quantificateurs, 237
 - d'hypothèses, 135
- Portée d'une fonction, 382
- Prédicat, 191, 322
- Prémisse, 81
- Produit cartésien, 355
- Projections
- d'un couple, 323, 339
 - d'un graphe, 356
- Proposition, 22
- atomique, 31
 - composée, 31
 - d'un langage du premier ordre, 53, 57
 - fausse, 98
 - fausse par excellence, 42, 47, 104, 139
 - indécidable, 99
 - vraie, 98
- Propositions
- duales, 159, 167
 - équivalentes, 155
- Prt, 382
- Pseudo-prédicat « existe », 343
- Quadruplet, 324
- Quantificateur
- d'unicité, 300
 - existentiel, 48
 - universel, 29, 48
 - typé, 50, 278
- Réduction à l'absurde
- J, 104, 109, 121
 - K, 104, 109, 127, 128
- Réflexivité
- de l'égalité, 282
 - de l'équivalence, 152
 - de l'implication, 141
- Relation
- d'appartenance, 186
 - d'inclusion, 188, 263
 - transitive, 266
- Réunion, 271, 351
- Réversibilité, 247
- Russel B., 241, 244
- Schéma de déduction, 83, 88
- booléen, 158
 - de récurrence, 90
 - de réunion, 353
 - de séparation, 363
 - élémentaire, 104, 192, 282, 301, 338, 353, 363
 - structurel, 106
- Sélecteur, 335, 338
- Sémantique, 17
- de l'implication, 46
 - d'un langage du premier ordre, 29
 - d'un langage formel, 20
- Simplification, 177
- Squelette d'une proposition, 251
- Substitution, 66, 67
- avec création de lien, 72
 - dans un schéma de déduction, 106
- Substitutions simultanées, 253
- Substitutivité
- de l'égalité, 282, 287, 289, 290, 349, 354, 374, 389
 - de l'équivalence, 155, 157, 229, 290, 309, 338, 348, 362, 374
- Sx, 335
- Symbole, 15
- d'un langage du premier ordre, 27
 - fonctionnel, 29, 30, 37

- fonctionnel binaire, 37
- fonctionnel unaire, 39
- primordial, 340
- relationnel, 29, 30, 40
- relationnel binaire, 40
- relationnel ternaire, 41
- relationnel unaire, 41

Symétrie

- de l'égalité, 286
- de l'équivalence, 152, 160

Syntaxe, 17, 18

- d'un langage du premier ordre, 53

Tautologie, 134

Terme, 22

- composé, 37
- conditionnel, 346
- descriptif, 336, 337
- d'un langage du premier ordre, 53, 56
- librement substituable, 71, 74

Théorème, 97

- de Russel, 368
- du barbier, 241
- sans hypothèses, 98

Théorie, 88

- contradictoire, 99
- logique, 90

Tiers exclu, 126, 147

Transformé d'un ensemble par un graphe, 364

Transitivité

- de l'égalité, 287
- de l'équivalence, 152
- de l'implication, 122, 141
- de l'inclusion, 266
- des jugements, 136, 211, 212

Triplet, 324

Type d'un opérateur, 330, 333

Unicité, 300

- d'une collection, 360

Univers

- de la théorie des ensembles, 185
- de l'interprétation d'un langage du premier ordre, 29
- du discours, 29

Variable individuelle, 29, 33

- du répertoire standard, 34
- globale, 61, 63

- libre, 61, 64, 65

- locale, 61

- réversiblement substituable, 246

Variables liées par un opérateur, 334

Vérité, 79

- *a fortiori*, 104
- d'une hypothèse, 104, 135