

Logique et Preuves

Pierre Castéran, Philippe Duchon, et d'autres
<http://www.labri.fr/~duchon/Enseignements/L-et-P/>

18 novembre 2019

Table des matières

1	Introduction	4
1.1	De la logique en général	4
1.2	La logique et l'informatique	4
1.3	Cadre de ce cours	5
1.3.1	Références	5
1.4	Divers	5
I	Le raisonnement logique	6
2	Le vocabulaire de la logique	7
2.1	Propositions	7
2.1.1	Notion intuitive et floue de proposition	7
2.1.2	Un cadre formel	7
2.2	L'importance du contexte : la notation de séquent	9
2.2.1	Séquents	10
2.3	Vérité et Preuve	11
2.3.1	Point de vue sémantique	11
2.3.2	Le point de vue syntaxique : notion de preuve	13
2.3.3	Notre point de vue (un rien subjectif)	13
3	La logique minimale	15
3.1	Les propositions	15
3.1.1	Remarques	16
3.2	Approche sémantique de la logique minimale	17
3.2.1	Définitions	17
3.3	Preuves en logique minimale	19
3.3.1	Règle d'hypothèse	19
3.3.2	Règle du modus ponens (élimination de l'implication)	20
3.3.3	Preuves en déduction naturelle	21
3.3.4	Règle d'abstraction (introduction de l'implication)	22
3.4	Simplifier l'écriture des preuves	26
3.4.1	Règles dérivées	26
3.4.2	Règle d'affaiblissement	31

3.4.3	Prouvabilité et règles dérivées	31
3.5	Preuves <i>vs</i> validité	32
3.5.1	Correction	32
3.5.2	Incomplétude	35
3.5.3	Théorème de substitution uniforme	36
3.6	Stratégies de démonstration	37
3.7	Logique et Programmation Fonctionnelle	39
3.7.1	Types et propositions	39
3.7.2	Termes et Preuves	39
4	Logique Propositionnelle (intuitionniste, puis classique)	42
4.1	La contradiction et la négation	42
4.1.1	Introduction	42
4.1.2	Syntaxe	43
4.1.3	Aspects sémantiques de la contradiction	44
4.1.4	Le principe d'explosion	45
4.1.5	Règles associées à la négation	45
4.1.6	Logique intuitionniste : définition	47
4.1.7	Peut-on prouver $\perp$?	47
4.1.8	Autour de la double négation	48
4.2	Les connecteurs binaires	48
4.2.1	La conjonction	48
4.2.2	La disjonction	50
4.2.3	L'équivalence logique	52
4.2.4	Équivalence entre propositions	53
4.3	Quelques méta-théorèmes	53
4.3.1	Cohérence et non-complétude	53
4.3.2	Relation avec la logique minimale	54
4.3.3	Remplacement d'une sous-formule	55
4.4	La logique propositionnelle classique	56
4.4.1	Règles dérivées pour la logique classique	56
4.5	Quelques méta-théorèmes pour la logique classique	57
4.5.1	Cohérence et complétude	57
4.5.2	Relation avec la logique intuitionniste	57
4.5.3	Quelques règles dérivées en logique classique	58
5	Calcul des prédicats (Logique du premier ordre)	60
5.1	Introduction	60
5.2	Formules du premier ordre	60
5.2.1	Types	60
5.2.2	Constantes et symboles de fonctions	61
5.2.3	Variables, Déclarations et Contextes	62
5.2.4	Termes	62
5.2.5	Règles de typage	63
5.2.6	Prédicats et Relations	63
5.2.7	Syntaxe des propositions	65

5.2.8	Variables libres et variables liées	65
5.2.9	Substitution	67
5.2.10	Substitution simultanée	68
5.2.11	Le coin du formaliste	68
5.3	Preuves en logique du premier ordre	69
5.3.1	Règles associées à l'égalité	69
5.3.2	Règles associées au quantificateur universel	70
5.3.3	Notation de bloc pour l'introduction du quantificateur universel	72
5.3.4	Quantifications imbriquées	72
5.3.5	Règles associées au quantificateur existentiel	73
5.4	Pratique de la logique du premier ordre	75
5.4.1	Quantificateurs et négation	75
5.4.2	Exemples	76
5.4.3	Règles dérivées	77
5.5	Sémantique	78
5.5.1	Structures d'interprétation	78
5.5.2	Valuations	78
5.5.3	Sémantique des formules	79
5.6	Méta-théorèmes principaux	79
II	Spécifier et prouver des programmes fonctionnels	80
6	Quelques types inductifs	82
6.1	Les valeurs booléennes	82
6.1.1	Règles spécifiques au type <code>bool</code>	82
6.1.2	Définitions de fonctions par cas	83
6.1.3	Exemples	83
6.2	Entiers naturels	83
III	Annexes	84
7	TO DO	85

Chapitre 1

Introduction

1.1 De la logique en général

La logique est une science dont l'objet est l'étude du raisonnement, abstraction faite du domaine auquel ce raisonnement s'applique : linguistique, mathématiques, propriétés des programmes informatiques, etc. Cette discipline est très ancienne, les premiers ouvrages sur la logique ayant été écrits par Aristote (384-322 avant J.C.). On pourra lire avec profit les ouvrages de Gilles Dowek [5] et de Jean-Pierre Belna [2].

Le côté obscur de la logique (la théorie de la mauvaise foi) est traité avec humour dans le petit livre d'Arthur Schopenhauer *L'Art d'avoir toujours raison* [7]. Trente-huit *stratagèmes* pour avoir raison dans une discussion y sont présentés. Un exercice intéressant est de voir lesquels de ces stratagèmes correspondent à des règles de raisonnement mal appliquées ou carrément dévoyées.

1.2 La logique et l'informatique

La logique, après avoir été au centre de la problématique des fondements des mathématiques [9], a pris une importance considérable avec le développement de l'informatique. En effet, la nécessité de *prouver* que tel logiciel ou composant matériel correspond bien à sa *spécification* conduit au développement d'outils informatiques permettant la vérification de programmes. On pourra trouver dans la préface de [3] un historique de ces outils.

Ces outils permettent de mener à bien des preuves qui, dans le cas de la vérification de programmes de taille courante, peuvent être très longues et complexes. Aucun être humain n'est en mesure d'en écrire et/ou en vérifier toutes les étapes.

Manipuler ces outils afin de *certifier* du logiciel ou du matériel (cartes, circuits) demande deux types de compétences :

- Savoir spécifier en langage non-ambigu (c.-à-d. mathématique) le comportement voulu du composant informatique considéré

- Savoir prouver de façon incontestable que le dit composant satisfait sa spécification.

La seconde de ces compétences ne peut être totalement automatisée dans la plupart des cas. La théorie de la calculabilité (qui sort très largement du cadre de ce cours !) permet de montrer cette impossibilité d’automatiser la preuve de programmes.

La preuve de correction d’un composant informatique peut donc exiger de la part des ingénieurs du logiciel la connaissance des règles de base du raisonnement, afin de guider l’outil en lui donnant des stratégies de démonstration (par exemple, suggérer un raisonnement par l’absurde ou par récurrence).

1.3 Cadre de ce cours

Ce cours est destiné aux étudiants de licence d’Informatique et de Mathématiques et Informatique. On privilégiera les exemples pris dans le domaine du logiciel, et plus particulièrement dans la sûreté du logiciel. Seules des connaissances de base de la programmation et de mathématiques discrètes seront nécessaires.

On étudiera des systèmes logiques de plus en plus expressifs, déterminés par les composants suivants :

- Un langage (formel) de *propositions* correspondant aux énoncés dont l’éventuelle vérité nous importe,
- Une *sémantique* permettant d’associer une *valeur de vérité* à toute proposition.
- Un langage (formel, aussi) permettant d’écrire des *preuves* de ces énoncés. On insistera sur la parenté entre langages de programmation et langages de description de preuves.
- Des *tactiques* et *stratégies* de démonstration, facilitant la recherche de la preuve d’un énoncé donné.

La plupart des séances se feront en salle de TD ; un petit nombre de séances seront consacrées à la présentation et l’apprentissage de l’*assistant de preuve* Coq[8, 3].

1.3.1 Références

Nous conseillons la consultation régulière du site du cours de Mathématiques pour l’Informatique de Christine Paulin-Mohring [6].

1.4 Divers

Merci de me signaler par courrier électronique les (certainement nombreuses) erreurs dans ce document.

Première partie

Le raisonnement logique

Chapitre 2

Le vocabulaire de la logique

2.1 Propositions

2.1.1 Notion intuitive et floue de proposition

Il est très difficile de définir de façon générale la notion de *proposition*. On trouve souvent des descriptions genre [1]

Une proposition est un énoncé (mathématique, informatique, mais pas que) susceptible d’être démontré ou réfuté, pour lequel il fait sens de parler de vérité.

En voici quelques exemples, pris dans des domaines divers : certaines sont vraies, d’autres fausses, et la dernière est encore une conjecture.

- “Mon mot de passe est ”password” ”
- “42 est un nombre premier”
- “41 et 43 sont des nombres premiers jumeaux”
- “Le langage C comporte plusieurs failles de sécurité” [4]
- “*Quicksort* est un algorithme correct de tri de tableaux”
- “Le programme de recherche binaire de la figure 2.2 page suivante est incorrect”
- “La fonction de tri de la figure 2.3 page 9 est correcte”
- “La suite de Syracuse (Figure 2.1) finit toujours par le cycle (1, 2, 4) quel que soit l’élément initial”

En revanche, voici des exemples de non-propositions :

- 42
- $b^2 - 4ac$
- L’algorithme *QuickSort*
- L’ensemble des entiers naturels

2.1.2 Un cadre formel

Pour ne pas nous disperser, la majorité des exemples que nous prendrons viendront des applications informatiques (spécification et vérification du logiciel)

initialisation On considère un nombre entier strictement positif.

calcul de l'élément suivant Soit n un élément de la suite : Par construction, il est strictement positif

- Si n est pair, l'élément suivant est $n/2$
- Sinon, l'élément suivant est $3n + 1$

Par exemple, si l'on démarre de 23, on obtient la suite suivante :

23, 70, 35, 106, 53, 160, 80, 40, 20, 10, 5, 16, 8, 4, 2, 1, 4, 2, 1, 4, 2, 1...

FIGURE 2.1 – La suite de Syracuse

```
int binary_search(long t[], int n, long v) {
    int low = 0, high = n - 1;
    while (low <= high) {
        int middle = (low + high) / 2;
        if (t[middle] < v)
            low = middle + 1;
        else if (t[middle] > v)
            high = middle - 1;
        else return middle;
    }
    return -1;
}
```

FIGURE 2.2 – Un programme buggé

et de mathématiques simples (inévitables en programmation).

De plus, les propositions que nous étudierons appartiendront à des *langages*, qui à l'instar du langage mathématique et des langages de programmation, disposent d'une *syntaxe* non-ambigüe. Il n'existe qu'une seule façon de lire une expression de ces langages. Comme pour les langages de programmation, des conventions (précédence et associativité des opérateurs) permettent de trouver un compromis entre lourdeur des notations et non-ambiguïté.

Il existe de nombreuses variantes de la notion de proposition, suivant les besoins à satisfaire. En voici une énumération non exhaustive.

- Logique *minimale* : permet de comprendre la notion d'implication (notée $\rightarrow$ dans ce cours).
- Logique *propositionnelle* : extension de la logique minimale par les connecteurs $\vee$, $\wedge$, $\neg$, $\leftrightarrow$, et les propositions $\perp$ (toujours fausse) et $\top$ (toujours vraie).
- Logique du premier ordre : définit les notions de prédicats, de relation (dont l'égalité), et la quantification sur des objets

```

let rec insert x l =
  match l with
  [] -> [x]
  | y::l' -> if x <= y
              then x::l
              else y::insert x l'

let rec sort l = match l with
  [] -> []
  | x::l' -> insert x (sort l')

let _ = sort [0; 2; 4; 3; 4; 5; 9]

```

FIGURE 2.3 – Un programme correct

- Logique d'ordre supérieur : on autorise la quantification sur des fonctions, des relations et des ensembles

Certaines logiques répondent à des besoins spécifiques :

- Logique de Hoare : preuve de programmes impératifs (contenant des affectations)
- Logique de séparation : extension de la logique de Hoare pour les programmes manipulant des pointeurs et des structures de données allouées dynamiquement (comme en C)
- Logiques temporelles : permettent de considérer des systèmes *réactifs* et de raisonner sur des suites d'*événements*.
- Logiques épistémiques : tel processeur peut n'avoir qu'une connaissance *partielle* de l'état d'un réseau.

Chacun de ces systèmes, à l'instar des langages de programmation, est muni d'une syntaxe précise et non ambiguë, que nous présenterons pour chaque logique étudiée dans ce cours.

2.2 L'importance du contexte : la notation de séquent

Une proposition comme “ $x = 2$ ”, prise isolément, ne peut être considérée en absolu comme vraie ou fausse. Tout dépend en effet de ce qu'on sait ou suppose sur x : est-ce un entier ? un nombre réel ? la valeur courante d'une variable d'un programme C ?

Dès qu'on veut être rigoureux, on se doit de rendre *explicite* le *contexte* dans lequel on se place. Cela consiste à préciser dans quel univers (nombres, rela-

tions, programmes, etc.) on se place, et quelles propriétés sont admises (analyse mathématique, théorie des ensembles, spécification d'un langage de programmation, etc.), mais aussi des hypothèses de travail spécifiques à une situation donnée (tel tableau est trié jusqu'à l'indice i , par exemple).

2.2.1 Séquents

Définition

Un *séquent* est une structure composée :

- d'un *contexte* formé d'un ensemble Γ de propositions appelées *prémisses* ou *hypothèses*,
- d'une proposition A appelée *conclusion* du séquent

Un tel séquent se note $\Gamma \vdash A$.

Il est d'usage d'écrire les hypothèses séparées par des virgules, et sans accolades, sous la forme $A_1, \dots, A_n \vdash A$ où les A_i sont les prémisses et A la *conclusion*.

Un tel séquent se lit de façon indifférente d'une des trois façons suivantes :

- “Sous les hypothèses $A_1, \dots, A_n$, la conclusion A est vraie” ;
- “La proposition A est une conséquence logique de $A_1, A_2, \dots, A_n$ ” ;
- Si $A_1, \dots, A_n$ sont toutes vraies, alors A est vraie”.

La notion de “validité” d'un séquent sera prise de façon intuitive dans ces exemples, et précisée dans les chapitres suivants.

Exemples

Voici quelques exemples de séquents, plaçant chacun la proposition $x = 2$ dans un contexte différent :

- Le premier séquent affirme que si x est un nombre réel vérifiant l'égalité $2x^2 - 5x + 2 = 0$, alors $x = 2$:

$$x \in \mathbb{R}, 2x^2 - 5x + 2 = 0 \vdash x = 2$$

Intuitivement, ce séquent ne peut être tenu pour « vrai ». En effet, si l'on prend pour x la valeur $1/2$, les hypothèses de ce séquent sont vérifiées, alors que la conclusion est fausse.

- En revanche, en *affaiblissant* la conclusion, on obtient un jugement valide :

$$x \in \mathbb{R}, 2x^2 - 5x + 2 = 0 \vdash x = 2 \vee x = 1/2$$

- Si une hypothèse contraint x à être un entier, nous avons le séquent valide suivant.

$$x \in \mathbb{Z}, 2x^2 - 5x + 2 = 0 \vdash x = 2$$

- Finalement, nous terminons par deux exemples de séquents valides *parce que leurs hypothèses sont incohérentes*.

$$x = 2, x = 3 \vdash x = 42$$

$$x \in \mathbb{R}, x^2 + x + 1 = 0 \vdash x = 42$$

Notations diverses

Traditionnellement, nous utiliserons la méta-variable Γ pour désigner un contexte arbitraire. L’adjonction d’une hypothèse A à un contexte Γ sera notée Γ, A au lieu du plus lourd $\Gamma \cup \{A\}$. L’union de deux contextes Γ et Δ sera notée simplement Γ, Δ .

Si l’ensemble des hypothèses est vide, on notera simplement $\vdash A$ le séquent $\emptyset \vdash A$.

Remarque

Un séquent sans prémisse, de la forme $\vdash A$ se lit “La proposition A est vraie” (sans condition).

Prononciation : Le symbole $\vdash$ se prononce “thèse” en français. Sa ressemblance avec un tourniquet vu de haut le fait appeler « turnstyle » en anglais. Certaines régions vinicoles y voient une forme abstraite de tire-bouchon.

Remarque

Dans le calcul des séquents dit *classique*, on peut avoir un nombre quelconque de propositions à droite du symbole $\vdash$. Nous n’étudierons pas ce calcul dans ce cours, car il s’agit d’une formalisation de la logique très différente.

2.3 Vérité et Preuve

Nous voulons donner un sens précis à la notion de “correction” d’un séquent, juste abordée de façon informelle dans les exemples de la section 2.2.1 page précédente. Nous présentons deux façons de voir.

2.3.1 Point de vue sémantique

Le cas des propositions

Il est communément admis qu’une proposition est vraie ou fausse. La “sémantique” d’une logique précise comment déterminer la “valeur de vérité” d’une proposition. Cette valeur dépend en général d’une *valuation*, c’est à dire une fonction qui attribue une valeur de vérité aux formules les plus élémentaires (*e.g.* les formules atomiques de la logique propositionnelle) à partir desquelles on calcule la valeur de vérité de n’importe quelle proposition (voir par exemple la section 3.2 page 17).

Dans le cas de la logique propositionnelle, des méthodes de calcul comme les tables de vérité ou d’autres méthodes plus sophistiquées (diagrammes binaires de décision) permettent de calculer cette sémantique de façon automatique.

Le cas des séquents

Considérons un séquent $A_1, \dots, A_n \vdash A$. Nous dirons qu'il est *valide* si, chaque fois que *toutes* les hypothèses sont vraies, alors la conclusion A est vraie. On notera $A_1, \dots, A_n \models A$ la propriété « le séquent $A_1, \dots, A_n \vdash A$ est valide ».

Dans la définition précédente, “chaque fois” doit être pris au sens de “dans chaque situation envisageable”, ou, quand on parle de valuation dans le cadre de la logique propositionnelle, “pour toute valuation telle que”. En particulier, s'il est *impossible* que les hypothèses soient toutes vraies (*i.e.* si les hypothèses sont *contradictoires* entre elles), alors le séquent est automatiquement valide.

Il faut bien distinguer un *jugement* de la forme $\Gamma \models A$ ou sa négation $\Gamma \not\models A$, du séquent (objet purement syntaxique) $\Gamma \vdash A$ dont les jugements ci-dessus portent sur la validité.

Ainsi, on a tout à fait le droit d'écrire le séquent $x = 2 \vdash x = 3$, bien qu'on ne puisse pas¹ énoncer le jugement $x = 2 \models x = 3$.

Remarques

1. Soit un séquent sans hypothèses, c'est à dire de la forme $\vdash A$. Alors ce séquent est valide (c'est à dire $\models A$) si et seulement si A est une tautologie.
2. Pour montrer qu'un séquent n'est pas valide, il suffit de donner une valuation telle que toutes les hypothèses sont vraies et la conclusion fausse.
3. Soit un séquent tel qu'il n'existe aucune valuation rendant *toutes* ses hypothèses vraies. Alors ce séquent est valide. *Autre formulation équivalente : pour toute valuation ν , il existe au moins une hypothèse fausse pour cette valuation.* C'est le cas du séquent $x=2, x=3 \vdash x=42$.

Limites de l'approche sémantique

Cette méthode fonctionne bien dans le cas de la logique propositionnelle : tables de vérité, diagrammes de décision binaire, etc.

Elle présente cependant quelques inconvénients :

- Elle est plus adaptée au calcul automatique qu'à la compréhension intuitive des énoncés, et du *sens* des formules. Qui peut lire une table de vérité sur plus de 5 variables propositionnelles ?
- En dehors de la logique propositionnelle, la détermination automatique de la valeur de vérité d'une proposition peut être soit difficile, soit impossible.
- Les algorithmes de calcul de la valeur de vérité d'une proposition peuvent utiliser des représentations qui risquent de masquer l'intuition d'un énoncé.
- La définition de la sémantique des connecteurs et quantificateurs montre une “circularité” déroutante. Par exemple :

1. Enfin, on peut, mais alors on énonce des choses manifestement fausses. L'erreur est humaine, après tout.

- “La proposition $A \wedge B$ est vraie si la proposition A et la proposition B sont vraies”
- “La proposition $\forall x, P(x)$ est vraie si pour toute valeur t , la proposition $P(t)$ est vraie”
- Il est difficile de bien saisir le *sens* de l’implication $\rightarrow$. Dire que $A \rightarrow B$ est une abréviation de $\neg A \vee B$ n’aide en rien à comprendre la relation de causalité liée à ce connecteur.
- Finalement, rappelons la présence de propositions dont on ne sait encore si elles sont vraies ou fausses. Exemple : la conjecture sur la suite de Syracuse.

2.3.2 Le point de vue syntaxique : notion de preuve

En complément de la vue sémantique précédente (liée au nom de Tarski), la *théorie de la démonstration* consiste à définir des objets appelées *preuves*. Étant donné un énoncé logique, on définira quelles sont ses preuves éventuelles. On remplace donc la question « Cet énoncé est-il vrai ? » par « Cet énoncé est-il prouvable ? », voire « Comment démontrer cet énoncé ? ».

Comme pour un langage de programmation où la notion de “programme bien écrit” est définie, nous allons définir quelles sont les preuves correctes d’un séquent. Les règles de construction de ces preuves donneront un sens aux connecteurs et aux quantificateurs, sens que nous espérons aussi intuitif que les méthodes sémantiques.

Il existe plusieurs notions de preuves en logique, notamment le *calcul des séquents* de Gentzen, les *systèmes à la Hilbert*, la *déduction naturelle*. C’est cette dernière que nous allons utiliser, car très proche de la notion intuitive de raisonnement.

Quelques remarques :

- Ce n’est pas parce qu’on n’a pas pu démontrer un énoncé qu’il est faux.
- L’écriture de preuves correctes est très proche de l’écriture de programmes informatiques. Les informaticien[ne]s devront se sentir en terrain familier.
- Contrairement aux tables de vérité et autres techniques similaires, l’approche par la notion de preuve est tout à fait adaptée aux logiques à fort pouvoir d’expression (premier ordre, ordre supérieur, par exemple)

2.3.3 Notre point de vue (un rien subjectif)

Il n’est pas question de prendre une position affirmée entre ces deux faces de la logique. Ce cours utilisera plutôt la notion de démonstration pour les raisons suivantes :

- Il est important de reconnaître et écrire des argumentations convaincantes.
- La similitude entre la structure des démonstrations et celle des programmes informatiques rend ce choix bien adapté à un public informaticien.

- Le fait d’avoir une notion proprement définie de preuve permet un raisonnement rigoureux sur les preuves elles-mêmes – que nous qualifierons de *méta-raisonnement*, car vivant “en dehors” de la logique étudiée. Nous verrons des exemples de tels raisonnements dès le chapitre suivant.

Chapitre 3

La logique minimale

Afin d'aider à mieux comprendre le sens des propositions logiques, des notions de démonstration, d'hypothèse, de théorème, etc., nous commençons par une logique très réduite et très abstraite, appelée *logique minimale*, et ne contenant qu'un connecteur, l'implication.

Les logiques propositionnelle ou du premier ordre seront des extensions de la logique minimale, où l'on introduira progressivement plus de connecteurs et des quantificateurs.

3.1 Les propositions

La description d'une logique commence par définir quels énoncés seront manipulés, et comment les écrire sans ambiguïté.

On considère un ensemble arbitraire V_P de *variables propositionnelles* souvent notées $P, Q, \dots$

L'ensemble des propositions sur V_P (ou plus simplement *propositions*) est défini de façon *inductive* par :

- Toute variable propositionnelle est une proposition (appelée également *proposition atomique* ou *atome*).
- Si A et B sont deux propositions, alors l'implication $(A \rightarrow B)$ est une proposition
- Sous-entendu : il n'y a pas d'autres propositions.

Une telle définition inductive doit être considérée comme très pratique : elle définit toutes les manières de “construire” une proposition. Elle permet également une analyse “par cas” de *toute* proposition : face à n'importe quelle proposition, on sait qu'elle est *soit* d'une forme (atomique), *soit* de l'autre (implication). Enfin, on notera la similitude de cette définition avec la forme des définitions de nouveaux types dans un langage de programmation comme `Ocaml`.

On utilisera également le terme “formule” pour désigner une proposition.

La proposition $(A \rightarrow B)$ se lit « A implique B », ou « si A , alors B ».

Intuitivement, elle « signifie » que, si l’on peut prouver A , alors on peut aussi prouver B . Cette signification intuitive sera précisée plus loin sous forme de *règles de déduction*.

3.1.1 Remarques

1. Par convention, les méta-variables P , Q , R , etc., seront utilisées pour désigner des variables propositionnelles et A , B , C , etc, des propositions quelconques.
2. Le fait d’utiliser, dans ce chapitre et le suivant, des noms de propositions “abstraites” nous permet à la fois de souligner l’universalité des schémas de raisonnement présentés, et de travailler avec des formules plus courtes et donc plus lisibles.

Rien n’interdit au lecteur de prendre des exemples comme `socrate_est_mortel` ou `mon_programme_boucle`, mais ces exemples seront bien mieux traités dans le cadre de la logique des prédicats¹.

3. On trouve fréquemment la notation $A \Rightarrow B$ pour l’implication. Nous avons choisi la simple flèche $\rightarrow$ par compatibilité avec le logiciel *Coq*.
4. Les propositions $((A \rightarrow B) \rightarrow C)$ et $(A \rightarrow (B \rightarrow C))$ sont *différentes*. La structure de la première est représentée en figure 3.1, et celle de la seconde en figure 3.2.

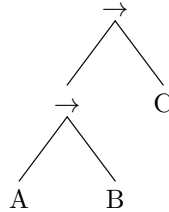


FIGURE 3.1 – structure de la proposition $((A \rightarrow B) \rightarrow C)$

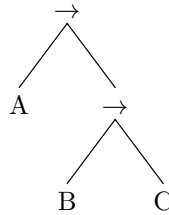


FIGURE 3.2 – structure de la proposition $(A \rightarrow (B \rightarrow C))$

1. Et l’expérience montre que les preuves sur papier prennent déjà bien assez de place sans allonger arbitrairement les noms des variables propositionnelles

- On peut enlever des parenthèses inutiles en considérant que l'opérateur $\rightarrow$ est *associatif à droite* (aussi dit *de droite à gauche*), c'est-à-dire que la formule $A \rightarrow B \rightarrow C$ est interprétée comme $A \rightarrow (B \rightarrow C)$ et non comme $(A \rightarrow B) \rightarrow C$. De nouveau, on pourra noter la similarité avec le sens de la notation “flèche” en `Ocaml`, où $A \rightarrow B \rightarrow C$ désigne le type des fonctions de A vers les fonctions de B vers C , et non celui des fonctions des fonctions de A vers B vers C ², lequel serait noté, justement, $(A \rightarrow B) \rightarrow C$.

Exercice 1. Dessiner la structure des formules suivantes :

- $P \rightarrow (P \rightarrow Q) \rightarrow Q$
- $(P \rightarrow Q) \rightarrow (Q \rightarrow R) \rightarrow P \rightarrow R$
- $(P \rightarrow Q \rightarrow R) \rightarrow Q \rightarrow P \rightarrow R$

3.2 Approche sémantique de la logique minimale

On considère toujours les propositions et séquents que l'on peut former à partir d'un ensemble V_P de variables propositionnelles. Dans la plupart de nos exemples, nous prendrons $V_P = \{P, Q, R, S, T, U\}$.

3.2.1 Définitions

- On considère l'ensemble $\mathbb{B}$ des *valeurs booléennes*. $\mathbb{B}$ ne contient que deux éléments : $\mathbf{v}$ (“vrai”) et $\mathbf{f}$ (“faux”). *On évitera l'assimilation de $\mathbf{f}$ au nombre 0 et de $\mathbf{v}$ à 1.*
- Une *valuation* est une fonction ν qui associe à toute variable propositionnelle une valeur dans $\mathbb{B} = \{\mathbf{v}, \mathbf{f}\}$. Par exemple, il y a $2^6 = 64$ valuations différentes sur l'ensemble V_P ci-dessus.
- Soit ν une valuation. On l'étend inductivement à l'ensemble de toutes les propositions en posant :

$$\begin{aligned} \nu((A \rightarrow B)) = \\ \mathbf{v} \text{ si } \nu(A) = \mathbf{f} \\ \nu(B) \text{ si } \nu(A) = \mathbf{v} \end{aligned}$$

- Une valuation ν *satisfait* un ensemble Γ de propositions si $\nu(A) = \mathbf{v}$ pour toute proposition A de Γ . Notation : $\nu \models \Gamma$.
- Un ensemble Γ de propositions est *satisfaisable*³ s'il existe une valuation ν qui satisfait $\nu \models \Gamma$.

2. Relire calmement cette phrase, en faisant des pauses mentales pour marquer les parenthèses implicites.

3. Pour le moment, cette notion n'est pas très pertinente ; elle prendra tout son sens une fois qu'on étendra ces définitions à des logiques plus riches.

6. Un séquent $A_1, A_2, \dots, A_n \vdash A$ est *valide* si pour toute valuation ν telle que $\nu(A_1) = \nu(A_2) = \dots = \nu(A_n) = \mathbf{v}$, alors $\nu(A) = \mathbf{v}$ – en d’autres termes, toutes valuation qui satisfait toutes les hypothèses satisfait également la conclusion. On exprime par le jugement $\Gamma \models A$ la validité du séquent $\Gamma \vdash A$.
7. Soit A une proposition. On dit que A est une *tautologie* si $\models A$.

La table de vérité de la figure 3.3 montre que le séquent $P \rightarrow Q, Q \rightarrow R \vdash P \rightarrow R$ est valide. En effet, toutes les valuations qui rendent vraies à la fois $P \rightarrow Q$ et $Q \rightarrow R$ rendent vraie la conclusion $P \rightarrow R$.

P	Q	R	$P \rightarrow Q$	$Q \rightarrow R$	$P \rightarrow R$
f	f	f	v	v	v
f	f	v	v	v	v
f	v	f	v	f	v
f	v	v	v	v	v
v	f	f	f	v	f
v	f	v	f	v	v
v	v	f	v	f	f
v	v	v	v	v	v

FIGURE 3.3 – Table de vérité du séquent $P \rightarrow Q, Q \rightarrow R \vdash P \rightarrow R$

En revanche, le séquent $P \rightarrow Q, Q \rightarrow R \vdash R \rightarrow P$ n’est pas valide (voir Figure 3.4, 2ème et 4ème lignes).

P	Q	R	$P \rightarrow Q$	$Q \rightarrow R$	$R \rightarrow P$
f	f	f	v	v	v
f	f	v	v	v	f
f	v	f	v	f	v
f	v	v	v	v	f
v	f	f	f	v	v
v	f	v	f	v	v
v	v	f	v	f	v
v	v	v	v	v	v

FIGURE 3.4 – Table de vérité du séquent $P \rightarrow Q, Q \rightarrow R \vdash R \rightarrow P$

Exercice 2. Revenons sur la notion d’ensemble d’hypothèses satisfaisable en logique minimale.

1. Montrer que la valuation qui attribue à chaque variable propositionnelle la valeur $\mathbf{v}$, attribue à toute proposition de la logique minimale la valeur $\mathbf{v}$.

2. En déduire que tout ensemble d'hypothèses Γ de la logique minimale est satisfaisable.

3.3 Preuves en logique minimale

Dans cette section, nous définissons quelles sont les éventuelles *preuves* d'un séquent, dans le cadre de la logique minimale.

Rappelons que le but est l'obtention d'une certitude que dans un séquent $\Gamma \vdash A$, A peut se déduire des hypothèses de Γ .

On définit une preuve de $\Gamma \vdash A$ comme une suite de séquents dont chacun s'obtient des précédents en appliquant une *règle d'inférence* parmi les trois règles de la logique minimale, définies ci-dessous, et dont le dernier élément est le séquent voulu $\Gamma \vdash A$.

Nous donnerons quelques exemples de preuves au fur et à mesure de la présentation des règles de la logique minimale.

On présente souvent chaque règle d'inférence de la façon suivante : la conclusion (le séquent que la règle permet de prouver) est placée sous une barre horizontale, le nom de la règle et d'éventuelles conditions d'application, à droite de cette barre. Enfin, les différents séquents qu'il est nécessaire de prouver pour utiliser proprement la règle, sont placés au-dessus de la barre.

Une telle règle se “lit” donc, dans le sens naturel du raisonnement déductif, *de haut en bas* : on *déduit* le séquent “du bas” à partir des séquents “du haut” (qu'on est censé avoir déjà prouvés) et de la règle de déduction utilisée.

Une telle preuve n'est complète que si chaque séquent est prouvé correctement à partir des précédents. Pour éviter un paradoxe de type “oeuf et poule” (ou tout simplement pour qu'il existe des preuves finies), il faut bien que certains séquents soient prouvables sans faire référence à d'autres : c'est ce que permet la règle d'hypothèse.

3.3.1 Règle d'hypothèse

La *règle d'hypothèse*⁴, comme son nom l'indique, précise le rôle joué par les hypothèses d'un séquent : si une proposition apparaît parmi les hypothèses d'un séquent, on considère qu'elle est bien une conséquence logique triviale de ces hypothèses.

Si $A \in \Gamma$, alors $\Gamma \vdash A$

Voici la présentation habituelle de la règle d'hypothèse.

$$\frac{}{\Gamma \vdash A} \text{ hyp } (A \in \Gamma)$$

Variante, signifiant la même chose :

$$\frac{}{\Gamma, A \vdash A} \text{ hyp}$$

4. « assumption » en anglais, parfois appelée “axiome” mais bof

Comme la règle d'hypothèse ne demande de prouver aucun autre séquent, il n'y a rien au-dessus de la barre.

Instances d'une règle

Dans la description ci-dessus, les symboles A et Γ sont des variables pouvant être respectivement remplacées par n'importe quelle proposition et n'importe quel contexte.

Cette opération de remplacement s'appelle *instantiation* de la règle.

Soient par exemple trois variables propositionnelles P , Q et R .

Alors, la règle d'hypothèse nous permet de déduire le séquent :

$$P \rightarrow Q, Q \rightarrow R, P \vdash P \rightarrow Q$$

On peut représenter la preuve associée comme une application de la règle d'hypothèse :

$$\frac{P \rightarrow Q, Q \rightarrow R, P}{P \vdash P \rightarrow Q} \text{ hyp}$$

En revanche, le séquent suivant *n'est pas* obtenu comme instance de la règle d'hypothèse (car la condition $A \in \Gamma$ n'est pas respectée). Le fait que la conclusion P soit une sous-formule de l'hypothèse $P \rightarrow Q$ ne fait pas de P une des hypothèses.

$$P \rightarrow Q \vdash P$$

3.3.2 Règle du modus ponens (élimination de l'implication)

Le sens de l'implication $A \rightarrow B$ est donné par la règle du *modus ponens* (également appelée *règle d'élimination de l'implication*). Avoir prouvé $A \rightarrow B$, permet, si l'on a également prouvé A , d'en déduire B . Pour référencer cette règle, on pourra utiliser indifféremment la notation $\rightarrow_e$ (élimination de l'implication) ou *mp* (pour *modus ponens*).

La présentation usuelle de cette règle consiste, comme pour la règle d'hypothèse, à placer la conclusion sous une barre horizontale. Au dessus de cette barre, on place les séquents dont la preuve préalable conditionne l'application de la règle.

$$\frac{\Gamma \vdash A \rightarrow B \quad \Gamma \vdash A}{\Gamma \vdash B} \text{ mp}$$

Cette règle se lit :

Si l'on peut prouver les séquents $\Gamma \vdash A \rightarrow B$ et $\Gamma \vdash A$, alors on peut prouver le séquent $\Gamma \vdash B$

En d'autres termes, cette règle permet d'obtenir une preuve de $\Gamma \vdash B$ à partir d'une preuve de $\Gamma \vdash A \rightarrow B$ et d'une preuve de $\Gamma \vdash A$.

Remarque

Il est important, lorsqu'on applique une règle d'inférence, de vérifier que l'on remplace **toutes** les occurrences d'une même méta-variable par la même formule ou le même contexte. Par exemple l'arbre suivant n'est pas une instance du modus-ponens :

$$\frac{P \rightarrow Q, R \rightarrow P \vdash P \rightarrow Q \quad R, R \rightarrow P \vdash P \rightarrow P}{P \rightarrow Q, R \rightarrow P \vdash Q} \text{ mp}$$

3.3.3 Preuves en déduction naturelle

Avec les deux règles : d'hypothèse et du modus-ponens, nous pouvons déjà écrire des preuves non triviales. Les preuves en déduction naturelle prennent alors la forme de *dérivations*. Une dérivation est un arbre dont la racine est le séquent prouvé, et chaque noeud correspond à l'application d'une règle d'inférence. Les feuilles de cet arbre sont les applications de la règle d'hypothèse (car c'est la seule règle qui ne demande pas de prouver d'autre séquent pour être appliquée).

Exemple

La preuve suivante utilise les deux règles d'hypothèse et du modus ponens. Afin d'alléger les notations nous posons $\Gamma = \{P \rightarrow Q, Q \rightarrow R, P\}$

$$\frac{\frac{\frac{}{\Gamma \vdash Q \rightarrow R} \text{ hyp} \quad \frac{\frac{\frac{}{\Gamma \vdash P \rightarrow Q} \text{ hyp} \quad \frac{}{\Gamma \vdash P} \text{ hyp}}{\Gamma \vdash Q} \text{ mp}}{\Gamma \vdash R} \text{ mp}}{\Gamma \vdash R} \text{ mp}$$

Présentation linéaire d'une preuve

Les arbres associés à des preuves complexes peuvent vite devenir peu lisibles ; en particulier, chaque utilisation de la règle du *modus ponens* introduit un noeud binaire, et les arbres peuvent facilement devenir "touffus". On pourra à la place représenter une preuve comme une suite de propositions. Les hypothèses seront signalées comme telles, et les propositions obtenues par l'application d'une règle d'inférence seront accompagnées d'une *justification* composée du nom de la règle appliquée, et de références permettant de retrouver à quelles formules ces règles sont appliquées.

De façon plus précise :

- Chaque ligne de la preuve est soit une hypothèse (qu'on introduira par le mot-clé "Supposons"), soit une proposition obtenue par l'application d'une règle d'inférence.
- À chaque ligne différente d'une hypothèse, on peut associer le séquent dont la conclusion est la formule courante, et le contexte l'ensemble de toutes les hypothèses précédant cette ligne dans la preuve.

— Il va de soi que les références présentes dans une justification doivent porter sur des étapes précédant la ligne considérée (pas de références aux lignes suivantes, sous peine d’introduire de la circularité dans les “preuves”).

Dans l’exemple ci-dessous, les lignes 1 à 3 sont des hypothèses. La ligne 4 est obtenue par modus ponens sur l’implication de la ligne 1 et l’hypothèse en ligne 3. De même pour la ligne 5, obtenue par modus ponens, à partir de la proposition $Q \rightarrow R$ et de la proposition Q prouvée en ligne 4.

Le séquent associé à la ligne 4 est

$$P \rightarrow Q, Q \rightarrow R, P \vdash Q$$

```

1  Supposons  $P \rightarrow Q$ 
2  Supposons  $Q \rightarrow R$ 
3  Supposons  $P$ 
4   $Q$  [mp, 1,3]
5   $R$  [mp, 2,4]
```

Exercice 3. Prouver le séquent suivant (dans les deux formats)

$$P \rightarrow Q \rightarrow R, P \rightarrow Q, P \vdash R$$

(attention à l’associativité droite de $\rightarrow$!)

3.3.4 Règle d’abstraction (introduction de l’implication)

Dans les exemples précédents, nous avons vu comment *utiliser* des propositions de la forme $A \rightarrow B$ présentes dans les hypothèses. La règle d’introduction de l’implication (notée $\rightarrow_i$) permet de prouver des séquents de la forme $\Gamma \vdash A \rightarrow B$.

Cette règle formalise le mode de raisonnement intuitif suivant.

Pour prouver $A \rightarrow B$:

1. **Supposons** A
2. *Sous cette hypothèse*, prouvons B

La règle $\rightarrow_i$, dite d’*introduction de l’implication* formalise ce mode de raisonnement dit *raisonnement hypothétique*. Le fait de prouver B *en supposant* A se traduit simplement par “prouver le séquent $\Gamma, A \vdash B$ ”.

Règle 1 (introduction de l’implication, $\rightarrow_i$).

$$\frac{\Gamma, A \vdash B}{\Gamma \vdash A \rightarrow B} \rightarrow_i$$

Remarquons que l’hypothèse sur A n’est “active” que lors de la preuve de B . Cette remarque est explicitée par la représentation formelle de la règle $\rightarrow_i$, où la formule A n’est plus dans les hypothèses de la conclusion de la règle.

Exemple

Soit $\Gamma = \{P \rightarrow Q, Q \rightarrow R\}$. La dérivation suivante utilise les trois règles de la logique minimale. *On remarquera que tous les séquents de cette preuve n'ont pas le même contexte.*

$$\frac{\frac{\frac{}{\Gamma, P \vdash Q \rightarrow R} \text{hyp}}{\Gamma, P \vdash Q} \text{mp} \quad \frac{\frac{\frac{}{\Gamma, P \vdash P \rightarrow Q} \text{hyp}}{\Gamma, P \vdash P} \text{hyp}}{\Gamma, P \vdash Q} \text{mp}}{\Gamma, P \vdash R} \rightarrow_i$$

Présentation “à la C” d’une preuve

Une autre présentation du raisonnement hypothétique se rapproche de la structure de bloc des langages de programmation. La notion d’hypothèse s’apparente alors à celle de *déclaration de variable locale*.

Prenons par exemple la fonction 4.1 page 43. La variable locale *i* est invisible depuis les lignes 9, 10 et 11 de la fonction.

Comme en **C**, une convention d’indentation permet de mieux visualiser la structure d’une dérivation. Nous adaptons donc la présentation de la section 3.3.3 page 21 en ajoutant une notation de bloc, délimité par des accolades. Chaque bloc peut contenir une hypothèse, dont la *portée* est limitée à ce bloc.

```

1  Supposons P → Q
2  Supposons Q → R
3  { Supposons P
4      Q [mp, 1, 3]
5      R [mp, 2, 4]
6  }
7  P → R [→i, 3, 5]
```

Il est important de comprendre la notion de *portée d’une hypothèse* et le rôle des blocs. Un bloc commençant par une ligne “**Supposons A**” contient des propositions prouvées en général à l’aide de l’hypothèse sur *A*. Afin de marquer cette dépendance, on considère qu’une proposition appartenant à ce bloc est *inaccessible* depuis l’extérieur de ce bloc – elle n’est peut-être pas prouvable sans cette hypothèse *A*.

Considérons par exemple la structure ci-dessous, comportant deux sous-dérivations :

```

1  { Supposons A
2      ...
3      B
4  }
5  ...
6  C
```

```

7   ...
8   { Supposons  $H$ 
9     ...
10     $D$ 
11  }
12   $E$ 

```

- Une première sous-dérivation, sous l’hypothèse A ; la preuve de la justification de la proposition B peut utiliser cette hypothèse.
- De même avec un bloc d’hypothèse H
- Les justifications des formules C et E ne peuvent pas utiliser les propositions A , B , H et D .
- Le seul moyen d’utiliser par exemple la proposition B prouvée en ligne 3, dans une justification de E est de l’exporter par la règle d’introduction de l’implication.

```

1   { Supposons  $A$ 
2     ...
3      $B$ 
4   }
5   ...
6    $C$ 
7   ...
8   { Supposons  $H$ 
9     ...
10     $D$ 
11  }
12   $A \rightarrow B \quad [\rightarrow_i, 1, 3]$ 
13   $E$ 

```

Dérivations imbriquées

Comme dans les langages de programmation, la structure de bloc associée au raisonnement hypothétique autorise l’imbrication de *sous-dérivations*. Chaque niveau d’imbrication pourra correspondre à autant d’applications de la règle d’introduction de l’implication.

La preuve suivante contient deux applications successives de la règle $\rightarrow_i$, traduites par deux blocs imbriqués.

$$\frac{\frac{\frac{}{P, Q \vdash P} \text{hyp}}{P \vdash Q \rightarrow P} \rightarrow_i}{\vdash P \rightarrow Q \rightarrow P} \rightarrow_i$$

```

1   { Supposons  $P$ 
2     { Supposons  $Q$ 

```

```

3     P [hyp, 1]
4   }
5   Q → P [→i, 2, 3]
6 }
7 P → Q → P [→i, 1, 5]

```

Erreurs communes

La dérivation suivante est un exemple de non-respect de la structure de bloc. À la ligne 8, on tente d’appliquer l’introduction de l’implication sur une formule de la ligne 5, non accessible de la ligne 8.

```

1  Supposons P → Q
2  { Supposons Q → R
3    { Supposons P
4      Q [mp, 1, 3]
5      R [mp, 2, 4]
6    }
7  }
8  (Q → R) → R [→i, 2, 5]

```

Cette dérivation incorrecte “prouve” le séquent (non valide) suivant :

$$P \rightarrow Q \vdash (Q \rightarrow R) \rightarrow R$$

Une dérivation correcte (mais ne prouvant pas le même séquent) serait la suivante, où la proposition de la ligne 5 est “exportée” en tenant compte de l’hypothèse sur Q .

```

1  Supposons P → Q
2  { Supposons Q → R
3    { Supposons P
4      Q [mp, 1, 3]
5      R [mp, 2, 4]
6    }
7    P → R [→i, 3, 5]
8  }
9  (Q → R) → P → R [→i, 2, 6]

```

Avec cette structuration par blocs des preuves, il nous faut redéfinir le séquent prouvé en chaque ligne d’une telle preuve : le contexte n’est plus défini par l’ensemble de toutes les hypothèses qui précèdent la ligne considérée, mais par l’ensemble des hypothèses *accessibles* qui précèdent cette ligne. Les hypothèses introduites dans des blocs qui ont été fermés ne sont plus accessibles, et donc ne font pas partie du contexte.

Définitions

Soit $\Gamma \vdash A$ un séquent. On dira que ce séquent est *prouvable* (en logique minimale) s'il existe une dérivation en logique minimale concluant par ce séquent. On utilisera la notation $\Gamma \vdash_M A$ pour exprimer le jugement “ $\Gamma \vdash A$ est prouvable en logique minimale”.

Notons que $\vdash_M$ *n'est pas un connecteur*. On ne confondra pas le *jugement* $A \vdash_M B$ avec la *proposition* $A \rightarrow B$.

Remarque Il faut bien distinguer la différence de statut linguistique entre un séquent $\Gamma \vdash A$ et un jugement $\Gamma \vdash_M A$. Le premier s'apparente à une structure de données (niveau “objet”) et le second à une propriété mathématique (niveau “méta”). On trouvera des exemples de cette seconde catégorie dans la suite de ce document : jugements de validité $\Gamma \models A$, prouvabilité dans d'autres logiques : $\Gamma \vdash_J A$, $\Gamma \vdash_K A$, etc.

Introduction, élimination

Les vocables “règle d'introduction” et “règle d'élimination” sont utilisés en logique pour classer les règles d'inférence associées aux connecteurs et quantificateurs.

- Une *règle d'introduction* permet de prouver un séquent de la forme $\Gamma \vdash A$ où l'opérateur principal de A est le connecteur ou quantificateur considéré.
- Une *règle d'élimination* permet de prouver un séquent de la forme $\Gamma \vdash A$, en utilisant un séquent $\Delta \vdash B$ où l'opérateur principal de B est le connecteur ou quantificateur considéré.

3.4 Simplifier l'écriture des preuves

3.4.1 Règles dérivées

Notion de règle dérivée

La preuve d'un séquent par les seules règles d'hypothèse, d'introduction et d'élimination peut se révéler très longue et fastidieuse. Une *règle dérivée* est par définition une règle qui n'apporte rien du point de vue de la prouvabilité, mais qui permet de remplacer plusieurs étapes d'application des règles de base. D'une certaine façon, l'usage de règles dérivées s'apparente à l'utilisation de fonctions dans un langage de programmation.

Avant d'utiliser une règle dérivée, il convient de montrer comment cette règle se traduit en termes de règles dont la validité est déjà connue (règles de base de la logique, ou règles dérivées déjà validées). En d'autres termes, une règle dérivée se *justifie* en donnant un algorithme permettant de traduire toute preuve utilisant cette nouvelle règle en une preuve n'utilisant que les règles de base ou les règles dérivées *déjà certifiées*.

Généralisation du modus-ponens

L'élimination de l'implication (modus ponens) permet d'éliminer une implication de la formule $A \rightarrow B$. On rencontre fréquemment dans une preuve des propositions de la forme $A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow B$ (rappelons que l'implication est associative à droite).

Supposons que l'on soit capable de prouver chacun des séquents $\Gamma \vdash A_i$.

On peut obtenir la dérivation suivante :

$$\begin{array}{c}
 \frac{\Gamma \vdash A_1 \rightarrow A_2 \rightarrow \dots A_{n-1} \rightarrow A_n \rightarrow B \quad \Gamma \vdash A_1}{\Gamma \vdash A_2 \rightarrow \dots A_{n-1} \rightarrow A_n \rightarrow B} \text{mp} \quad \Gamma \vdash A_2 \quad \text{mp} \\
 \frac{\Gamma \vdash A_3 \rightarrow \dots A_{n-1} \rightarrow A_n \rightarrow B}{\vdots} \\
 \frac{\Gamma \vdash A_{n-1} \rightarrow A_n \rightarrow B \quad \Gamma \vdash A_{n-1}}{\Gamma \vdash A_n \rightarrow B} \text{mp} \quad \Gamma \vdash A_n \quad \text{mp} \\
 \hline
 \Gamma \vdash B
 \end{array}$$

Il est plus simple de court-circuiter toutes ces étapes en une seule, formalisée par une généralisation du modus-ponens :

$$\frac{\Gamma \vdash A_1 \rightarrow A_2 \rightarrow \dots A_{n-1} \rightarrow A_n \rightarrow B \quad \Gamma \vdash A_1 \quad \Gamma \vdash A_2 \quad \dots \quad \Gamma \vdash A_{n-1} \quad \Gamma \vdash A_n}{\Gamma \vdash B} \text{mp}'$$

Exemple

Soit $\Gamma = \{P \rightarrow Q \rightarrow R, P, Q\}$. La dérivation suivante est correcte :

$$\frac{\overline{\Gamma \vdash P \rightarrow Q \rightarrow R} \text{ hyp} \quad \overline{\Gamma \vdash P} \text{ hyp} \quad \overline{\Gamma \vdash Q} \text{ hyp}}{\Gamma \vdash R} \text{mp}'$$

- 1 $P \rightarrow Q \rightarrow R$
- 2 P
- 3 Q
- 4 R [mp, 1,2,3]

Exemple

Cet autre exemple montre une application *partielle* du modus ponens généralisé :

$$\frac{\Gamma \vdash P \rightarrow Q \rightarrow R \rightarrow S \rightarrow T \quad \Gamma \vdash P \quad \Gamma \vdash Q}{\Gamma \vdash R \rightarrow S \rightarrow T}$$

Généralisation de l'introduction de l'implication

De façon symétrique au *modus ponens*, on peut dériver la règle suivante :

$$\frac{\Gamma, A_1, A_2, \dots, A_n \vdash A}{\Gamma \vdash A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow A} \rightarrow_i$$

En voici un exemple d'instance :

$$\frac{\overline{P, Q \vdash P} \text{ hyp}}{\vdash P \rightarrow Q \rightarrow P} \rightarrow_i$$

Exercice 4. Montrer que cette généralisation de $\rightarrow_i$ est bien une règle dérivée de la logique minimale.

Et les blocs ?

Dans la présentation structurées des preuves, il suffit de procéder aux modifications suivantes :

- Admettre l'introduction de plusieurs hypothèses $A_1, A_2, \dots, A_n$ en tête de bloc (une par ligne : il est important de pouvoir faire référence séparément à chacune d'entre elles)
- Une proposition B située à l'intérieur d'un tel bloc s'exporte sous la forme $A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow B$.

```

1 { Supposons P
2   Supposons Q
3   P [hyp, 1]
4 }
5 P → Q → P [→i, 1-2, 3]
```

Remarque

Il est important d'exporter *toutes* les hypothèses d'un bloc. L'exemple suivant montre que le non-respect de cette règle peut conduire à la « preuve » de séquents non valides.

```

1 { Supposons P
2   Supposons Q
3   P [hyp, 1]
4 }
5 Q → P [→i, 2, 3]
```

Exercice 5. Prouver les séquents suivants (si possible dans les deux formats de preuve).

$\vdash P \rightarrow P$
 $\vdash (P \rightarrow P) \rightarrow (P \rightarrow P)$
 $\vdash (P \rightarrow Q) \rightarrow (Q \rightarrow R) \rightarrow P \rightarrow R$
 $P \rightarrow Q \rightarrow R \vdash Q \rightarrow P \rightarrow R$
 $P \rightarrow R \vdash P \rightarrow Q \rightarrow R$
 $P \rightarrow P \rightarrow Q \vdash P \rightarrow Q$
 $P \rightarrow Q, P \rightarrow R, Q \rightarrow R \rightarrow T \vdash P \rightarrow T$

Règle de coupure

Supposons qu'on veuille prouver un séquent $\Gamma \vdash A$, et que l'on trouve cette tâche difficile. En revanche, on trouve une proposition B telle que :

- On se sent capable de prouver le séquent $\Gamma \vdash B$
- On pense qu'une fois prouvée B , il est plus facile de prouver A .

Une stratégie possible peut consister à travailler en trois étapes :

1. Prouver d'abord un séquent $\Gamma \vdash B$, où B est une proposition appelée *lemme*.
2. Prouver l'implication $\Gamma \vdash B \rightarrow A$.
3. Appliquer le modus ponens pour en déduire $\Gamma \vdash A$.

En bref :

$$\frac{\Gamma \vdash B \quad \frac{\Gamma, B \vdash A}{\Gamma \vdash B \rightarrow A} \rightarrow_i}{\Gamma \vdash A} \text{mp}$$

Nous pouvons associer une nouvelle règle dérivée, la *règle de coupure* à cette stratégie, de façon à court-circuiter la succession formée d'une introduction, immédiatement suivie d'une élimination de l'implication.

Nous utilisons (provisoirement) des couleurs pour distinguer les deux composantes d'une coupure.

$$\frac{\Gamma \vdash B \quad \Gamma, B \vdash A}{\Gamma \vdash A} \text{cut}$$

Voici un premier exemple d'application de cette règle.

Soit $\Gamma = \{P \rightarrow Q \rightarrow R, P \rightarrow Q, P\}$.

$$\frac{\frac{\overline{\Gamma \vdash P \rightarrow Q} \text{ hyp} \quad \overline{\Gamma \vdash P} \text{ hyp}}{\Gamma \vdash Q} \text{ mp} \quad \frac{\frac{\overline{\Gamma, Q \vdash P \rightarrow Q \rightarrow R} \text{ hyp} \quad \overline{\Gamma, Q \vdash P} \text{ hyp}}{\Gamma, Q \vdash Q \rightarrow R} \text{ mp} \quad \overline{\Gamma, Q \vdash Q} \text{ hyp}}{\Gamma, Q \vdash R} \text{ mp} \quad \frac{\Gamma \vdash Q \quad \Gamma, Q \vdash R}{\Gamma \vdash R} \text{ cut}$$

Dans le format structuré des preuves, la règle de coupure peut se représenter très simplement. On commence par prouver la proposition B , dans un bloc dédié à cette preuve. Une fois prouvée, B peut être utilisée dans le reste de la démonstration.

Reprenons notre exemple :

```

1  Supposons  $P \rightarrow Q \rightarrow R$ 
2  Supposons  $P \rightarrow Q$ 
3  Supposons  $P$ 
4  Prouvons  $Q$  {
5     $Q$  [mp, 2, 3]
6  }
7   $R$  [mp, 1, 3, 4]
```

Stricto sensu, rien n'oblige à même à faire apparaître cette coupure : on peut tout à fait prouver le lemme comme faisant partie du cours naturel de la preuve.

L'exemple suivant, un peu plus complexe, montre comment la règle de coupure permet de prouver une seule fois la proposition $P \rightarrow R$ et de l'appliquer deux fois (en lignes 13 et 16). L'utilisation de lemmes (et donc de la règle de coupure) permet souvent de réduire la taille des preuves. En ce sens, la règle de coupure s'apparente aux techniques pour éviter la duplication de code en programmation. Là où le programmeur utilise des fonctions auxiliaires pour "factoriser" son code, le "logicien" invente des lemmes.

```

1  Supposons  $P \rightarrow Q$ 
2  Supposons  $Q \rightarrow R$ 
3  Supposons  $(P \rightarrow R) \rightarrow T \rightarrow Q$ 
4  Supposons  $(P \rightarrow R) \rightarrow T$ 
5  Prouvons  $P \rightarrow R$  {
6    { Supposons  $P$ 
7       $Q$  [mp, 1, 6]
8       $R$  [mp, 2, 7]
9    }
10    $P \rightarrow R$  [ $\rightarrow_i$ , 6, 8]
11 }
12 Prouvons  $T$  {
13    $T$  [mp, 4, 5]
14 }
15
16  $Q$  [mp, 3, 5, 12]
```

3.4.2 Règle d'affaiblissement

Si l'on peut prouver un séquent $\Gamma \vdash A$, on peut *a fortiori* prouver sa conclusion en ajoutant un ensemble arbitraire d'hypothèses.

La *règle d'affaiblissement* tient compte de ce fait :

$$\frac{\Gamma \vdash A}{\Gamma, \Delta \vdash A} \text{ aff}$$

Le nom d'affaiblissement est souvent mal compris. Il ne vient pas du fait qu'on semble retirer des hypothèses quand on lit l'arbre depuis sa racine (ce qui est souvent l'ordre d'écriture d'un arbre de preuve) ; il faut plutôt voir chaque séquent qu'on prouve dans le cadre d'une preuve comme un petit théorème. Ajouter des hypothèses à un théorème, c'est l'affaiblir : le théorème obtenu est moins puissant puisqu'il s'applique plus difficilement. C'est en ce sens (et donc bien dans le sens de lecture “vers la racine”) qu'il faut comprendre ce terme d'affaiblissement.

En pratique, la règle d'affaiblissement est surtout utilisée pour réduire l'encombrement des preuves présentées sous forme d'arbres, ou pour la justification des règles dérivées.

Exercice 6. *Montrer que la règle d'affaiblissement est bien une règle dérivée. On pourra procéder par récurrence sur [la taille de] la preuve de $\Gamma \vdash A$*

Exercice 7. *Montrer que la “règle” suivante est visiblement fausse.*

$$\frac{\Gamma, \Delta \vdash A}{\Gamma \vdash A} \text{ renforcement}$$

Exercice 8. *Le but de cet exercice est de comprendre comment créer une règle dérivée à partir de la preuve d'un séquent.*

Soient A et B deux propositions quelconques. On suppose qu'on a une preuve en logique minimale⁵ du séquent $A \vdash B$. Montrer qu'on peut alors dériver la règle suivante :

$$\frac{\Gamma \vdash A}{\Gamma \vdash B}$$

3.4.3 Prouvabilité et règles dérivées

Les règles dérivées que nous avons vues (ainsi, en général, que toutes celles que nous verrons par la suite), présentent plusieurs avantages :

- Elles permettent de réduire notablement la taille des preuves, et donc d'accroître leur lisibilité.
- Elles facilitent la recherche d'une preuve d'un séquent donné (notamment la règle de coupure).

5. Ce résultat s'étendra aux logiques présentées dans les chapitres suivants.

En revanche, elles n'ajoutent rien à l'ensemble des séquents prouvables : si la preuve d'un séquent utilise des règles dérivées, alors on peut construire une preuve du même séquent n'utilisant que les trois règles de base de la logique minimale.

Les règles dérivées : mode d'emploi

Les remarques qui suivent sont bien sûr immédiatement applicables à la logique minimale, mais seront utiles dans tout le document.

- Les règles dérivées sont introduites dans le texte du document, mais aussi dans les exercices.
- La preuve de validité de ces règles est parfois laissée en exercice. Dans d'autres cas, cette preuve est trop longue ou complexe pour ce cours, et on pourra admettre le résultat ou chercher de la documentation sur Internet.
- Si par malheur, vous ne pouvez pas faire un tel exercice, vous pouvez cependant admettre la règle en question dans la suite du document.
- Une fois qu'une règle est présentée et prouvée ou admise, son utilisation permet de simplifier la recherche ou l'écriture de démonstrations.
- On peut trouver que de nombreux exercices se ressemblent ; cependant nous encourageons le lecteur à essayer de trouver les preuves les plus simples et courtes ; l'emploi de règles dérivées est un moyen d'obtenir cette simplicité.

3.5 Preuves *vs* validité

Nous avons deux notions différentes de vérité pour un séquent $\Gamma \vdash A$:

- une vérité “sémantique”, exprimée par le jugement $\Gamma \models A$,
- une vérité “syntaxique” exprimée par le jugement $\Gamma \vdash_M A$.

Il est temps de comparer ces deux approches.

3.5.1 Correction

Le (méta-)théorème de correction établit que les règles de déduction de la logique minimale ne permettent de prouver que des séquents valides.

Méta-théorème 1 (Correction). *Si $\Gamma \vdash_M A$, alors $\Gamma \models A$*

Démonstration. Pour prouver ce théorème, on nous allons procéder par *récur-
rence* sur la structure des preuves. Pour tout arbre correct de preuve, appelons *taille* de cet arbre son nombre de noeuds, c'est-à-dire le nombre total de règles d'inférences utilisées dans la preuve. C'est donc un entier strictement positif : les plus “petites” preuves sont celles qui se réduisent à la règle d'hypothèse (taille 1).

Soit, pour $n > 0$, H_n la propriété suivante : “Tout séquent de la logique propositionnelle qui est prouvable par un arbre de preuve de taille au plus n ,

est un séquent valide”. Démontrer le théorème revient à prouver que, pour tout $n > 0$, H_n est vraie. Nous allons faire cela par récurrence sur n , ce qui veut dire que nous avons deux sous-preuves à effectuer :

- H_1 est vraie.
- Si $n \geq 1$ est tel que H_n est vraie, alors H_{n+1} est vraie.

(Notre récurrence commence avec $n = 1$ car les arbres de preuves ne peuvent pas être de taille 0)

Commençons par discuter de la propriété H_1 . Un arbre de preuve de taille au plus 1 est forcément de taille 1, et c’est forcément une preuve par la règle d’hypothèse (c’est la seule règle qui ne demande pas de prouver un autre séquent). La propriété H_1 sera donc une conséquence du lemme 1, énoncé et prouvé ci-dessous.

Passons maintenant à la seconde sous-preuve : si H_n est vraie, alors H_{n+1} est vraie. Soit donc un entier n tel que H_n soit vraie, et considérons un séquent $\Gamma \vdash B$ qui admet une preuve de taille au plus $n + 1$; notons $\mathcal{A}$ cet arbre de preuve. Deux cas peuvent se produire :

- $\mathcal{A}$ est de taille au plus n : dans ce cas, puisque H_n a été supposée vraie, le séquent $\Gamma \vdash B$ est valide.
- $\mathcal{A}$ est de taille $n + 1$: l’hypothèse de récurrence H_n ne peut donc pas être appliquée directement à $\mathcal{A}$.

Considérons alors, dans l’arbre de preuve $\mathcal{A}$, les noeuds qui sont des fils de la racine. Ils sont soit au nombre de 2 (si la racine de $\mathcal{A}$ est une utilisation de la règle du modus ponens) ou de 1 (si la racine de $\mathcal{A}$ est une utilisation de la règle d’introduction de l’implication). (À ce stade, la racine de $\mathcal{A}$ ne peut pas être une feuille, car l’arbre serait alors de taille 1 ; or, on a supposé $n \geq 1$ et $\mathcal{A}$ est de taille $n + 1$). Notons $\mathcal{A}_i$ les sous-arbres enracinés aux fils de la racine de $\mathcal{A}$ (soit seulement pour $i = 1$, soit pour $i \in \{1, 2\}$).

Chacun des sous-arbres $\mathcal{A}_i$ est de taille au plus n (car il y “manque” au moins la racine de $\mathcal{A}$). Donc on peut leur appliquer l’hypothèse de récurrence H_n : chacun des séquents qui se trouvent à la racine des arbres $\mathcal{A}_i$ est un séquent valide. Le lemme 1 ci-dessous permet donc de conclure que le séquent $\Gamma \vdash B$ est valide.

□

Dans la preuve qui précède, nous avons par deux fois éludé une partie du raisonnement, correspondant au lemme suivant.

Lemme 1. *Pour chacune des règles d’inférence de la logique minimale, dans toute instanciation correcte de la règle, si chacun des séquents au-dessus de la barre (“obligations de preuve”) est valide, alors le séquent en-dessous de la barre est également valide.*

Démonstration. La logique minimale n’utilise que trois règles : nous pouvons donc nous contenter d’examiner chacune des règles individuellement.

1. Soit une preuve obtenue par la règle d’hypothèse :

$$\frac{}{\Gamma \vdash A} \text{ hyp}$$

Par construction, la proposition A appartient au contexte Γ . Donc, toute valuation ν qui satisfait toutes les hypothèses de Γ satisfait forcément A , et le séquent $\Gamma \vdash A$ est donc valide.

2. Soit une instanciation de la règle du *modus ponens*. Elle a la forme suivante :

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \Gamma \vdash A \rightarrow B \end{array} \quad \begin{array}{c} \vdots \pi_2 \\ \Gamma \vdash A \end{array}}{\Gamma \vdash B} \text{ mp}$$

On suppose $\Gamma \models A \rightarrow B$ et $\Gamma \models A$; montrons le jugement $\Gamma \models B$:

- (a) Soit ν une valuation qui rend vraies toutes les hypothèses de Γ . Par hypothèse ($\Gamma \models A$), on a $\nu(A) = \mathbf{v}$, et (d'après $\Gamma \models A \rightarrow B$) $\nu(A \rightarrow B) = \mathbf{v}$.
 - (b) On a donc $\nu(A) = \mathbf{v}$ et $\nu(A \rightarrow B) = \mathbf{v}$. Par conséquent, on a $\nu(B) = \mathbf{v}$. En d'autres termes, ν satisfait la conclusion de $\Gamma \vdash B$.
 - (c) On vient de montrer que toute valuation qui satisfait les hypothèses du séquent $\Gamma \vdash B$, satisfait également sa conclusion ; en d'autres termes, on vient de montrer $\Gamma \models B$.
3. Soit une instanciation correcte de l'introduction de l'implication. Elle a la forme suivante :

$$\frac{\begin{array}{c} \vdots \pi \\ \Gamma, A \vdash B \end{array}}{\Gamma \vdash A \rightarrow B} \rightarrow_i$$

On suppose $\Gamma, A \models B$; montrons le jugement $\Gamma \models A \rightarrow B$:

- (a) Soit ν une valuation satisfaisant toutes les hypothèses de Γ . Deux cas sont possibles :
 - i. Si $\nu(A) = \mathbf{v}$, alors ν satisfait toutes les hypothèses de Γ ainsi que l'hypothèse A , et par conséquent $\nu(B) = \mathbf{v}$, et $\nu(A \rightarrow B) = \mathbf{v}$.
 - ii. Si $\nu(A) = \mathbf{f}$, alors $\nu(A \rightarrow B) = \mathbf{v}$.
- (b) Donc, dans les deux cas, on a bien $\nu(A \rightarrow B) = \mathbf{v}$; on a donc montré que toute valuation qui satisfait les hypothèses de $\Gamma \vdash B$ satisfait également sa conclusion, *i.e.* on a montré $\Gamma \models B$.

□

Si l'on revient sur la preuve du théorème, on s'aperçoit que tout repose sur le lemme, lequel se ramène à une vérification locale. En particulier, chaque fois qu'on étendra la logique en ajoutant des connecteurs et des règles d'inférence, on obtiendra un nouveau méta-théorème de correction (analogue du théorème 1) en étendant le lemme (lemme 1) aux nouvelles règles.

Remarque sur la preuve par récurrence La preuve par récurrence que nous avons donnée du théorème 1 a été énoncée sous la forme d’une preuve sur la taille des arbres de preuve. Il est plus usuel en informatique de faire ce genre de preuve sur la *structure* de l’arbre ; on parle de récurrence structurelle.

Le principe de récurrence sur les arbres s’énoncerait de la manière suivante : si une propriété d’arbres P est vraie sur tous les arbres réduits à une feuille, et est vraie sur un arbre dès lors qu’elle est vraie sur chacun de ses sous-arbres (au sens : pour chacun des sous-arbres enracinés en les fils de la racine), alors cette propriété est vraie pour tous les arbres.

Dans la preuve du théorème de correction, les arbres sont des arbres de preuve, et la propriété P considérée est celle d’avoir à sa racine un séquent valide. “La propriété est vraie pour les arbres réduits à une feuille” correspond au cas “règle d’hypothèse” du lemme ; les deux autres cas du lemme (modus ponens, introduction de l’implication) couvrent la partie “si la propriété est vraie pour les sous-arbres, elle est vraie pour l’arbre”.

3.5.2 Incomplétude

On aurait aimé prouver la réciproque du théorème de correction :

Tout séquent valide est prouvable en logique minimale.

Eh bien non ! Au contraire on peut prouver le résultat suivant :

Méta-théorème 2 (Incomplétude de la logique minimale). *Il existe au moins un séquent valide, et non prouvable en logique minimale.*

Idée de démonstration

On considère la *Formule de Peirce* : $((P \rightarrow Q) \rightarrow P) \rightarrow P$.

Exercice 9. *Montrer que la formule de Peirce est une tautologie ; autrement dit, le séquent $\vdash ((P \rightarrow Q) \rightarrow P) \rightarrow P$ est valide.*

Une analyse des propriétés de la logique minimale (que nous ne détaillons pas ici) permet de démontrer qu’aucune preuve de la logique minimale ne permet de prouver le séquent $\vdash ((P \rightarrow Q) \rightarrow P) \rightarrow P$.

Commentaires sur le résultat d’incomplétude

Bon, ce dernier résultat semble montrer que la logique minimale est un peu faible, puisqu’elle est incapable de démontrer certaines tautologies. Par ailleurs, elle nous a permis de nous familiariser avec les notions de règles d’inférence, de preuve, de règle dérivée, ...

Dans les chapitres suivants, nous procéderons à des extensions successives de cette logique, en enrichissant la syntaxe et l’ensemble des règles d’inférence.

3.5.3 Théorème de substitution uniforme

La logique propositionnelle permet, une fois qu'on a prouvé un théorème de la forme $\Gamma \vdash A$ d'en déduire une infinité d'autres sans effort ©. En fait, cela revient à transformer tout séquent en une règle. Pour ce faire, il suffit de remplacer *uniformément* des variables propositionnelles par des propositions à la fois dans Γ et dans A .

Plus précisément :

Définition 1 (Substitution). *On considère un ensemble de n variables propositionnelles $\{v_1, \dots, v_n\} \subseteq V_p$, et n formules $B_1, \dots, B_n$. On définit récursivement la substitution dans une proposition A les variables v_i par des formules B_i (pour $i \in 1..n$), notée $A[v_1/B_1, \dots, v_n/B_n]$ par*

- $v_i[v_1/B_1, \dots, v_n/B_n] = B_i$
- $v[v_1/B_1, \dots, v_n/B_n] = v$ ($v \in V_p \setminus \{v_1, \dots, v_n\}$)
- $(A \rightarrow B)[v_1/B_1, \dots, v_n/B_n] = (A[v_1/B_1, \dots, v_n/B_n] \rightarrow B[v_1/B_1, \dots, v_n/B_n])$

Exercice 10. *Étendre la notion de substitution aux séquents, de façon à définir la notation $(\Gamma \vdash A)[v_1/B_1, \dots, v_n/B_n]$*

Méta-théorème 3 (Théorème de substitution uniforme). *Soit un séquent prouvable en logique minimale $\Gamma \vdash A$. Alors le séquent $(\Gamma \vdash A)[v_1/B_1, \dots, v_n/B_n]$ (avec les notations ci-dessus) est aussi prouvable en logique minimale.*

Exemple

On considère le jugement $A \rightarrow B, B \rightarrow C \vdash_M A \rightarrow C$. Si l'on remplace *uniformément* A par $P \rightarrow Q$, B par R , et C par $S \rightarrow T$, on obtient le jugement suivant : $(P \rightarrow Q) \rightarrow R, R \rightarrow S \rightarrow T \vdash_M (P \rightarrow Q) \rightarrow S \rightarrow T$.

Remarques

Attention aux parenthèses !

Le jugement $P \rightarrow Q \rightarrow R, Q \rightarrow R \rightarrow S \vdash_M P \rightarrow S$ n'est pas une instance de $A \rightarrow B, B \rightarrow C \vdash_M A \rightarrow C$.

En effet si l'on veut substituer P à A , $Q \rightarrow R$ à B et S à C , nous obtenons plutôt le jugement ci-dessous, de structure différente :

$$P \rightarrow (Q \rightarrow R), (Q \rightarrow R) \rightarrow S \vdash_M P \rightarrow S$$

On se rappellera les problèmes liés à l'usage du `#define` du langage C, où une méconnaissance des priorités et l'associativité des opérateurs peuvent conduire à des substitutions malheureuses. En cas de doute, il vaut mieux définir les substitutions par des expressions bien parenthésées, quitte à enlever les parenthèses inutiles après avoir effectué la substitution.

Dans notre exemple, la substitution correcte est :

substituer P à A , $(Q \rightarrow R)$ à B et S à C

On obtient bien le jugement correct.

Par ailleurs, pour que le théorème de substitution uniforme soit correctement appliqué, il est important que la substitution soit *uniformément* appliquée dans le contexte et la conclusion du séquent : *toute* occurrence des variables propositionnelles sujettes à substitution doit être remplacée.

Prenons comme exemple à ne pas suivre le séquent (prouvable)

$$P \rightarrow Q \vdash (Q \rightarrow R) \rightarrow P \rightarrow R$$

Si l'on remplace Q par S *seulement dans la conclusion du séquent* on obtient le séquent non valide (et donc non prouvable)

$$P \rightarrow Q \vdash (S \rightarrow R) \rightarrow P \rightarrow R$$

3.6 Stratégies de démonstration

Supposons que l'on cherche à écrire (si possible) une preuve d'un séquent $\Gamma \vdash A$.

Une démarche possible est de considérer la preuve de ce séquent comme un *but* à atteindre ; nous pourrions utiliser la notation $\Gamma \vdash^? A$ pour désigner un tel but, c'est à dire une preuve incomplète.

Par exemple, considérons le but $\Gamma \vdash^? P \rightarrow R$, avec $\Gamma = \{P \rightarrow Q, Q \rightarrow R\}$.

Au départ de notre recherche de preuve, nous avons la situation suivante :

$$\Gamma \vdash^? P \rightarrow R$$

La conclusion du but courant étant une implication, nous pouvons tenter d'appliquer $\rightarrow_i$, ce qui engendre un nouveau *sous-but*.

$$\frac{\Gamma, P \vdash^? R}{\Gamma \vdash P \rightarrow R} \rightarrow_i$$

On pense résoudre le but $\Gamma, P \vdash^? R$ en appliquant le modus ponens, ce qui nous donne deux sous-buts :

$$\frac{\Gamma, P \vdash^? Q \rightarrow R \quad \Gamma, P \vdash^? Q}{\Gamma, P \vdash R} \text{ mp}$$

$$\frac{\Gamma, P \vdash R}{\Gamma \vdash P \rightarrow R} \rightarrow_i$$

Le sous-but de gauche se résout immédiatement avec la règle d'hypothèse, celui de droite peut utiliser encore une fois le modus-ponens.

$$\frac{\frac{\Gamma, P \vdash Q \rightarrow R}{\Gamma, P \vdash Q \rightarrow R} \text{ hyp} \quad \frac{\Gamma, P \vdash^? P \rightarrow Q \quad \Gamma, P \vdash^? P}{\Gamma, P \vdash Q} \text{ mp}}{\Gamma, P \vdash R} \text{ mp}$$

$$\frac{\Gamma, P \vdash R}{\Gamma \vdash P \rightarrow R} \rightarrow_i$$

On obtient finalement une preuve complète (sans “?”) !

$$\frac{\frac{\frac{\Gamma, P \vdash Q \rightarrow R}{\Gamma, P \vdash Q \rightarrow R} \text{hyp} \quad \frac{\frac{\Gamma, P \vdash P \rightarrow Q}{\Gamma, P \vdash P \rightarrow Q} \text{hyp} \quad \frac{\Gamma, P \vdash P}{\Gamma, P \vdash P} \text{hyp}}{\Gamma, P \vdash Q} \text{mp} \quad \frac{\Gamma, P \vdash R}{\Gamma \vdash P \rightarrow R} \rightarrow_i}{\Gamma \vdash P \rightarrow R} \text{mp}$$

Si l’on veut généraliser cet exemple, on peut proposer l’heuristique suivante pour prouver un séquent $\Gamma \vdash A$:

1. Créer un but $\Gamma \vdash^? A$.
2. tant qu’il reste un but non résolu de la forme $\Gamma \vdash^? A$, appliquer l’une des tactiques suivantes :
 - Si $A \in \Gamma$, le but se résout immédiatement à l’aide de la règle d’hypothèse.
 - Si A est de la forme $A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow A$, appliquer la règle d’introduction de l’implication, et engendrer un nouveau but $\Gamma, A_1, A_2, \dots, A_n \vdash^? A$
 - Si l’on peut prouver $\Gamma \vdash A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow A$ ⁶, alors appliquer la règle d’élimination de l’implication, et engendrer les sous-buts $\Gamma \vdash^? A_1, \Gamma \vdash^? A_2, \dots, \Gamma \vdash^? A_n$.

Remarque

Appliquer une telle stratégie de démonstration revient à dessiner un arbre de preuve en commençant par la racine (le séquent qu’on veut prouver) et en terminant par les feuilles (résolution de sous-buts triviaux).

Cette stratégie est celle utilisée par la tactique **auto** de l’assistant de preuve *Coq*.

Remarque

On peut bien sûr établir des stratégies plus élaborées (par exemple en appliquant des règles dérivées). On remarquera cependant le cas particulier de la règle de coupure :

Soit un but $\Gamma \vdash^? A$. On peut essayer de résoudre ce but en prouvant un lemme :

$$\frac{\Gamma \vdash^? B \quad \Gamma, B \vdash^? A}{\Gamma \vdash A} \text{cut}$$

Il reste que la proposition B peut très bien ne pas être une sous-formule de $\Gamma \cup \{A\}$. C’est à l’auteur de la démonstration de trouver quelle formule B aidera à compléter la preuve en cours.

Exercice 11. Reprendre les exercices 3 et 5, en utilisant la stratégie de recherche systématique de preuve.

6. en particulier, si la proposition $A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow A$ fait partie des hypothèses de Γ

Exercice 12. *Il est difficile de trouver à la main une preuve du séquent suivant :*

$$(((P \rightarrow Q) \rightarrow P) \rightarrow P) \rightarrow Q \vdash Q$$

En revanche, il ne posera pas de problème en séance sur machine.

3.7 Logique et Programmation Fonctionnelle

Dans cette section, nous commençons à explorer la correspondance entre la programmation fonctionnelle et la logique. Cette correspondance, dite *de Curry-Howard* est à la base de la réalisation d’assistants de preuves tels que *Coq*, dont nous allons utiliser les notations, très voisines de celles du langage *OCaML*.

3.7.1 Types et propositions

Considérons un ensemble de propositions atomiques P, Q , etc. Nous avons vu comment former les propositions de la logique minimale à l’aide du connecteur $\rightarrow$. Si nous considérons P, Q comme des types de base du lambda-calcul simplement typé, nous pouvons associer à toute proposition de la logique minimale un type du lambda-calcul simplement typé.

- Si P est une proposition atomique, on lui associe le type de base P
- Soit une implication de la forme $A \rightarrow B$, on lui associe le type $A \rightarrow B$ des fonctions totales du type A vers le type B .

On remarque que jusqu’à présent, cette “correspondance” n’est que l’identité. Le plus important est de voir dans cette correspondance une assimilation entre deux notions de nature apparemment différente : l’implication (intuitionniste) logique, et le typage des fonctions dans un langage de programmation.

3.7.2 Termes et Preuves

Nous proposons une représentation des preuves de la logique minimale par des expressions (aussi appelés *termes*) de notre langage fonctionnel. Cette représentation est basée sur l’assimilation entre deux jugements : *t est une preuve de A* et *t est un terme de type de A* .

Plus précisément, nous devons prendre en compte les notions d’hypothèse et de séquent, et nous nous intéresserons aux jugements “*dans le contexte Γ , t est une preuve de A* ” et “*dans le contexte Γ , t est un terme de type de A* ”.

Dans le premier jugement, Γ est un ensemble de propositions considérées comme les prémisses du jugement, dans le second, Γ est un ensemble de déclarations de variables. Afin de rendre compatibles ces deux visions du contexte, on conviendra que chaque hypothèse d’un contexte doit avoir un nom (distinct de celui des autres hypothèses).

Du coup, un contexte de la logique minimale aura la forme $h_1 : A_1, h_2 : A_2, \dots, h_n : A_n$, où les h_i sont des identificateurs deux à deux distincts. Chaque déclaration $h_i : A_i$ peut se lire dans un cadre comme “*L’hypothèse h_i d’énoncé A_i* ”, dans l’autre comme “*La déclaration de la variable h_i de type A_i* ”.

Reprenons les trois règles de la logique minimale, en les mettant en correspondance avec les règles de typage de la programmation fonctionnelle.

Un jugement pourra s'écrire $\Gamma \vdash_M t : A$ et se lire “dans le contexte Γ , t est une preuve⁷ de A ”, mais aussi “dans le contexte Γ , t est un terme de type A ”.

Règle d'hypothèse

On peut amender la règle telle que présentée en section 3.3.1 page 19 :

Si le contexte Γ contient la déclaration $a : A$, alors $\Gamma \vdash_M a : A$

Règle du modus ponens

La règle d'élimination de l'implication (modus-ponens) correspond à la règle de typage associée à l'application d'une fonction en lambda-calcul simplement typé.

$$\frac{\Gamma \vdash f : A \rightarrow B \quad \Gamma \vdash a : A}{\Gamma \vdash f a : B} \text{ mp}$$

Autrement dit : Une preuve de $A \rightarrow B$ est une fonction f qui, appliquée à une preuve a de A , retourne une preuve de B . Cette dernière preuve est juste l'application de la fonction f à l'argument a .

Par exemple, soit un contexte Γ contenant les déclarations $f : A \rightarrow B \rightarrow C$, $b : B$, et $c : C$. Alors le terme $f a$ aura pour type $B \rightarrow C$, et le terme $f a b$ ⁸ aura pour type C . Donc, le terme $f a b$ est une preuve de la proposition C dans le contexte Γ .

Remarque

La généralisation du modus-ponens vue en section 3.4.1 page 27 s'interprète très facilement en programmation fonctionnelle. *Rappelons que dans les langages fonctionnels comme OCaml et Coq, l'application d'une fonction à son argument est considérée comme une opération associant à gauche : l'écriture $f a_1 a_2 \dots a_{n-1} a_n$ est une abréviation de $((f a_1) a_2) \dots a_{n-1} a_n$.*

$$\frac{\Gamma \vdash f : A_1 \rightarrow A_2 \rightarrow \dots A_{n-1} \rightarrow A_n \rightarrow B \quad \Gamma \vdash a_1 : A_1 \quad \Gamma \vdash a_2 : A_2 \quad \dots \quad \Gamma \vdash a_n : A_n}{\Gamma \vdash f a_1 a_2 \dots a_n : B} \text{ mp}$$

Règle d'abstraction

Puisque la preuve d'une implication $A \rightarrow B$ s'interprète comme une fonction du type fonctionnel $A \rightarrow B$, l'introduction de l'implication s'apparente à la règle de typage des fonctions en lambda-calcul simplement typé.

$$\frac{\Gamma, a : A \vdash t : B}{\Gamma \vdash \text{fun } a : A \Rightarrow t : A \rightarrow B} \rightarrow_i$$

7. On dit aussi “est un terme de preuve de A ”

8. avec les conventions syntaxiques de *Coq* et *OCaml* sur l'application de fonction

Exemple

Par exemple, considérons le but $\Gamma \vdash^2 P \rightarrow R$, avec $\Gamma = \{P \rightarrow Q, Q \rightarrow R\}$.

On commence par nommer les hypothèses de Γ : $\Gamma = \{H : P \rightarrow Q, H0 : Q \rightarrow R\}$.
Pour construire un terme de type $P \rightarrow R$ dans Γ , on cherche à construire un terme de type R dans le contexte étendu $\Gamma' = \{H : P \rightarrow Q, H0 : Q \rightarrow R, p : P\}$.

L'application $H0(H\ p)$ est clairement de type R dans Γ' ,

et donc l'abstraction $\text{fun } p : P \Rightarrow H0(H\ p)$ est bien une preuve de $P \rightarrow R$ dans Γ .

Résumé

On peut donc convenir qu'une preuve en logique minimale du séquent $h_1 : A_1, h_2 : A_2, \dots, h_n : A_n \vdash A$ est un terme du lambda-calcul simplement typé de type A dans le contexte $h_1 : A_1, h_2 : A_2, \dots, h_n : A_n$.

Exercice 13. Reprendre les exercices 3 et 5, en donnant le terme du lambda-calcul simplement typé pour chacune des preuves.

Chapitre 4

Logique Propositionnelle (intuitionniste, puis classique)

Dans ce chapitre, nous étudions une première extension de la logique minimale, afin d’étudier les schémas de raisonnement par l’absurde, puis nous introduirons les connecteurs $\wedge$, $\vee$, et $\Leftrightarrow$.

Finalement, nous considérons deux versions de la logique propositionnelle, selon que nous incluons ou pas le *principe du tiers exclu*.

4.1 La contradiction et la négation

4.1.1 Introduction

La *négation* est omniprésente dans le langage courant : en voici quelques exemples :

“Socrate n’est pas immortel.”

“Je ne te hais point.”

“Il ne pleuvra pas au sud d’une ligne Bordeaux Lyon.”

En Mathématiques et en Informatique, elle apparaît dans des situations très courantes, parfois masquée par des abréviations :

- La proposition $i \neq j$ est une abréviation de $\neg(i = j)$
- La proposition $i < j$ est une abréviation de $i \leq j \wedge i \neq j$, etc
- L’instruction **assert** de **C**, permet de stopper l’exécution d’un programme si son argument est une expression fausse (voir figure 4.1 page suivante)

```

1 double harmonic(long t[], int n) {
2     double s = 0.0;
3     assert (n > 0);
4     { int i = 0;
5         for(i=0, i < n ; i++){
6             assert(t[i] != 0.0);
7             s += 1.0/t[i];
8         }
9         ...
10    }
11    return n/s ;
12 }

```

FIGURE 4.1 – Utilisation d'`assert` en C

La notion de contradiction, plus simplement de proposition fausse, est tellement ancrée dans nos neurones qu'il semble difficile d'en donner une *définition* précise. Des exemples de propositions fausses seraient :

- “ $2 = 3$ ”
- “Bordeaux est une ville de montagne”
- “C est le seul langage de programmation existant”

4.1.2 Syntaxe

Pour des raisons de simplicité, nous procédons en deux étapes :

1. Nous introduisons une proposition “toujours fausse” : la *contradiction*, notée $\perp$.
2. Nous en dérivons la *négation* : La négation notée $\sim A$ d'une proposition A est juste une abréviation de la proposition $A \rightarrow \perp$.

La contradiction

Définition 2. On ajoute une clause à la définition de l'ensemble des propositions donnée en section 3.1 page 15 :

- toute variable propositionnelle est une proposition,
 - la contradiction $\perp$ est une proposition,
 - si A et B sont deux propositions, alors l'implication $(A \rightarrow B)$ est une proposition
- sous-entendu : Il n'y a pas d'autres propositions

La négation

En suivant la tradition des *macros* des langages de programmation, nous ne considérons pas la *négation* comme une nouvelle construction mais comme une abréviation.

Soit A une proposition ; la proposition $\sim A$ est une abréviation de $(A \rightarrow \perp)$.

L'opérateur $\sim$ se prononce “non”. La notation $\neg A$ est aussi très usitée.

Par exemple, la proposition $(P \rightarrow Q) \rightarrow \sim Q \rightarrow \sim P$ est une abréviation de $(P \rightarrow Q) \rightarrow (Q \rightarrow \perp) \rightarrow P \rightarrow \perp$ et donc une instance de $(A \rightarrow B) \rightarrow (B \rightarrow C) \rightarrow A \rightarrow C$

En reprenant un exemple antérieur, la proposition $i \neq j$ est donc une abréviation de $i = j \rightarrow \perp$

Une proposition toujours vraie ?

Nous avons introduit, avec le symbole $\perp$, une proposition qui est “toujours fausse”. Pourquoi ne pas faire de même avec une proposition “toujours vraie” ?

Nous n'avons pas, *stricto sensu*, besoin d'un tel ajout : n'importe quelle tautologie prouvable sans hypothèse fera l'affaire, par exemple $P \rightarrow P$ où P est l'une quelconque de nos variables propositionnelles, ou $\perp \rightarrow \perp$.

Nous pourrions aussi introduire un nouveau symbole (typiquement $\top$) qui jouerait le rôle d'une telle proposition atomique “triviale” ; cela nécessiterait d'ajouter à notre arsenal une ou plusieurs règles permettant de raisonner sur les formules faisant intervenir ce nouveau symbole.

Exercice 14. *Supposons que l'on souhaite introduire un nouveau symbole $\top$ représentant une nouvelle proposition atomique “toujours vraie”, et prouvable dans n'importe quel contexte (y compris le contexte vide). Quelle(s) règle(s) de raisonnement (d'introduction, d'élimination) faudrait-il ajouter ?*

4.1.3 Aspects sémantiques de la contradiction

Pour définir la sémantique de propositions et de séquents contenant la proposition $\perp$, et, par conséquent, la négation $\sim$, il suffit de poser que pour n'importe quelle valuation ν , on a $\nu(\perp) = \mathbf{f}$. Par conséquent, pour toute proposition A , nous avons

$$\nu(\sim A) = \nu(A \rightarrow \perp) \quad (4.1)$$

$$= \mathbf{f} \text{ si } \nu(A) = \mathbf{v} \quad (4.2)$$

$$= \mathbf{v} \text{ si } \nu(A) = \mathbf{f} \quad (4.3)$$

Considérons maintenant un jugement de la forme $\Gamma \models \perp$. Nous pouvons en inférer que toute valuation ν qui satisfait (simultanément) toutes les hypothèses de Γ satisfait également $\perp$, ce qui est impossible : donc *aucune* valuation ne satisfait toutes les hypothèses ; on dit que ces hypothèses sont *contradictoires*.

Soit maintenant une proposition quelconque A . Il n'existe aucune valuation ν satisfaisant toutes les hypothèses de Γ et ne satisfaisant pas la conclusion A , donc il n'existe aucun contre-exemple possible à la validité de $\Gamma \vdash A$. On peut donc en conclure le jugement $\Gamma \models A$.

Nous venons donc de prouver le résultat suivant :

Si $\Gamma \models \perp$, alors $\Gamma \models A$, pour n'importe quelle proposition A .

4.1.4 Le principe d'explosion

La règle d'inférence ci-dessous dite “*principe d'explosion*”, mais aussi “*ex falso (sequitur) quodlibet*”, ou “*élimination de la contradiction*” reflète le raisonnement du paragraphe précédent.

$$\frac{\Gamma \vdash \perp}{\Gamma \vdash A} \perp_e$$

Voici un exemple de preuve utilisant ce principe.

$$\frac{\overline{\perp \vdash \perp} \text{ hyp}}{\perp \vdash P} \perp_e$$

Cette règle est la base du raisonnement par l'absurde.

Remarque

La règle $\perp_e$ peut sembler artificielle au premier abord. Remarquons cependant que le principe d'explosion (sous sa dénomination latine) est connu depuis le moyen âge (pour ne pas dire l'antiquité). Nous verrons plus loin (Section 4.1.5 page suivante) comment cette règle sert à valider le “raisonnement par l'absurde”.

Une autre façon de concevoir la règle $\perp_e$ serait de considérer la contradiction $\perp$ comme une façon de dire “toute proposition est vraie”. Si toute proposition est vraie, alors $2 = 2$, $2 \neq 2$ sont vraies, et la notion de vérité perd tout sens. Nous ne nous attarderons pas sur cet aspect, car cela demanderait d'étudier la *logique du second ordre* qui sort des ambitions de ce cours ☹.

Exercice 15. Faire des recherches sur Internet sur les termes “*ex falso quodlibet sequitur*” et “*principe of explosion*”.

4.1.5 Règles associées à la négation

Rappelons que toute proposition de la forme $\sim A$ est juste une abréviation de $A \rightarrow \perp$. Les règles d'inférence associées à la négation seront donc des règles dérivées.

Macro-expansion

Dans toute preuve, on peut remplacer n'importe quelle sous-formule de la forme $\sim A$ par $(A \rightarrow \perp)$ et réciproquement.

Introduction de la négation

Pour prouver un séquent $\Gamma \vdash \sim A$ il suffit d'expanser la négation et d'appliquer la règle d'introduction.

$$\frac{\frac{\Gamma, A \vdash \perp}{\Gamma \vdash A \rightarrow \perp} \rightarrow_i}{\Gamma \vdash \sim A} \sim\text{-expansion}$$

D'où la règle dérivée :

$$\frac{\Gamma, A \vdash \perp}{\Gamma \vdash \sim A} \sim_i$$

Contraposée

$$\frac{\Gamma \vdash A \rightarrow B}{\Gamma \vdash \sim B \rightarrow \sim A} \text{ contraposée}$$

En voici la justification :

$$\frac{\frac{\frac{\Gamma \vdash A \rightarrow B}{\Gamma, \sim B, A \vdash A \rightarrow B} \text{ aff} \quad \frac{}{\Gamma, \sim B, A \vdash A} \text{ hyp} \quad \frac{}{\Gamma, \sim B, A \vdash \sim B} \text{ hyp}}{\Gamma, \sim B, A \vdash B} \text{ mp} \quad \frac{}{\Gamma, \sim B, A \vdash B \rightarrow \perp} \sim\text{-expansion}}{\frac{\Gamma, \sim B, A \vdash \perp}{\Gamma, \sim B \vdash \sim A} \sim_i} \rightarrow_i$$

Raisonnement par l'absurde (intuitionniste)

A l'aide de la règle $\perp_e$, nous pouvons implanter un premier mode de raisonnement par l'absurde : Pour prouver $\Gamma \vdash A$ il *suffit* de trouver une proposition B et de prouver les séquents $\Gamma \vdash \sim B$ et $\Gamma \vdash B$. En visualisant $\sim B$ comme l'implication $B \rightarrow \perp$, ce n'est rien d'autre que le raisonnement par *modus ponens* !

$$\frac{\Gamma \vdash \sim B \quad \Gamma \vdash B}{\Gamma \vdash A} \text{ absurde}$$

Exercice 16. Justifier (par une dérivation) cette nouvelle règle dérivée.

Exercice 17. En utilisant la règle *absurde*, prouver le séquent

$$P \rightarrow Q, P \rightarrow \sim Q \vdash P \rightarrow S$$

Remarque

Du point de vue de la preuve interactive de théorèmes, la règle **absurde** peut s'utiliser comme une “tactique d'élimination de la négation” : Si l'on cherche à prouver un séquent $\Gamma \vdash A$ et qu'on dispose déjà d'une preuve de $\Gamma \vdash \sim B$, alors il *suffit* de prouver aussi le séquent $\Gamma \vdash B$.

Il est remarquable dans ce cas que la proposition A “disparaît” du nouveau but. C'est encore une conséquence du principe d'explosion. Dans le cas considéré, la preuve de $\Gamma \vdash A$ revient plus à chercher une contradiction dans Γ qu'à chercher une preuve de A dans ce contexte.

Exercice 18. *Prouver les règles dérivées suivantes (une fois prouvées, on pourra les utiliser dans tout raisonnement ultérieur)*

$$\frac{\Gamma \vdash A}{\Gamma \vdash \sim \sim A} \text{ double négation}_i$$

$$\frac{\Gamma \vdash \sim \sim \sim A}{\Gamma \vdash \sim A} \text{ triple négation}_e$$

4.1.6 Logique intuitionniste : définition

Un séquent $\Gamma \vdash A$ est prouvable en logique propositionnelle intuitionniste si l'on peut en construire une preuve utilisant les règles de la logique minimale ainsi que les règles associées à la contradiction et à la négation (ainsi que les règles dérivées).

On note $\Gamma \vdash_J A$ le jugement “ $\Gamma \vdash A$ est prouvable en logique propositionnelle intuitionniste”.

Notons que $\vdash_J$, à l'instar de $\vdash_M$, *n'est pas un connecteur*. On ne confondra pas le *jugement* $A \vdash_J B$ avec la *proposition* $A \rightarrow B$.

4.1.7 Peut-on prouver $\perp$?

Il n'y a pas de règle d'introduction pour la contradiction.

Un séquent de la forme $\Gamma \vdash \perp$ ne peut être prouvable que si les hypothèses de Γ contiennent cette contradiction, soit directement, soit par la présence d'une négation.

Exemple

$$\frac{\frac{\frac{}{P \rightarrow Q \rightarrow \perp, P, Q \vdash P \rightarrow Q \rightarrow \perp} \text{hyp}}{\frac{P \rightarrow Q \rightarrow \perp, P, Q \vdash \perp}{P \rightarrow \sim Q, P, Q \vdash \perp} \sim\text{-expansion}} \text{hyp} \quad \frac{\frac{}{P \rightarrow Q \rightarrow \perp, P, Q \vdash Q} \text{hyp}}{\frac{P \rightarrow Q \rightarrow \perp, P, Q \vdash P}{P \rightarrow Q \rightarrow \perp, P, Q \vdash P} \text{hyp}} \text{mp}$$

Voici la même preuve, dans un autre format :

```

1  Supposons  $P \rightarrow \sim Q$ 
2  Supposons  $P$ 
3  Supposons  $Q$ 
4   $P \rightarrow Q \rightarrow \perp$  [expansion de  $\sim$ , 1]
5   $\perp$  [mp, 4, 2, 3]

```

4.1.8 Autour de la double négation

On remarquera que dans l'exercice 18 page précédente, on ne demande pas de prouver de règle de la forme :

$$\frac{\Gamma \vdash \sim\sim A}{\Gamma \vdash A} \text{ double négation}_e$$

La raison en est encore un problème d'incomplétude. En effet, bien que pour toute proposition A et tout contexte Γ , le jugement $\Gamma \models \sim\sim A$ implique le jugement $\Gamma \models A$, la règle ci-dessus n'est pas dérivable en logique intuitionniste. On verra dans la section 4.4 que cette règle est dérivable dans le cadre de la logique dite *classique*.

4.2 Les connecteurs binaires

4.2.1 La conjonction

La conjonction permet, à partir de deux propositions A et B de former une seule proposition qui est vraie si et seulement si A et B sont vraies.

Syntaxe

On étend de nouveau la définition de l'ensemble des propositions (Section 2 page 43), en ajoutant la clause :

Définition 3. Si A et B sont deux propositions, alors $(A \wedge B)$ est une proposition, prononcée " $(A \text{ et } B)$ "

On considèrera que $\wedge$ est un opérateur associatif à droite (comme $\rightarrow$) et de priorité supérieure à $\rightarrow$ et inférieure à $\sim$.

Par exemple, la proposition $\sim P \wedge Q \rightarrow R \wedge S$ doit se lire comme $((\sim P) \wedge Q) \rightarrow (R \wedge S)$. Une proposition de la forme $A \wedge B \wedge C$ se lira comme $A \wedge (B \wedge C)$.

Sémantique

On étend de nouveau la définition de la valuation d'une formule en ajoutant une clause permettant d'obtenir $\nu(A \wedge B)$ à partir de $\nu(A)$ et $\nu(B)$, ce qui se résume par la table de vérité suivante.

A	B	$A \wedge B$
f	f	f
f	v	f
v	f	f
v	v	v

Règle d'introduction

Dans un raisonnement oral, on ne fait pas vraiment la différence entre “prouver $A \wedge B$ ” et “prouver A , et prouver B ”. La règle d'introduction de la conjonction reflète exactement cette remarque.

Pour prouver un séquent $\Gamma \vdash A \wedge B$, il suffit de prouver séparément les séquents $\Gamma \vdash A$ et $\Gamma \vdash B$:

$$\frac{\Gamma \vdash A \quad \Gamma \vdash B}{\Gamma \vdash A \wedge B} \wedge_i$$

Règle d'élimination

La règle d'élimination de la conjonction peut paraître un peu déroutante, mais la pratique de nombreux exemples en fera saisir l'intuition. Elle exprime que, pour prouver une proposition C à partir de $A \wedge B$, il suffit de prouver C à partir de A et B *prises séparément*.

$$\frac{\Gamma \vdash A \wedge B \quad \Gamma, A, B \vdash C}{\Gamma \vdash C} \wedge_e$$

Dérouté ? Prenons comme exemple la validation de deux règles dérivées plus intuitives.

$$\frac{\Gamma \vdash A \wedge B}{\Gamma \vdash A} \wedge_1$$

$$\frac{\Gamma \vdash A \wedge B}{\Gamma \vdash B} \wedge_2$$

Montrons que $\wedge_1$ est bien une règle dérivée. La preuve pour $\wedge_2$ est laissée en exercice.

$$\frac{\Gamma \vdash A \wedge B \quad \overline{\Gamma, A, B \vdash A} \text{ hyp}}{\Gamma \vdash A} \wedge_e$$

Exercice 19. Montrer les règles suivantes :

$$\overline{A \wedge B \vdash B \wedge A} \text{ commutativité de } \wedge$$

$$\overline{(A \wedge B) \wedge C \vdash A \wedge (B \wedge C)} \text{ associativité de } \wedge$$

$$\overline{A \wedge \perp \vdash \perp} \text{ loi d'absorption de } \wedge$$

4.2.2 La disjonction

Reprenons le même plan que pour la conjonction (nom savant du “ou”, toujours considéré comme inclusif).

On prendra garde au fait que le langage courant est souvent ambigu sur le statut sémantique du “ou”. Il est parfois inclusif (“Pour discuter de votre orientation, prenez un rendez-vous avec votre responsable pédagogique en discutant avec lui à la sortie d’un cours, ou envoyez-lui un courrier électronique” : le fait d’engager les deux démarches n’empêche pas d’obtenir un rendez-vous), et parfois exclusif (le restaurateur qui affiche sur son menu “fromage ou dessert” a généralement l’intention de facturer un supplément au client qui demande les deux). Cette ambiguïté n’est pas acceptable dans un cadre formel, et suivant la tradition mathématique comme informatique, le “ou” sera toujours pris dans le sens inclusif.

Syntaxe

On étend de nouveau la définition de l’ensemble des propositions (Section 2 page 43), en ajoutant la clause :

Définition 4. Si A et B sont deux propositions, alors $(A \vee B)$ est une proposition, prononcée “ $(A \text{ ou } B)$ ”

On considérera que $\vee$ est un opérateur associant à droite et de priorité supérieure à $\rightarrow$ et inférieure à $\wedge$.

Par exemple, la proposition $\sim P \wedge Q \rightarrow R \vee S \wedge T$ doit se lire comme $((\sim P) \wedge Q) \rightarrow (R \vee (S \wedge T))$

Sémantique

Une fois de plus, on étend la définition d’une valuation en ajoutant une façon de déterminer $\nu(A \vee B)$ à partir de $\nu(A)$ et de $\nu(B)$, résumée par la table de vérité :

A	B	$A \vee B$
f	f	f
f	v	v
v	f	v
v	v	v

Règles d’introduction

À cause de la fréquente ambiguïté du langage courant, mentionnée précédemment, il est difficile de faire appel à l’intuition pour définir ce que c’est que “prouver $A \vee B$ ”.

La disjonction possède *deux* règles d’introduction : Pour prouver un séquent $\Gamma \vdash A \vee B$, on peut soit prouver le séquent $\Gamma \vdash A$ soit prouver le séquent $\Gamma \vdash B$

$$\frac{\Gamma \vdash A}{\Gamma \vdash A \vee B} \vee_{i,1}$$

$$\frac{\Gamma \vdash B}{\Gamma \vdash A \vee B} \vee_{i,2}$$

Règle d'élimination

La règle d'élimination de la disjonction formalise la notion de *preuve par cas*.
Si l'on veut prouver $\Gamma \vdash C$ et que l'on sait prouver $\Gamma \vdash A \vee B$ alors il suffit de prouver C en considérant les deux cas possibles

- C peut être prouvée en supposant A
- C peut être prouvée en supposant B

$$\frac{\Gamma \vdash A \vee B \quad \Gamma, A \vdash C \quad \Gamma, B \vdash C}{\Gamma \vdash C} \vee_e$$

Sous un autre format, l'élimination de la disjonction se présente ainsi :

```

1  ...
2  A ∨ B
3  ...
4  { Supposons A
5      ...
6      C
7  }
8  { Supposons B
9      ...
10     C
11 }
12 C [Ve, 2, 4 – 6, 8 – 10]
```

Exemple

Voici une preuve, dans les deux formats, du séquent $A \vee \perp \vdash A$:

$$\frac{\frac{}{A \vee \perp \vdash A \vee \perp} \text{hyp} \quad \frac{}{A \vee \perp, A \vdash A} \text{hyp} \quad \frac{\frac{}{A \vee \perp, \perp \vdash \perp} \text{hyp} \quad \frac{}{A \vee \perp, \perp \vdash A} \perp_e}{A \vee \perp \vdash A} \vee_e$$

```

1  Supposons A ∨ ⊥
2  { Supposons A
3      A [hyp, 2]
4  }
5  { Supposons ⊥
6      A [⊥e, 5]
7  }
8  A [Ve, 1, 2-3, 5-6]
```

Exemple

Voici la preuve de la règle de “commutativité de $\vee$ ”

$$\frac{\frac{}{A \vee B \vdash A \vee B} \text{hyp} \quad \frac{\frac{}{A \vee B, A \vdash A} \text{hyp} \quad \frac{}{A \vee B, A \vdash B \vee A} \vee i, 2}{A \vee B \vdash B \vee A} \vee e}{\frac{}{A \vee B, B \vdash B} \text{hyp} \quad \frac{}{A \vee B, B \vdash B \vee A} \vee i, 1}{A \vee B \vdash B \vee A} \vee e}$$

Voici la même preuve sous un autre format :

```

1  Supposons  $A \vee B$ 
2  { Supposons  $A$ 
3     $A$  [hyp,2]
4     $B \vee A$  [ $\vee i, 2, 3$ ]
5  }
6  { Supposons  $B$ 
7     $B$  [hyp,6]
8     $B \vee A$  [ $\vee i, 1, 7$ ]
9  }
10  $B \vee A$  [ $\vee e, 1, 2-4, 6-8$ ]

```

Exercice 20. Prouver les règles dérivées suivantes :

$$\frac{\Gamma \vdash A \vee B \quad \Gamma \vdash \sim B}{\Gamma \vdash A}$$

$$\frac{}{\sim A \vee B \vdash A \rightarrow B}$$

$$\frac{}{\sim(A \vee B) \vdash \sim A \wedge \sim B} \text{ De Morgan}_1$$

4.2.3 L'équivalence logique

De même que pour la négation, nous allons considérer l'équivalence logique comme une abréviation plutôt qu'un nouveau connecteur : une équivalence se réduit à deux implications.

Définition 5. Si A et B sont deux propositions, alors $(A \leftrightarrow B)$ est une abréviation de $((A \rightarrow B) \wedge (B \rightarrow A))$.

L'opérateur $\leftrightarrow$ sera d'une priorité équivalente à celle de $\rightarrow$. Il conviendra de mettre suffisamment de parenthèses dans toute formule comportant à la fois ces deux connecteurs.

On peut dériver les règles suivantes :

$$\frac{\Gamma \vdash A \rightarrow B \quad \Gamma \vdash B \rightarrow A}{\Gamma \vdash A \leftrightarrow B} \leftrightarrow_i$$

$$\frac{\Gamma \vdash A \leftrightarrow B}{\Gamma \vdash A \rightarrow B} \leftrightarrow_e$$

$$\frac{\Gamma \vdash A \leftrightarrow B}{\Gamma \vdash B \rightarrow A} \leftrightarrow_e$$

$$\frac{\Gamma \vdash A \leftrightarrow B}{\Gamma \vdash B \leftrightarrow A} \leftrightarrow\text{-sym}$$

L'équivalence étant une simple abréviation comme la négation, elle ne fait pas l'objet de règles d'inférence spécifiques.

4.2.4 Équivalence entre propositions

Soient A et B deux formules. On notera $A \equiv_J B$ le *jugement* $\vdash_J A \leftrightarrow B$.

Par exemple, la loi de De Morgan de l'exercice précédent s'exprime par le jugement $\sim(A \vee B) \equiv_J \sim A \wedge \sim B$.

Notons que $\equiv$ n'est pas un connecteur. On ne confondra pas le *jugement* $A \equiv_J B$ avec la *proposition* $A \leftrightarrow B$.

Exercice 21. Montrer les jugements suivants (qu'on pourra par la suite utiliser comme règles dérivées) :

$$A \wedge \perp \equiv_J \perp \quad (4.4)$$

$$A \vee \perp \equiv_J A \quad (4.5)$$

$$A \vee B \equiv_J B \vee A \quad (4.6)$$

$$A \wedge B \equiv_J B \wedge A \quad (4.7)$$

$$A \vee (B \vee C) \equiv_J A \vee B \vee C \quad (4.8)$$

$$A \wedge (B \wedge C) \equiv_J A \wedge B \wedge C \quad (4.9)$$

$$A \wedge (B \vee C) \equiv_J A \wedge B \vee A \wedge C \quad (4.10)$$

$$A \vee B \wedge C \equiv_J (A \vee B) \wedge (A \vee C) \quad (4.11)$$

$$\sim(A \rightarrow B) \equiv_J (A \wedge \sim B) \quad (4.12)$$

$$\sim(A \vee B) \equiv_J \sim A \wedge \sim B \text{ [loi de De Morgan]} \quad (4.13)$$

$$A \rightarrow \sim A \vdash_J \sim A \text{ [reductio ad absurdum]} \quad (4.14)$$

$$\sim A \rightarrow A \vdash_J \sim \sim A \quad (4.15)$$

4.3 Quelques méta-théorèmes

4.3.1 Cohérence et non-complétude

Les propriétés suivantes (que nous admettrons pour la plupart) permettent d'exprimer la cohérence de la logique propositionnelle intuitionniste.

Tout d'abord, nous avons un résultat similaire au théorème 1 page 32.

Méta-théorème 4 (Correction). *Si $\Gamma \vdash_J A$, alors $\Gamma \models A$.*

La preuve en est laissée en exercice ; en s’inspirant de la preuve donnée dans le chapitre sur la logique minimale, le squelette de la preuve reste le même et il suffit de faire des vérifications mineures correspondant aux connecteurs ajoutés.

Nous en déduisons (par contraposition) le corollaire suivant :

Méta-théorème 5 (non-contradiction). *Le séquent $\vdash \perp$ n’est pas démontrable en logique propositionnelle intuitionniste.*

En effet, le séquent $\vdash \perp$ n’est clairement pas valide, ce qui entraîne la non-prouvabilité de ce séquent.

Ce méta-théorème entraîne l’impossibilité de prouver une proposition *et* sa négation.

Méta-théorème 6. *Il n’existe aucune proposition A telle que $\vdash_J A$ et $\vdash_J \sim A$.*
Preuve laissée en exercice

Exercice 22. *L’énoncé suivant est-il vrai ou faux ? Pourquoi ?*

“ Il n’existe aucun séquent $\Gamma \vdash A$ tel qu’on ait à la fois $\Gamma \vdash_J A$ et $\Gamma \vdash_J \sim A$.”

Remarque

Le théorème de correction nous donne un moyen de vérifier qu’un séquent $\Gamma \vdash A$ n’est pas prouvable : il suffit de montrer qu’il n’est pas valide.

Non-complétude

Nous donnons ci-dessous une liste de séquents valides (au sens sémantique) mais non prouvables en logique propositionnelle intuitionniste (nous admettrons ce dernier résultat).

$$\vdash A \vee \sim A \text{ [tiers exclu]} \quad (4.16)$$

$$\sim \sim A \vdash A \text{ [élimination de la double négation]} \quad (4.17)$$

$$\vdash \sim(A \wedge B) \leftrightarrow \sim A \vee \sim B \text{ [loi de De Morgan]} \quad (4.18)$$

$$\vdash \sim(A \rightarrow B) \leftrightarrow (A \wedge \sim B) \quad (4.19)$$

$$\vdash (A \rightarrow B) \leftrightarrow (\sim B \rightarrow \sim A) \quad (4.20)$$

4.3.2 Relation avec la logique minimale

La logique propositionnelle intuitionniste est une extension de la logique minimale, dans la mesure où toute règle d’inférence de la logique minimale est aussi une règle de la logique propositionnelle intuitionniste. Toute preuve en logique minimale peut alors être convertie (“castée”) en une preuve de la logique propositionnelle intuitionniste :

Méta-théorème 7. Si $\Gamma \vdash_M A$, alors $\Gamma \vdash_J A$.

Méta-théorème 8. Toutes les règles dérivées de la logique minimale restent applicables dans la logique propositionnelle intuitionniste.

Exercice 23. Donner une idée de preuve de l'énoncé ci-dessus.

4.3.3 Remplacement d'une sous-formule

L'équivalence entre deux propositions a une conséquence intéressante, capturée par le théorème ci-dessous.

Méta-théorème 9. Si $\Gamma \vdash_J A \leftrightarrow B$ et $\Gamma \vdash_J C$, alors $\Gamma \vdash_J D$, où D est obtenue en remplaçant dans C une ou plusieurs occurrences (comme sous-formules) de A par B .

Exercice 24. Donner une idée de preuve de l'énoncé ci-dessus.

Remarque

Le théorème précédent permet d'affirmer que des séquents sont prouvables, sans en écrire explicitement de preuves formelles, et ce, de manière bien plus "chirurgicale" que le théorème de substitution : d'une part, A peut être une sous-formule arbitraire, et pas seulement une formule atomique ; d'autre part, on peut choisir, pour *chaque* occurrence de A comme sous-formule, de la remplacer ou non par B .

Exemple

$$\frac{\frac{\frac{}{\vdash (P \vee Q) \leftrightarrow (Q \vee P)}}{P \vee Q, \sim P \vdash (P \vee Q) \leftrightarrow (Q \vee P)} \text{ aff} \quad P \vee Q, \sim P \vdash Q}{Q \vee P, \sim P \vdash Q}$$

Remarques

Dans l'énoncé du théorème 9, le contexte des trois séquents est le même. En fait, on peut relâcher un peu cette condition ; il faut simplement s'assurer que l'équivalence qu'on utilise reste prouvable dans tous les contextes considérés.

Méta-théorème 10. Si $\Gamma \vdash_J A \leftrightarrow B$ et $\Delta \vdash_J C$ avec $\Gamma \subseteq \Delta$, alors $\Delta' \vdash_J D$, où Δ' et D sont obtenues en remplaçant dans Δ et C une ou plusieurs occurrences (comme sous-formules) de A par B , pour peu que l'on ait $\Gamma \subseteq \Delta'$.

Méta-théorème 11. En particulier, si $A \equiv_J B$, et si $\Gamma \vdash_J C$, alors $\Gamma' \vdash_J D$, où $\Gamma' \vdash D$ est obtenu en remplaçant dans $\Gamma \vdash C$ une ou plusieurs occurrences de A par B .

Exercice 25. Pouvez-vous justifier ces variantes ?

Exercice 26. *Inventez un exemple où le non-respect de la condition $\Gamma \subseteq \Delta'$ entraîne la dérivation d'un séquent absurde.*

4.4 La logique propositionnelle classique

Nous avons vu en 4.3.1 page 54 une liste de schémas “classiques” que la logique propositionnelle intuitionniste ne permet pas de prouver. À chaque fois, on a un séquent sémantiquement valide qui n’admet pas de preuve en logique propositionnelle intuitionniste.

Définition 6 (Logique classique). *La logique propositionnelle classique s'obtient en ajoutant à la logique propositionnelle intuitionniste la règle du tiers exclu¹ :*

$$\frac{}{\Gamma \vdash A \vee \sim A} \text{ } exm$$

Nous noterons $\Gamma \vdash_K A$ le jugement “Le séquent $\Gamma \vdash A$ est prouvable en logique propositionnelle classique”.

4.4.1 Règles dérivées pour la logique classique

Tous les séquents de la liste donnée en 4.3.1 page 54 deviennent prouvables en logique classique. Voici par exemple comment dériver l'élimination de la double négation :

$$\frac{\frac{}{\sim\sim A \vdash A \vee \sim A} \text{ exm + aff} \quad \frac{}{\sim\sim A, A \vdash A} \text{ hyp}}{\sim\sim A \vdash A} \quad \frac{\frac{\frac{}{\sim\sim A, \sim A \vdash \sim A} \text{ hyp} \quad \frac{}{\sim\sim A, \sim A \vdash \sim A} \text{ hyp}}{\sim\sim A, \sim A \vdash \perp} \text{ hyp} \quad \frac{}{\sim\sim A, \sim A \vdash A} \text{ absurde}}{\sim\sim A \vdash A} \text{ }_{\vee_e}$$

```

1  Supposons  $\sim \sim A$ 
2   $A \vee \sim A$  [exm]
3  { supposons  $A$ 
4     $A$  [hyp,3]
5  }
6  { supposons  $\sim A$ 
7     $A$  [absurde, 1, 6]
8  }
9   $A$  [ $\vee_e$ ,2,3-4,6-7]

```

Exercice 27. *Prouver les autres règles dérivées de 4.3.1 page 54.*

1. “Excluded middle” en anglais

4.5 Quelques méta-théorèmes pour la logique classique

4.5.1 Cohérence et complétude

Méta-théorème 12 (Correction). *Si $\Gamma \vdash_K A$, alors $\Gamma \models A$.*

Comme pour le théorème de correction de la logique intuitionniste, ce théorème peut être prouvé en reprenant le squelette de la preuve énoncée pour la logique minimale, en ajoutant les vérifications correspondant à la règle d'introduction du tiers exclu.

Nous en déduisons (par contraposition) le corollaire suivant :

Méta-théorème 13 (non-contradiction). *Le séquent $\vdash \perp$ n'est pas démontrable en logique classique.*

En effet, le séquent $\vdash \perp$ n'est clairement pas valide, ce qui entraîne la non-prouvabilité de ce séquent.

Ce méta-théorème entraîne l'impossibilité de prouver une proposition *et* sa négation.

Méta-théorème 14. *Il n'existe aucune proposition A telle que $\vdash_K A$ et $\vdash_K \sim A$.*
Preuve laissée en exercice

Remarque

Le théorème de correction nous donne un moyen de vérifier qu'un séquent $\Gamma \vdash A$ n'est pas prouvable : il suffit de montrer qu'il n'est pas valide.

Méta-théorème 15 (Complétude). *Si $\Gamma \models A$ alors $\Gamma \vdash_K A$.*

Preuve longue, assez technique, et ennuyeuse. On en trouve plein de versions sur Internet.

Équivalence en logique classique

On notera $A \equiv_K B$ le jugement $\vdash_K A \leftrightarrow B$.

Il est clair que si l'on a $A \equiv_J B$, alors $A \equiv_K B$.

Dans la mesure où le contexte – intuitionniste ou classique – est clair, on pourra simplifier en la notation $A \equiv B$.

4.5.2 Relation avec la logique intuitionniste

La logique propositionnelle classique est une extension de la logique intuitionniste. Toute preuve en logique propositionnelle intuitionniste peut être convertie (“castée”) en une preuve de la logique propositionnelle classique :

Méta-théorème 16. *Si $\Gamma \vdash_J A$, alors $\Gamma \vdash_K A$.*

Méta-théorème 17. *Toutes les règles dérivées de la logique intuitionniste restent applicables dans la logique propositionnelle classique.*

Exercice 28. Donner une idée de preuve de l'énoncé ci-dessus.

Méta-théorème 18. *Le théorème de substitution uniforme reste applicable en logique classique.*

4.5.3 Quelques règles dérivées en logique classique

Attention ! Les règles de cette section ne sont correctes qu'en logique classique et sont donc à éviter en logique propositionnelle intuitionniste.

Raisonnement par l'absurde classique

$$\frac{\Gamma, \sim A \vdash \perp}{\Gamma \vdash A} \text{ absurde classique}$$

Exercice 29. Montrer que cette règle est valide.

Elimination du tiers-exclu

$$\frac{\Gamma, B \vdash A \quad \Gamma, \sim B \vdash A}{\Gamma \vdash A} \text{ exm}_e$$

Exercice 30. Montrer que cette règle est valide.

Exercice 31. Valider la règle dérivée ci-dessous :

$$\frac{\Gamma, B \vdash A \quad \Gamma \vdash \sim A}{\Gamma \vdash \sim B}$$

Exercice 32. Montrer le théorème suivant :

Méta-théorème 19. Soient Γ un contexte et A une proposition, alors $\Gamma \vdash_K A$ si et seulement si $\Gamma \cup \{\sim A\}$ n'est pas satisfaisable.

Exercice 33. Prouver le jugement suivant

$$\sim A \rightarrow A \vdash_K A$$

Aide : On pourra d'abord prouver le jugement

$$\sim A \rightarrow A \vdash_J \sim \sim A$$

puis utiliser l'élimination de la double négation.

Exercice 34. Prouver le jugement suivant

$$\vdash_K (A \rightarrow B) \vee (B \rightarrow A)$$

Remarque : Nous avons là un exemple de jugement prouvable en logique classique, donc valide, mais dont l'intuition nous échappe.

Que veut dire, et surtout comment utiliser l'instance suivante ?

$$\vdash_K (x = 2 \rightarrow x = 3) \vee (x = 3 \rightarrow x = 2)$$

Exercice 35. Soient P et Q deux propositions atomiques.

1. Montrer l'équivalence $\sim P \rightarrow P \equiv_J \sim \sim P$.
2. En déduire l'équivalence $(\sim P \rightarrow P) \rightarrow P \equiv_J \sim \sim P \rightarrow P$.
3. (*) En déduire que si la formule de Peirce était prouvable en logique intuitionniste, alors on pourrait prouver l'élimination de la double négation c'est à dire prouver tout jugement de la forme $\sim \sim A \vdash_J A$. On pourra utiliser le théorème de substitution uniforme.
4. Donner une preuve en logique classique de la formule de Peirce :

$$((P \rightarrow Q) \rightarrow P) \rightarrow P$$

5. Donner une preuve en logique intuitionniste du séquent suivant :

$$P \vee \sim P \vdash ((P \rightarrow Q) \rightarrow P) \rightarrow P$$

6. (*) Montrer que si l'on ajoute à la logique intuitionniste la règle d'élimination de la double négation :

$$\frac{\Gamma \vdash \sim \sim A}{\Gamma \vdash A} \text{ double negation}_e$$

Alors, on peut en dériver la règle du tiers exclu.

7. Que peut-on en conclure sur le tiers-exclu, la loi de la double négation et la formule de Peirce ?

Chapitre 5

Calcul des prédicats (Logique du premier ordre)

5.1 Introduction

Dans les chapitres précédents, nous n'avons vu que des formules très abstraites, engendrées à partir de simples variables propositionnelles. Cette abstraction nous a permis d'étudier les règles et structures de raisonnement sans être influencés par l'application à des domaines précis.

Nous allons procéder à une extension des logiques propositionnelles (intuitionniste et classique) dans deux directions, qui nous autoriseront à présenter des exemples plus réalistes :

- Remplacement des variables propositionnelles par des formules plus complexes composées à partir d'objets, de fonctions, de prédicats et de relations
- Introduction des deux *quantificateurs* “pour tout” et “il existe au moins”

Les règles (de base et dérivées) de la logique propositionnelle seront encore applicables. Quant aux méta-théorèmes, nous examinerons un par un s'il sont encore applicables dans ce nouveau cadre.

5.2 Formules du premier ordre

5.2.1 Types

L'étudiant en informatique a déjà rencontré la notion de type dans le cadre des langages de programmation. Cette notion permet de contrôler que les expressions apparaissant dans un programme sont bien formées : dans le sens où une fonction ne peut être appliquée qu'à des arguments du type prévu. Par exemple la fonction `C itoa` ne doit être appliquée qu'à un argument de type `int` et renvoie un résultat de type `char *`. On dispose d'une déclaration similaire pour `strlen`.

On en déduit que si i est une variable de type `int`, alors l'expression `strlen (atoi (i))` est bien formée et de type `int`.

Syntaxe des types

On pourra considérer suivant les exemples des *types atomiques* qui seront des symboles déclarés comme tels. En voici quelques exemples, que nous reprendrons dans nos exemples :

- bool** Le type des valeurs booléennes
- nat** Le type des entiers naturels
- Z** Le type des nombres entiers
- int** Le type des entiers sur 32 bits (`int31` en *Coq*)
- char** Le type des caractères
- unit** Le type singleton (contenant une seule valeur)
- Empty_set** Le type vide
- individu** Un type pour les syllogismes à la Aristote

5.2.2 Constantes et symboles de fonctions

Les constantes sont des identificateurs à la base de la construction des *termes*, les expressions dont le calcul des prédicats permet d'étudier les propriétés. À chaque constante c sera associé un type unique. On utilisera la notation de *typage* $c : A$ pour exprimer que A est le type de c .

Les constantes sont en fait un cas particulier de *symboles de fonctions*. À chaque symbole de fonction est associé un type de la forme $A_1 \times A_2 \times \dots \times A_k \rightarrow A$ ou les A_i et A sont des types atomiques.

Ce type se veut le type des fonctions totales qui prennent leurs arguments dans les types $A_1, \dots, A_n$ et retournent un résultat de type A . Si $k = 0$, on retrouve la notion de constante.

Remarque

Dans nombre d'ouvrages sur la logique mathématique, on ne considère qu'un seul type de base. Dans ce cas, un symbole de fonction est caractérisé par son *arité*, c'est à dire son nombre d'arguments. Nous avons choisi de donner une définition plus proche des langages de programmation.

Exemples

Nous donnons quelques exemples de constantes et symboles de fonctions qui seront utilisés plus loin :

- `true` : `bool`
- `false` : `bool`
- `xor` : `bool × bool → bool`

- $O : \text{nat}$ *le nombre 0*
- $S : \text{nat} \rightarrow \text{nat}$ *(la fonction successeur)*
- $+$: $\text{nat} \times \text{nat} \rightarrow \text{nat}$
- $*$: $\text{nat} \times \text{nat} \rightarrow \text{nat}$
- $'a'$: char
- Socrate : individu

Remarques

1. On ne confondra pas la constante false avec la valeur de vérité f , ni avec la contradiction $\perp$. Ce sont des objets de nature totalement différente :
 - false est une constante associée au type bool
 - f est une valeur de vérité, du domaine de la sémantique
 - $\perp$ est une proposition logique
 De même pour true à ne pas confondre avec v .
2. **À déplacer** On s'autorisera les abréviations suivantes : 0 pour O , 1 pour $S\ O$, 2 pour $S\ (S\ O)$, etc.

5.2.3 Variables, Déclarations et Contextes

Afin de rester le plus compatibles possible avec *Coq*, nous allons légèrement modifier les notations de contextes et de séquents vus précédemment, le but étant d'unifier le traitement des hypothèses et des déclarations de variables.

Nous considérons dans toute la suite un ensemble infini de symboles appelés *variables*. En utilisant une politique de nommage appropriée (comme dans les bonnes règles de programmation) on s'arrangera à éviter toute confusion avec les noms de constantes, de fonctions ou de prédicats.

Définition 7 (Contexte). *Un contexte est un ensemble de déclarations d'une des deux formes suivantes :*

- Une déclaration de variable $v : A$ où A est un type.
- Une hypothèse $H : A$ où A est une proposition et H un identificateur.

On suppose qu'un contexte ne contient pas deux déclarations différentes de la même variable ou du même nom d'hypothèse.

Remarque

Afin de rester compatible avec les notations usuelles et les chapitres précédents, on pourra continuer à ne pas nommer les hypothèses. En revanche, comme on l'a vu sur machine, le logiciel *Coq* demande que chaque hypothèse soit nommée.

5.2.4 Termes

Comme dans tout langage de programmation typé, l'écriture d'expressions correctes est guidée par des conventions syntaxiques et l'application de *règles*

de typage. Ces règles permettent de construire des *jugements de typage*, de la forme $\Gamma \vdash t : A$ qui se lisent “ dans le contexte Γ , le terme t a pour type A ”.

5.2.5 Règles de typage

Les deux premières règles permettent de construire des termes de base, soit à partir d’une constante, soit à partir d’une variable déclarée dans le contexte considéré. La troisième permet de construire des termes plus complexes, en appliquant des fonctions à des arguments du type approprié.

Constantes

$$\frac{}{\Gamma \vdash c : A} \text{ (} c \text{ est une constante de type } A \text{)}$$

Variables

$$\frac{}{\Gamma \vdash v : A} ((v : A) \in \Gamma)$$

Application d’une fonction

$$\frac{\Gamma \vdash f : A_1 \times A_2 \times \dots \times A_k \rightarrow A \quad \Gamma \vdash t_1 : A_1 \quad \Gamma \vdash t_2 : A_2 \quad \dots \quad \Gamma \vdash t_k : A_k}{\Gamma \vdash f(t_1, t_2, \dots, t_k) : A}$$

Exemple

Reprenons les symboles de fonctions de la section 5.2.2 page 61. Soit un contexte Γ contenant deux variables $i : \text{nat}$ et $j : \text{nat}$. Alors nous obtenons le jugement de typage $\Gamma \vdash *(i, (+ (j, S(S(0))))) : \text{nat}$.

Exercice 36. Vérifier l’exemple précédent règle par règle. Construire d’autres exemples simples.

Remarque

Afin d’alléger les notations, nous conviendrons d’écrire $i+j$ au lieu de $+(i,j)$, $i*j$ au lieu de $*(i,j)$, 1 au lieu de $S(0)$, 2 au lieu de $S(S(0))$, etc.

Ainsi notre exemple précédent peut-il s’écrire $\Gamma \vdash i * (j + 2) : \text{nat}$.

5.2.6 Prédicats et Relations

Les symboles de prédicats et de relations nous permettent de construire des propositions atomiques énonçant des propriétés des termes.

Un *symbole de prédicat* est un identificateur R (différent des symboles de constantes et de fonction), auquel on associe un type d’argument A .

Les symboles de prédicats sont des cas particuliers de *symboles de relation*, auxquels on associe un type de la forme $A_1 \times A_2 \times \dots \times A_k$.

Notation

Nous proposons les notations suivantes (cohérentes avec *Coq*) pour déclarer le type des arguments d'un symbole de prédicat ou relation :

$$R : A \rightarrow \mathbf{Prop}$$

et

$$R : A_1 \times A_2 \times \dots \times A_k \rightarrow \mathbf{Prop}$$

Cette notation revient à attribuer un type spécial aux propositions, ce qui permet de simplifier l'écriture de la phrase “ P est une proposition” en “ $P : \mathbf{Prop}$ ”.

Exemples

Reprenons les types utilisés dans l'exemple 5.2.2 page 61.

Nous pouvons considérer des symboles de prédicat associés aux propriétés sur les entiers naturels : “être nul”, “être positif”, et un symbole de relation binaire “être strictement inférieur à” :

- $\text{nul} : \text{nat} \rightarrow \mathbf{Prop}$
- $\text{positif} : \text{nat} \rightarrow \mathbf{Prop}$
- $< : \text{nat} \times \text{nat} \rightarrow \mathbf{Prop}$

Dans le même contexte que l'exemple 5.2.5 page précédente, nous pouvons former les propositions suivantes :

- $i < i + 1$ (autrement dit $< (i, i + 1)$)
- $\text{nul}(i) \vee \text{positif}(i)$
- $0 < i \leftrightarrow \text{positif}(i)$

L'égalité

On considère un symbole de relation binaire $=$ *polymorphe* : Pour tout type atomique A , et tout couple de termes t et t' *du même type* A , la proposition $t = t'$ est bien formée.

On considérera aussi $t \neq t'$ comme une abréviation de $\sim(t = t')$.

Notons que “ $1 \neq \text{true}$ ” n'est pas une proposition, car l'égalité n'a pas de sens entre deux termes de type différent. Or autoriser l'écriture de “ $1 \neq \text{true}$ ” reviendrait à autoriser celle de “ $\sim(1 = \text{true})$ ” qui contiendrait comme sous-formule “ $1 = \text{true}$ ”.

Définition 8 (Signature). *En logique du premier ordre (calcul des prédicats) on appelle **signature** la donnée d'un ensemble de types, de symboles de constantes, de fonctions, et de relations (avec leur type). On supposera toujours que les symboles utilisés pour désigner les types, constantes, fonctions, relations sont deux à deux disjoints, afin d'éviter toute ambiguïté.*

5.2.7 Syntaxe des propositions

Les constructions de propositions vues en logique propositionnelle ont besoin d'être étendues.

Tout d'abord, les propositions atomiques, au lieu d'être réduites aux variables propositionnelles, sont maintenant formées à partir des termes et des symboles de relations (y compris l'égalité).

Les constructions à partir de connecteurs des chapitres précédents restent valables. Il nous reste à introduire les *quantificateurs* pour obtenir la logique des prédicats.

Définition 9 (Quantifications). *Soit v une variable, A un type, et P une proposition. Les formules $\forall v : A, P$ et $\exists v : A, P$ sont des propositions. Les symboles $\forall$ et $\exists$ sont appelés respectivement *quantificateur universel* et *quantificateur existentiel*.*

On considère que les quantifications ont une priorité plus faible que tous les connecteurs. Par exemple la proposition “ $\forall i : \text{nat}, \text{positif}(i) \vee i = 0$ ” se lit “ $\forall i : \text{nat}, (\text{positif}(i) \vee i = 0)$ ”.

On peut utiliser des parenthèses pour modifier cette priorité comme dans la formule suivante :

$$\forall i : \text{nat}, (\exists j : \text{nat}, j < i) \leftrightarrow 0 < i$$

Remarques

Dans le cas d'une signature à un seul type A , ou si le type des variables est rendu clair par le contexte, on pourra omettre l'indication de type et noter “ $\forall v, P$ ” et “ $\exists v, P$ ” au lieu de “ $\forall v : A, P$ ” et “ $\exists v : A, P$ ”.

De même, on pourra abréger des quantifications imbriquées de même nature et portant sur le même type en écrivant par exemple “ $\forall v w x : A, P$ ” au lieu de “ $\forall v : A, \forall w : A, \forall x : A, P$ ”.

5.2.8 Variables libres et variables liées

Dans une formule de la forme “ $\forall v : A, P$ ” ou “ $\exists v : A, P$ ”, on appelle P la *portée* de la quantification sur x . Par exemple, dans la formule vue précédemment :

$$\forall i : \text{nat}, (\exists j : \text{nat}, j < i) \leftrightarrow 0 < i$$

la portée de la quantification sur i est l'expression “ $(\exists j : \text{nat}, j < i) \leftrightarrow 0 < i$ ” et la portée de la quantification sur j est l'expression “ $j < i$ ”.

On dit qu'une variable est *liée* dans une formule si elle apparaît dans la portée d'un quantificateur de cette formule, *libre* sinon.

Par exemple, soit $i : \text{nat}$, et la formule “ $(\exists j : \text{nat}, j < i) \leftrightarrow 0 < i$ ”. La variable i est libre dans cette proposition alors que j est liée (par le quantificateur existentiel).

La figure 5.1 page suivante visualise par une flèche la liaison entre chaque occurrence de variable liée et *son* quantificateur.

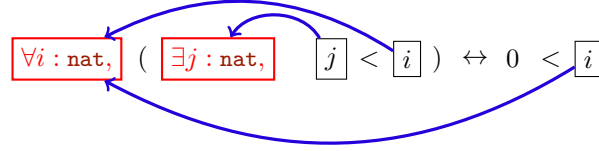


FIGURE 5.1 – géométrie des liaisons(1)

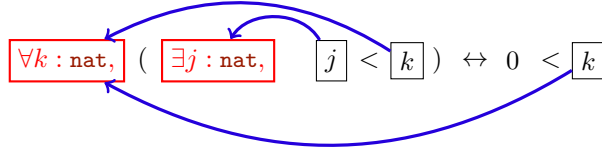


FIGURE 5.2 – géométrie des liaisons : résultat d'une α -conversion

On appelle α -conversion la transformation consistant à renommer toutes les variables liées à un même quantificateur *sans modifier la géométrie des liaisons*.

La figure 5.2 montre le résultat d'une α -conversion sur la formule précédente : la variable liée i a été renommée en k . On voit clairement que la structure des liaisons n'a pas été modifiée par rapport à la figure 5.1.

En revanche si l'on renomme la même variable en j , on voit que la structure de la formule a considérablement changé (figure 5.3). En effet la deuxième occurrence de la variable j a été "capturée" par la quantification existentielle.

On considère que des formules égales à α -conversion près sont identiques. C'est le cas entre les formules des figures 5.1 et 5.2. En revanche, les figures 5.2 et 5.3 sont vraiment différentes.

On évitera d'avoir des variables à la fois libres et liées dans une même formule, par exemple comme dans " $(\exists i : \text{nat}, i < j) \vee i = 2$ ". Dans ce cas, il est conseillé de procéder à une α -conversion, ce qui donnerait par exemple " $(\exists k : \text{nat}, k < j) \vee i = 2$ ". On appelle *polie* une formule dans laquelle aucune variable n'est à la fois libre et liée, et dans laquelle aucune variable liée n'est liée dans deux quantifications différentes.

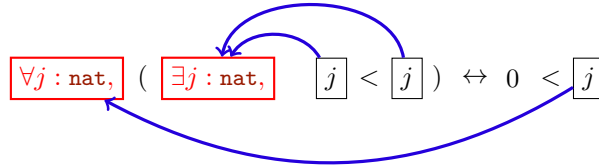


FIGURE 5.3 – géométrie des liaisons : mauvaise α -conversion



FIGURE 5.4 – phénomène de capture : avant

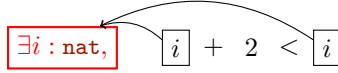


FIGURE 5.5 – phénomène de capture : après

On qualifie de *close* toute formule qui ne contient aucune variable libre. Par exemple " $\forall i : \text{nat}, \exists j : \text{nat}, i < j$ " est close.

En revanche, une formule qui contient une ou plusieurs occurrences d'une variable v exprime une propriété de v . Par exemple soit $n : \text{nat}$. La proposition " $\exists p : \text{nat}, n = p \times 2$ " exprime la propriété " n est pair".

De même pour une formule avec plusieurs variables libres : la formule " $\exists d : \text{nat}, n = d + p$ " exprime la relation " $n \leq p$ ".

5.2.9 Substitution

Soit F une formule, une variable $v : A$ et un terme $t : A$ (dans un contexte donné). On note $F[v \leftarrow t]$ le résultat de la substitution dans F de chaque occurrence libre de v par le terme t , à condition qu'aucune variable de t ne devienne liée dans F (sinon, la notation $F[v \leftarrow t]$ n'est pas définie).

Remarques

Dans la définition précédente, on exige que le terme t ne contienne aucune variable qui pourrait être "capturée" par une quantification.

Par exemple, prenons la proposition $\exists i : \text{nat}, j < i$. Cette proposition "parle" d'un entier naturel, noté j , et énonce qu'il existe un entier plus grand.

Si nous avons dans le contexte considéré, la déclaration $i : \text{nat}$, le remplacement de j par $i + 2$ nous donnerait $\exists i : \text{nat}, i + 2 < i$: dans cette formule, toutes les variables sont liées par le quantificateur existentiel – et on voit bien que cette nouvelle formule a un sens bien différent de la précédente. D'une part, elle ne parle plus d'un entier (non défini, *i.e.* il n'y a plus de variable libre) ; mais surtout, elle énonce qu'il existe un entier plus grand que lui-même augmenté de deux.

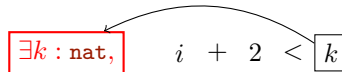


FIGURE 5.6 – substitution correcte de j par $i + 2$ dans la formule de la figure 5.4

Dans cette situation, on commencera par renommer la variable liée i , par exemple en k , ce qui donne $\exists k : \text{nat}, j < k$, puis on opère le remplacement de j par $i + 2$, ce qui donne $\exists k : \text{nat}, i + 2 < k$ (voir Figure 5.6). On obtient bien, conformément à ce qu'on souhaitait, la proposition précédente (“il existe un entier plus grand que”), appliquée à l'entier $i + 2$ (“il existe un entier plus grand que $i + 2$ ”).

Cette précaution préalable (remplacer chaque variable liée de la formule “accueillant” une substitution, par une variable “fraîche”, qui n'apparaît nulle part ailleurs) pourrait être prise systématiquement avant toute tentative de substitution. Ce n'est pas l'usage en pratique : nous avons, depuis quelques siècles que les notations symboliques existent, pris l'habitude d'utiliser des lettres comme noms de variables, et, qui plus est, d'utiliser des noms de variables “parlants” (les entiers s'appellent souvent n, m, k ; les réels, x, y, z , etc). Multiplier les substitutions conduirait rapidement à sortir des alphabets usuels, et à utiliser soit des noms plus longs, soit des indices ; or, des formules $\exists bbc : \text{nat}, bcc + 2 < bbc$ ou $\exists v_{42} : \text{nat}, v_{53} + 2 < v_{42}$ s'avèrent plus difficiles à déchiffrer pour les humains.

Résultat, il reste indispensable de faire attention avant d'opérer toute substitution d'une variable par un terme dans une formule, sous peine de faire des bêtises.

5.2.10 Substitution simultanée

On étend la notion de substitution au cas de plusieurs variables libres distinctes 2 à 2 : Notation $F[v_1 \leftarrow t_1; v_2 \leftarrow t_2; \dots; v_n \leftarrow t_n]$

Exercice 37. Donner une définition formelle de cette opération.

5.2.11 Le coin du formaliste

On peut trouver une définition précise de la substitution, par exemple dans <https://www.lri.fr/~paulin/Logique/html/cours005.html>

1. On commence par définir la substitution de v par t dans un terme t' par récurrence sur ce terme.
 - Si c est une constante, alors $c[v \leftarrow t] = c$
 - Si w est une variable distincte de v , alors $w[v \leftarrow t] = w$
 - $v[v \leftarrow t] = t$
2. Si la formule P est atomique, la substitution se définit sans problème :

$$R(t_1, \dots, t_n)[v \leftarrow t] = R(t_1[v \leftarrow t], \dots, t_n[v \leftarrow t])$$
3. Les connecteurs propositionnels ne posent aucun problème :
 - $\perp[v \leftarrow t] = \perp$
 - $(\sim P)[v \leftarrow t] = \sim(P[v \leftarrow t])$
 - $(P \vee Q)[v \leftarrow t] = P[v \leftarrow t] \vee Q[v \leftarrow t]$
 - etc.
4. Les règles de substitution pour les quantificateurs doivent tenir compte des problèmes de capture de variables :

Soit P une formule et w une variable :

- Si $w = v$, alors $(\forall w : A, P)[v \leftarrow t] = (\forall w : A, P)$ et $(\exists w : A, P)[v \leftarrow t] = (\exists w : A, P)$
- Si $w \neq v$ et w n'a pas d'occurrence dans t , alors

$$(\forall w : B, P)[v \leftarrow t] = \forall w : B, P[v \leftarrow t]$$

et

$$(\exists w : B, P)[v \leftarrow t] = \exists w : B, P[v \leftarrow t]$$

5.3 Preuves en logique du premier ordre

On a tous les outils pour définir ce qu'est une preuve dans le calcul des prédicats. Il suffit d'ajouter aux règles de la logique propositionnelle des règles pour l'égalité et les deux quantificateurs. Suivant qu'on admet ou non la règle du tiers-exclu, on obtient une logique du premier ordre classique ou intuitionniste.

5.3.1 Règles associées à l'égalité

Règle d'introduction de l'égalité

Soit t un terme de type A . La règle suivante applique la réflexivité de l'égalité au terme t .

$$\frac{}{\Gamma \vdash t = t} =_i \quad (\Gamma \vdash t : A)$$

Par exemple nous pouvons montrer le théorème suivant :

$$\frac{}{\vdash 36 = 36} =_i$$

On admet une légère extension de $=_i$, utile dans le raisonnement sur des calculs (par exemple les expressions d'un langage de programmation) :

$$\frac{}{\Gamma \vdash t = t'} =_i \quad (\Gamma \vdash t : A) \quad t \text{ et } t' \text{ ont la même valeur}$$

Ce qui nous permet d'obtenir le résultat suivant :

$$\frac{}{\vdash 6 \times 6 = 9 \times 4} =_i$$

Règle d'élimination de l'égalité

$$\frac{\Gamma \vdash t = t' \quad \Gamma \vdash P[x \leftarrow t]}{\Gamma \vdash P[x \leftarrow t']} =_e$$

Commentaires La variable x (libre dans P) dans cette règle sert uniquement à *marquer* les emplacements où l'on remplace le terme t par t' . Bien sûr le nom x utilisé dans la présentation de la règle peut être remplacé par n'importe quel nom non-utilisé par ailleurs. Cette variable est souvent laissée implicite dans les preuves.

Exemple

Voici une preuve de symétrie de l'égalité (qui devient donc une règle dérivée) : (a et b sont des termes d'un type A).

$$\frac{\Gamma \vdash a = b \quad \overline{\Gamma \vdash a = a}}{\Gamma \vdash b = a} \begin{matrix} =_i \\ =_e \end{matrix}$$

Si l'on veut détailler l'application de la règle $=_e$, il suffit d'écrire $a = a$ comme $(x = a)[x \leftarrow a]$. Du coup, $(x = a)[x \leftarrow b]$ est bien la proposition $b = a$.

Exercice 38. *En utilisant les règles associées à l'égalité, prouver la règle de transitivité de l'égalité.*

$$\frac{\Gamma \vdash a = b \quad \Gamma \vdash b = c}{\Gamma \vdash a = c}$$

5.3.2 Règles associées au quantificateur universel

Règle d'élimination

Soit A un type ; dans un contexte où l'on peut prouver la proposition $\forall x : A, P$ on peut en déduire une *instance* en instanciant la variable liée v par un terme de type A .

$$\frac{\Gamma \vdash t : A \quad \Gamma \vdash \forall x : A, P}{\Gamma \vdash P[x \leftarrow t]} \forall_e, t$$

Exemple

$$\frac{\overline{\Gamma \vdash 23 : \text{nat}} \quad \Gamma \vdash \forall i : \text{nat}, i < i + 1}{\Gamma \vdash 23 < 23 + 1} \forall_e, 23$$

Ajouter l'exemple sur Socrate

Donner une version complète, puis l'élimination du quantificateur et de l'implication en un seul raccourci

Attention !

La règle d'élimination du quantificateur universel utilise l'opération de substitution. Ne pas respecter les précautions destinées à empêcher la capture de variables peut conduire à des déductions absurdes :

$$\frac{\Gamma \vdash j : \text{nat} \quad \Gamma \vdash \forall i : \text{nat}, \exists j : \text{nat}, i \neq j}{\Gamma \vdash \exists j : \text{nat}, j \neq j} \forall_e, j$$

En revanche, si l'on renomme la variable liée j , tout se passe bien :

$$\frac{\Gamma \vdash j : \text{nat} \quad \frac{\Gamma \vdash \forall i : \text{nat}, \exists j : \text{nat}, i \neq j}{\Gamma \vdash \forall i : \text{nat}, \exists k : \text{nat}, i \neq k} \alpha\text{-conversion}}{\Gamma \vdash \exists k : \text{nat}, j \neq k} \forall_e, j$$

Règle d'introduction

Soit x une variable non-libre dans le contexte Γ (*i.e.*, qui n'apparaît comme variable libre dans **aucune** des hypothèses de Γ).

$$\frac{\Gamma, x : A \vdash P}{\Gamma \vdash \forall x : A, P} \forall_i$$

Exemple

$$\frac{\overline{\Gamma, x : A \vdash x = x} \quad =_i}{\Gamma \vdash \forall x : A, x = x} \forall_i$$

Attention !

Ne pas respecter la condition ci-dessus “ x non-libre dans le contexte Γ ” conduit encore à des raisonnements absurdes :

$$\frac{\frac{\overline{i : \text{nat}, i = 2 \vdash i = 2}}{i : \text{nat}, i = 2 \vdash \forall i : \text{nat}, i = 2} \forall_i}{i : \text{nat}, i = 2 \vdash 23 = 2} \forall_e, 23$$

Dans ce cas, on choisit une variable “*fraîche*” pour l'introduction du quantificateur universel :

$$\frac{\overline{i : \text{nat}, i = 2 \vdash i = 2}}{i : \text{nat}, i = 2 \vdash \forall k : \text{nat}, i = 2} \forall_i$$

5.3.3 Notation de bloc pour l'introduction du quantificateur universel

Comme pour les hypothèses, les variables gérées dans la règle $\forall_i$ ont un comportement très similaire aux variables locales des langages tels que **C**. On créera donc un bloc pour chaque introduction de $\forall$ avec une mention “**Soit** $x : A$ ”. La variable x est bien sûr *locale* à ce bloc. Une proposition P prouvée à l'intérieur du bloc est *exportée* sous la forme $\forall x : A, P$.

```

1 { Soit  $x : A$ 
2    $x = x$   $[=_i]$ 
3 }
4  $\forall x : A, x = x$   $[\forall_i, 2, x]$ 

```

5.3.4 Quantifications imbriquées

Comme pour l'implication, nous généralisons l'introduction et l'élimination du quantificateur à un nombre quelconque de variables.

Introduction du quantificateur universel (variante)

Soient $x_1, \dots, x_n$ n variables non-libres dans le contexte Γ .

$$\frac{\Gamma, x_1 : A_1, \dots, x_n : A_n \vdash P}{\Gamma \vdash \forall x_1 : A_1, \dots, x_n : A_n, P} \forall_i$$

Élimination du quantificateur universel (variante)

$$\frac{\Gamma \vdash t_1 : A_1 \quad \dots \quad \Gamma \vdash t_n : A_n \quad \Gamma \vdash \forall x_1 : A_1, \dots, x_n : A_n, P}{\Gamma \vdash P[x_1 \leftarrow t_1; \dots; x_n \leftarrow t_n]} \forall_e, t_1, \dots, t_n$$

Exemples

$$\frac{\frac{\frac{x : A, y : A, x = y \vdash x = y}{x : A, y : A, x = y \vdash y = x} \rightarrow_i}{x : A, y : A \vdash x = y \rightarrow y = x} \forall_i, x, y}{\vdash \forall x y : A, x = y \rightarrow y = x} =_e$$

```

1 { Soient  $x, y : A$ 
2   { Supposons  $x = y$ 
3      $y = y$   $[=_i]$ 
4      $y = x$   $[=_e, 3, 2]$ 
5   }

```

```

6       $x = y \rightarrow y = x$  [ $\rightarrow_i$ , 2, 4]
7    }
8     $\forall x y : A, x = y \rightarrow y = x$  [ $\forall_i$ , 6, x,y]

```

Voici un exemple d'utilisation de $\forall_e$ sur deux quantifications imbriquées :

```

1  { Supposons  $\forall x y : A, P(x, y)$ 
2    { Soit  $x : A$ 
3       $P(x, x)$  [ $\forall_e$ , 2, x, x]
4    }
5     $\forall x : A, P(x, x)$  [ $\forall_i$ , 3, x]
6  }
7   $(\forall x y : A, P(x, y)) \rightarrow (\forall x A, P(x, x))$  [ $\rightarrow_i$ , 1, 4]

```

5.3.5 Règles associées au quantificateur existentiel

Règle d'introduction

La règle ci-dessous exprime que pour prouver une formule $\exists x : A, P$, il suffit de fournir un terme $t : A$ tel que l'on puisse prouver la proposition $P[x \leftarrow t]$.

Le terme t est appelé le *témoin* de l'application de cette règle.

$$\frac{\Gamma \vdash t : A \quad \Gamma \vdash P[x \leftarrow t]}{\Gamma \vdash \exists x : A, P} \exists_i, t$$

Exemple

```

1  Supposons  $\forall x, \text{humain}(x) \rightarrow \text{mortel}(x)$ 
2  Supposons  $\text{humain}(\text{Socrate})$ 
3  ...
4   $\text{mortel}(\text{Socrate})$ 
5   $\exists h : \text{individu}, \text{mortel}(h)$  [ $\exists_i$ , 4, Socrate]

```

Exemple

```

1  { Soit  $x : A$ 
2     $x = x$  [ $=_i$ ]
3     $\exists y : A, x = y$  [ $\exists_i$ , 2, x]
4  }
5   $\forall x : A, \exists y : A, x = y$  [ $\forall_i$ , 3]

```

Attention !

La règle d'introduction du quantificateur utilise l'opération de substitution. Il est donc nécessaire de vérifier que les précautions d'usage de cette opération (voir Section 5.2.9 page 67) sont bien respectées.

Voici un exemple de mauvaise application de cette règle.

```
1  Supposons  $\forall y : \text{nat}, y \leq 2y$ 
2   $\exists x : \text{nat}, \forall y : \text{nat}, y \leq x$   $[\exists_i, 1, 2y]$ 
```

Règle d'élimination

On considère une variable x , non libre dans $\Gamma \cup \{Q\}$.

$$\frac{\Gamma, x : A, P \vdash Q \quad \Gamma \vdash \exists x : A, P}{\Gamma \vdash Q}$$

Ce schéma peut s'écrire facilement en considérant un bloc contenant la variable x et l'hypothèse P :

```
1  ...
2   $\exists x : A, P$ 
3  ...
4  { Soit  $x : A$  tel que  $P$ 
5      ...
6       $Q$ 
7  }
8   $Q$   $[\exists_e, 2, 6]$ 
```

Cela paraît-il compliqué ? Intuitivement, le schéma d'élimination du quantificateur existentiel exprime le raisonnement suivant :

1. On considère qu'il existe un x vérifiant la propriété P
2. Prenons donc un élément quelconque x vérifiant P ¹.
(a) Dans ce contexte, nous arrivons à prouver Q
3. Puisque x était quelconque, Q est donc prouvé.

Exemple

```
1  Supposons  $\exists x : A, \forall y : A, x = y$ 
2  { Soit  $x : A$  tel que  $\forall y : A, x = y$ 
3      { soient  $a, b : A$ 
4           $x = a$   $[\forall_e, 2, a]$ 
5           $x = b$   $[\forall_e, 2, b]$ 
6           $a = b$   $[=, 5, 4]$ 
7      }
```

1. Cette quelconquitude est exprimée par la condition " x , non libre dans $\Gamma \cup \{Q\}$ "

```

8       $\forall a b : A, a = b$  [ $\forall_i, 6, a, b$ ]
9    }
10    $\forall a b : A, a = b$  [ $\exists_e, 1, 8$ ]

```

Attention !

Qu'arrive-t-il si on ne respecte pas la condition x , non libre dans $\Gamma \cup \{Q\}$?
Voici deux exemples de faux raisonnements :

```

1  Supposons  $x = 2$ 
2  Prouvons  $\exists x : \text{nat}, x + x = x$  {
3     $0 = 0 + 0$  [ $=_i$ ]
4     $\exists x : \text{nat}, x + x = x$  [ $\exists_i, 3, 0$ ]
5  }
6  { Soit  $x : \text{nat}$  tel que  $x + x = x$ 
7     $2 + 2 = 2$  [ $=_e, 1, 6$ ]
8     $4 = 2$  [calcul sur 7]
9  }
10  $4 = 2$  [ $\exists_e, 2, 7$ ]

```

```

1   $\exists x : \text{nat}, x = 0$  [exercice]
2  { soit  $x : \text{nat}$  tel que  $x = 0$ 
3     $x + x = x$  [exercice]
4  }
5   $x + x = x$  [ $\exists_e, 2, 7$ ]

```

5.4 Pratique de la logique du premier ordre

Nous pouvons voir quelques exemples de preuve en logique du premier ordre, utilisant les règles de la logique propositionnelle et celles associées à l'égalité et aux deux quantificateurs. Suivant que l'on admet ou pas la règle du tiers exclu, on aura des théorèmes "classiques" ou "intuitionnistes". Nous étendons au calcul des prédicats l'usage des symboles $\vdash_K$ et $\vdash_J$.

5.4.1 Quantificateurs et négation

Soient A un type quelconque et P un prédicat sur A . Voici une preuve *en logique intuitionniste* du séquent $\sim(\exists x : A, P(x)) \vdash \forall x : A, \sim P(x)$

```

1  Supposons  $\sim(\exists x : A, P(x))$ 
2  { Soit  $x : A$ 
3    { Supposons  $P(x)$ 
4       $\exists x : A, P(x)$  [ $\exists_i, 4, x$ ]
5       $\perp$  [mp, 1, 4]

```

```

6      }
7       $\sim P(x)$  [ $\rightarrow_i$ , 3,5]
8    }
9  }
10  $\forall x : A, \sim P(x)$  [ $\forall_i$ , 2, 7]

```

La preuve suivante est une preuve *en logique classique* du séquent $\sim(\forall x : A, \sim P(x)) \vdash \exists x : A, P(x)$

```

1  Supposons  $\sim(\forall x : A, \sim P(x))$ 
2  { Supposons  $\sim(\exists x : A, P(x))$ 
3    { Soit  $x : A$ 
4      { Supposons  $P(x)$ 
5         $\exists x : A, P(x)$  [ $\exists_i$ ,5,x]
6         $\perp$  [ $\rightarrow_e$ ,2,5]
7      }
8       $\sim P(x)$  [ $\rightarrow_i$ , 4, 6]
9    }
10    $\forall x : A, \sim P(x)$ , [ $\forall_i$ ,3, 8]
11    $\perp$ , [ $\rightarrow_e$ ,1,10]
12 }
13  $\exists x : A, P(x)$  [absurde classique]

```

On aurait pu aussi utiliser le tiers-exclu pour la même preuve :

```

1  Supposons  $\sim(\forall x : A, \sim P(x))$ 
2   $(\exists x : A, P(x)) \vee \sim(\exists x : A, P(x))$  [tiers exclu]
3  { Supposons  $\exists x : A, P(x)$ 
4     $\exists x : A, P(x)$  [hypothèse]
5  }
6  { Supposons  $\sim(\exists x : A, P(x))$ 
7    même dérivation que ci-dessus
8     $\exists x : A, P(x)$ 
9  }
10  $\exists x : A, P(x)$  [ $\vee_e$ ,2,3-5,6-8]

```

5.4.2 Exemples

Exercice 39. Prouver les jugements suivants :

```

 $\vdash_J \forall x : A, P(x) \rightarrow \exists y : A, P(y)$ 
 $\sim(\forall x : A, P(x)) \equiv_K \exists x : A, \sim P(x)$ 
 $\sim(\exists x : A, P(x)) \equiv_J \forall x : A, \sim P(x)$ 
 $(\exists x : A, P(x)) \rightarrow Q \equiv_J \forall x : A, P(x) \rightarrow Q$ 
 $\exists x : A, \forall y : B, Q(x,y) \vdash_J \forall y : B, \exists x : A, Q(x,y)$ 

```

5.4.3 Règles dérivées

Les règles dérivées applicables en logique propositionnelle intuitionniste [resp. classique] restent applicables en logique du premier ordre intuitionniste [resp. classique].

Il est toutefois intéressant de modifier le théorème de substitution uniforme afin de réutiliser des jugements similaires à ceux de l'exercice 39 page précédente.

Prenons par exemple le jugement suivant :

$$\exists x : A, \forall y : B, Q(x, y) \vdash_J \forall y : B, \exists x : A, Q(x, y)$$

Dans ce jugement, Q est un symbole de relation sur le type $A \times B$.

On peut obtenir une instance de ce jugement en substituant à A un type X et à B un type Y , puis à Q une relation sur $X \times Y$.

Par exemple, substituons nat à A et B , et associons à Q le prédicat défini par $Q(n, p) = n < p$. L'opération de substitution consiste alors à remplacer toute sous-formule de la forme $Q(t_1, t_2)$ par $(n \leq p)[n \leftarrow t_1; p \leftarrow t_2]$ soit $t_1 \leq t_2$.

Nous obtenons donc le jugement :

$$\exists x : \text{nat}, \forall y : \text{nat}, x \leq y \vdash_J \forall y : \text{nat}, \exists x, x \leq y$$

Nous admettons que, en logique des prédicats, la substitution uniforme peut s'appliquer à des symboles de prédicats et de relations, à condition que l'on remplace ces symboles en respectant les contraintes de type.

Ce résultat permet d'utiliser les jugements de l'exercice 39 comme des règles dérivées.

Notons que l'utilisation de l'opération de substitution doit être correctement appliquée, notamment en ce qui concerne les captures de variables.

Par exemple, considérons une variable $y : \text{nat}$, et le jugement $\vdash_J \forall x : \text{nat}, P(x) \rightarrow \exists y : \text{nat}, P(y)$. Si nous associons au symbole P le prédicat défini par $P(\alpha) =_{\text{def}} y < \alpha$, une substitution appliquée sans précaution donnerait le jugement $\vdash_J \forall x : \text{nat}, y < x \rightarrow \exists y : \text{nat}, y < y$ au lieu du (correct) $\vdash_J \forall x : \text{nat}, y < x \rightarrow \exists z : \text{nat}, y < z$.

Exercice 40. Démontrer en logique classique le paradoxe du buveur :

“Soit un café avec au moins une personne (par exemple le patron).
Alors, il existe une personne dans ce café qui, si elle boit, alors tout le monde boit”

On considère un type resents pour les personnes présentes dans le café, et une constante $\text{patron} : \text{resents}$

Il s'agit alors de montrer le jugement $\vdash_K \exists c : \text{resents}, \text{boit}(c) \rightarrow \forall x : \text{resents}, \text{boit}(x)$.

Aide : On pourra utiliser le tiers-exclu avec la proposition $\forall x : \text{resents}, \text{boit}(x)$.

Exercice 41. Soit A un type. On peut exprimer que A est vide par la proposition $\text{vide}(A) =_{\text{def}} \forall x : A, x \neq x$.

Montrer les jugements :

$$\text{vide}(A) \vdash_J \forall x : A, P(x) \quad (5.1)$$

$$\sim \text{vide}(A), \forall x : A, P(x) \vdash_K \exists x : A, P(x) \quad (5.2)$$

$$\forall x y : A, x \neq y \vdash_J \text{vide}(A) \quad (5.3)$$

5.5 Sémantique

5.5.1 Structures d'interprétation

Définition 10. Soit σ une signature pour la logique du premier ordre. Une **interprétation** pour σ est une structure I constituée des champs suivants :

- Pour chaque type A , un domaine I_A ²
- Pour chaque constante $c : A$, un élément $I_c \in I_A$,
- Pour chaque symbole de fonction $f : A_1 \times A_2 \times \cdots \times A_k \rightarrow A$, une fonction totale I_f de $I_{A_1} \times I_{A_2} \times \cdots \times I_{A_k}$ dans I_A ,
- Pour chaque symbole de relation R sur $A_1 \times A_2 \times \cdots \times A_k$, un sous-ensemble I_R du produit cartésien $I_{A_1} \times I_{A_2} \times \cdots \times I_{A_k}$.

Exemple

Prenons la signature vue en 5.2 page 60

Nous pouvons lui associer une interprétation I de la façon suivante :

- $I_{\text{nat}} =_{\text{def}} \mathbb{N}$
- $I_0 =_{\text{def}} 0$
- $I_S =_{\text{def}} \{(n, n+1) | n \in \mathbb{N}\}$
- $I_+ =_{\text{def}} +$
- $I_* =_{\text{def}} \times$
- $I_{\text{nul}} =_{\text{def}} \{0\}$
- $I_{\text{positif}} =_{\text{def}} \mathbb{N} \setminus \{0\}$
- $I_{<} =_{\text{def}} \{(n, p) | n, p \in \mathbb{N} \wedge n < p\}$

5.5.2 Valuations

On suppose dans ce paragraphe que pour chaque type A de la signature considérée, nous disposons d'une infinité de variables de ce type.

Définition 11. Une **valuation** est une application ν qui à toute variable de type A associe une valeur $\nu(v) \in I_A$.

Soit V un ensemble de variables ; on dit que deux valuations ν et ν' sont **congruentes sur V** (noté $\nu \equiv_V \nu'$) si pour toute variable v dans V , $\nu(v) = \nu'(v)$.

Soit v une variable de type A et $a \in I_A$. On note $\nu\{v := a\}$ la valuation ν'' définie par

- $\nu'(v) = a$

2. En théorie des modèles, on impose $I_A \neq \emptyset$. Nous n'imposerons pas cette hypothèse dans ce cours.

— $\nu'(w) = \nu(w)$ pour toute variable w différente de v .

5.5.3 Sémantique des formules

Soit ν une valuation.

Pour tout terme t , notons $[t]_\nu$ la valeur de t dans la valuation ν .

Exercice : En donner la définition.

La *valeur de vérité* d'une formule φ dans ν , notée $[\varphi]_\nu$, se définit par récurrence sur la structure de φ .

- $[R(t_1, \dots, t_n)]_\nu = \mathbf{v}$ si $([t_1]_\nu, \dots, [t_n]_\nu) \in I_R$
- $[t_1 = t_2]_\nu = \mathbf{v}$ si $[t_1]_\nu = [t_2]_\nu$
- $[\varphi \rightarrow \varphi']_\nu =$
 - $\mathbf{v}$ si $[\varphi]_\nu = \mathbf{f}$ ou $[\varphi']_\nu = \mathbf{v}$
 - $\mathbf{f}$ si $[\varphi]_\nu = \mathbf{v}$ ou $[\varphi']_\nu = \mathbf{f}$
- ... **exercice :** compléter la définition pour chaque connecteur
- $[\forall v : A, \varphi]_\nu = \mathbf{v}$ si pour toute valeur $a \in I_A$, on a $[\varphi]_{\nu\{v:=a\}} = \mathbf{v}$
- $[\exists v : A, \varphi]_\nu = \mathbf{v}$ s'il existe une valeur $a \in I_A$, telle que $[\varphi]_{\nu\{v:=a\}} = \mathbf{v}$

Définition 12. On dit qu'un séquent est valide (notation $\Gamma \models A$) si pour toute valuation ν qui satisfait toutes les hypothèses de Γ , alors $[A]_\nu = \mathbf{v}$.

5.6 Méta-théorèmes principaux

Nous donnons les deux résultats suivants sans démonstration. Le lecteur est invité à consulter la nombreuse bibliographie sur ce sujet.

Méta-théorème 20 (Cohérence). Si $\Gamma \vdash_J A$, alors $\Gamma \models A$.

A fortiori, si $\Gamma \vdash_K A$, alors $\Gamma \models A$.

Méta-théorème 21 (Complétude de la logique classique). Si $\Gamma \models A$, alors $\Gamma \vdash_K A$.

Exercice 42. Montrer que le séquent $\forall x : A, P(x) \vdash \exists x : A, P(x)$ n'est pas prouvable en logique classique, ni en logique intuitionniste. **Aide :** Revoir l'exercice 41 page 77.

Exercice 43. On considère le contexte Γ formé des axiomes suivants :

$$\begin{aligned} &\forall x : A, \sim x < x \\ &\forall xyz : A, x < y \rightarrow y < z \rightarrow x < z \\ &\forall xy : A, x \leq y \leftrightarrow x < y \vee x = y \end{aligned}$$

Montrer que de Γ on peut déduire les propositions suivantes :

$$\begin{aligned} &\forall x : A, x \leq x \\ &\forall xy : A, x \leq y \rightarrow y \leq x \rightarrow x = y \\ &\forall xyz : A, x \leq y \rightarrow y \leq z \rightarrow x \leq z \end{aligned}$$

Deuxième partie

Spécifier et prouver des programmes fonctionnels

Dans cette partie, nous montrons quelques techniques de base pour spécifier et prouver des programmes fonctionnels très simples. Pour des preuves de programmes plus réalistes, voir par exemple [3].

Nous nous limitons à quelques types de données inductifs simples : booléens, entiers naturels, listes et arbres (monomorphes). Pour chaque type, nous énumérerons ses constructeurs, des axiomes ou schémas d'axiomes spécifiques à ce type, et présenterons des techniques de preuves de fonctions opérant sur ce type.

Chapitre 6

Quelques types inductifs

6.1 Les valeurs booléennes

Le type `bool`, présent en *OCaML* et en *Coq* ne contient que deux valeurs appelées *constructeurs* : `true` et `false`.

6.1.1 Règles spécifiques au type `bool`

Une première règle stipule que ces deux constructeurs sont différents.

$$\frac{}{\text{true} \neq \text{false}} \text{ bool-diff}$$

Comment exprimer que `bool` ne contient *que* `true` et `false` ? La solution la plus commode est d'énoncer le *schéma de règle* suivant :

Soit P un prédicat défini sur `bool`.

$$\frac{P(\text{true}) \quad P(\text{false})}{\forall b : \text{bool}, P(b)} \text{ bool-cases}$$

En combinant cette règle avec une élimination du quantificateur universel, nous obtenons la variante suivante.

$$\frac{P(\text{true}) \quad P(\text{false})}{b : \text{bool} \vdash P(b)} \text{ bool-cases}$$

Du coup, nous pouvons prouver que tout booléen est, soit `true`, soit `false`.

$$\frac{\frac{\frac{}{\text{true} = \text{true}}{=i}}{\text{true} = \text{true} \vee \text{true} = \text{false}} \vee_{i,1} \quad \frac{\frac{\frac{}{\text{false} = \text{false}}{=i}}{\text{false} = \text{true} \vee \text{false} = \text{false}} \vee_{i,2}}{\frac{b : \text{bool} \vdash b = \text{true} \vee b = \text{false}}{\forall b : \text{bool}, b = \text{true} \vee b = \text{false}}} \text{ bool-cases} \vee_i$$

6.1.2 Définitions de fonctions par cas

L'utilisation du type `bool` dans des programmes se fait simplement, soit en plaçant un des deux constructeurs, `true` ou `false` dans une expression du langage, soit en utilisant la construction `if b then e1 else e2`, où b est une expression de type `bool`, et e_1 et e_2 deux expressions *d'un même type* A . Alors, l'expression conditionnelle a aussi le type A .

Par exemple, en *OCaML*, voici les définitions de la négation, la conjonction et la disjonction booléennes.

```
let notb b = if b then false else true
```

```
let andb b b' = if b then b' else false
```

```
let orb b b' = if b then true else b'
```

Pour raisonner sur ces programmes, on utilise les règles de calcul suivante

Toute expression de la forme `if true then e1 else e2` se simplifie en e_1

Toute expression de la forme `if false then e1 else e2` se simplifie en e_2

Par exemple, l'expression `notb false`, équivalente à `if false then false else true`, se simplifie en `true`.

L'expression `notb (not b false)`, équivalente à `if (notb false) then false else true` va se simplifier en `if true then false else true`, puis en `false`.

Exercice 44. Démontrer le théorème $\forall b : \text{bool}, \text{notb}(\text{notb } b) = b$.

Exercice 45. Écrire les règles de simplifications associées aux fonctions `orb` et `andb`.

Montrer l'égalité $\forall b : \text{bool}, \text{orb } b (\text{notb } b) = \text{true}$.

Expressions conditionnelles

Syntaxe, règles de simplification.

6.1.3 Exemples

négation, conjonction et disjonction booléennes.

6.2 Entiers naturels

Même plan avec la récursion en plus.

Troisième partie

Annexes

Chapitre 7

TO DO

Les trucs à rédiger et placer correctement.

1. Axiomes vs hypothèses

Bibliographie

- [1] apprendre-en-ligne.net. Site Web sur l’informatique et les maths au lycée.
- [2] J.P. Belna. *Histoire de la logique*. L’esprit des sciences. Ellipses, 2005.
- [3] Yves Bertot and Pierre Castéran. *Le Coq’Art*. Version française : <http://www.labri.fr/~casteran/CoqArt/coqartF.pdf>.
- [4] www.cgsecurity.org/wiki/Articles.
- [5] G. Dowek. *La logique*. Flammarion, coll. “Dominos”, 1995.
- [6] Christine Paulin-Mohring. Mathématiques pour l’informatique 2. www.lri.fr/~paulin/MathInfo2/.
- [7] Arthur Schopenhauer. *l’Art d’avoir toujours raison*. Éditions Mille et une Nuits. Voir aussi [https://fr.wikisource.org/wiki/L’Art d’avoir toujours raison](https://fr.wikisource.org/wiki/L’Art_d’avoir_toujours_raison).
- [8] Coq Development Team. The Coq Proof Assistant. <https://coq.inria.fr/>.
- [9] Wikipedia. Crise des fondements. https://fr.wikipedia.org/wiki/Crise_des_fondements.