

et, plus généralement, les 2^k théorèmes suivants, pour tout $k \in \{0, 1, \dots, n\}$:

$$\{p'_1, \dots, p'_k\} \vdash A,$$

et donc, en particulier ($k = 0$) le théorème

$$\vdash A,$$

ce qui achève la démonstration.

5 Logique prédicative : syntaxe et sémantique

5.1 Introduction

Nous avons vu d'emblée que la principale limitation de la logique propositionnelle est l'impossibilité de modéliser adéquatement des propositions “paramétriques”, dont la valeur de vérité dépend de la signification de termes contenus dans la proposition. Reconsidérons les exemples suivants :

- Les entiers naturels x et $x + 2$ sont premiers.
- Il existe un naturel x tel que x et $x + 2$ sont premiers.
- Il existe une infinité de nombres premiers x tels que $x + 2$ est aussi premier.
- Il fera beau à tel endroit, à tel instant.
- Il fera beau à Liège le 29 avril de l'an 2021.
- $x^2 + y^2 = z^2$.
- $x^3 + y^3 = z^3$.
- Il existe des entiers strictement positifs x, y, z tels que $x^2 + y^2 = z^2$.
- Il existe des entiers strictement positifs x, y, z tels que $x^3 + y^3 = z^3$.

On observe d'abord que, faute de pouvoir analyser des propositions paramétriques telles que “Les entiers naturels x et $x + 2$ sont premiers” et “ $x^2 + y^2 = z^2$ ”, on ne peut pas non plus analyser des propositions non paramétriques telles que “Il existe un naturel x tel que x et $x + 2$ sont premiers” et “Il existe des entiers strictement positifs x, y, z tels que $x^2 + y^2 = z^2$ ”. En effet, il est fréquent que des propositions non paramétriques admettent comme composants (dans un sens à préciser) des propositions paramétriques. Dans la mesure où l'approche compositionnelle semble incontournable, on voit que la logique propositionnelle ne pourra pas à elle seule rendre compte des mécanismes de raisonnement des mathématiciens.

En fait, même les raisonnements courants impliquent des propositions paramétriques comme le montre l'exemple classique suivant :

- Tous les hommes sont mortels.
- Or, Socrate est un homme.
- Donc, Socrate est mortel.

On voit que la troisième proposition est conséquence logique des deux premières, mais une analyse purement propositionnelle ne rendra pas compte de ce fait. Nous avons donc besoin d'un langage formel plus riche que le calcul des propositions, qui permettra d'exprimer des propriétés vraies pour certains individus pris dans un ensemble, c'est-à-dire des relations.

En mathématique, on définit une *relation* $\mathcal{R}$ d'arité n sur les ensembles $D_1, D_2, \dots, D_n$ comme un sous-ensemble du produit cartésien $D_1 \times D_2 \times \dots \times D_n$. Voici à titre d'exemple la description de quelques relations importantes de l'arithmétique :

$$\begin{aligned} PPQ(x, y) &= \{(x, y) \in (\mathbb{N} \times \mathbb{N}) : x < y\} \\ &= \{(0, 1), (0, 2), (0, 3), \dots, (1, 2), (1, 3), \dots, (2, 3), \dots\} \end{aligned}$$

$$CARRE(x, y) = \{(x, y) \in (\mathbb{N} \times \mathbb{N}) : y = x^2\} = \{(0, 0), (1, 1), (2, 4), (3, 9), \dots\}$$

$$PR(x) = \{x \in \mathbb{N} : x \text{ est un nombre premier}\} = \{2, 3, 5, 7, 11, 13, \dots\}$$

Définitions. Soit D un ensemble. $\mathcal{R}$ est une *relation* d'arité n sur le domaine D si $\mathcal{R}$ est une relation sur D^n . Le *prédicat* R associé à $\mathcal{R}$ est défini par

$$R(d_1, \dots, d_n) = \mathbf{V} \text{ si et seulement si } (d_1, \dots, d_n) \in \mathcal{R}.$$

On aura donc

$$PPQ(0, 1) = \mathbf{V}, PPQ(8, 4) = \mathbf{F}, PPQ(3, 6) = \mathbf{V}, \dots$$

$$CARRE(0, 0) = \mathbf{V}, CARRE(0, 2) = \mathbf{F}, CARRE(2, 4) = \mathbf{V} \dots$$

$$PR(3) = \mathbf{V}, PR(8) = \mathbf{F} \dots$$

On voit qu'un prédicat est une proposition paramétrique, vraie pour certains éléments d'un domaine et fausse pour les autres.

Il faut souligner d'emblée que le principal apport du calcul des prédicats ne sera pas l'étude des formules paramétriques pour elles-mêmes, mais plutôt l'étude de formules non paramétriques dont certaines composantes sont paramétriques. Pour prendre un exemple célèbre, la question de Fermat n'est pas de savoir si la formule

$$x^n + y^n = z^n \wedge x, y, z \neq 0 \wedge n > 2 \quad (1)$$

est vraie pour des entiers x, y, z, n donnés, mais bien de savoir si, oui ou non, il existe un quadruplet d'entiers tel que la formule soit vraie. La première question, du ressort du calcul élémentaire, est clairement paramétrique ; le fait que $2^3 + 3^3 = 35 \neq 64 = 4^3$ établit clairement que la formule est fausse pour le quadruplet $(2, 3, 4, 3)$ mais ne détermine pas la valeur de vérité pour le quadruplet $(12, 13, 14, 15)$ par exemple. En revanche, le fait que la formule

$$\exists x, y, z, n \in \mathbb{Z} [x^n + y^n = z^n \wedge x, y, z \neq 0 \wedge n > 2] \quad (2)$$

soit fausse (ce fait a été — avec beaucoup de difficulté — démontré récemment) établit bien que la formule paramétrique précédente est fausse pour tous les quadruplets, et notamment pour $(12, 13, 14, 15)$.

Remarque. Insistons sur le fait que la formule 1 est paramétrique (et sans grand intérêt) alors que la formule 2 ne l'est pas ; il s'agit d'une “honnête” proposition, qui ne peut être que vraie ou

fausse, indépendamment de tout contexte.⁵¹ Cela n’a pas empêché les mathématiciens d’étudier cette formule pendant plus de trois siècles.

Notre introduction à la logique des prédicats se limitera à l’essentiel. On verra d’abord que l’interprétation d’une formule prédicative, quantifiée ou non, implique un domaine de référence D et l’association, à chaque prédicat, d’une relation sur ce domaine. Les constantes individuelles et les variables libres s’interprètent en des éléments de D . On peut aussi introduire des constantes fonctionnelles, dont l’interprétation sera naturellement une fonction qui, à tout n -uplet d’éléments de D , associe un élément de D .

On étudiera ensuite comment les procédures de décision introduites pour la logique propositionnelle s’adaptent à la logique prédicative ; nous verrons que ces techniques (tableaux sémantiques de Beth et Hintikka, séquents de Gentzen, systèmes axiomatiques de Hilbert et résolution de Davis, Putnam et Robinson) donnent lieu à des “semi-procédures” de décision.

5.2 Syntaxe du calcul des prédicats simplifié

Dans un premier temps, nous introduisons les prédicats, les variables et les constantes individuelles, mais pas les fonctions.

5.2.1 Lexique, termes et formules

Soit

- $\mathcal{P} = \{p, q, r, \dots\}$, un ensemble de symboles arbitraires appelés *symboles de prédicats* (chacun ayant une arité).

NB : Les propositions atomiques sont des symboles de prédicats d’arité 0.

- $\mathcal{A} = \{a, a_1, a_2, \dots, b, c, \dots\}$, un ensemble de symboles arbitraires appelés *constantes* (ou *constantes individuelles*).
- $\mathcal{X} = \{x, x_1, x_2, x', \dots, y, z, \dots\}$, un ensemble de symboles arbitraires appelés *variables* (ou *variables individuelles*).

Un *terme* est une constante $a \in \mathcal{A}$ ou une variable $x \in \mathcal{X}$. Une *formule atomique* (ou un *atome*) est une expression $p(t_1, \dots, t_n)$, où $p \in \mathcal{P}$ est un symbole prédicatif d’arité n et $t_1, \dots, t_n$ sont des termes. Le concept de *formule* est défini récursivement comme suit.

- Une formule atomique est une formule.
- *true*, *false* sont des formules.
- Si A_1 et A_2 sont des formules, alors $\neg A_1$, $(A_1 \vee A_2)$, $(A_1 \wedge A_2)$, $(A_1 \Rightarrow A_2)$ et $(A_1 \equiv A_2)$ sont des formules.
- Si A est une formule et x une variable, alors $\forall x A$ et $\exists x A$ sont des formules.

Comme dans le cadre propositionnel, on peut justifier l’omission de certaines parenthèses par des règles de précedence. Dans l’ordre décroissant, on a la négation et les quantificateurs, puis la conjonction, la disjonction, l’implication et enfin l’équivalence. Par exemple, la formule $\forall x((\neg \exists y p(x, y)) \vee (\neg \exists y p(y, x)))$ peut se récrire plus simplement en $\forall x(\neg \exists y p(x, y) \vee \neg \exists y p(y, x))$. Néanmoins, l’excès de concision peut nuire à la clarté. Dans la suite, nous utiliserons seulement le fait que la négation et les quantificateurs ont une précedence plus forte

⁵¹A condition d’accorder aux symboles qui composent la formule leur signification mathématique habituelle.

que les connecteurs binaires. Notons aussi que, dans le cas d’une formule dont l’opérateur principal est un connecteur binaire, il est d’usage d’omettre les parenthèses extérieures.⁵²

5.2.2 Portée des quantificateurs, variable libre, variable liée

- La *portée* d’une quantification (d’un quantificateur, d’une variable quantifiée) est la formule à laquelle la quantification s’applique. Dans $\forall x A$ ou dans $\exists x A$, la portée de x (de $\forall x$) est A .⁵³
- L’occurrence de la variable x dans la quantification $\forall x$ ou $\exists x$ est dite *quantifiée*.
- Toute occurrence de x dans la portée d’une quantification est dite *liée*.
- Une variable est *libre* si elle n’est ni quantifiée, ni liée.
- Les portées de deux variables x et y sont disjointes ou l’une est incluse dans l’autre. Ces notions existent aussi en programmation. Considérons l’exemple suivant :

```
program Principal ;
var x : integer ;

procedure p ;
var x : integer ;
begin x := 1 ; writeln(x + x) end ;

procedure q ;
var y : integer ;
begin y := 1 ; writeln(x + y) end ;

begin x := 5 ; p ; q end .
```

On remarque que les portées de x local et de y sont disjointes et incluses dans la portée de x global. Dans la procédure q , on se réfère au x global. De même, dans

$(+ x ((\text{lambda } (x) (+ x y)) x))$

les première et dernière occurrences de x sont libres, de même que la variable y , tandis que les deuxième et troisième occurrences de x sont liées. (On pourrait dire, plus justement, que la deuxième occurrence est “liante” et que la troisième est liée.)

Ces notions apparaissent également en mathématique, et notamment en algèbre et en analyse. Dans l’expression

$$C_{ij} = \sum_{k=1}^n A_{ik} B_{kj} ,$$

⁵²Selon la syntaxe adoptée ici, les formules quantifiées et les négations ne comportent pas de paire de parenthèses extérieures.

⁵³Les parenthèses extérieures d’une formule dont le connecteur principal est binaire peuvent être omises, mais il n’en découle pas, pour $A =_{\text{def}} p(x) \vee q(x)$, que la portée de $\forall x$ dans $\forall x p(x) \vee q(x)$ soit $p(x) \vee q(x)$; cette portée est $p(x)$. La formule $\forall x A$ doit s’écrire $\forall x (p(x) \vee q(x))$.

les variables i et j sont libres, la variable k est liée.⁵⁴ Dans l'expression

$$y(x) = y(x_0) + \int_{x_0}^x f(t, y(t)) dt,$$

la variable x est libre, la variable t est liée.

La distinction entre variable libre et variable liée se fait par simple inspection de la formule considérée. En revanche, la distinction entre constante et variable libre est moins immédiate. Le critère est qu'une variable libre est susceptible d'être quantifiée (et de devenir liée), tandis qu'une constante n'est jamais quantifiable. La distinction se fait par le contexte, ou par des conventions plus ou moins contraignantes et plus ou moins explicites. Dans la formule

$$S = \pi R^2,$$

il est "naturel" de considérer π comme la constante bien connue 3.14... , parce que l'égalité évoque la relation existant entre la surface d'un cercle et son rayon. En revanche, l'égalité

$$V = hb^2$$

évoque la relation entre le volume d'un parallélépipède à base carrée et ses dimensions b et h ; il sera alors tout aussi naturel de considérer h comme une variable libre. La confusion provient des libertés de notation que se permettent les mathématiciens. Les deux formules ci-dessus peuvent se récrire

$$\forall C \in \mathcal{C} [S(C) = \pi(R(C))^2],$$

et

$$\forall P \in \mathcal{P} [V(P) = h(P)(b(P))^2],$$

ce qui évite toute ambiguïté. Notons cependant que, parfois, le mathématicien est moins laxiste que le logicien. En analyse, on évitera d'écrire

$$y(x) = y(x_0) + \int_{x_0}^x f(x, y(x)) dx,$$

alors qu'en logique il n'est pas interdit d'écrire

$$P(x) \wedge \forall x Q(x),$$

même si nous préférons

$$P(x) \wedge \forall u Q(u).$$

Considérons quelques exemples.

$$1. \varphi_1 =_{def} \forall x (p(x, a) \Rightarrow \exists x q(x)).$$

On préférera éviter d'imbriquer plusieurs quantifications sur la même variable, sans toutefois l'interdire. En fait, on verra que la sémantique de φ_1 est exactement celle de

$$\forall x (p(x, a) \Rightarrow \exists y q(y)) \text{ ou encore de } \forall y (p(y, a) \Rightarrow \exists x q(x)).$$

⁵⁴Selon le contexte, A , B , C et n sont des constantes ou des variables libres.

$$2. \varphi_2 =_{def} \exists x \forall x A.$$

Ici aussi, deux quantifications sur x sont imbriquées. La sémantique de φ_2 est celle de $\exists y \forall x A$, pour n'importe quelle variable y sans occurrence (libre) dans A ; la quantification sur y étant inutile, la formule équivaut à $\forall x A$.

$$3. \varphi_3 =_{def} \forall x p(x, a) \Rightarrow \exists x q(x).$$

Deux variables liées ont le même nom, mais les portées sont disjointes ; il n'y a donc pas de problème.

$$4. \varphi_4 =_{def} \forall x p(x, a) \Rightarrow q(x).$$

Une variable libre et une variable liée ont le même nom x . C'est acceptable, mais il est préférable de *renommer* la variable liée et d'écrire, par exemple, $\forall y p(y, a) \Rightarrow q(x)$.

5.2.3 Fermatures universelle et existentielle

Une formule est *fermée* ou *close* si elle ne contient aucune variable libre. Lorsqu'une formule A contient les variables libres $x_1, x_2, \dots, x_n$, on la notera aussi $A(x_1, x_2, \dots, x_n)$.

Si $x_1, x_2, \dots, x_n$ sont toutes les variables libres d'une formule A ,

– $\forall x_1 \forall x_2 \dots \forall x_n A$ est la *fermeture universelle* de A .

– $\exists x_1 \exists x_2 \dots \exists x_n A$ est la *fermeture existentielle* de A .

Exemple. La fermeture universelle de la formule $p(x) \Rightarrow q(x)$ est $\forall x (p(x) \Rightarrow q(x))$ et non $\forall x p(x) \Rightarrow q(x)$.

5.3 Sémantique du calcul des prédicats

5.3.1 Interprétations

Une *interprétation* $\mathcal{I}$ est un triplet (D, I_c, I_v) tel que :

– D est un ensemble non vide, appelé *domaine d'interprétation* ;

– I_c est une fonction qui associe

– à toute *constante* a , un objet $I_c[a]$ appartenant à D ,

– à tout symbole prédicatif p (arité n), une relation (arité n) sur D , c'est-à-dire une fonction de D^n dans $\{\mathbf{V}, \mathbf{F}\}$;

– I_v est une fonction qui associe à toute variable x un élément $I_v[x]$ de D .

Voici quatre exemples d'interprétations pour la formule $\forall x p(a, x)$:

– $\mathcal{I}_1 = (\mathbb{N}, I_{1c}[p] = \leq, I_{1c}[a] = 0)$;

– $\mathcal{I}_2 = (\mathbb{N}, I_{2c}[p] = \leq, I_{2c}[a] = 1)$;

– $\mathcal{I}_3 = (\mathbb{Z}, I_{3c}[p] = \leq, I_{3c}[a] = 0)$;

– $\mathcal{I}_4 = (\mathcal{S}, I_{4c}[p] = \sqsubseteq, I_{4c}[a] = \lambda)$,

où $\mathcal{S}$ est l'ensemble des mots sur un alphabet donné ; $w_1 \sqsubseteq w_2$ signifie que w_1 est un préfixe de w_2 ; λ représente le mot vide.

5.3.2 Règles d'interprétation

Des règles sémantiques, appelées *règles d'interprétation*, permettent d'étendre une interprétation à l'ensemble des formules. En fait, une interprétation $\mathcal{I} = (D, I_c, I_v)$ associe

une valeur de vérité à toute formule A et associe un élément de D à tout terme t . En ce qui concerne les termes, on a

- Si x est une variable libre, $\mathcal{I}[x] = I_v[x]$.
- Si a est une constante, $\mathcal{I}[a] = I_c[a]$.

En ce qui concerne les formules, on a

- Si p est un symbole prédicatif d'arité n et si $t_1, \dots, t_n$ sont des termes, alors $\mathcal{I}[p(t_1, \dots, t_n)] = (I_c[p])(\mathcal{I}[t_1], \dots, \mathcal{I}[t_n])$.
- $\mathcal{I}[true] = \mathbf{V}$ et $\mathcal{I}[false] = \mathbf{F}$.
- Si A est une formule, alors $\neg A$ s'interprète comme dans le calcul des propositions, c'est-à-dire $\mathcal{I}[\neg A] = \mathbf{V}$ si $\mathcal{I}[A] = \mathbf{F}$ et $\mathcal{I}[\neg A] = \mathbf{F}$ si $\mathcal{I}[A] = \mathbf{V}$.
- Si A_1 et A_2 sont des formules, alors $(A_1 \vee A_2)$, $(A_1 \wedge A_2)$, $(A_1 \Rightarrow A_2)$, $(A_1 \equiv A_2)$ s'interprètent comme dans le calcul des propositions.
 $\mathcal{I}[(A_1 \wedge A_2)]$ vaut $\mathbf{V}$ si $\mathcal{I}[A_1] = \mathbf{V}$ et $\mathcal{I}[A_2] = \mathbf{V}$, et vaut $\mathbf{F}$ sinon.
 $\mathcal{I}[(A_1 \vee A_2)]$ vaut $\mathbf{V}$ si $\mathcal{I}[A_1] = \mathbf{V}$ ou $\mathcal{I}[A_2] = \mathbf{V}$, et vaut $\mathbf{F}$ sinon.
 $\mathcal{I}[(A_1 \Rightarrow A_2)]$ vaut $\mathbf{V}$ si $\mathcal{I}[A_1] = \mathbf{F}$ ou $\mathcal{I}[A_2] = \mathbf{V}$, et vaut $\mathbf{F}$ sinon.
 $\mathcal{I}[(A_1 \equiv A_2)]$ vaut $\mathbf{V}$ si $\mathcal{I}[A_1] = \mathcal{I}[A_2]$, et vaut $\mathbf{F}$ sinon.

Notation. Si $\mathcal{I} = (D_{\mathcal{I}}, I_c, I_v)$ est une interprétation, si x est une variable et si d est un élément de $D_{\mathcal{I}}$, alors $\mathcal{I}_{x/d}$ désigne l'interprétation $\mathcal{J} = (D_{\mathcal{J}}, J_c, J_v)$ telle que $D_{\mathcal{J}} = D_{\mathcal{I}}$, $J_c = I_c$, $J_v[x] = d$ et $J_v[y] = I_v[y]$ pour toute variable y distincte de x .

- Si A est une formule et x une variable, $\mathcal{I}[\forall x A]$ vaut $\mathbf{V}$ si $\mathcal{I}_{x/d}[A] = \mathbf{V}$ pour tout élément d de D , et vaut $\mathbf{F}$ sinon.
- Si A est une formule et x une variable, $\mathcal{I}[\exists x A]$ vaut $\mathbf{V}$ si $\mathcal{I}_{x/d}[A] = \mathbf{V}$ pour au moins un élément d de D , et vaut $\mathbf{F}$ sinon.

5.3.3 Capture de variable

Les règles d'interprétation des quantifications sont conformes à l'intuition traduite par les noms des quantificateurs. Il faut quand même souligner deux points importants, que l'emploi d'un même lexique pour les variables libres et les variables liées rend délicats :

- La valeur de $\mathcal{I}[\forall x A(x)]$ ne dépend pas de $\mathcal{I}[x]$.
- Si $\mathcal{I}[\forall x A(x)] = \mathbf{V}$, alors $\mathcal{I}[A(t)] = \mathbf{V}$, pour tout terme t ne donnant lieu à *aucune capture de variable*.

Exemple. Si $\forall x \exists y p(x, y)$ est vrai, alors les instances $\exists y p(a, y)$, $\exists y p(x, y)$ et $\exists y p(z, y)$ sont nécessairement vraies, mais l'instance $\exists y p(y, y)$ peut être fausse. Dans cette instance, l'occurrence y premier argument de p a été capturée et est devenue liée.

Conclusion. On ne peut pas dire que $\exists y p(y, y)$ est une instance licite de $\forall x \exists y p(x, y)$; on ne peut pas substituer y à x dans $\exists y p(x, y)$. Si on veut quand même effectuer cette substitution ou instantiation, on commencera par renommer la variable liée pour éviter la *capture*. On pourra dire, par exemple, que $\exists z p(y, z)$ est une instance (après renommage) de $\forall x \exists y p(x, y)$, ou le résultat de la substitution (après renommage) de y à x dans $\exists y p(x, y)$.

On omettra souvent de rappeler que les instantiations et substitutions donnant lieu à capture sont interdites ... tout en signalant une fois pour toutes que le phénomène de capture est à la source de nombreuses erreurs !

5.3.4 Satisfaction, modèle

Une formule A est vraie pour une interprétation $\mathcal{I}$ ou A est satisfaite par une interprétation $\mathcal{I}$ ou $\mathcal{I}$ est un modèle de A si $\mathcal{I}[A] = \mathbf{V}$. Cela se note $\models_{\mathcal{I}} A$.

Remarque. On rencontre parfois l'écriture $\mathcal{I} \models A$, mais nous ne l'emploierons pas dans ce cours, pour éviter tout risque de confusion avec l'écriture $U \models A$, introduite au paragraphe suivant.

Exemples. Soit A la formule $\forall x p(a, x)$. Les quatre interprétations introduites plus haut attribuent à A une valeur de vérité :

- $D_{\mathcal{I}_1} = \mathbb{N}$, $I_{1c}[p] = \leq$, $I_{1c}[a] = 0$; on a $\models_{\mathcal{I}_1} A$.
- $D_{\mathcal{I}_2} = \mathbb{N}$, $I_{2c}[p] = \leq$, $I_{2c}[a] = 1$; on a $\not\models_{\mathcal{I}_2} A$.
- $D_{\mathcal{I}_3} = \mathbb{Z}$, $I_{3c}[p] = \leq$, $I_{3c}[a] = 0$; on a $\not\models_{\mathcal{I}_3} A$.
- $D_{\mathcal{I}_4} = \mathcal{S}$, $I_{4c}[p] = \sqsubseteq$, $I_{4c}[a] = \lambda$; on a $\models_{\mathcal{I}_4} A$.

Définitions. Soit A une formule du calcul des prédicats.

- A est satisfaisable ou consistante si A a au moins un modèle.
- A est valide (cela se note $\models A$) si $\mathcal{I}[A] = \mathbf{V}$ pour toute interprétation $\mathcal{I}$.
- A est insatisfaisable ou inconsistant si A n'est pas satisfaisable, donc si $\mathcal{I}[A] = \mathbf{F}$ pour toute interprétation $\mathcal{I}$.
- A est simplement consistante ou contingente si A est consistante mais non valide.

Théorème (dualité validité – consistance). La formule A est valide si et seulement si $\neg A$ est inconsistant.

Exemples.

- $\forall x p(a, x)$ est consistante mais non valide.
 $D_{\mathcal{I}_1} = \mathbb{N}$, $I_{1c}[p] = \leq$, $I_{1c}[a] = 0$: $\models_{\mathcal{I}_1} A$.
 $D_{\mathcal{I}_3} = \mathbb{Z}$, $I_{3c}[p] = \leq$, $I_{3c}[a] = 0$: $\not\models_{\mathcal{I}_3} A$.
- $\forall x p(x) \Rightarrow p(a)$ est valide.
- $\exists x p(x) \Rightarrow p(a)$ est simplement consistante.

Remarque. Tout schéma propositionnel valide est aussi un schéma prédicatif valide. Par exemple, du schéma propositionnel valide $\neg \neg A \equiv A$, on peut déduire $\neg \neg(p \wedge q) \equiv (p \wedge q)$, mais aussi $\neg \neg \forall x p(x) \equiv \forall x p(x)$.

5.3.5 Quelques formules valides importantes

- $(\forall x A \wedge \forall x B) \equiv \forall x (A \wedge B)$
- $(\forall x A \vee \forall x B) \Rightarrow \forall x (A \vee B)$
- $\forall x (A \Rightarrow B) \Rightarrow (\forall x A \Rightarrow \forall x B)$
- $\forall x (A \equiv B) \Rightarrow (\forall x A \equiv \forall x B)$
- $\exists x (A \vee B) \equiv (\exists x A \vee \exists x B)$
- $\exists x (A \wedge B) \Rightarrow (\exists x A \wedge \exists x B)$
- $\exists x (A \Rightarrow B) \equiv (\forall x A \Rightarrow \exists x B)$
- $\forall x A \equiv \neg(\exists x \neg A)$
- $\forall x A \Rightarrow \exists x A$
- $\forall x \forall y A \equiv \forall y \forall x A$

- $\exists x \exists y A \equiv \exists y \exists x A$
- $\exists x \forall y A \Rightarrow \forall y \exists x A$

On observera que le remplacement d’une implication par une équivalence produit, dans chaque cas, une formule non valide. Considérons par exemple le cas de la formule $\exists x(A \wedge B) \Rightarrow (\exists x A \wedge \exists x B)$. Il est évident, vu la règle sémantique se rapportant à l’existentielle, que, si $C \Rightarrow D$ est valide, alors $\exists x C \Rightarrow \exists x D$ est valide. En conséquence, les deux formules $\exists x(A \wedge B) \Rightarrow \exists x A$ et $\exists x(A \wedge B) \Rightarrow \exists x B$ sont valides. D’autre part, si $C \Rightarrow D$ et $C \Rightarrow E$ sont vraies ou valides, alors $C \Rightarrow (D \wedge E)$ est vraie ou valide. Il en découle que $\exists x(A \wedge B) \Rightarrow (\exists x A \wedge \exists x B)$ est valide.

Pour montrer que l’implication inverse (ou réciproque, ou converse) n’est pas valide, il suffit d’en donner un antimodèle. On prend pour domaine l’ensemble $\mathbb{N}$; $A(x)$ est interprété en “ x est pair” et $B(x)$ en “ x est impair”. La formule $\exists x A \wedge \exists x B$ est vraie : elle signifie qu’il existe au moins un entier naturel pair, et au moins un entier naturel impair. La formule $\exists x(A \wedge B)$ est fausse : elle signifie qu’il existerait au moins un entier naturel à la fois pair et impair.

Notons enfin que le passage des quantifications informelles aux quantifications formelles (et réciproquement) est un exercice important, souvent facile, mais parfois délicat. Considérons un exemple :

Toutes les licornes sont dangereuses, donc il existe une licorne dangereuse.

Une modélisation hâtive telle que

$$\forall x LD(x) \Rightarrow \exists x LD(x)$$

pourrait laisser croire à la validité du raisonnement informel, ce qui serait incorrect. En effet, les licornes n’existent pas ; on peut donc les qualifier sans erreur de dangereuses (ou d’inoffensives), mais on ne peut pas affirmer qu’il existe une licorne, dangereuse ou non. Le paradoxe apparent disparaît si l’on utilise un modèle formel correct, à savoir

$$\forall x[L(x) \Rightarrow D(x)] \Rightarrow \exists x[L(x) \wedge D(x)]$$

Cette dernière formule est consistante mais n’est pas valide.

5.3.6 Conséquence logique, équivalence logique

Définitions. Soit U un ensemble de formules et soient A et B deux formules.

- A est une *conséquence logique* de U (cela se note $U \models A$) si A est vrai dans tous les modèles de U .

Remarque. En pratique, l’ensemble U sera souvent un ensemble de formules fermées.

On a alors $U \models A$ si et seulement si $U \models \forall x A$.

- A et B sont *logiquement équivalentes* (cela se note $A \leftrightarrow B$) si $\mathcal{I}[A] = \mathcal{I}[B]$ pour toutes les interprétations $\mathcal{I}$.

Comme dans le calcul des propositions, on a $\models A$ si et seulement si $\emptyset \models A$.

Théorème. Une formule est valide si et seulement si sa fermeture universelle est valide ; une formule est consistante si et seulement si sa fermeture existentielle est consistante.

Théorème. Deux formules A et B sont logiquement équivalentes si et seulement si la formule $A \equiv B$ est valide.

Remarque. Ces théorèmes découlent immédiatement des définitions et des règles d’interprétation. On notera que, si x est la seule variable libre de $A(x)$, alors $A(x)$ et $A(y)$ ne

sont en général pas logiquement équivalentes, mais $\forall x A(x)$ et $\forall y A(y)$ le sont. On évite des complications sans perdre d’expressivité en considérant les problèmes de validité, de consistance et de conséquence logique seulement pour les formules fermées. Lors de la fermeture d’une formule, l’ordre des quantifications n’a pas d’importance (c’est pourquoi on parle de “la” fermeture universelle ou existentielle d’une formule).

Théorème de l’échange. Soit A une sous-formule d’une formule B et soit A' une formule telle que $A \leftrightarrow A'$. Soit B' la formule résultant du remplacement de A par A' dans B . On a $B \leftrightarrow B'$.

Démonstration. Comme dans le cas propositionnel, on procède par induction structurale. Les seuls cas inductifs nouveaux sont liés à la quantification. Pour la quantification universelle, on doit seulement montrer que si $B(x) \leftrightarrow B'(x)$, on a aussi $\forall x B(x) \leftrightarrow \forall x B'(x)$, ce qui est évident.

5.4 Le théorème de compacité

Le théorème de compacité subsiste en logique prédicative, et un ensemble de formules est consistant si et seulement si tous ses sous-ensembles finis sont consistants. On peut prouver ce résultat important en adaptant la preuve donnée dans le cadre propositionnel, mais nous verrons un moyen plus rapide plus loin.

6 Analyse des formules prédicatives

6.1 Méthode simple pour formules simples

En logique propositionnelle, l'application directe des règles sémantiques permet toujours d'analyser une formule, c'est-à-dire de déterminer si elle est valide, contingente ou inconsistante. La méthode des tables de vérité concrétise cette approche fondamentalement simple. En logique prédicative, la situation est moins favorable parce qu'une formule consistante admet souvent une infinité de modèles très différents. Néanmoins, si on accepte certaines restrictions sur l'emploi des quantificateurs, l'approche sémantique directe reste possible.

6.1.1 Formules sans quantification

Une formule sans quantification est une combinaison booléenne de formules atomiques. De telles formules peuvent s'analyser par la méthode des tables de vérité, si on assimile tout atome à une proposition élémentaire. Considérons par exemple la formule Φ :

$$P(a, a) \wedge \neg P(a, x) \wedge Q(a, b) \wedge (Q(a, a) \Rightarrow P(a, x)).$$

La formule Φ comporte quatre atomes syntaxiquement distincts qui, par ordre d'occurrence, sont $P(a, a)$, $P(a, x)$, $Q(a, b)$ et $Q(a, a)$. La *version propositionnelle* de Φ s'obtient en substituant uniformément à ces quatre atomes les propositions élémentaires distinctes, par exemple p_1, p_2, p_3 et p_4 , respectivement, ce qui donne

$$p_1 \wedge \neg p_2 \wedge p_3 \wedge (p_4 \Rightarrow p_2)$$

On note immédiatement les lemmes suivants :

Lemme 1. Toute formule sans quantification admet une version propositionnelle unique.

Lemme 2. Si Φ est une formule sans quantification, la version propositionnelle de $\neg\Phi$ est la négation de la version propositionnelle de Φ .

Remarque. Deux formules sans quantification distinctes peuvent avoir la même version propositionnelle.

Définition. Une formule sans quantification est *p-valide* (resp. *p-consistante*, *p-contingente*) si sa version propositionnelle est valide (resp. consistante, contingente).

Exemple. La formule Φ donnée plus haut est *p-contingente*, puisque sa version propositionnelle est contingente.⁵⁵

Lemme 3. Une formule sans quantification est valide (resp. consistante, contingente) si et seulement si elle est *p-valide* (resp. *p-consistante*, *p-contingente*).

Démonstration. Vu le lemme 2, il suffit de démontrer qu'une formule sans quantification Φ admet un modèle si et seulement si sa version propositionnelle admet un modèle.

La condition est nécessaire. Soit I un modèle de Φ . L'interprétation I attribue une valeur de vérité à chacun des atomes de Φ . Soit J l'interprétation telle que $J(p_k)$ est la valeur associée par I au k ème atome de Φ ; l'interprétation J est un modèle de la version propositionnelle de Φ .

⁵⁵Cette version propositionnelle admet en fait un modèle et quinze anti-modèles.

La condition est suffisante. Soit J un modèle de la version propositionnelle de Φ . On construit un modèle I de Φ comme suit. Le domaine d'interprétation se compose des constantes et des variables de Φ . La fonction d'interprétation I applique chaque terme sur lui-même. Il reste à définir $I(P)$, pour tout prédicat P intervenant dans Φ . Si P est, par exemple, d'arité 2, il faut définir, vu le choix que nous avons fait pour D , $I(P(d_1, d_2))$ pour tous $d_1, d_2 \in D$. On distingue deux cas : si $P(d_1, d_2)$ est le k ème atome de Φ , on pose $I(P(d_1, d_2)) = J(p_k)$, sinon on choisit (arbitrairement) $I(P(d_1, d_2)) = \mathbf{V}$.

6.2 Méthode des tableaux sémantiques

Dans le cadre prédicatif comme dans le cadre propositionnel, la méthode des tableaux sémantiques consiste en une recherche systématique des modèles. Pour déterminer si la formule A est valide, on recherche un modèle de $\neg A$; si un tel modèle n'existe pas, A est valide ; s'il en existe (au moins) un, A n'est pas valide. De plus, la méthode des tableaux sémantiques est analytique : elle réduit une formule à ses composants. En ce sens, les composants d'une formule universelle $\forall x A(x)$ seront les formules $A(c)$, où c est n'importe quel terme ; le composant d'une formule universelle $\exists x A(x)$ sera la formule $A(a)$, où a est une constante inédite, appelée paramètre. Le traitement de la quantification étant délicat, nous en illustrerons d'abord les dangers. On se limite à l'étude des formules fermées, ce qui n'est pas une réelle restriction.

6.2.1 Quelques exemples

Exemple 1 (naïf). Test de validité de $\forall x (p(x) \Rightarrow q(x)) \Rightarrow (\forall x p(x) \Rightarrow \forall x q(x))$.

Dans le tableau 50, on a d'abord instancié $\neg \forall x q(x)$ (formule existentielle, équivalente à $\exists x \neg q(x)$), en $\neg q(a)$. On a ensuite instancié les formules universelles $\forall x p(x)$ et $\forall x (p(x) \Rightarrow q(x))$ en $p(a)$ et $p(a) \Rightarrow q(a)$. Intuitivement, une existentielle sera instanciée une seule fois (par branche), en utilisant une constante spécifique ; au contraire, une universelle pourra être instanciée plusieurs fois, au moyen de toutes les constantes disponibles. Pour rendre ceci rigoureux, il faudra préciser les mots "pourra", "spécifique" et "disponible".

Exemple 2 (incorrect !). Test de validité de $\forall x (p(x) \vee q(x)) \Rightarrow (\forall x p(x) \vee \forall x q(x))$.

Le tableau 51 montre simplement que sa racine n'admet pas de modèle à un élément. En revanche, elle admet un modèle à deux éléments (ce que le tableau ne montre pas ; il est donc incorrect !) et la formule testée n'est donc pas valide. Le problème est lié au choix de la même constante a dans l'instantiation de $\neg \forall x p(x)$ et de $\neg \forall x q(x)$. Cette identité est abusive : le fait que les formules $p(x)$ et $q(x)$ admettent chacune des "contre-exemples" n'impliquent pas qu'elles admettent des contre-exemples communs.

Exemple 2 (version corrigée). Test de validité de $\forall x (p(x) \vee q(x)) \Rightarrow (\forall x p(x) \vee \forall x q(x))$. On recommence donc la dérivation, en utilisant pour les existentielles deux constantes distinctes a et b . Cela implique naturellement que l'universelle soit instanciée au moyen de a et de b . Le tableau de la figure 52 comporte une branche ouverte à laquelle correspond un modèle $\mathcal{I}$ de la racine, tel que $\mathcal{I}[p(a)] = \mathcal{I}[q(b)] = \mathbf{V}$ et $\mathcal{I}[p(b)] = \mathcal{I}[q(a)] = \mathbf{F}$. Cette interprétation montre que la formule testée n'est pas valide (cf. fig. 52).

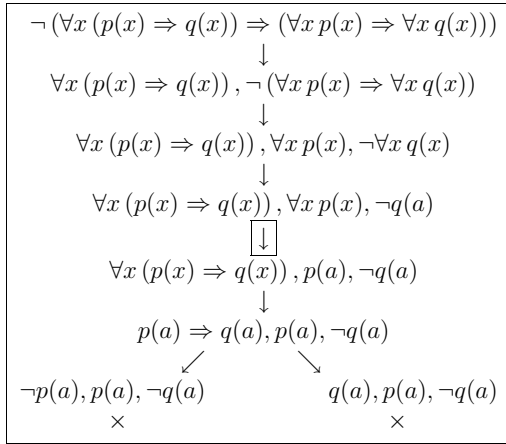


FIG. 50 – Exemple 1 (naïf) ; l'étape encadrée est fautive.

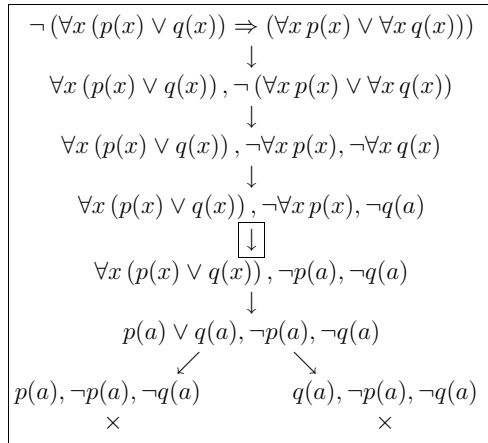


FIG. 51 – Exemple 2, tableau incorrect ; l'étape encadrée est fautive.

Exemple 3. Test de $\forall x \exists y p(x, y) \wedge \forall x \neg p(x, x) \wedge \forall x \forall y \forall z (p(x, y) \wedge p(y, z) \Rightarrow p(x, z))$.
Le tableau sémantique de la figure 53 est infini. Son unique branche doit être considérée comme ouverte car elle définit un modèle (nécessairement infini) de la formule testée.

Exemple 4. Test de validité pour la formule
 $\forall x \exists y p(x, y) \wedge \forall x \neg p(x, x) \wedge \forall x \forall y \forall z (p(x, y) \wedge p(y, z) \Rightarrow p(x, z)) \wedge \forall x (q(x) \wedge \neg q(x))$.
Si, dans le tableau 54, on instanciat indéfiniment $\forall x \exists y p(x, y)$, en négligeant à tort les autres formules, la branche ne se fermerait pas et serait infinie.

Les exemples 1 et 2 suggèrent que l'instantiation des existentielles, ou *exemplification*,

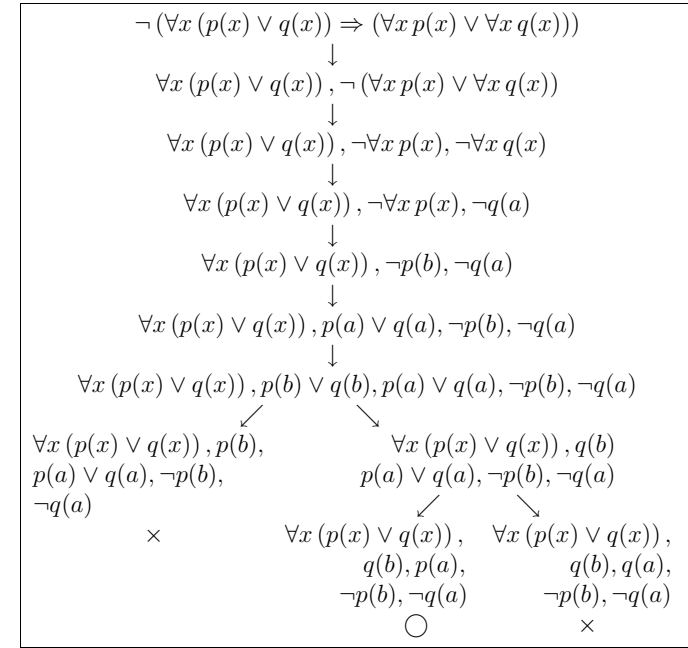


FIG. 52 – Exemple 2, tableau correct.

se fasse au moyen de constantes *inédites*, appelées aussi *paramètres*. Cela n'exclut pas les formules à “petits” modèles : en l'absence du prédicat spécial d'égalité, si $\{a, b\}$ par exemple est le domaine d'un modèle de A , $\{a, a_1, \dots, b, b_1, \dots\}$ donnera aussi lieu à un modèle, si les a_i et b_j sont des “clones” de a et b , c'est-à-dire tels que $\varphi, \varphi[a/a_i]$ et $\varphi[b/b_j]$ aient même valeur de vérité, pour toute formule φ .

L'exemple 3 montre que la construction d'un tableau sémantique peut ne pas se terminer, en particulier si la formule étudiée est consistante mais n'admet que des modèles infinis. On espère néanmoins que la méthode permettra toujours de reconnaître les formules inconsistantes, négations de formules valides.

L'exemple 4 indique enfin que cette inconsistance pourrait n'être pas reconnue si les règles de décomposition n'étaient pas appliquées de manière “équitable” ; il faut notamment se méfier de la règle *généralisatrice* d'instantiation des universelles, qui peut s'appliquer indéfiniment. On doit l'appliquer à toute constante introduite par la règle d'exemplification (sauf si la branche se ferme).

6.2.2 Règles de décomposition

Aux règles propositionnelles α et β s'ajoutent les règles prédictives γ et δ .

– Règles de prolongation (type α) et de ramification (type β)

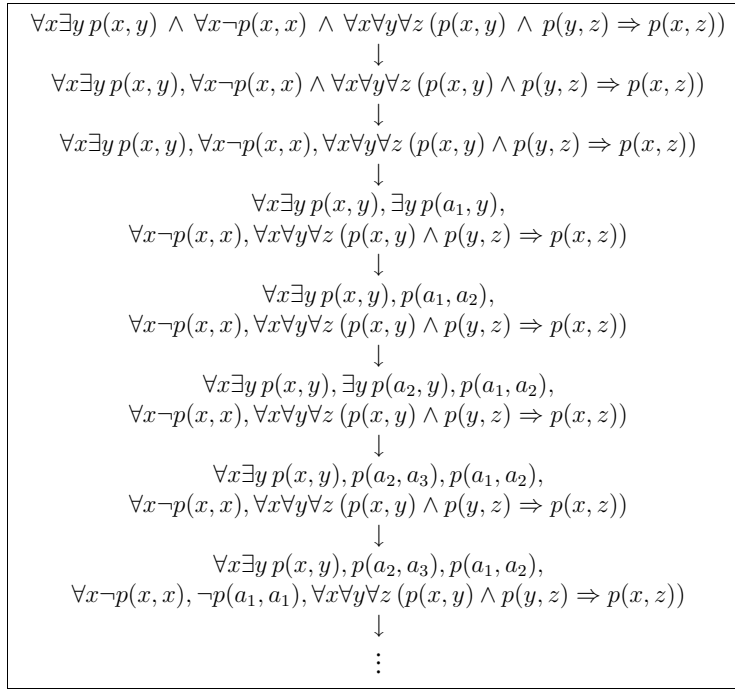


FIG. 53 – Exemple 3, tableau infini.

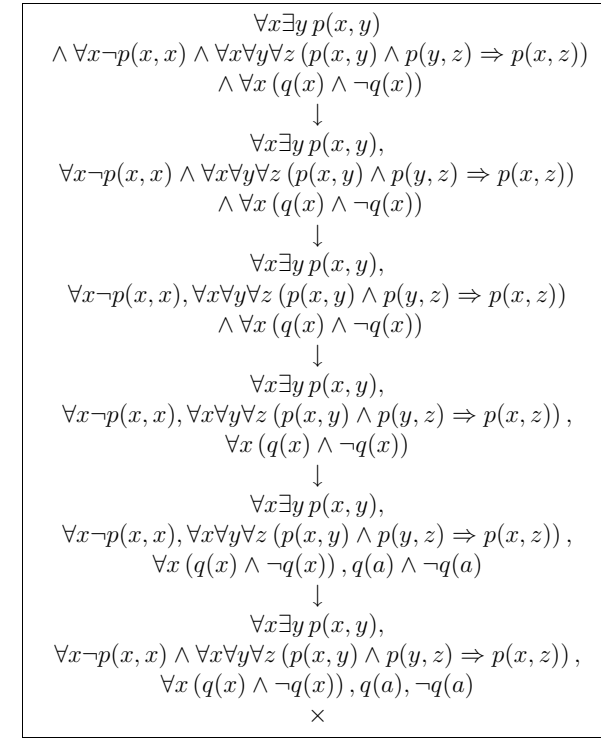


FIG. 54 – Exemple 4.

α	α_1	α_2
$A_1 \wedge A_2$	A_1	A_2
$\neg(A_1 \vee A_2)$	$\neg A_1$	$\neg A_2$
$\neg(A_1 \Rightarrow A_2)$	A_1	$\neg A_2$
$\neg(A_1 \Leftarrow A_2)$	$\neg A_1$	A_2

β	β_1	β_2
$B_1 \vee B_2$	B_1	B_2
$\neg(B_1 \wedge B_2)$	$\neg B_1$	$\neg B_2$
$B_1 \Rightarrow B_2$	$\neg B_1$	B_2
$B_1 \Leftarrow B_2$	B_1	$\neg B_2$

– Règles génératives (type γ) et exemplatives (type δ)

γ	$\gamma(c)$
$\forall x A(x)$	$A(c)$
$\neg \exists x A(x)$	$\neg A(c)$

(constante c quelconque)
(constante a inédite)

δ	$\delta(a)$
$\exists x A(x)$	$A(a)$
$\neg \forall x A(x)$	$\neg A(a)$

Rappelons aussi la règle d'élimination des doubles négations.

6.2.3 Construction d'un tableau sémantique

On présente d'abord l'algorithme, qui est non déterministe, puis des restrictions à ce non-déterminisme ; ces restrictions sont nécessaires pour assurer la terminaison (dans certains cas) et la complétude.

Algorithme de construction.

Initialisation : une racine étiquetée $\{A\}$.

Etape inductive : sélectionner une feuille non marquée ℓ ; soit $U(\ell)$ son étiquette.

- Si $U(\ell)$ contient une paire complémentaire, alors marquer ℓ comme *fermée* ' $\times$ ' ;
- Si $U(\ell)$ ne contient que des littéraux (sans paire complémentaire), alors marquer ℓ comme *ouverte* ' $\circ$ ' ;
- Si $U(\ell)$ n'est pas un ensemble de littéraux, sélectionner une formule dans $U(\ell)$:
 - si c'est une α -formule A , créer un nouveau nœud ℓ' , descendant de ℓ , et étiqueter ℓ' avec $U(\ell') = (U(\ell) \setminus \{A\}) \cup \{\alpha_1, \alpha_2\}$;
 - si c'est une β -formule B , créer deux nouveaux nœuds ℓ' et ℓ'' , descendants de ℓ , et étiqueter ℓ' avec $U(\ell') = (U(\ell) \setminus \{B\}) \cup \{\beta_1\}$ et étiqueter ℓ'' avec $U(\ell'') = (U(\ell) \setminus \{B\}) \cup \{\beta_2\}$;
 - si c'est une γ -formule C , créer un nouveau nœud ℓ' , descendant de ℓ , et étiqueter ℓ' avec $U(\ell') = U(\ell) \cup \{\gamma(c)\}$;
 - si c'est une δ -formule D , créer un nouveau nœud ℓ' , descendant de ℓ , et étiqueter ℓ' avec $U(\ell') = (U(\ell) \setminus \{D\}) \cup \{\delta(a)\}$, où a est une constante qui n'apparaît pas dans

$U(\ell)$.⁵⁶

Terminaison : survient quand toutes les feuilles sont marquées.

Règles additionnelles de construction.

Le non-déterminisme de l'algorithme de construction intervient

1. lors du choix du nœud à développer ;
2. lors du choix de la formule à décomposer dans ce nœud ;
3. lors du choix du terme c lors d'une γ -réduction.⁵⁷

Il faut adopter une stratégie qui garantisse les deux conditions suivantes.

- Toute formule qui apparaît sur une branche ouverte de l'arbre se voit appliquer une règle de décomposition quelque part sur cette branche.
Autrement dit, toute formule décomposable est décomposée, à moins que la branche se ferme.
- Pour toute γ -formule A et toute constante a qui apparaissent sur une branche ouverte, une règle d'instanciation est appliquée à la formule A avec la constante a quelque part sur cette branche.
Toute constante apparaissant sur une branche est utilisée à un moment donné pour instancier les γ -formules sur cette branche, à moins qu'elle se ferme.

Un moyen simple et classique d'assurer le respect des conditions d'équité est d'étiqueter les nœuds par des *listes* de formules. Le(s) nœud(s) successeur(s) de n est (sont) obtenus par "décomposition" de la première formule de la liste $U(n)$ non réduite à un littéral ; la liste $U(n')$ (et $U(n'')$, s'il y a lieu) est obtenue en supprimant de $U(n)$ la formule traitée, et en ajoutant en fin de liste le(s) "composant(s)" adéquats. Dans le cas d'une formule générative, la formule supprimée en tête de liste est réinsérée en queue de liste.

Une autre méthode appropriée est la suivante. Lorsqu'une règle générative est activée, on construit immédiatement les instances correspondant à toutes les constantes introduites jusque là dans la branche. De même, quand une exemplification est faite, ce qui provoque l'adjonction dans la branche d'une constante inédite, on "réactive" les γ -réductions déjà accomplies, pour insérer les instances correspondant à cette nouvelle constante. Ceci nécessite une gestion organisée de l'ensemble des constantes et des activations de règles génératives.

La stratégie n'est pas nécessaire pour obtenir l'adéquation, mais elle l'est pour obtenir la complétude. En effet, la stratégie ne vise qu'à éviter le report définitif de réductions susceptibles de fermer une branche. Le point est d'ailleurs délicat, puisque la construction d'un tableau sémantique peut ne pas se terminer.

⁵⁶On voit que cette constante n'apparaît pas non plus dans l'étiquette d'un ancêtre de ℓ ; cette contrainte devrait être introduite explicitement si on convenait de ne pas récrire les littéraux étiquetant un nœud dans l'étiquette de ses successeurs (convention que l'on adopte parfois pour alléger la construction). Dans ce cas, il convient de préciser que la recherche de paires complémentaires se fait dans toute la branche, et non seulement dans son dernier nœud.

⁵⁷Lors d'une δ -réduction, la constante choisie doit être inédite ; on a vu que le non-respect de cette condition rendait la méthode inadéquate (exemple 1). En revanche, le choix du nom de cette constante inédite est clairement sans importance ; les δ -réductions, au contraire des γ -réductions, n'introduisent donc pas de vrai non-déterminisme.

Rappelons enfin que certaines règles de priorité permettent souvent d'accélérer la construction du tableau. En particulier, on effectuera les α -réductions avant les β -réductions, pour limiter le nombre de branchements. On évitera d'instancier une γ -formule par une constante inédite (c'est inutile) sauf naturellement dans le cas où aucune δ -réduction n'a pu être effectuée. L'exemple 5 (fig. 55) illustre certaines de ces règles. Il illustre aussi un point délicat. La formule $\forall x \exists y r(x, y) \Rightarrow \exists y \forall x r(x, y)$, où r est un prédicat binaire, est non valide ; on en déduit naturellement que, si $R(x, y)$ est une formule quelconque admettant x et y comme variables libres, la formule $\forall x \exists y R(x, y) \Rightarrow \exists y \forall x R(x, y)$ est *généralement* non valide. Pour certains choix de R , la formule peut cependant être valide ; c'est le cas notamment si $R(x, y)$ est $p(x) \Rightarrow q(y)$.

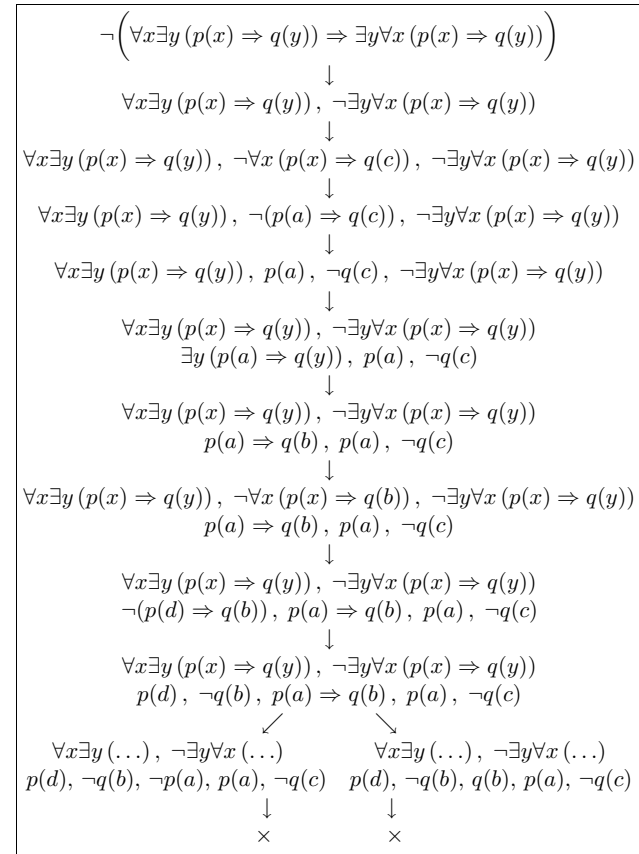


FIG. 55 – Exemple 5.

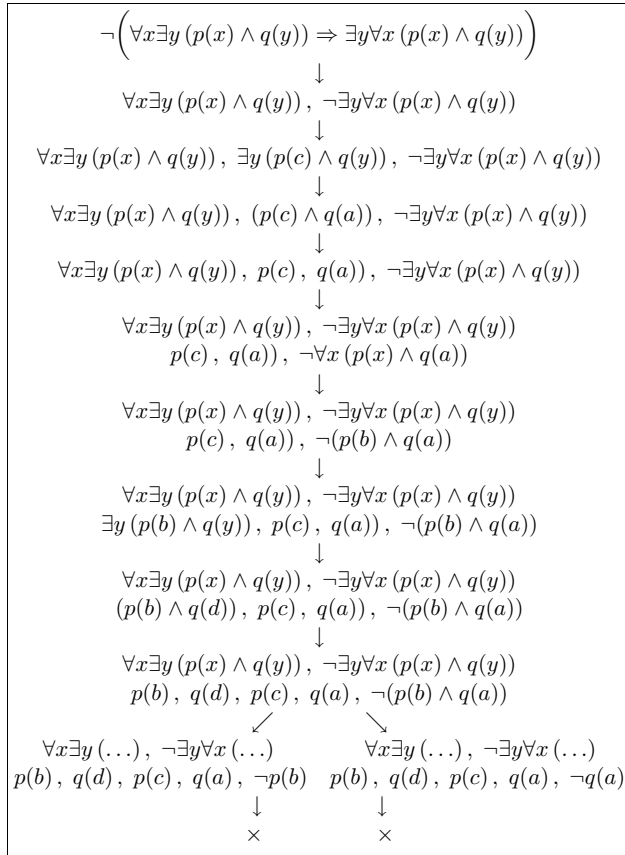


FIG. 56 – Exemple 6.

6.2.4 Adéquation de la méthode des tableaux sémantiques

Théorème. Soit $T(A)$ un tableau sémantique dont la racine est A . Si $T(A)$ est fermé,⁵⁸ alors la formule A est inconsistante.

Remarque. Vu que tout sous-arbre d'un tableau fermé est aussi un tableau fermé, on prouvera un résultat apparemment plus fort, à savoir que les étiquettes de tous les nœuds (pas seulement la racine) d'un tableau fermé sont inconsistantes.

Démonstration. On prouve par induction sur la hauteur h du nœud n dans $T(A)$ que l'étiquette de n est un ensemble inconsistent. Ce sera vrai en particulier pour la racine de l'arbre, dont l'étiquette est le singleton $\{A\}$.

⁵⁸c'est-à-dire si toutes les branches de $T(A)$ sont fermées.

- Cas de base, $h = 0$: le nœud n est une feuille, nécessairement fermée, donc $U(n)$ contient une paire complémentaire et est inconsistent.
- Cas inductif, $h > 0$: une règle α, β, γ ou δ a été utilisée pour créer le(s) descendant(s) du nœud n . Les cas α et β sont les mêmes que dans la démonstration de la version propositionnelle du théorème. On considère successivement les cas γ et δ .

– **Règle γ** : $n : \{ \forall x A(x) \} \cup U_0$

$$\downarrow$$

$$n' : \{ \forall x A(x), A(c) \} \cup U_0$$

$U(n')$ est inconsistent par hypothèse inductive, donc $U(n)$ est inconsistent; en effet, tout modèle de $U(n)$ serait aussi un modèle de $U(n')$, ou s'étendrait immédiatement en un tel modèle, au cas où la constante c n'interviendrait pas dans U_0 .

– **Règle δ** : $n : \{ \exists x A(x) \} \cup U_0$

$$\downarrow$$

$$n' : \{ A(a) \} \cup U_0$$

où a est une constante qui n'apparaît pas dans $U(n)$. Si $U(n)$ était consistant, il existerait une interprétation $\mathcal{I} = (D, I_c, I_v)$ telle que $\mathcal{I}[\exists x A(x)] = \mathbf{V}$, donc il existerait $d \in D$ tel que $\mathcal{I}_{x/d}[A(x)] = \mathbf{V}$.

Définissons $\mathcal{J} = (D, J_c, I_v)$ avec J_c obtenu en étendant⁵⁹ I_c de sorte que $J_c[a] = d$. Alors, $\mathcal{J}[A(a)] = \mathbf{V}$ et $\mathcal{J}[U_0] = \mathcal{I}[U_0] = \mathbf{V}$, donc $\mathcal{J}$ satisfait $U(n')$, une contradiction.

6.2.5 Complétude de la méthode des tableaux sémantiques

On abordera la complétude comme dans le cas propositionnel, via la notion d'ensemble de Hintikka.

Ensembles de Hintikka. *Définition.* Soit U un ensemble de formules fermées, et C_U l'ensemble des constantes individuelles ayant au moins une occurrence dans U . L'ensemble U est un *ensemble de Hintikka* si les cinq conditions suivantes sont satisfaites :

1. Si A est une formule atomique, on a $A \notin U$ ou $\neg A \notin U$.
2. Si $\alpha \in U$ est une α -formule, alors $\alpha_1 \in U$ et $\alpha_2 \in U$.
3. Si $\beta \in U$ est une β -formule, alors $\beta_1 \in U$ ou $\beta_2 \in U$.
4. Si γ est une γ -formule, alors pour tout $a \in C_U$ on a $\gamma(a) \in U$.
5. Si δ est une δ -formule, alors il existe $a \in C_U$ tel que $\delta(a) \in U$.

Théorème. Soit b une branche ouverte d'un tableau T construit en respectant les conditions d'équité. L'ensemble $U = \bigcup_{n \in b} U(n)$ est de Hintikka.

Démonstration. La condition d'ouverture assure le respect par U de la condition 1. Les règles α, β, γ et δ permettent l'insertion dans U des éléments requis par les conditions 2, 3, 4 et 5, respectivement. La stratégie de construction impose que tout élément ajoutable soit effectivement ajouté.

⁵⁹On est sûr de pouvoir procéder à l'extension puisque a est une constante inédite, telle que $I_c[a]$ n'existe pas.

Remarque. La branche b peut être infinie ; dans ce cas, elle est nécessairement ouverte et elle définit un modèle infini.

Lemme de Hintikka. Tout ensemble de Hintikka est consistant.

Démonstration. Soit U un ensemble de Hintikka. Le modèle canonique $\mathcal{I}_U = (D, I_c, I_v)$ associé à U est défini comme suit :

- $D = \{a, b, \dots\}$ est l'ensemble des constantes apparaissant dans les formules de U ;
- On construit la fonction d'interprétation I_c comme suit :
 - Pour toute constante $d \in D$, on pose $I_c[d] = d$.
 - Pour tout symbole prédicatif p (arité m) apparaissant dans U , on pose
$$I_c[p](I_c[a_1], \dots, I_c[a_m]) = \mathbf{V}, \text{ si } p(a_1, \dots, a_m) \in U,$$

$$I_c[p](I_c[a_1], \dots, I_c[a_m]) = \mathbf{F}, \text{ si } \neg p(a_1, \dots, a_m) \in U.$$

$$I_c[p](I_c[a_1], \dots, I_c[a_m]) \text{ est arbitraire si } \{p(a_1, \dots, a_m), \neg p(a_1, \dots, a_m)\} \cap U = \emptyset.$$
- I_v est quelconque, puisqu'il n'y a pas de variables libres.

Il reste à montrer que pour toute formule (fermée) $A \in U$, on a $\mathcal{I}[A] = \mathbf{V}$. Cela se fait par induction sur la structure de A . (Exercice.)

Complétude. Comme dans le cas propositionnel, on peut montrer que si un tableau sémantique (respectant la stratégie de construction) est ouvert, alors la formule étiquetant la racine est consistante. Un modèle est le modèle canonique de Hintikka associé à une branche ouverte. On en déduit que si A est une formule inconsistante, tout tableau $T(A)$ (respectant la stratégie de construction) est fermé. Rappelons que, dans le cadre prédicatif, une branche peut être infinie. Cependant, si on respecte la stratégie de construction, une branche infinie est nécessairement ouverte.

Pour analyser une formule (fermée) A , on peut construire les tableaux sémantiques $T(A)$ et $T(\neg A)$. Si $T(A)$ est fermé (et donc fini), A est inconsistent. Si $T(\neg A)$ est fermé (et donc fini), A est valide. Si $T(A)$ et $T(\neg A)$ sont ouverts, A et $\neg A$ sont simplement consistants.

Dans le cas propositionnel, l'analyse se termine toujours. Dans le cas prédicatif, l'analyse peut ne pas se terminer si A et $\neg A$ sont simplement consistants. Cela n'a rien d'étonnant ; contrairement au calcul des propositions, le calcul des prédicats n'est que semi-décidable.

6.3 Méthode des séquents

6.3.1 Dualité entre séquents et tableaux

Comme dans le cas propositionnel, on peut “par dualité” obtenir une dérivation de séquent au départ d'un tableau sémantique.

Un exemple suffira à rappeler le procédé. La validité de la formule

$$(\forall x p(x) \vee \forall x q(x)) \Rightarrow \forall x (p(x) \vee q(x))$$

est démontrée par la méthode des tableaux (figure 57) puis par celle des séquents sans antécédent (figure 58). D'une figure à l'autre, l'arbre est retourné et chaque formule est remplacée par son complément. Les feuilles fermées deviennent des séquents valides ou *axiomes* ; les feuilles ouvertes deviennent des séquents non valides ou *hypotheses*. Dans les tableaux sémantiques, les ensembles de formules sont conjonctifs et la virgule a donc valeur

conjonctive. Dans les séquents, la virgule a valeur disjonctive quand elle se trouve à droite de la flèche, dans le succédent. Un séquent peut aussi avoir un antécédent, dans lequel la virgule a valeur conjonctive. On ne change pas la sémantique d'un séquent en faisant passer l'une de ses formules du succédent vers l'antécédent ou réciproquement, à condition de changer sa polarité. Par exemple, les quatre séquents ci-dessous sont équivalents :

$$\rightarrow A, \neg B \quad B \rightarrow A \quad \neg A \rightarrow \neg B \quad \neg A, B \rightarrow$$

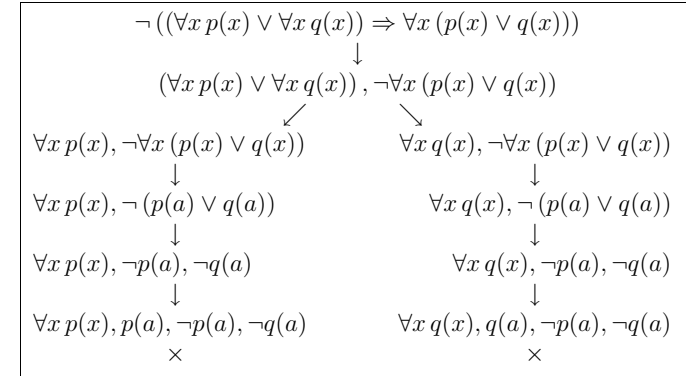


FIG. 57 – Un tableau sémantique ...

$\mathbf{A}$	$\mathbf{A}$
$\rightarrow \neg \forall x p(x), \neg p(a), p(a), q(a)$	$\rightarrow \neg \forall x q(x), \neg q(a), p(a), q(a)$
$\rightarrow \neg \forall x p(x), p(a), q(a)$	$\rightarrow \neg \forall x q(x), p(a), q(a)$
$\rightarrow \neg \forall x p(x), (p(a) \vee q(a))$	$\rightarrow \neg \forall x q(x), (p(a) \vee q(a))$
$\rightarrow \neg \forall x p(x), \forall x (p(x) \vee q(x))$	$\rightarrow \neg \forall x q(x), \forall x (p(x) \vee q(x))$
$\rightarrow \neg (\forall x p(x) \vee \forall x q(x)), \forall x (p(x) \vee q(x))$	
$\rightarrow (\forall x p(x) \vee \forall x q(x)) \Rightarrow \forall x (p(x) \vee q(x))$	

FIG. 58 – ... et la dérivation de séquent duale.

6.3.2 Règles du système de Gentzen

Comme dans le cas propositionnel, un séquent est un axiome si le même atome (variante : la même formule) apparaît dans l'antécédent et dans le succédent. En outre, les règles α et β

A	A
$\forall x p(x), p(a) \rightarrow p(a), q(a)$	$\forall x q(x), q(a) \rightarrow p(a), q(a)$
$\forall x p(x) \rightarrow p(a), q(a)$	$\forall x q(x) \rightarrow p(a), q(a)$
$\forall x p(x) \rightarrow (p(a) \vee q(a))$	$\forall x q(x) \rightarrow (p(a) \vee q(a))$
$\forall x p(x) \rightarrow \forall x (p(x) \vee q(x))$	$\forall x q(x) \rightarrow \forall x (p(x) \vee q(x))$
$\frac{(\forall x p(x) \vee \forall x q(x)) \rightarrow \forall x (p(x) \vee q(x))}{\rightarrow (\forall x p(x) \vee \forall x q(x)) \Rightarrow \forall x (p(x) \vee q(x))}$	

FIG. 59 – Dérivation, séquents avec antécédent.

sont les mêmes que dans le cadre propositionnel. Nous rappelons seulement celles relatives à l'implication.

- *règle α* :
- $$\frac{U, A \rightarrow V, B}{U \rightarrow V, (A \Rightarrow B)}$$
- *règle β* :
- $$\frac{U \rightarrow V, A \quad U, B \rightarrow V}{U, (A \Rightarrow B) \rightarrow V}$$

On ajoute des règles génératives (règles γ) et les règles d'exemplification (règles δ), pour traiter les formules quantifiées :

- *règles γ* :
- $$\exists : \frac{U \rightarrow V, \exists x A(x), A(c)}{U \rightarrow V, \exists x A(x)}$$
- $$\forall : \frac{U, \forall x A(x), A(c) \rightarrow V}{U, \forall x A(x) \rightarrow V}$$
- *règles δ* :
- $$\forall : \frac{U \rightarrow V, A(a)}{U \rightarrow V, \forall x A(x)} \text{ si } a \text{ n'apparaît pas dans la conclusion.}$$
- $$\exists : \frac{U, A(a) \rightarrow V}{U, \exists x A(x) \rightarrow V} \text{ si } a \text{ n'apparaît pas dans la conclusion.}$$

La dérivation de la figure 58 (séquents sans antécédent) est reprise dans le cadre général à la figure 59. La figure 60 montre la dérivation relative à une formule importante. On voit à la figure 61 comment l'analyse d'une formule non valide peut conduire à une séquence infinie. La figure 62 illustre le danger du non-respect de la restriction attachée à la règle δ . C'est cette restriction qui empêcherait ici la fermeture (incorrecte). La dérivation de la figure 62 montre que la formule est vraie dans un domaine réduit à un élément, sans mettre en évidence le fait

— essentiel — que la même formule est le plus souvent fausse dans un domaine comportant plusieurs éléments.

A	
$\forall y p(a, y), p(a, b), p(a, a) \rightarrow \exists x p(x, b), p(a, b)$	(règle $\gamma, \forall$)
$\forall y p(a, y), p(a, a) \rightarrow \exists x p(x, b), p(a, b)$	(règle $\gamma, \forall$)
$\forall y p(a, y) \rightarrow \exists x p(x, b), p(a, b)$	(règle $\gamma, \exists$)
$\forall y p(a, y) \rightarrow \exists x p(x, b)$	(règle $\delta, \forall$)
$\forall y p(a, y) \rightarrow \forall y \exists x p(x, y)$	(règle $\delta, \exists$)
$\exists x \forall y p(x, y) \rightarrow \forall y \exists x p(x, y)$	(règle $\alpha, \Rightarrow$)
$\rightarrow \exists x \forall y p(x, y) \Rightarrow \forall y \exists x p(x, y)$	

FIG. 60 – Validité de $\exists x \forall y p(x, y) \Rightarrow \forall y \exists x p(x, y)$.

H?	
$\forall \exists, p(d, b), p(e, c), p(c, a) \rightarrow \exists \forall, p(b, f), p(c, g), p(a, b)$	δ
$\forall \exists, \exists x p(x, b), \exists x p(x, c), p(c, a) \rightarrow \exists \forall, \forall y p(b, y), \forall y p(c, y), p(a, b)$	γ
$\forall y \exists x p(x, y), p(c, a) \rightarrow \exists x \forall y p(x, y), p(a, b)$	δ
$\forall y \exists x p(x, y), \exists x p(x, a) \rightarrow \exists x \forall y p(x, y), \forall y p(a, y)$	γ
$\forall y \exists x p(x, y) \rightarrow \exists x \forall y p(x, y)$	$\alpha \Rightarrow$
$\rightarrow \forall y \exists x p(x, y) \Rightarrow \exists x \forall y p(x, y)$	

FIG. 61 – Non-validité de $\forall y \exists x p(x, y) \Rightarrow \exists x \forall y p(x, y)$.

6.3.3 Propriétés du système de Gentzen

Adéquation et complétude. Une formule A est valide si et seulement si elle est racine d'une dérivation de séquent finie dont toutes les feuilles sont des axiomes.

A	
$!! \forall y \exists x p(x, y), p(a, a) \rightarrow \exists x \forall y p(x, y), p(a, a) !!$	
$\forall y \exists x p(x, y), \exists x p(x, a) \rightarrow \exists x \forall y p(x, y), \forall y p(a, y)$	$!! \delta !!$
$\forall y \exists x p(x, y) \rightarrow \exists x \forall y p(x, y)$	γ
$\rightarrow \forall y \exists x p(x, y) \Rightarrow \exists x \forall y p(x, y)$	$\alpha \Rightarrow$

FIG. 62 – Dérivation incorrecte de $\forall y \exists x p(x, y) \Rightarrow \exists x \forall y p(x, y)$.

Terminaison. L'obtention d'une dérivation adéquate exige le respect d'une stratégie équitable, comme pour les tableaux sémantiques. Malgré cela, l'analyse d'une formule non valide peut donner lieu à une dérivation infinie.

Analyticité. Une règle est *analytique* si tous les composants (formules et sous-formules) des prémisses apparaissent dans la conclusion. Les règles α et β sont analytiques. L'idée sous-jacente est que la découverte de prémisses appropriées au départ de la conclusion doit être triviale. En ce sens, on peut considérer que les règles δ sont analytiques. Pour les règles γ , le choix de la constante c devient critique s'il peut être effectué d'une infinité de manières. Ce sera le cas pour le calcul des prédicats avec symboles fonctionnels (pour l'instant, c ne peut être qu'une constante individuelle).

Réversibilité. Tout modèle des prémisses d'une règle (correcte) est aussi un modèle de sa conclusion. Une règle est *réversible* si la réciproque est également vraie. Les règles α , β et γ sont réversibles. Les règles δ sont "quasi réversibles" : tout modèle de la conclusion peut être étendu en un modèle de la prémisse. Dans tous les cas, la validité de la conclusion implique celle de la ou des prémisses(s).

Remarque. La correction des règles δ n'est pas évidente ; elle dépend crucialement de la condition imposée à la constante a .

6.4 Système axiomatique de Hilbert

6.4.1 Définition du système

Le système formel $\mathcal{H}$ a déjà été présenté dans sa version propositionnelle, cadre où il présentait peu d'intérêt, vu l'existence de procédures de décision relativement efficaces. La version prédicative, plus intéressante, est constituée de

– cinq schémas d'axiomes :

1. $\vdash A \Rightarrow (B \Rightarrow A)$
2. $\vdash (A \Rightarrow (B \Rightarrow C)) \Rightarrow ((A \Rightarrow B) \Rightarrow (A \Rightarrow C))$
3. $\vdash (\neg B \Rightarrow \neg A) \Rightarrow (A \Rightarrow B)$
4. $\vdash \forall x A(x) \Rightarrow A(t)$ (sauf capture)
5. $\vdash \forall x (A \Rightarrow B(x)) \Rightarrow (A \Rightarrow \forall x B(x))$ où x n'est pas libre dans A

– la règle d'inférence *Modus Ponens* :

$$\frac{\vdash A \quad \vdash A \Rightarrow B}{\vdash B}$$

– la règle de *Généralisation* :

$$\frac{\vdash A(x)}{\vdash \forall x A(x)}$$

Remarque. La restriction relative à la capture est naturellement essentielle ; l'instance $\forall x \exists y p(x, y) \Rightarrow \exists y p(y, y)$ ne peut pas être un axiome, parce que ce n'est pas une formule valide. De même, l'instance $\forall x (p(x) \Rightarrow p(x)) \Rightarrow (p(x) \Rightarrow \forall x p(x))$ du schéma 5 ne peut pas être un axiome.

Remarques. L'expression

$$A \Rightarrow (B \Rightarrow A)$$

est un schéma d'axiome ; cela implique, notamment, que la formule

$$(p \Rightarrow q) \Rightarrow ((\neg p \Rightarrow r) \Rightarrow (p \Rightarrow q))$$

est un axiome. L'expression

$$\vdash A \Rightarrow (B \Rightarrow A)$$

est une assertion ; cette assertion exprime que $A \Rightarrow (B \Rightarrow A)$ est un (schéma de) théorème, c'est-à-dire, est dérivable dans le système de Hilbert. Notons enfin que les notions de preuve, de dérivation et de théorème sont les mêmes que dans le cadre propositionnel.

6.4.2 Règle de déduction

La règle de déduction est une règle d'inférence dérivée, déjà introduite dans le cadre propositionnel :

$$\frac{U, A \vdash B}{U \vdash A \Rightarrow B}$$

Cette règle reste correcte dans le cadre prédicatif, à condition de respecter lors de son emploi une restriction essentielle : dans la déduction de B à partir de $U \cup \{A\}$, on n'utilise pas la règle de généralisation sur une variable ayant une occurrence libre dans A . Le non-respect de cette restriction conduit aisément à des "théorèmes" non valides. Par exemple, $p(x) \vdash \forall x p(x)$ est licite (il suffit de généraliser $p(x) \vdash p(x)$)⁶⁰ mais $p(x) \Rightarrow \forall x p(x)$ ne peut être un théorème puisque ce n'est pas une formule valide.

Remarque. Il est préférable d'interdire les variables libres à gauche du symbole $\vdash$, comme d'ailleurs à gauche du symbole $\models$. Cette restriction n'est pas réellement gênante.

Pour justifier la règle de déduction, on doit montrer que toute conclusion obtenue en l'utilisant aurait pu aussi être obtenue sans l'utiliser. Dans l'optique constructive que nous adoptons, cette preuve sera fournie par une technique de conversion d'une dérivation du type $U, A \vdash B$ en une dérivation (nettement plus longue) de $U \vdash (A \Rightarrow B)$. Cela se fait en adaptant

⁶⁰Dans certaines variantes du système axiomatique de Hilbert, ce genre de dérivation indésirable est illicite, ce qui est intéressant. Le système que nous présentons ici est néanmoins plus simple, globalement.

la technique donnée dans le cas propositionnel. Le seul point délicat est la conversion des fragments du type

$$\begin{aligned} U, A \vdash C(x), \\ U, A \vdash \forall x C(x). \end{aligned}$$

La conversion est

$$\begin{aligned} U \vdash A \Rightarrow C(x), \\ U \vdash \forall x (A \Rightarrow C(x)), \\ U \vdash \forall x (A \Rightarrow C(x)) \Rightarrow (A \Rightarrow \forall x C(x)), \\ U \vdash A \Rightarrow \forall x C(x). \end{aligned}$$

La troisième ligne n'est correcte que si x n'intervient pas dans A , d'où la restriction concernant l'emploi de la règle de généralisation dans les dérivations.

6.4.3 Substitution uniforme, échange

Le *principe de substitution uniforme* reste valable dans le cadre prédicatif. Cela revient à dire qu'un théorème propositionnel donne lieu à un schéma de théorème. Par exemple,

$$\neg p \Rightarrow (p \Rightarrow q)$$

est un théorème, donc

$$\neg A \Rightarrow (A \Rightarrow B)$$

est un schéma de théorème, et

$$\neg \forall x P(x) \Rightarrow (\forall x P(x) \Rightarrow \forall y (R(y) \Rightarrow Q(z)))$$

est un théorème. L'adaptation de la démonstration donnée en logique propositionnelle est immédiate.

La *règle de l'échange* reste valable aussi. Par exemple, si $A \equiv B$ est un théorème et si C est un théorème, toute formule obtenue en remplaçant une ou plusieurs occurrences de A par B dans C sera aussi un théorème.

Exercice. Justifier formellement les règles !

Mentionnons encore deux règles dérivées élémentaires mais utiles. Dans le cadre prédicatif, on appelle souvent *tautologie* une formule qui s'obtient par substitutions uniformes au départ d'une tautologie propositionnelle. La règle *propositionnelle* (notée PC) affirme que toute tautologie est un théorème. C'est un corollaire immédiat de la propriété de complétude du système de Hilbert propositionnel et du principe de substitution uniforme (dans sa version prédicative).

Remarques. Vu la règle propositionnelle, on autorise dans la suite les connecteurs $\vee$, $\wedge$ et $\equiv$, jusqu'ici exclus. La justification "PC" couvrira la règle et aussi toutes les règles utilisables dans le cadre propositionnel. On (ré)introduit aussi le quantificateur existentiel, en admettant que $\exists x \phi$ est une variante notationnelle de $\neg \forall x \neg \phi$. La règle PC est souvent implicitement couplée avec une version élémentaire de la règle de l'échange. Si l'équivalence $A \equiv B$ est une tautologie, on s'autorise à déduire immédiatement $U \vdash B$ de $U \vdash A$.

6.4.4 Quelques dérivations

Théorème. $\vdash p(a) \Rightarrow \exists x p(x)$

Démonstration.

1. $\vdash \forall x \neg p(x) \Rightarrow \neg p(a)$ (Axiome 4)
2. $\vdash p(a) \Rightarrow \neg \forall x \neg p(x)$ (PC, 1)
3. $\vdash p(a) \Rightarrow \exists x p(x)$ (Définition $\exists$)

Théorème. $\vdash (A \Rightarrow \forall x C(x)) \Rightarrow \forall x (A \Rightarrow C(x))$, si x n'a pas d'occurrence libre dans A .

Démonstration.

1. $A, A \Rightarrow \forall x C(x) \vdash \forall x C(x)$ (Hypothèse, MP)
2. $A, A \Rightarrow \forall x C(x) \vdash C(x)$ (Axiome 4, 1)
3. $A \Rightarrow \forall x C(x) \vdash (A \Rightarrow C(x))$ (Déduction, 2)
4. $A \Rightarrow \forall x C(x) \vdash \forall x (A \Rightarrow C(x))$ (Généralisation, 3)
5. $\vdash (A \Rightarrow \forall x C(x)) \Rightarrow \forall x (A \Rightarrow C(x))$ (Déduction, 4)

Remarque. Quelle(s) étape(s) de la démonstration serai(en)t illicite(s) si la restriction n'était pas respectée ?

Théorème. $\vdash \forall x (p(x) \Rightarrow q) \equiv (\exists x p(x) \Rightarrow q)$, si x n'a pas d'occurrence libre dans q .

Démonstration.

1. $\forall x (p(x) \Rightarrow q) \vdash \forall x (p(x) \Rightarrow q)$ (Hypothèse)
2. $\forall x (p(x) \Rightarrow q) \vdash \forall x (\neg q \Rightarrow \neg p(x))$ (PC, échange, 1)
3. $\forall x (p(x) \Rightarrow q) \vdash \neg q \Rightarrow \forall x \neg p(x)$ (Axiome 5, 2)
4. $\forall x (p(x) \Rightarrow q) \vdash \exists x p(x) \Rightarrow q$ (PC, $\exists$, 3)
5. $\exists x p(x) \Rightarrow q \vdash \exists x p(x) \Rightarrow q$ (Hypothèse)
6. $\exists x p(x) \Rightarrow q \vdash \neg q \Rightarrow \forall x \neg p(x)$ (PC, $\exists$, 5)
7. $\exists x p(x) \Rightarrow q \vdash \forall x (\neg q \Rightarrow \neg p(x))$ (Théorème, 6)
8. $\exists x p(x) \Rightarrow q \vdash \forall x (p(x) \Rightarrow q)$ (PC, 7)
9. $\vdash \forall x (p(x) \Rightarrow q) \equiv (\exists x p(x) \Rightarrow q)$ (Déduction, 4, 8)

Remarque. Quelle(s) étape(s) de la démonstration serai(en)t illicite(s) si la restriction n'était pas respectée ?

Il est permis d'utiliser une existentielle $\exists x p(x)$ en posant "soit x tel que $p(x)$ ", ou "soit a tel que $p(a)$ ". La *règle des constantes* formalise ce mode de raisonnement.

Théorème. (Règle C). Si $U \vdash \exists x p(x)$, si x n'a pas d'occurrence libre dans A et si on peut établir $U, p(x) \vdash A$ sans généralisation sur x , alors $U \vdash A$.

Démonstration.

1. $U, p(x) \vdash A$ (Hypothèse)
2. $U \vdash p(x) \Rightarrow A$ (Déduction, 1)
3. $U \vdash \forall x (p(x) \Rightarrow A)$ (Généralisation, 2)
4. $U \vdash \exists x p(x) \Rightarrow A$ (Théorème)
5. $U \vdash \exists x p(x)$ (Hypothèse)
6. $U \vdash A$ (PC, 4, 5)

Remarque. L'usage de la règle C est soumis à deux restrictions importantes, dont le non-respect conduit naturellement à des erreurs.

1. Le “blocage” de la généralisation sur x est nécessaire ; il permet l’application de la règle de déduction (ligne 2 de la démonstration). Négliger ce blocage permettrait de prouver par exemple $\exists x p(x) \vdash \forall x p(x)$. En posant $U = \{\exists x p(x)\}$ et $A = \forall x p(x)$, on aurait

1. $\exists x p(x) \vdash \exists x p(x)$ (Hypothèse)
2. $\exists x p(x), p(x) \vdash p(x)$ (Hypothèse)
3. $\exists x p(x), p(x) \vdash \forall x p(x)$ (Généralisation)
4. $\exists x p(x) \vdash \forall x p(x)$ (Règle C [usage incorrect])

2. L’absence d’occurrence libre de x dans A permet d’utiliser le (méta)théorème $\forall x(p(x) \Rightarrow A) \vdash (\exists x p(x) \Rightarrow A)$ à la ligne 4 de la démonstration. Négliger cette exigence permettrait aussi de prouver $\exists x p(x) \vdash \forall x p(x)$. En posant cette fois $U = \{\exists x p(x)\}$ et $A = p(x)$, on aurait

1. $\exists x p(x) \vdash \exists x p(x)$ (hypothèse)
2. $\exists x p(x), p(x) \vdash p(x)$ (hypothèse)
3. $\exists x p(x) \vdash p(x)$ (Règle C [usage incorrect])
4. $\exists x p(x) \vdash \forall x p(x)$ (Généralisation)

Remarque. Nous avons signalé qu’un moyen simple et radical d’éviter les risques de généralisation abusive était de proscrire toute variable libre à gauche du symbole $\vdash$. L’emploi de la règle C est une entorse temporaire à cette pratique. En particulier, même si ce n’est pas formellement requis, l’ensemble U des hypothèses devrait ne contenir que des formules fermées.

La règle C n’est pas indispensable mais elle a le mérite de rendre plus intuitives certaines preuves, et de formaliser une démarche fréquente en mathématique. Observons par exemple qu’une dérivation directe de

$$\exists x \forall y p(x, y) \Rightarrow \forall y \exists x p(x, y)$$

peut être laborieuse mais, d’après la règle C, il est suffisant de prouver

$$\forall y p(a, y) \Rightarrow \forall y \exists x p(x, y)$$

ou encore de prouver

$$p(a, y) \Rightarrow \exists x p(x, y)$$

ce qui est évident.

Remarque. Certains auteurs utilisent la règle “naturelle”

$$\frac{U \vdash \exists x p(x)}{U \vdash p(a)} \quad (a \text{ inédit})$$

Cette règle a un sens intuitif clair : on donne un nom (inédit) à un objet dont l’existence est prouvée. Si on accepte cette règle (ce que nous ne faisons pas), on doit nuancer le fait que, en l’absence de variables libres, $U \vdash A$ équivaut à $U \models A$. Dans le même ordre d’idée, on pourrait refuser (ce que nous ne faisons pas non plus) toute généralisation de variable libre présente dans une hypothèse, et en fait se restreindre aux hypothèses sans variable libre. Cela

bloquerait une déduction du type $p(x) \vdash \forall x p(x)$. En fait, plusieurs variantes existent pour le système de Hilbert, chacune ayant ses avantages et ses inconvénients.

Remarque. Une *théorie du premier ordre* est définie par une collection d’axiomes utilisant un lexique spécial, pouvant comporter des constantes individuelles. Par exemple, $\forall x [x * i(x) = e]$ est un axiome de la théorie des groupes, où e dénote l’élément neutre. La constante e ne sera jamais “inédite” au sens où ce mot est utilisé ici, même si elle n’a qu’une seule occurrence dans les hypothèses d’une dérivation. En particulier, la “dérivation”

1. $U \vdash \forall x [x * i(x) = e]$ (hypothèse)
2. $U \vdash \forall y \forall x [x * i(x) = y]$ (Généralisation 1)

est naturellement incorrecte ; elle illustre les dangers de la règle “naturelle” que nous venons d’évoquer.

6.4.5 Adéquation et complétude du système de Hilbert

La preuve d’adéquation consiste, comme d’habitude, à montrer que les axiomes sont des formules valides, et que les règles d’inférence préservent la validité. On a les résultats suivants.

Lemme. Soient A une formule, x une variable et t un terme tels que l’instantiation $[x/t]$ ne provoque pas de capture de variable dans A , alors $\forall x A \Rightarrow A[x/t]$ est une formule valide.

Lemme. Si A et B sont des formules et si x est une variable sans occurrence libre dans A , alors $\forall x (A \Rightarrow B(x)) \Rightarrow (A \Rightarrow \forall x B(x))$ est une formule valide.

Lemme. Si A est une formule valide, alors $\forall x A$ est une formule valide.

Théorème. Le système de Hilbert est adéquat.

Les démonstrations sont laissées au lecteur.

La technique de Kalmar, utilisée pour prouver la complétude du système de Hilbert dans le cadre propositionnel, n’est plus applicable dans le cadre prédicatif, puisque la notion de table de vérité n’existe plus. L’idée de base de cette technique était de montrer que toute conclusion, à laquelle la méthode des tables de vérité pouvait conduire, restait accessible à la méthode de Hilbert. Cette idée reste valable ici, il suffit de prendre un autre point de départ et de montrer, par exemple, que toute dérivation effectuée dans le système des séquents de Gentzen peut être simulée dans le système de Hilbert. Concrètement, cela implique de démontrer le lemme suivant :

Lemme. Tout usage des règles α , β , γ et δ dans le système de Gentzen peut être simulé dans le système de Hilbert.

Démonstration. Nous considérons successivement une règle α , une règle β , une règle γ et une règle δ . Les autres règles sont laissées au lecteur.

Lemme α . La règle de Gentzen

$$\frac{\rightarrow V, \neg A, B}{\rightarrow V, (A \Rightarrow B)}$$

est simulée dans le système de Hilbert par la règle

$$\frac{\vdash W \vee \neg A \vee B}{\vdash W \vee (A \Rightarrow B)}$$

Lemme β . La règle de Gentzen

$$\frac{\rightarrow V, A \quad \rightarrow V, \neg B}{\rightarrow V, \neg(A \Rightarrow B)}$$

est simulée dans le système de Hilbert par la règle

$$\frac{\vdash W \vee A \quad \vdash W \vee \neg B}{\vdash W \vee \neg(A \Rightarrow B)}$$

Lemme γ . La règle de Gentzen

$$\frac{\rightarrow V, \exists x A(x), A(c)}{\rightarrow V, \exists x A(x)}$$

est simulée dans le système de Hilbert par la règle

$$\frac{\vdash W \vee \exists x A(x) \vee A(c)}{\vdash W \vee \exists x A(x)}$$

Démonstration.

1. $\vdash \forall x \neg A(x) \Rightarrow \neg A(c)$ (Axiome 4)
2. $\vdash \neg \forall x \neg A(x) \vee \neg A(c)$ (PC 1)
3. $\vdash V \vee \neg \forall x \neg A(x) \vee \neg A(c)$ (PC 2)
4. $\vdash V \vee \exists x A(x) \vee \neg A(c)$ ($\exists$)
5. $\vdash V \vee \exists x A(x) \vee A(c)$ (Hypothèse)
6. $\vdash V \vee \exists x A(x)$ (PC 4, 5)

Lemme δ . La règle de Gentzen

$$\frac{\rightarrow V, A(a)}{\rightarrow V, \forall x A(x)}$$

est simulée dans le système de Hilbert par la règle

$$\frac{\vdash W \vee A(x)}{\vdash W \vee \forall x A(x)}$$

Démonstration.

1. $\vdash V \vee A(x)$ (Hypothèse)
2. $\vdash \neg V \Rightarrow A(x)$ (PC 1)
3. $\vdash \forall x (\neg V \Rightarrow A(x))$ (Généralisation 2)
4. $\vdash \neg V \Rightarrow \forall x A(x)$ (Axiome 5, PC 4)
5. $\vdash V \vee \forall x A(x)$ (PC 4)

Corollaire. Le système de Hilbert est complet.

Corollaire. Si U et A sont sans variables libres on a $U \vdash A$ si et seulement si $U \models A$.

6.4.6 Preuve indirecte du théorème de compacité

On doit prouver que tout ensemble inconsistant admet un sous-ensemble fini inconsistant ; on peut se limiter aux ensembles de formules fermées. Soit U un ensemble inconsistant de formules fermées, donc tel que $U \models A$ pour toute formule A , et en particulier pour $A = \neg(p \Rightarrow p)$. Vu la complétude du système de Hilbert, on a $U \vdash A$, donc il existe une dérivation dont la dernière ligne est $U \vdash A$. Cette dérivation est nécessairement finie (par définition) et ne peut donc évoquer qu'un nombre fini d'hypothèses. Ces hypothèses forment un sous-ensemble fini V de U , tel que $V \vdash A$. Le système de Hilbert étant adéquat, on a nécessairement $V \models A$, ce qui montre que V est inconsistant.

7 Logique prédicative avec fonctions

En mathématique, on rencontre souvent des formules du type

$$x > y \Rightarrow (x + 1) > (y + 1),$$

ou encore, en notation préfixée,

$$> (x, y) \Rightarrow > (+ (x, 1), + (y, 1)).$$

Ce sont des instances du schéma

$$p(x, y) \Rightarrow p(f(x, a), f(y, a)).$$

Il est naturel et utile de compléter le langage des prédicats par des *symboles fonctionnels* qui représenteront des *fonctions* sur le domaine d'interprétation.

On introduit donc

– $\mathcal{F} = \{f, g, h, \dots\}$: un ensemble de symboles arbitraires appelés *symboles fonctionnels* (chacun ayant une arité),

en plus des *symboles prédictatifs*, des *constantes* et des *variables*.

7.1 Syntaxe du calcul des prédicats

La syntaxe des *termes* est généralisée mais les règles syntaxiques définissant les formules sont inchangées. Le concept de *terme* est défini récursivement :

- Une *variable* est un terme.
- Une *constante* est un terme.
- Si f est un *symbole fonctionnel* (arité m) et si $t_1, t_2, \dots, t_m$ sont des *termes*, alors $f(t_1, \dots, t_m)$ est un terme.
- Rien d'autre n'est un terme.

Remarques. Les constantes sont des symboles fonctionnels d'arité 0. Un terme est *clos* s'il ne contient aucune variable.

Exemples de termes :

$$a \quad x \quad f(a, x) \quad g(f(a)) \quad f(g(x, h(y))) .$$

Exemples de formules atomiques :

$$p(a, b) \quad p(x, f(a, x)) \quad p(f(a, b), f(g(x), g(x))) .$$

7.2 Sémantique du calcul des prédicats

Une *interprétation* $\mathcal{I}$ est un triplet (D, I_c, I_v) tel que

- D est un ensemble non vide, appelé *domaine d'interprétation* ;
- I_c est une fonction qui associe
 - à toute *constante* a , un objet $I_c[a]$ appartenant à D ,
 - à tout *symbole fonctionnel* f d'arité m , une fonction $I_c[f]$ de D^m dans D ;
 - à tout *symbole prédicatif* p d'arité n , un *prédicat* d'arité n sur D , c'est-à-dire une fonction $I_c[p]$ de D^n dans $\{\mathbf{V}, \mathbf{F}\}$;
- I_v est une fonction qui associe à toute variable x un élément $I_v[x]$ de D .

Les règles d'interprétation permettent d'associer un objet de D à chaque terme et une valeur de vérité à chaque formule. Soit $\mathcal{I} = (D, I_c, I_v)$ une interprétation. On a

- Si x est une variable libre, alors $\mathcal{I}[x] = I_v[x]$.
- Si a est une constante, alors $\mathcal{I}[a] = I_c[a]$.
- Si f est un symbole fonctionnel d'arité m et si $t_1, t_2, \dots, t_m$ sont des termes, alors $\mathcal{I}[f(t_1, t_2, \dots, t_m)] = I_c[f](\mathcal{I}[t_1], \mathcal{I}[t_2], \dots, \mathcal{I}[t_m])$.

Les règles d'interprétation des formules sont inchangées.

Exemple : La formule

$$\forall x \forall y (p(x, y) \Rightarrow p(f(x, a), f(y, a)))$$

est satisfaite par l'interprétation

$$\mathcal{I}_1 = (\mathbb{Z}, I_c, I_v) : I_c[a] = 1, I_c[f] = +, I_c[p] = \leq ,$$

mais pas par l'interprétation

$$\mathcal{I}_2 = (\mathbb{Z}, I_c, I_v) : I_c[a] = -1, I_c[f] = *, I_c[p] = > .$$

7.3 Formes normales

Nous avons vu au paragraphe 3.6.1 l'intérêt de définir des formes normales, ou canoniques, pour les formules et les objets formels en général. Les notions de formes normales disjonctives et conjonctives se sont révélées fructueuses en logique propositionnelle et il en ira de même en logique prédicative.

On peut raisonnablement espérer que les notions propositionnelles continuent à s'étendre au cadre prédicatif, moyennant quelques restrictions et quelques complications liées notamment à la quantification. Par rapport à celle-ci, on peut envisager deux types de normalisation. D'une part, on peut s'efforcer de restreindre la quantification à des formules aussi "simples" que possible, et montrer que toute formule peut se ramener à une combinaison booléenne de formules quantifiées "simples". On peut d'autre part imposer que la portée des quantifications soit toujours maximale, et que tout atome soit dans la portée de toutes les quantifications présentes dans la formule. Il n'est pas évident a priori de savoir s'il est préférable de minimiser ou de maximiser la portée des quantifications. L'étude qui suit montre que les deux techniques ont leur utilité.

7.3.1 Lois de passage

Si on envisage de modifier la portée d'une quantification sans altérer la sémantique de la formule concernée, il faut connaître les relations existant entre les quantifications et les atomes, entre les quantifications et les connecteurs et entre les quantifications entre elles.

Si Φ est un atome ne comportant pas la variable x ou, plus généralement, une formule ne comportant pas d'occurrence libre de x , alors les trois formules Φ , $\forall x \Phi$ et $\exists x \Phi$ sont logiquement équivalentes. Concrètement, cela signifie que toute quantification portant sur une formule ne comportant pas d'occurrence libre de la variable quantifiée est superflue et donc, en pratique, supprimée. Une formule telle que $\forall x \forall y \exists x P(x, y)$ peut donc se simplifier en $\forall y \exists x P(x, y)$; la formule $\forall z A(x, y)$ se simplifie en $A(x, y)$ mais $\forall x A(x, y)$ ne se simplifie pas.

Les formules valides introduites au paragraphe 5.3.5 constituent un bon point de départ pour déterminer les relations entre quantificateurs et connecteurs. Un raisonnement sémantique direct permet de vérifier les équivalences logiques suivantes, concernant les connecteurs de négation, de conjonction et de disjonction :

$\forall x \neg \Phi$	$\leftrightarrow$	$\neg \exists x \Phi$	$\exists x \neg \Phi$	$\leftrightarrow$	$\neg \forall x \Phi$
$\forall x (\Phi \wedge \Psi)$	$\leftrightarrow$	$\forall x \Phi \wedge \forall x \Psi$	$\exists x (\Phi \vee \Psi)$	$\leftrightarrow$	$\exists x \Phi \vee \exists x \Psi$
$\forall x (\Phi \vee \Xi)$	$\leftrightarrow$	$\forall x \Phi \vee \Xi$	$\exists x (\Phi \wedge \Xi)$	$\leftrightarrow$	$\exists x \Phi \wedge \Xi$

Dans ce tableau,
 Φ et Ψ désignent des formules quelconques,
 Ξ désigne une formule sans occurrence libre de x

On dit parfois que la quantification universelle distribue la conjonction et que la quantification existentielle distribue la disjonction. Rappelons, comme cela a été mentionné au paragraphe 5.3.5, que la formule $\forall x (\Phi \vee \Psi)$ est en général strictement plus faible que la formule $\forall x \Phi \vee \forall x \Psi$, tandis que la formule $\exists x (\Phi \wedge \Psi)$ est en général strictement plus forte que la formule $\exists x \Phi \wedge \exists x \Psi$.⁶¹ Notons enfin que, l'implication étant de nature disjonctive, on a aussi

$$\exists x (\Phi \Rightarrow \Psi) \leftrightarrow \forall x \Phi \Rightarrow \exists x \Psi$$

7.3.2 Forme pure

Une formule Φ est *pure* si toute sous-formule de Φ se trouvant dans la portée d'une quantification sur une variable x comporte une occurrence libre de x . Les lois de passage permettent de réduire une formule quelconque à la forme pure, c'est-à-dire de construire une forme pure qui soit logiquement équivalente à la formule initiale.

⁶¹Tout nombre naturel est pair ou impair, mais tous les naturels ne sont pas pairs, et tous les naturels ne sont pas impairs ; de la même manière, il n'existe pas de nombre naturel simultanément pair et impair, mais il existe des naturels pairs et des naturels impairs.

7.3.3 Forme prénexe

Une formule est en *forme prénexe* si elle est de la forme

$$\underbrace{Q_1 x_1 \cdots Q_n x_n}_{\text{préfixe}} \underbrace{M}_{\text{matrice}}$$

où chaque Q_i désigne soit $\forall$, soit $\exists$, pour $i = 1, \dots, n$ et où la *matrice* M est une formule sans quantification. La portée du préfixe doit être la matrice tout entière.

Remarque. On peut supposer (sans restriction) que seules les variables apparaissant (libres) dans la matrice sont quantifiées dans le préfixe.

Théorème. Pour toute formule du calcul des prédicats, il existe (au moins) une formule en forme prénexe qui lui est équivalente.

7.3.4 Réduction à la forme prénexe

Exemple : $\forall x (p(x) \wedge \neg \exists y \forall x \neg (\neg q(x, y) \Rightarrow \forall z r(a, x, y)))$.

1. Eliminer tous les connecteurs autres que $\neg$, $\vee$, $\wedge$.
Ex. : $\forall x (p(x) \wedge \neg \exists y \forall x \neg (\neg \neg q(x, y) \vee \forall z r(a, x, y)))$.
2. Renommer des variables liées (si nécessaire) de manière à ce qu'aucune variable n'ait simultanément des occurrences libres et liées dans la formule ou une de ses sous-formules.
Ex. : $\forall x (p(x) \wedge \neg \exists y \forall u \neg (\neg \neg q(u, y) \vee \forall z r(a, u, y)))$.
3. Supprimer les quantifications dont la portée ne contient pas la variable quantifiée.
Ex. : $\forall x (p(x) \wedge \neg \exists y \forall u \neg (\neg \neg q(u, y) \vee r(a, u, y)))$.
4. Propager les occurrences de $\neg$ vers l'intérieur et éliminer les doubles négations.

$$\begin{aligned} \neg \forall x A &\rightarrow \exists x \neg A, \\ \neg \exists x A &\rightarrow \forall x \neg A, \\ \neg \neg C &\rightarrow C \end{aligned}$$

Ex. : $\forall x (p(x) \wedge \forall y \exists u (q(u, y) \vee r(a, u, y)))$.

5. Propager les quantifications vers l'extérieur.

$$\begin{aligned} \forall x A \wedge \forall x B &\rightarrow \forall x (A \wedge B) & \exists x A \vee \exists x B &\rightarrow \exists x (A \vee B) \\ \text{si } x \text{ n'a pas d'occurrence dans } B : & & & \\ \forall x A \wedge B &\rightarrow \forall x (A \wedge B) & \exists x A \vee B &\rightarrow \exists x (A \vee B) \\ \forall x A \vee B &\rightarrow \forall x (A \vee B) & \exists x A \wedge B &\rightarrow \exists x (A \wedge B) \end{aligned}$$

Renommer si nécessaire :

$$\exists x p(x) \wedge \forall x q(x) \rightarrow \exists x p(x) \wedge \forall y q(y) \rightarrow \exists x \forall y (p(x) \wedge q(y)).$$

Ex. : $\forall x \forall y \exists u (p(x) \wedge (q(u, y) \vee r(a, u, y)))$.

7.3.5 Forme de Skolem

Certaines définitions du cadre propositionnel restent valables dans le cadre prédicatif.

- Un *littéral* est un atome ou la négation d'un atome.
- Une *clause* (un *cube*) est une disjonction (une conjonction) de littéraux.
- Une *forme conjonctive (disjonctive) normale* est une conjonction (disjonction) de clauses (de cubes).
- Une forme prénexe est *conjonctive (disjonctive)* si sa matrice est en forme conjonctive (disjonctive) normale.

Une *forme de Skolem* est une forme prénexe sans quantifications existentielles. A toute forme prénexe, on associe une forme de Skolem au moyen de l'algorithme suivant.

Pour chaque quantification existentielle $\exists x$ se trouvant dans la portée de $k \geq 0$ quantifications universelles ($\forall x_1 \cdots \forall x_k$),

1. remplacer chaque occurrence de x dans la matrice par $f(x_1, \dots, x_k)$ où f est un *nouveau* symbole fonctionnel d'arité k ($k = 0$ n'est pas exclu).
2. supprimer la quantification $\exists x$.

Exemples :

- $\forall x \forall y \exists u (q(u, y) \Rightarrow r(a, u, y, z))$
se transforme en
 $\forall x \forall y (q(f(x, y), y) \Rightarrow r(a, f(x, y), y, z)).$
- $\forall x \exists u \forall v \exists w \forall x \forall y \exists z M(u, v, w, x, y, z)$
se simplifie en
 $\exists u \forall v \exists w \forall x \forall y \exists z M(u, v, w, x, y, z)$
qui se transforme en
 $\forall v \forall x \forall y M(a, v, f(v), x, y, g(v, x, y)).$

La formule $A =_{def} \forall x \forall y \exists u [q(u, y) \Rightarrow r(a, u, y, x)]$ affirme l'existence d'un certain u , dépendant de x et y , tel que la formule $q(u, y) \Rightarrow r(a, u, y, x)$ soit vraie. Le passage à la forme de Skolem associée S_A consiste simplement à *nommer* ce u ; le nom $f(x, y)$ rappelle la dépendance. Le symbole f représente une *fonction de choix*.

Remarque. Il n'est pas indispensable de passer par la forme prénexe pour obtenir une forme de Skolem. Il suffit de reconnaître les quantifications *sémantiquement* existentielles et de nommer l'objet dont l'existence est assertée. Par exemple,

$$\begin{aligned} &\exists z \forall x (p(x) \Rightarrow \neg \forall y [q(x, y) \Rightarrow p(z)]) \\ \text{devient} &\quad \forall x (p(x) \Rightarrow \neg [q(x, f(x)) \Rightarrow p(a)]) \end{aligned}$$

Motivation. Pour déterminer la consistance d'une formule quelconque φ , il suffira d'étudier la consistance de la forme de Skolem associée à une forme prénexe logiquement équivalente à φ . On rappelle aussi qu'une formule est consistante si et seulement si sa fermeture existentielle est consistante.

Théorème de Skolem. La forme de Skolem S_A associée à la forme prénexe A est consistante si et seulement si A est consistante.

La démonstration n'est pas difficile mais les notations sont lourdes. Un exemple simple suffira à illustrer les points essentiels : tout modèle de S_A est un modèle de A , et tout modèle de A s'étend en un modèle de S_A si on donne une interprétation adéquate aux symboles de Skolem.

Exemple. Soit $A : \forall x \exists y p(x, y)$ et $S_A : \forall x p(x, f(x))$.

On se donne d'abord un modèle de A , soit $\mathcal{I} = (\{1, 2\}, I_c, I_v)$, où $I_c[p]$ est vrai pour $(1, 1)$, $(1, 2)$ et $(2, 1)$, et faux pour $(2, 2)$.

On obtient un modèle $\mathcal{J} = (\{1, 2\}, J_c, J_v)$ de S_A en étendant I_c en J_c ; $J_c[f]$ appliquera naturellement 1 sur 1 ou 2 (au choix) et 2 sur 1 (obligatoirement).

La sémantique de $\exists$ garantit la possibilité de construire la fonction $J_c[f]$ (totale sur D).

Réciproquement, on peut obtenir $\mathcal{I}$ à partir de $\mathcal{J}$ en “oubliant” $J_c[f]$. La totalité de $J_c[f]$ garantit le respect de la sémantique de $\exists$.

On a $\models_{\mathcal{I}} A$, $\models_{\mathcal{J}} A$ et $\models_{\mathcal{J}} S_A$, mais pas $\models_{\mathcal{I}} S_A$ ($\mathcal{I}$ n'est pas une interprétation pour S_A).

Remarque. Peut-on dire qu'une forme prénexe A est valide si et seulement si la forme de Skolem associée S_A est valide? Peut-on dire que A et S_A sont logiquement équivalentes?

7.3.6 Forme clausale

Une formule est dite *en forme clausale* si elle est en forme de Skolem et si la matrice est en forme conjonctive normale.

Exemple de mise en forme clausale.

- $\exists x \forall y p(x, y) \Rightarrow \forall y \exists x p(x, y)$
- $\neg \exists x \forall y p(x, y) \vee \forall y \exists x p(x, y)$
- $\forall x \exists y \neg p(x, y) \vee \forall y \exists x p(x, y)$
- $\forall x \exists y \neg p(x, y) \vee \forall w \exists z p(z, w)$
- $\forall x \exists y \forall w \exists z (\neg p(x, y) \vee p(z, w))$
- $\forall x \forall w (\neg p(x, f(x)) \vee p(g(x, w), w))$

Remarque. Lorsqu'une formule est en forme de Skolem, on omet souvent le préfixe (si l'on peut distinguer les constantes des variables). Dans ce cas, une formule en forme clausale est simplement un ensemble de clauses; la formule

$$\forall x \forall y \forall z [(p(x, y) \vee \neg q(a)) \wedge (q(x) \vee \neg r(b, z))]$$

est représentée par l'ensemble

$$\{ p(x, y) \vee \neg q(a), q(x) \vee \neg r(b, z) \}$$

équivalent à l'ensemble

$$\{ p(x, y) \vee \neg q(a), q(u) \vee \neg r(b, v) \}.$$

7.4 Théorie de Herbrand

Idée : Définir un ensemble d'interprétations canoniques telles que si une formule φ en forme de Skolem est consistante, alors elle a au moins un modèle canonique.

Les interprétations canoniques, dites *de Herbrand*, sont basées sur un domaine particulier, le *domaine de Herbrand* ou l'*univers de Herbrand*. Tout terme clos (construit avec le lexique de φ) doit être interprété en une valeur d'un domaine D . Le domaine générique sera simplement l'ensemble des termes clos.

7.4.1 Domaines de Herbrand

Soit S une forme de Skolem dont les constantes et les symboles fonctionnels forment les ensembles $\mathcal{A}$ et $\mathcal{F}$. Le *domaine de Herbrand* H_S (ou *univers de Herbrand*) de S est défini récursivement de la manière suivante.

- Si $a \in \mathcal{A}$, alors $a \in H_S$. (Si $\mathcal{A} = \emptyset$, créer une constante arbitraire $a \in H_S$; un domaine ne peut être vide.)
- Si $f \in \mathcal{F}$ (f d'arité m) et $t_1, \dots, t_m \in H_S$, alors $f(t_1, \dots, t_m) \in H_S$.

Les éléments du domaine de Herbrand sont des objets syntaxiques, sans signification particulière : ce sont tous les termes clos que l'on peut construire à l'aide de $\mathcal{A}$ et $\mathcal{F}$.

Exemples d'univers de Herbrand associés à une matrice clausale.

- Pour $S_1 = (p(a) \vee \neg p(b) \vee q(z)) \wedge (\neg q(z) \vee \neg p(b) \vee q(z))$, on a $H_{S_1} = \{a, b\}$.
- Pour $S_2 = (\neg p(x, f(y))) \wedge (p(w, g(w)))$, on a $H_{S_2} = \{a, f(a), g(a), f(f(a)), g(f(a)), f(g(a)), g(g(a)), f(f(f(a))), g(f(f(a))), f(g(f(a))), g(g(f(a))), \dots\}$.
- Pour $S_3 = \neg p(a, f(x, y)) \vee p(b, f(x, y))$, on a $H_{S_3} = \{a, b, f(a, a), f(a, b), f(b, a), f(b, b), f(f(a, a), a), f(f(a, a), b), \dots, f(f(b, a), f(a, b)), \dots\}$.
- Pour $S_4 = (p(x) \vee q(x)) \wedge \neg q(x)$, on a $H_{S_4} = \{a\}$.

7.4.2 Interprétations, bases et modèles de Herbrand

Une *interprétation de Herbrand* d'une formule en forme de Skolem S est une interprétation $\mathcal{H}$ de S qui satisfait les conditions suivantes.

- Le domaine d'interprétation de $\mathcal{H}$ est le domaine de Herbrand H_S .
- Pour toute constante a dans S : $\mathcal{H}[a] = a$.
- Pour tout symbole fonctionnel f (arité m) dans S et pour tous termes $t_1, \dots, t_m$: $\mathcal{H}(f(t_1, \dots, t_m)) = f(\mathcal{H}[t_1], \dots, \mathcal{H}[t_m])$.

Remarques. L'interprétation des symboles prédicatifs et des variables est libre.

Si t est un terme clos, on a $\mathcal{H}[t] = t$.

Si $f \in \mathcal{F}$, $H_c[f]$ est la fonction de H^m dans H qui applique le m -uplet $(t_1, \dots, t_m)$ de termes clos sur le terme clos $f(t_1, \dots, t_m)$.

Un terme (un atome, un littéral, une clause, une matrice de forme de Skolem) est dit *clos* ou *complètement instancié* s'il ne contient aucune variable. Les éléments de l'univers de Herbrand H_S sont des termes clos. Ces termes et les atomes, littéraux et clauses qu'ils permettent de construire (au moyen des symboles prédicatifs de S) sont dits *fondamentaux*. La *base de Herbrand* B_S est l'ensemble des atomes fondamentaux. Une interprétation de Herbrand attribue une valeur de vérité à tout élément de B_S . Un *modèle de Herbrand* d'une formule en forme de Skolem S est une interprétation de Herbrand qui satisfait S (qui rend S vrai).

7.4.3 Simplification de Herbrand

Soit $\varphi =_{\text{def}} \forall x [p(x) \Rightarrow p(f(x))]$. Soit $\mathcal{I}$ l'interprétation de domaine $\mathbb{N}$, telle que $I_c[f](n) = 4 * n$ et $I_c[p](n) = \mathbf{V}$ si n est un carré. On voit immédiatement que $\mathcal{I}[\varphi] = \mathbf{V}$. Informellement, on le justifie en notant que l'on a $\mathcal{I}[p(n) \Rightarrow p(f(n))] = \mathbf{V}$, ou $p(n) \Rightarrow p(4*n)$, pour tout $n \in \mathbb{N}$. Cette écriture est abusive, parce qu'elle mêle des objets syntaxiques (p, f, x) et des objets sémantiques ($0, 4, n, *$). L'écriture correcte est $\mathcal{I}_{x/n}[p(x) \Rightarrow p(f(x))] = \mathbf{V}$.

Plus généralement, $\mathcal{I}_{x/d}[A(x)]$ ou $\mathcal{I}_{x/d}[B]$ est correct (si d est un élément du domaine d'interprétation), tandis que $\mathcal{I}[A(d)]$ ou $\mathcal{I}[B(x/d)]$ est abusif.

Soit alors $\mathcal{H}$ une interprétation de Herbrand. Le domaine est $H = \{a, f(a), f(f(a)), \dots\}$. On a $\mathcal{H}[\varphi] = \mathbf{V}$ si et seulement si $\mathcal{H}_{x/h}[p(x) \Rightarrow p(f(x))] = \mathbf{V}$ pour tout $h \in H$. On observe qu'ici, la notation simplifiée n'est plus abusive : on a

$$\mathcal{H}_{x/h}[p(x) \Rightarrow p(f(x))] = \mathcal{H}[p(h) \Rightarrow p(f(h))]$$

puisque l'objet h a le double statut syntaxique et sémantique.

Théorème. Si $\mathcal{H}$ est une interprétation de Herbrand pour la matrice $A(x_1, \dots, x_n)$, on a $\mathcal{H}[\forall x_1 \dots \forall x_n A(x_1, \dots, x_n)] = \mathbf{V}$ si et seulement si $\mathcal{H}[A(h_1, \dots, h_n)] = \mathbf{V}$ pour tous $h_1, \dots, h_n \in H$.

Définition. Les formules $A(h, h')$ sont les *instances fondamentales* de la matrice $A(x, y)$, ou de la forme de Skolem correspondante.

Corollaire. Une forme de Skolem est vraie pour une interprétation de Herbrand si et seulement si toutes ses instances fondamentales sont vraies pour cette interprétation.

Simplification. Les interprétations de Herbrand s'identifient aux fonctions (totales) de la base de Herbrand B_H sur $\{\mathbf{V}, \mathbf{F}\}$, ou encore aux sous-ensembles de B_H , c'est-à-dire aux interprétations propositionnelles de lexique $\Pi = B_H$.

Exemples d'interprétations de Herbrand,

Pour $S_3 = \neg p(a, f(x, y)) \vee p(b, f(x, y))$, on a

$$H_{S_3} = \{a, b, f(a, a), f(a, b), f(b, a), f(b, b), \dots, f(f(a, b), f(a, a)), \dots, f(f(f(a, a), a), f(a, b)), \dots\}.$$

$$B_{S_3} = \{p(a, a), p(a, b), p(b, a), p(b, b), \dots, p(f(a, b), f(a, a)), \dots, p(f(f(a, a), a), f(a, b)), \dots\}.$$

Une interprétation de Herbrand $\mathcal{I}$ de S_3 est (définie par) un sous-ensemble $\mathcal{I}$ (fini ou non) de B_{S_3} ; on identifie donc $A \in \mathcal{I}$ et $\mathcal{I}[A] = \mathbf{V}$, pour tout atome fondamental A .

Remarque. La théorie de Herbrand permet de réduire quasiment le calcul des prédicats au calcul des propositions. La seule différence (importante !) est que le lexique "propositionnel" (la base de Herbrand) est généralement infini.

7.4.4 Théorèmes de Herbrand

Premier théorème de Herbrand. Une formule en forme de Skolem S est consistante si et seulement si elle admet un modèle de Herbrand.

Démonstration. La condition est visiblement suffisante. On montre qu'elle est nécessaire en donnant une technique de transformation d'un modèle quelconque $\mathcal{I}$ (de domaine D quelconque) en un modèle de Herbrand $\mathcal{H}$ (de domaine $H = H_S$).

1. On commence par donner une fonction w qui à tout élément $h \in H$ du domaine de Herbrand H associe un élément $w(h) \in D$.

(a) Si au moins une constante apparaît dans S , toutes ces constantes sont interprétées par $\mathcal{I}$ et on pose $w(c_i) = \mathcal{I}[c_i] = I_c[c_i] \in D$; sinon, la constante arbitraire a est interprétée en un élément $d = w(a) \in D$ quelconque.

(b) Pour tout terme composé $h = f(h_1, \dots, h_m) \in H$, on pose

$$w(h) = I_c[f](w(h_1), \dots, w(h_m)) \in D. \text{ (Cette expression est simplement } \mathcal{I}[h], \text{ sauf si on a ajouté une constante arbitraire.)}$$

2. Pour donner une interprétation de Herbrand $\mathcal{H}$, il faut spécifier l'ensemble des atomes fondamentaux qui seront vrais dans $\mathcal{H}$.

Soient $h_1, \dots, h_n \in H$ et p un symbole prédicatif d'arité n . Pour interpréter l'atome fondamental $p(h_1, \dots, h_n)$, on pose $\mathcal{H}[p(h_1, \dots, h_n)] = I_c[p](w(h_1), \dots, w(h_n))$ ($I_c[p]$ est une fonction de D^n dans $\{\mathbf{V}, \mathbf{F}\}$).

On a donc

$$\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[p(x_1, \dots, x_n)] = \mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[p(x_1, \dots, x_n)]$$

3. Soit $\varphi(x_1, \dots, x_n)$ une matrice ne contenant aucune variable libre autre que $x_1, \dots, x_n$. On a $\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[\varphi(x_1, \dots, x_n)] = \mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[\varphi(x_1, \dots, x_n)]$

4. Toute formule de la forme $\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)$ satisfaite par $\mathcal{I}$ est aussi satisfaite par $\mathcal{H}$. On a successivement

$$\mathcal{I}[\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)] = \mathbf{V} \text{ (hypothèse),}$$

$$\mathcal{I}_{x_1/d_1, \dots, x_n/d_n}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } d_1, \dots, d_n \in D,$$

$$\mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } h_1, \dots, h_n \in H.$$

$$\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } h_1, \dots, h_n \in H.$$

$$\mathcal{H}[\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)] = \mathbf{V}.$$

Remarque. La théorie de Herbrand s'applique seulement aux formes de Skolem. Par exemple, la formule

$$p(a) \wedge \exists x \neg p(x)$$

est consistante, mais n'a pas de modèle de Herbrand : l'univers de Herbrand (si on le considère comme défini) serait le singleton $\{a\}$, et la formule n'admet que des modèles à deux éléments au moins. En revanche, la forme de Skolem correspondante

$$p(a) \wedge \neg p(b)$$

admet le modèle de Herbrand $\{p(a)\}$, tel que $p(a) = \mathbf{V}$ et $p(b) = \mathbf{F}$.

Remarque. Dans le cas particulier où la matrice d'une forme de Skolem est une conjonction, on peut tenir compte de la relation existant entre $\forall$ et $\wedge$. Par exemple, les formules $\forall x \forall y [\varphi(x) \wedge \psi(y)]$, $\forall x \varphi(x) \wedge \forall y \psi(y)$ et $\forall x [\varphi(x) \wedge \psi(x)]$ sont équivalentes et représentables indifféremment par $\{\varphi(x), \psi(y)\}$ ou $\{\varphi(x), \psi(x)\}$.

Second théorème de Herbrand. Une formule S en forme de Skolem est inconsistante si et seulement s'il existe une conjonction finie inconsistante d'instances fondamentales de sa matrice M .

Démonstration.

- La formule S est consistante si et seulement si elle admet un modèle de Herbrand.
- L'interprétation de Herbrand $\mathcal{H}$ est un modèle de S si et seulement si $\mathcal{H}'[M] = \mathbf{V}$ pour toute variante $\mathcal{H}'$ de $\mathcal{H}$ attribuant aux variables de M des valeurs quelconques prise dans l'univers de Herbrand.
- On a $\mathcal{H}'[M] = \mathcal{H}[M']$, où M' est l'instance fondamentale de M obtenue en remplaçant dans M les variables par les valeurs qui leur sont attribuées par $\mathcal{H}'$.

En conclusion, S est (in)consistant si et seulement si l'ensemble de ses instances fondamentales est (in)consistant. Vu le théorème de compacité (logique des propositions), S est inconsistant si et seulement s'il existe un ensemble fini inconsistant d'instances fondamentales de M .

Corollaire. Une formule en forme clausale S est inconsistante si et seulement s'il existe une conjonction finie inconsistante de clauses fondamentales.

Remarque. Soit $\forall x \forall y \forall z [C_1(x, y) \wedge C_2(y, z)]$ une forme clausale et H son domaine de Herbrand. L'ensemble des instances fondamentales de la matrice est

$$\{[C_1(h, h') \wedge C_2(h', h'')] : h, h', h'' \in H\};$$

il est logiquement équivalent à

$$\{C_1(h, h') : h, h' \in H\} \cup \{C_2(h', h'') : h', h'' \in H\},$$

lui-même équivalent à

$$\{C_1(h, h'), C_2(h, h') : h, h' \in H\},$$

qui est l'ensemble des *clauses* fondamentales.

7.4.5 Analyse de formes clausales

Premier exemple.

Soit $A = \neg (\forall x (p(x) \Rightarrow q(x)) \Rightarrow (\forall x p(x) \Rightarrow \forall x q(x)))$.

On a $S_A = \forall x [(\neg p(x) \vee q(x)) \wedge p(x) \wedge \neg q(a)]$.

La forme clausale est $\{\neg p(x) \vee q(x), p(x), \neg q(a)\}$.

Le domaine de Herbrand est $\{a\}$ et l'ensemble des clauses fondamentales est

$$\{\neg p(a) \vee q(a), p(a), \neg q(a)\};$$

cet ensemble est inconsistant donc la formule A est inconsistante.

Deuxième exemple.

Soit $S = \{p(f(x), a) \vee p(y, g(a)), \neg p(f(f(a)), z)\}$.

Le domaine de Herbrand et l'ensemble des clauses fondamentales sont infinis. Trois clauses fondamentales intéressantes sont

$$C_1 =_{\text{def}} p(f(f(a)), a) \vee p(f(f(a)), g(a)),$$

$$C_2 =_{\text{def}} \neg p(f(f(a)), a),$$

$$C_3 =_{\text{def}} \neg p(f(f(a)), g(a)).$$

$\{C_1, C_2, C_3\}$ est inconsistant, donc S est inconsistant.

Rappel. Attention aux quantifications implicites. Les deux formules éléments de S sont en fait $\forall x \forall y [p(f(x), a) \vee p(y, g(a))]$ et $\forall z \neg p(f(f(a)), z)$.

Semi-procédure de décision. Le théorème de Herbrand suggère une semi-procédure de décision pour la validité des formules du calcul des prédicats :

1. Considérer la négation de la formule donnée.
2. La mettre en forme clausale.
3. Générer un ensemble fini de clauses fondamentales.
4. Vérifier si cet ensemble de clauses fondamentales est inconsistant.

Les points 1 et 2 sont triviaux, même si la mise en forme normale est une procédure parfois longue et fastidieuse. Le point 4 se fait dans le cadre propositionnel ; les atomes fondamentaux se traitent en effet comme des atomes de logique des propositions. Seul le point 3 pose un réel problème ; produire des instances fondamentales est facile, produire "les bonnes" est plus délicat.

7.4.6 Analyse de règles d'inférence

Premier exemple.

Soient

$$H_1 : \forall x (p(x) \Rightarrow q(x)),$$

$$H_2 : \forall x (q(x) \Rightarrow r(x)),$$

$$C : \forall x (p(x) \Rightarrow r(x)).$$

On voudrait montrer que

$$\frac{H_1, H_2}{C}$$

est une règle correcte, c'est-à-dire que

$$H_1, H_2 \models C,$$

ou encore que

$$A =_{\text{def}} H_1 \wedge H_2 \wedge \neg C$$

est une formule inconsistante.

Transformons A en forme clausale :

$$\forall x ((p(x) \Rightarrow q(x)) \wedge (q(x) \Rightarrow r(x))) \wedge \exists y (p(y) \wedge \neg r(y))$$

$$\exists y \forall x ((p(x) \Rightarrow q(x)) \wedge (q(x) \Rightarrow r(x)) \wedge (p(y) \wedge \neg r(y)))$$

$$\exists y \forall x ((\neg p(x) \vee q(x)) \wedge (\neg q(x) \vee r(x)) \wedge p(y) \wedge \neg r(y))$$

$$\forall x ((\neg p(x) \vee q(x)) \wedge (\neg q(x) \vee r(x)) \wedge p(c) \wedge \neg r(c))$$

Le domaine de Herbrand est le singleton $\{c\}$. Les quatre clauses fondamentales sont

$$\neg p(c) \vee q(c), \neg q(c) \vee r(c), p(c) \text{ et } \neg r(c).$$

Ces clauses forment un ensemble inconsistant, donc la formule A est inconsistante et la règle est correcte.

Remarque. La règle reste correcte si $p(x)$, $q(x)$ et $r(x)$ sont des formules quelconques dont x est l'unique variable libre.

Deuxième exemple.

Soient

$$\begin{aligned} H_1 &: p(a), \\ H_2 &: \forall x (p(x) \Rightarrow p(f(x))), \\ C' &: \forall x p(x). \end{aligned}$$

On voudrait montrer que

$$\frac{H_1, H_2}{C}$$

est une règle correcte, c'est-à-dire que

$$A =_{\text{def}} H_1 \wedge H_2 \wedge \neg C$$

est une formule inconsistante.

Transformons A en forme clausale :

$$\begin{aligned} &p(a) \wedge \forall x (p(x) \Rightarrow p(f(x))) \wedge \neg \forall x p(x), \\ &p(a) \wedge \forall x (\neg p(x) \vee p(f(x))) \wedge \exists x \neg p(x), \\ &p(a) \wedge \forall x (\neg p(x) \vee p(f(x))) \wedge \neg p(b). \end{aligned}$$

Le domaine de Herbrand est $H = \{a, b, f(a), f(b), \dots\} = \{f^n(a), f^n(b) : n \in \mathbb{N}\}$.

La base de Herbrand est $B = \{p(f^n(a)), p(f^n(b)) : n \in \mathbb{N}\}$.

Une interprétation de Herbrand intéressante est $\mathcal{H} = \{p(f^n(a)) : n \in \mathbb{N}\}$.

Cette interprétation rend vraies les clauses $p(a)$ et $\neg p(b)$, ainsi que toutes les instances fondamentales de $\neg p(x) \vee p(f(x))$. C'est donc un modèle de l'ensemble des clauses fondamentales; cela montre que A est une formule consistante, et aussi que la règle est *incorrecte*.

7.5 Résolution fondamentale

La résolution fondamentale est simplement la méthode de résolution vue dans le cadre propositionnel, appliquée à des ensembles de clauses fondamentales.

Règle de résolution fondamentale.

Soit S un ensemble de clauses fondamentales et soient $C_1 = (C'_1 \vee \ell)$ et $C_2 = (C'_2 \vee \neg \ell)$ deux clauses fondamentales de S . La règle est

$$S \vdash \text{res}(C_1, C_2) = C'_1 \vee C'_2.$$

La clause $\text{res}(C_1, C_2)$ est la *résolvante* des clauses C_1 et C_2 .

Tant que $\square \notin S$, répéter :
 choisir $C_1 = (C'_1 \vee \ell)$, $C_2 = (C'_2 \vee \neg \ell) \in S$
 redéfinir $S := S \cup \{\text{res}(C_1, C_2)\}$

FIG. 63 – Procédure de résolution fondamentale

7.5.1 Procédure de résolution fondamentale

L'algorithme d'analyse d'un ensemble S de clauses fondamentales est donné à la figure 63. Soit

$$S = \{\neg p(c) \vee q(c), \neg q(c) \vee r(c), p(c), \neg r(c)\}.$$

Voici une réfutation de S par la méthode de résolution :

1. $\neg p(c) \vee q(c)$ (clause de S)
2. $p(c)$ (clause de S)
3. $q(c)$ (résolution 1, 2)
4. $\neg q(c) \vee r(c)$ (clause de S)
5. $r(c)$ (résolution 3, 4)
6. $\neg r(c)$ (clause de S)
7. $\square$ (résolution 5, 6)

L'inconvénient de la méthode de résolution fondamentale est qu'une formule prédictive donne souvent lieu à un ensemble infini de clauses fondamentales. L'obtention d'une réfutation passe par la détermination d'un sous-ensemble fini inconsistant. Diverses techniques sont développées pour permettre une meilleure mécanisation de l'analyse d'une formule prédictive par la méthode de résolution. La principale d'entre elles est abordée au cours de *Représentation de la connaissance*. Elle est à la base de la programmation logique et du langage PROLOG. Elle permet d'apporter une solution simple à de nombreux problèmes informatiques, en particulier dans le domaine de l'intelligence artificielle.

7.5.2 Preuve du théorème de compacité

Il faut montrer que tout ensemble inconsistant U admet un sous-ensemble fini inconsistant. On peut sans restriction supposer que les éléments de U sont des formes clausales. Si U est inconsistant, il est possible de déduire la clause vide par résolution fondamentale. Les clauses fondamentales utilisées sont en nombre fini, et donc sont des instances d'un nombre fini d'éléments de U ; ces éléments forment un sous-ensemble fini inconsistant de U .

8 Logiques prédictives décidables

Le calcul des prédicats est indécidable, mais admet des fragments intéressants décidables. Nous en considérons ici deux exemples, la logique des prédicats monadiques et la logique de Bernays et Schönfinkel.

8.1 Calcul des prédicats monadiques

Cette logique est le plus connu des fragments décidables de la logique prédictive, car elle généralise la théorie classique du syllogisme catégorique.

8.1.1 Brève introduction à la théorie du syllogisme catégorique

Pendant des siècles, l'enseignement de la logique s'est limité à la théorie du syllogisme catégorique, inventée par Aristote et systématisée par les Scolastiques, des logiciens du Moyen-Age.

La formule de base et ses variantes Etant donnés deux prédicats unaires⁶² P et Q (dans cet ordre), on appelle *formule de base* la formule $\forall x (P(x) \Rightarrow Q(x))$. Les *variantes* s'obtiennent en introduisant des négations, portant sur le conséquent de l'implication ou sur toute la formule. On a donc quatre possibilités, souvent identifiées par les quatre lettres **A**, **E**, **I** et **O** :⁶³

$\forall x (P(x) \Rightarrow Q(x))$	A universelle affirmative	$\neg \exists x (P(x) \wedge \neg Q(x))$
$\forall x (P(x) \Rightarrow \neg Q(x))$	E universelle négative	$\neg \exists x (P(x) \wedge Q(x))$
$\neg \forall x (P(x) \Rightarrow \neg Q(x))$	I particulière affirmative	$\exists x (P(x) \wedge Q(x))$
$\neg \forall x (P(x) \Rightarrow Q(x))$	O particulière négative	$\exists x (P(x) \wedge \neg Q(x))$

Ces quatre formules sont dites "de type PQ ". Une formule dont le type est PQ ou QP est une $\{P, Q\}$ -formule ; il y a donc huit $\{P, Q\}$ -formules.

Le syllogisme catégorique. Le "jeu du syllogisme catégorique" consiste à choisir une $\{P, Q\}$ -formule, une $\{Q, R\}$ -formule et une $\{P, R\}$ -formule, puis à déterminer si la troisième formule (la *conclusion*) est conséquence logique des deux premières (les *prémisses*). Il y a donc $8^3 = 512$ possibilités. Dans la mesure où les rôles de P et de R sont interchangeables, on peut imposer que la conclusion soit de type PR (et non de type RP), ce qui élimine la moitié des possibilités. On appelle alors *mineure* la $\{P, Q\}$ -prémisse, et *majeure* la $\{Q, R\}$ -prémisse ; les

⁶²ou monadiques, c'est-à-dire d'arité 1.

⁶³Pour "AffIrmo" et "nEgO".

termes $P(x)$, $Q(x)$ et $R(x)$ sont dits respectivement *mineur*, *moyen* et *majeur*.⁶⁴ Le syllogisme catégorique est la règle d'inférence

$$\frac{\text{majeure} \quad \text{mineure}}{\text{conclusion}}$$

Etant donnés les trois prédicats P , Q et R , il existe donc 256 syllogismes catégoriques. Le problème est de déterminer lesquels sont valides, c'est-à-dire tels que la conclusion soit conséquence logique des prémisses. Une grande partie des raisonnements courants (y compris beaucoup de raisonnements incorrects) se formalisent naturellement en des enchaînements de syllogismes, ce qui justifie l'intérêt particulier que cette notion a suscité dans le passé. Le point de vue moderne accorde peu d'importance à la théorie du syllogisme, simple fragment de la logique monadique ; la raison essentielle en est que, le nombre de syllogismes étant fini, leur étude est triviale : il suffit de les passer en revue un à un et de tester, pour chacun d'eux, sa validité. Cela se fait aisément, par exemple en utilisant la méthode des tableaux sémantiques ou celles de Herbrand.⁶⁵

Esquisse de l'approche classique du problème. Une approche systématique consiste à organiser les 256 possibilités selon divers critères. Classiquement, on utilise les notions de *figure* et de *mode*. La figure est fonction du type des prémisses et, plus précisément, de l'ordre des termes dans les prémisses. Il y a donc quatre figures, reprises dans le tableau suivant :

Figure :	première	deuxième	troisième	quatrième
Majeure	QR	RQ	QR	RQ
Mineure	PQ	PQ	QP	QP
Conclusion	PR	PR	PR	PR

Le mode d'un syllogisme est déterminé par la nature des prémisses et de la conclusion. Par exemple, le mode AEI désigne le cas où la *majeure* est universelle affirmative (A), la *mineure* est universelle négative (E) et la *conclusion* est particulière affirmative (I). Il y a donc $4^3 = 64$ modes possibles, chacun pouvant exister dans les quatre figures. Cependant, on peut vérifier que 12 modes seulement peuvent donner lieu à des syllogismes valides. Ce sont

AAA, AAI, AEE, AEO, AII, AOO, EAE, EAO, EIO, IAI, IEO, OAO.

Les Anciens avaient synthétisé ce résultat en cinq règles mnémotechniques :

1. Si les deux prémisses sont négatives, le syllogisme n'est pas valide.
2. Si les deux prémisses sont particulières, le syllogisme n'est pas valide.
3. Si une prémisse est particulière, la conclusion doit être particulière.

⁶⁴La conclusion comporte le mineur puis le majeur. La prémisse mineure comporte le mineur et le moyen, la prémisse majeure comporte le majeur et le moyen ; dans les prémisses, l'ordre des deux termes n'est pas imposé. Notons aussi l'emploi classique du mot "terme" ; dans la terminologie moderne, $P(x)$, $Q(x)$ et $R(x)$ sont en fait des atomes, ou formules atomiques.

⁶⁵Le lecteur est invité à construire quelques uns des 256 syllogismes et à tester leur validité par l'une ou l'autre méthode.

4. Si une prémisses est négative, la conclusion doit être négative.
5. Si les deux prémisses sont affirmatives, la conclusion doit être affirmative.

Ces règles, dont l'exactitude avait été reconnue empiriquement, permettent de rejeter les 52 modes "stériles". Par exemple, la première des règles permet d'éliminer des modes tels que *EEE* et *OEO* ; la deuxième permet d'éliminer *III* (entre autres) ; la troisième provoque notamment le rejet de *AIA* ; des modes tels que *AOI* et *AIO* contreviennent respectivement aux quatrième et cinquième règles.

Il ne reste donc que $12 \times 4 = 48$ syllogismes potentiellement valides ; parmi ceux-ci, nous allons voir que 15 syllogismes seulement sont valides.

Les diagrammes de Venn. C'est par une démarche informelle qu'Aristote et les Scolastiques ont déterminé quels syllogismes étaient valides et lesquels ne l'étaient pas. Les syllogismes valides ont reçu des noms conventionnels, dont les voyelles rappellent le mode. Par exemple, le raisonnement classique "Tous les humains sont mortels, tous les Grecs sont des humains, donc tous les Grecs sont mortels" est un syllogisme dont on détecte aisément la structure :⁶⁶

- (M) *Tous les humains sont mortels.*
- (m) *Tous les Grecs sont des humains.*
- (C) *Tous les Grecs sont mortels.*

Ce syllogisme appartient à la première figure et au mode *AAA* ; cette combinaison se note *AAA-1* ou, de manière plus classique (et plus poétique), *BARBARA*. Les trois "A" rappellent que les prémisses et la conclusion sont toutes trois des universelles affirmatives. De même, le syllogisme

- (M) *Tous les étudiants sont intelligents.*
- (m) *Certains humains ne sont pas intelligents.*
- (C) *Certains humains ne sont pas des étudiants.*

appartient à la deuxième figure ; c'est un exemple de *AOO-2* ou *BAROCO*, la majeure étant universelle affirmative, la mineure et la conclusion étant particulières négatives.

Un moyen simple et concret d'appréhender la validité d'un syllogisme consiste à utiliser un diagramme de Venn à trois composants (figure 64). Le cercle de gauche représente le mineur $P(x)$, celui de droite le majeur $R(x)$, le cercle du haut correspondant au moyen $Q(x)$. Ces cercles déterminent huit zones numérotées de 0 à 7. Tout objet appartient à l'une de ces zones, selon la valeur de vérité qu'il attribue au mineur, au majeur et au moyen. Par exemple, la zone 4 est intérieure aux cercles mineur et moyen, mais extérieure au cercle majeur ; elle regroupe donc les objets rendant vrais le mineur et le moyen, mais faux le majeur.

Les prémisses et conclusions des syllogismes correspondent à des assertions de vacuité ou de non-vacuité de certaines zones ; un syllogisme sera valide si son "interprétation graphique" est correcte. Nous illustrons cette technique de vérification par quelques exemples et contre-exemples.

Le syllogisme *AAA-1* que nous venons d'évoquer s'analyse aisément. La prémisses majeure $\forall x (Q(x) \Rightarrow R(x))$ signifie que tout objet vérifiant le moyen vérifie aussi le majeur, donc que les zones 1 et 4 sont vides, ce que nous notons $1 \cup 4 = \emptyset$; de même, la prémisses mineure

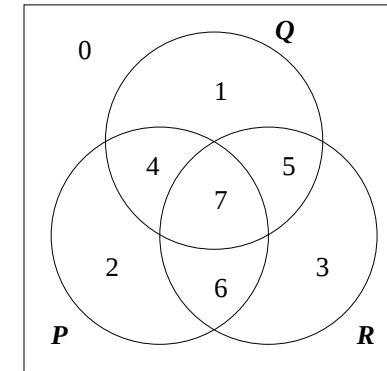


FIG. 64 – Diagramme de Venn

$\forall x (P(x) \Rightarrow Q(x))$ devient $2 \cup 6 = \emptyset$. A la conclusion $\forall x (P(x) \Rightarrow R(x))$ correspond l'assertion $2 \cup 4 = \emptyset$. Il est clair que, si les zones 1, 4, 2 et 6 sont vides, alors les zones 2 et 4 sont vides ; on en déduit la validité de *BARBARA*.

Considérons encore le cas *AOO-2*. La prémisses majeure $\forall x (R(x) \Rightarrow Q(x))$ devient $3 \cup 6 = \emptyset$ et la prémisses mineure $\exists x (P(x) \wedge \neg Q(x))$ devient $2 \cup 6 \neq \emptyset$; la conclusion $\exists x (P(x) \wedge \neg R(x))$ devient $2 \cup 4 \neq \emptyset$. A nouveau, il est clair que si les zones 3 et 6 sont toutes deux vides et que les zones 2 et 6 ne sont pas toutes deux vides, alors la zone 2 n'est pas vide, et donc les zones 2 et 4 ne sont pas toutes deux vides, ce qui établit la validité de *BAROCO*.

Le diagramme de Venn permet aussi de voir pourquoi un syllogisme n'est pas valide. Considérons par exemple *AIO-4*. La prémisses majeure $\forall x (R(x) \Rightarrow Q(x))$ devient $3 \cup 6 = \emptyset$ et la prémisses mineure $\exists x (Q(x) \wedge P(x))$ devient $4 \cup 7 \neq \emptyset$; on ne peut pas déduire de cela $2 \cup 4 \neq \emptyset$, correspondant à la conclusion $\exists x (P(x) \wedge \neg R(x))$. En effet, la situation où 2, 3, 4 et 6 sont vides tandis que 7 ne l'est pas vérifie les deux prémisses mais pas la conclusion.

Les diagrammes de Venn peuvent aussi être utilisés pour l'étude des cinq règles intuitives données plus haut. Considérons par exemple la deuxième règle : "si les deux prémisses sont particulières, le syllogisme n'est pas valide". En effet, une prémisses particulière se traduit par l'assertion "les zones x et y ne sont pas vides toutes les deux". De deux assertions de ce type, on ne peut rien déduire d'intéressant.

Taxonomie des syllogismes catégoriques. Il suffit de passer en revue les 48 syllogismes "potentiellement valides" et d'appliquer la technique du diagramme de Venn à chacun d'eux pour isoler les quinze syllogismes valides. La plupart des auteurs mentionnent cependant plus de quinze syllogismes valides, parce qu'ils acceptent comme valide, au moins dans certains cas, le mécanisme de *subalternation*. La *subalterne* d'une PQ -formule universelle est la

⁶⁶Nous convenons d'énoncer systématiquement la majeure avant la mineure, quoique l'ordre des prémisses soit sans influence sur la validité d'un raisonnement.

PQ -formule particulière (ou existentielle) correspondante.⁶⁷ Le mécanisme de subalternation consiste à déduire la subalterne de la formule universelle correspondante.

Ce mécanisme n'est pas valide stricto sensu, puisque la subalterne n'est pas conséquence logique de l'universelle ; on a

$$\forall x[L(x) \Rightarrow D(x)] \not\models \exists x[L(x) \wedge D(x)].$$

Cependant, on peut, au moyen d'une prémisse additionnelle, obtenir une version correcte de la subalternation :

$$\{\exists x L(x), \forall x[L(x) \Rightarrow D(x)]\} \models \exists x[L(x) \wedge D(x)].$$

Dans le langage naturel, on peut parfois considérer que la prémisse manquante est implicite. Nous qualifierons de *quasi-valide* un syllogisme dont la validité dépend de l'emploi de la subalternation, et donc de l'existence d'une prémisse implicite. Cette prémisse sera toujours de la même nature : elle affirme l'existence d'un objet au moins vérifiant le majeur, le moyen ou le mineur. Dans le cadre des diagrammes de Venn, cette prémisse prend donc l'une des trois formes suivantes :

- (Moyen) $1 \cup 4 \cup 5 \cup 7 \neq \emptyset$;
- (Mineur) $2 \cup 4 \cup 6 \cup 7 \neq \emptyset$;
- (Majeur) $3 \cup 5 \cup 6 \cup 7 \neq \emptyset$.

Le tableau récapitulatif des syllogismes valides ou quasi-valides est représenté à la figure 65.⁶⁸

Les quinze syllogismes valides sont AAA-1, EAE-1, AII-1 et EIO-1 pour la première figure, AEE-2, EAE-2, AOO-2 et EIO-2 pour la deuxième figure, AII-3, IAI-3, EIO-3 et OAO-3 pour la troisième figure, et AEE-4, IAI-4 et EIO-4, pour la quatrième figure. Dans le tableau, les noms anciens ont été utilisés ; on obtient la nomenclature moderne en ne retenant que les voyelles.⁶⁹

Cinq syllogismes valides ont une conclusion universelle ; en remplaçant celle-ci par sa subalterne, on obtient cinq syllogismes quasi-valides. Les dix autres syllogismes valides, dont la conclusion est particulière, ont également une prémisse particulière ;⁷⁰ en remplaçant celle-ci par sa superalterne, on obtient aussi des syllogismes quasi-valides, dont quatre distincts des précédents. On a donc en tout neuf syllogismes quasi-valides.

8.1.2 Schémas monadiques sur une variable

Introduction. Dans la mesure où les syllogismes catégoriques sont en nombre fini, leur théorie est nécessairement décidable et une procédure de décision triviale serait un tableau récapitulatif à 256 entrées, mentionnant pour chaque syllogisme s'il est valide ou non. La méthode des diagrammes de Venn est une procédure de décision plus intéressante, dont

⁶⁷ On dit parfois aussi que la formule universelle est la *superalterne* de la formule particulière.

⁶⁸ La dénomination "quasi-valide" n'appartient pas à la terminologie usuelle ; nous l'entendons du raisonnement privé de sa prémisse implicite existentielle. En effet, si nous tenons compte de celle-ci, le raisonnement n'est plus, stricto sensu, un syllogisme ; en revanche, il devient valide.

⁶⁹ Les consonnes ont également une signification, liée aux règles qu'utilisaient les logiciens médiévaux pour convertir en syllogismes de la première figure ceux des autres figures, la première figure étant jugée plus fondamentale.

⁷⁰ jamais les deux (règle 2 du paragraphe 8.1.1).

Première figure

<u>BARBARA</u> $\frac{\forall x [Q(x) \Rightarrow R(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\forall x [P(x) \Rightarrow R(x)]}$	<u>CELARENT</u> $\frac{\forall x [Q(x) \Rightarrow \neg R(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\forall x [P(x) \Rightarrow \neg R(x)]}$
<u>DARII</u> $\frac{\forall x [Q(x) \Rightarrow R(x)] \quad \exists x [P(x) \wedge Q(x)]}{\exists x [P(x) \wedge R(x)]}$	<u>FERIO</u> $\frac{\forall x [Q(x) \Rightarrow \neg R(x)] \quad \exists x [P(x) \wedge Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$
<u>BARBARI</u> $\frac{\exists x P(x) \quad \forall x [Q(x) \Rightarrow R(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\exists x [P(x) \wedge R(x)]}$	<u>CELARO</u> $\frac{\exists x P(x) \quad \forall x [Q(x) \Rightarrow \neg R(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$

Deuxième figure

<u>CAMESTRES</u> $\frac{\forall x [R(x) \Rightarrow Q(x)] \quad \forall x [P(x) \Rightarrow \neg Q(x)]}{\forall x [P(x) \Rightarrow \neg R(x)]}$	<u>CESARE</u> $\frac{\forall x [R(x) \Rightarrow \neg Q(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\forall x [P(x) \Rightarrow \neg R(x)]}$
<u>BAROCO</u> $\frac{\forall x [R(x) \Rightarrow Q(x)] \quad \exists x [P(x) \wedge \neg Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$	<u>FESTINO</u> $\frac{\forall x [R(x) \Rightarrow \neg Q(x)] \quad \exists x [P(x) \wedge Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$
<u>CAMESTROS</u> $\frac{\exists x P(x) \quad \forall x [R(x) \Rightarrow Q(x)] \quad \forall x [P(x) \Rightarrow \neg Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$	<u>CESARO</u> $\frac{\exists x P(x) \quad \forall x [R(x) \Rightarrow \neg Q(x)] \quad \forall x [P(x) \Rightarrow Q(x)]}{\exists x [P(x) \wedge \neg R(x)]}$

Troisième figure

<u>DATISI</u> $\frac{\forall x [Q(x) \Rightarrow R(x)] \quad \exists x [Q(x) \wedge P(x)]}{\exists x [P(x) \wedge R(x)]}$	<u>DISAMIS</u> $\frac{\exists x [Q(x) \wedge R(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge R(x)]}$
<u>BOCARDO</u> $\frac{\exists x [Q(x) \wedge \neg R(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$	<u>FERISON</u> $\frac{\forall x [Q(x) \Rightarrow \neg R(x)] \quad \exists x [Q(x) \wedge P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$
<u>DARAPTI</u> $\frac{\exists x Q(x) \quad \forall x [Q(x) \Rightarrow R(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge R(x)]}$	<u>FELAPTON</u> $\frac{\exists x Q(x) \quad \forall x [Q(x) \Rightarrow \neg R(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$

Quatrième figure

<u>CAMENES</u> $\frac{\forall x [R(x) \Rightarrow Q(x)] \quad \forall x [Q(x) \Rightarrow \neg P(x)]}{\forall x [P(x) \Rightarrow \neg R(x)]}$	<u>DIMARIS</u> $\frac{\exists x [R(x) \wedge Q(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge R(x)]}$
<u>CAMENOS</u> $\frac{\exists x P(x) \quad \forall x [R(x) \Rightarrow Q(x)] \quad \forall x [Q(x) \Rightarrow \neg P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$	<u>FRESISON</u> $\frac{\forall x [R(x) \Rightarrow \neg Q(x)] \quad \forall x [Q(x) \wedge P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$
<u>BRAMANTIP</u> $\frac{\exists x R(x) \quad \forall x [R(x) \Rightarrow Q(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge R(x)]}$	<u>FESAPO</u> $\frac{\exists x Q(x) \quad \forall x [R(x) \Rightarrow \neg Q(x)] \quad \forall x [Q(x) \Rightarrow P(x)]}{\exists x [P(x) \wedge \neg R(x)]}$

FIG. 65 – Tableau des syllogismes valides ou quasi-valides

on imagine aisément qu'elle reste applicable pour une classe infinie de formules.⁷¹ Pour déterminer une telle classe, aussi grande que possible, nous reconsidérons des syllogismes et essayons de les généraliser.

⁷¹ Il est plus commode de parler de formules que de raisonnement ; rappelons que le raisonnement dont les prémisses sont $P_1, \dots, P_n$ et dont la conclusion est C est valide (ou correct) si et seulement si la formule $(P_1 \wedge \dots \wedge P_n) \Rightarrow C$ est valide.

Concrètement, valider le syllogisme BARBARA revient à valider la disjonction

$$\exists x [Q(x) \wedge \neg R(x)] \vee \exists x [P(x) \wedge \neg Q(x)] \vee \forall x [P(x) \Rightarrow R(x)];$$

de même, valider FERIO revient à valider la disjonction

$$\exists x [Q(x) \wedge R(x)] \vee \forall x [P(x) \Rightarrow \neg Q(x)] \vee \exists x [P(x) \wedge \neg R(x)].$$

Il semble clair que la méthode des diagrammes de Venn restera applicable si les disjonctions comportent plus de trois éléments et si matrices des formules impliquées sont des fonctions booléennes quelconques des formes $P(x)$, $Q(x)$ et $R(x)$. De même, rien ne devrait empêcher l'introduction d'une quatrième forme $S(x)$: on pourrait alors maintenir la forme graphique élégante de la méthode de Venn en passant dans l'espace à trois dimensions, les formes étant représentées par quatre sphères, délimitant en tout $2^4 = 16$ zones. On pourrait aussi renoncer à l'aspect graphique de la méthode et autoriser n prédicats distincts.⁷²

Dans la suite de ce chapitre, on précise ces extensions et leur traitement, ce qui conduit à une procédure de décision pour le calcul des prédicats monadiques.

Schémas monadiques booléens sur une variable. Un schéma monadique booléen (SMB) sur la variable x est une combinaison booléenne (finie) de formes telles que $P(x)$, $Q(x)$, ... résultant de l'application à la variable x d'un prédicat monadique.

Dans la méthode de Venn, on se préoccupe seulement de savoir si une zone est vide ou non, sans distinguer les cas où une zone contient un ou plusieurs élément(s) ; c'est justifié par le théorème suivant :

Théorème. Si un SMB $\Phi(x)$ admet un modèle, alors il admet un modèle à un seul élément.

Démonstration. Soit I une interprétation de $\Phi(x)$ de domaine D et soit $a \in D$ tel que $I[x] = a$. L'interprétation J de domaine $\{a\}$ telle que $J[x] = a$ et, pour tout prédicat (monadique) P , telle que $J[P]$ est la restriction à $\{a\}$ de $I[P]$ est telle que $J(\Phi) = I(\Phi)$.

Définition. Une interprétation *fondamentale* d'un SMB $\Phi(x)$ est une interprétation de $\Phi(x)$ dont le domaine comporte un seul élément.

Remarque. Les interprétations dont le domaine est un singleton ont une propriété intéressante dépassant le cadre monadique. Une telle interprétation attribue toujours la même valeur de vérité à une formule (quelconque), à sa fermeture universelle et à sa fermeture existentielle.

Remarque. Etant donné le lexique $\Pi = \{P_1, \dots, P_n\}$, un SMB $\Phi(x)$ admet 2^n interprétations fondamentales distinctes.

Corollaire. Un SMB est valide s'il est vrai pour toutes les interprétations fondamentales ; il est consistant s'il est vrai pour une interprétation fondamentale au moins.

Remarque. Les schémas monadiques booléens peuvent être assimilés aux formules propositionnelles ; ils ne sont donc pas intéressants en soi.

Schémas monadiques quantifiés sur une variable. La fermeture (existentielle ou universelle) d'un schéma monadique booléen sur x est un *schéma monadique quantifié* (existantiel ou universel) sur la variable x . Les prémisses et la conclusion d'un syllogisme sont des schémas monadiques quantifiés (SMQ).

⁷²Cela correspondrait à n hypersphères de l'espace à $n - 1$ dimensions, déterminant 2^n zones.

Remarque. La matrice d'un SMQ ne contient pas d'autre variable que la variable quantifiée (unique) ; un SMQ est donc une formule fermée.

Théorème. Un schéma monadique quantifié est valide (resp. consistant, contingent) si et seulement si sa matrice est valide (resp. consistante, contingente).

Démonstration. Il suffit de démontrer que, pour tout schéma monadique booléen $\Phi(x)$, la validité de $\exists x \Phi(x)$ entraîne celle de $\Phi(x)$. On procède par l'absurde. Si $\Phi(x)$ admet un antimodèle, il existe un antimodèle à un seul élément ; celui-ci est nécessairement un antimodèle de $\exists x \Phi(x)$.

Corollaire. Si $\Phi(x)$ est un SMB, les trois formules $\Phi(x)$, $\forall x \Phi(x)$ et $\exists x \Phi(x)$ sont simultanément valides, contingentes ou inconsistantes.⁷³

Remarque. Ce corollaire montre que les schémas monadiques quantifiés, pris isolément, ne sont pas non plus très intéressants. En revanche, les combinaisons booléennes de tels schémas le sont, comme nous le voyons au paragraphe suivant. La formule correspondant à un syllogisme peut toujours s'écrire comme une disjonction de trois SMQ.

8.1.3 Formules monadiques sur une variable

Combinaisons booléennes de SMQ sur une variable. Les syllogismes correspondent à des disjonctions de SMQ. Cela suggère l'étude systématique de ces disjonctions ou, plus généralement, des combinaisons booléennes de SMQ. Cette "généralisation" n'est d'ailleurs pas très profonde, car on a le résultat suivant :

Théorème. Toute combinaison booléenne B de SMQ est logiquement équivalente à une disjonction D de conjonctions de SMQ, et aussi à une conjonction C de disjonctions de SMQ.

Démonstration. Pour obtenir par exemple C à partir de B , on réduit d'abord B à la forme normale conjonctive, chaque SMQ de B étant assimilé à un atome ; les "littéraux négatifs" éventuels, c'est-à-dire des négations de SMQ, sont alors remplacés par des SMQ.⁷⁴

Si on peut tester la validité de disjonctions de schémas monadiques sur une variable, on pourra donc tester la validité de la conjonction de ces disjonctions et donc d'une combinaison booléenne quelconque ; de même, si on peut tester la consistance d'une conjonction de schémas monadiques sur une variable, on pourra tester la consistance d'une combinaison booléenne quelconque de tels schémas. Notons aussi que les deux problèmes sont équivalents.⁷⁵ Signalons enfin que, vu les règles de renommage, on peut toujours supposer que seule la variable x est utilisée.

Les deux résultats qui suivent simplifient grandement le problème.

Théorème. Une conjonction de schémas existentiels $E_1 \wedge \dots \wedge E_n$ est consistante si et seulement si chacun des schémas E_j est consistant.

Démonstration. La condition est visiblement nécessaire : un modèle d'une conjonction est aussi un modèle de chacun de ses éléments. D'autre part, supposons que tous les $E_k = \exists x M_k$ soient consistants. Les matrices M_k sont consistantes et admettent un modèle fondamental I_k sur un

⁷³En cas de validité ou d'inconsistance, les trois formules sont évidemment logiquement équivalentes, mais en cas de contingence, $\Phi(x)$ n'est jamais logiquement équivalente à l'une de ses fermetures.

⁷⁴La négation d'un SM universel est un SM existentiel, la négation d'un SM existentiel est un SM universel.

⁷⁵Rappelons qu'une conjonction est valide si et seulement si tous ses éléments le sont ; une disjonction est consistante si et seulement si tous ses éléments le sont. De plus, la négation d'une conjonction (disjonction) de SMQ est une disjonction (conjonction) de SMQ.

singleton $\{a_k\}$. On définit une interprétation I dont le domaine est $\{a_1, \dots, a_n\}$; pour tout prédicat P et pour tout $k \in \{1, \dots, n\}$, on pose $I[P](a_k) = I_k[P](a_k)$ si P intervient dans E_k et $I[P](a_k) = \mathbf{V}$ (par exemple) sinon. L'interprétation I est un modèle commun à tous les E_k et donc un modèle de la conjonction car, par construction, $I_{x/a_k}(M_k) = \mathbf{V}$ donc $I(E_k) = \mathbf{V}$.
Remarque. On ne peut pas déduire de ce qui précède qu'en logique monadique, la formule $\exists x [\Phi(x) \wedge \Psi(x)]$ serait logiquement équivalente à la formule $\exists x \Phi(x) \wedge \exists x \Psi(x)$. En outre, le résultat selon lequel un schéma existentiel consistant admet un modèle fondamental ne s'étend pas à une conjonction de tels schémas. Un contre-exemple commun évident est fourni par les deux schémas $\exists x P(x)$ et $\exists x \neg P(x)$.

Théorème. Un schéma existentiel $\exists x \Psi(x)$ est conséquence logique d'un schéma existentiel $\exists x \Phi(x)$ si et seulement si la matrice $\Psi(x)$ est conséquence logique de la matrice $\Phi(x)$.

Démonstration. La condition est visiblement suffisante. Soit I un modèle fondamental de $\exists x \Phi(x)$; c'est aussi un modèle (fondamental) de $\Phi(x)$, de $\Psi(x)$ et de $\exists x \Psi(x)$. La condition est aussi nécessaire. Si $\Psi(x)$ n'est pas conséquence logique de $\Phi(x)$, alors le SMB $\Phi(x) \wedge \neg \Psi(x)$ admet un modèle fondamental, qui est aussi un modèle de $\exists x \Phi(x)$, mais un antimodèle fondamental de $\Psi(x)$ et donc de $\exists x \Psi(x)$.

Corollaire. La conjonction $(\exists x \Phi(x) \wedge \forall x \Psi(x))$ est (in)consistante si et seulement si le SMB $\Phi(x) \wedge \neg \Psi(x)$ est (in)consistant.

Corollaire. Si $\exists x \Psi(x)$ n'est pas conséquence logique de $\exists x \Phi(x)$, la conjonction $\exists x \Phi(x) \wedge \forall x \neg \Psi(x)$ admet un modèle fondamental.

Il reste à étudier le cas où la conjonction de SMQ comporte aussi un schéma universel ou plusieurs; on peut se limiter à un seul, puisque la conjonction de schémas universels est un schéma universel.⁷⁶

On note d'abord que la conjonction d'un schéma existentiel E et d'un schéma universel U est consistante si et seulement si le schéma existentiel $\neg U$ n'est pas conséquence logique du schéma existentiel E , ce que l'on peut tester par le théorème précédent.

On a enfin le théorème suivant.

Théorème. La conjonction $E_1 \wedge \dots \wedge E_n \wedge U$ est consistante si et seulement si chacune des conjonctions $E_k \wedge U$ est consistante.

Démonstration. La condition est visiblement nécessaire. Elle est aussi suffisante. Vu le corollaire précédent, si les conjonctions $E_k \wedge U$ sont consistantes, elles admettent respectivement les modèles fondamentaux I_k , de domaine $\{a_k\}$. On définit alors l'interprétation I de domaine est $\{a_1, \dots, a_n\}$; pour tout prédicat P et pour tout $k \in \{1, \dots, n\}$, on pose $I[P](a_k) = I_k[P](a_k)$ si P intervient dans E_k ou dans U et $I[P](a_k) = \mathbf{V}$ (par exemple) sinon. L'interprétation I est un modèle commun à U et à tous les E_k et donc un modèle de la conjonction car, par construction, on a $I_{x/a_k}(M_k) = I_{x/a_k}(M)\mathbf{V}$ donc $I(E_k \wedge U) = \mathbf{V}$, où M_k et M sont les matrices de E_k et U .

Corollaire. La disjonction $U_1 \vee \dots \vee U_n \vee E$ est valide si et seulement si au moins une des disjonctions $U_k \vee E$ est valide.

Remarque. On teste la validité d'une combinaison booléenne de SMQ en la transformant en une conjonction de disjonctions de SMQ, et en traitant séparément chaque disjonction. Le test d'une disjonction de n SMQ se ramène à au plus n tests de conséquence logique entre deux

SMQ existentiels ou, plus simplement, entre deux SMB sur une même variable x et donc, plus simplement encore, au test de validité de n SMB sur x . Enfin, un SMB est valide si et seulement si le schéma propositionnel correspondant (obtenu en remplaçant chaque occurrence de $P_i(x)$ par l'atome p_i) est valide.

Remarque. La technique vue ici s'étend immédiatement aux combinaisons booléennes comportant non seulement des SMQ mais aussi des propositions élémentaires.

Exemples. Nous reconsidérons d'abord les cas des syllogismes BARBARA et FERIO évoqués au paragraphe 8.1.2. La disjonction

$$\exists x [Q(x) \wedge \neg R(x)] \vee \exists x [P(x) \wedge \neg Q(x)] \vee \forall x [P(x) \Rightarrow R(x)],$$

associée à BARBARA, se réécrit d'abord en

$$\exists x ([Q(x) \wedge \neg R(x)] \vee [P(x) \wedge \neg Q(x)]) \vee \forall x [P(x) \Rightarrow R(x)];$$

elle est valide si son premier élément est conséquence logique de son second, ou encore si

$$\neg[P(x) \Rightarrow R(x)] \models [Q(x) \wedge \neg R(x)] \vee [P(x) \wedge \neg Q(x)],$$

c'est-à-dire si

$$\neg[p \Rightarrow r] \models [q \wedge \neg r] \vee [p \wedge \neg q],$$

ce qui est visiblement le cas. De même, la disjonction

$$\exists x [Q(x) \wedge R(x)] \vee \forall x [P(x) \Rightarrow \neg Q(x)] \vee \exists x [P(x) \wedge \neg R(x)].$$

associée à FERIO, se réécrit d'abord en

$$\exists x ([Q(x) \wedge R(x)] \vee [P(x) \wedge \neg R(x)]) \vee \forall x [P(x) \Rightarrow \neg Q(x)];$$

elle est valide si son premier élément est conséquence logique de son second, ou encore si

$$\neg[P(x) \Rightarrow \neg Q(x)] \models [Q(x) \wedge R(x)] \vee [P(x) \wedge \neg R(x)],$$

c'est-à-dire si

$$\neg[p \Rightarrow \neg q] \models [q \wedge r] \vee [p \wedge \neg r],$$

ce qui est visiblement le cas.

On étudie ensuite DARAPTI. Sans le présupposé d'existence, la disjonction correspondante est

$$\exists x [Q(x) \wedge \neg R(x)] \vee \exists x [Q(x) \wedge \neg P(x)] \vee \exists x [P(x) \wedge R(x)],$$

qui se réécrit en

$$\exists x ([Q(x) \wedge \neg R(x)] \vee [Q(x) \wedge \neg P(x)] \vee [P(x) \wedge R(x)]).$$

Ce SBQ est valide si la formule

$$[q \wedge \neg r] \vee [q \wedge \neg p] \vee [p \wedge r],$$

⁷⁶Quelles que soient les formules A_i et la variable x , les formules $\forall x \bigwedge_i A_i$ et $\bigwedge_i \forall x A_i$ sont logiquement équivalentes.

ce qui n'est pas le cas. En revanche, avec le présupposé d'existence, la disjonction correspondante devient

$$\forall x \neg Q(x) \vee \exists x [Q(x) \wedge \neg R(x)] \vee \exists x [Q(x) \wedge \neg P(x)] \vee \exists x [P(x) \wedge R(x)],$$

qui se réécrit en

$$\forall x \neg Q(x) \vee \exists x ([Q(x) \wedge \neg R(x)] \vee [Q(x) \wedge \neg P(x)] \vee [P(x) \wedge R(x)]).$$

Cette disjonction est valide si son second élément est conséquence logique de la négation de son premier, ou encore si la formule

$$q \Rightarrow ([q \wedge \neg r] \vee [q \wedge \neg p] \vee [p \wedge r])$$

est valide, ce qui est le cas.

Les diagrammes de Venn. Reconsidérons DARPTI par la méthode des diagrammes de Venn. Dans le diagramme standard du syllogisme (Fig. 64), la prémisse majeure exprime que les zones 1 et 4 sont vides, ce que l'on peut écrire $Q \subset R$ ou, pourquoi pas, $q \Rightarrow r$; la prémisse mineure exprime de même $q \Rightarrow p$ et la conclusion exprime que l'une au moins des zones 6 et 7 n'est pas vide, ce que l'on écrit $p \wedge r$.

D'après la méthode de Venn, la validité de DARPTI (sans prémisse supplémentaire) correspond à l'énoncé

$$\{q \Rightarrow r, q \Rightarrow p\} \models p \wedge r;$$

la validité de DARPTI avec présupposé d'existence correspond à l'énoncé

$$\{q, q \Rightarrow r, q \Rightarrow p\} \models p \wedge r,$$

ou encore à l'énoncé

$$\{q, q \Rightarrow r, q \Rightarrow p, \neg(p \wedge r)\} \text{ est inconsistent.}$$

On note que la méthode de Venn n'est autre qu'une version graphique de la méthode introduite et justifiée dans les paragraphes précédents; elle est donc correcte.

Remarque. Le traitement de DARPTI par la méthode de Herbrand revient à déterminer l'inconsistance de l'ensemble

$$\{Q(a), Q(a) \Rightarrow R(a), Q(a) \Rightarrow P(a), \neg(P(a) \wedge R(a))\};$$

la méthode de Venn est donc dans ce cas une version graphique de la méthode de Herbrand.

8.1.4 La logique des prédicats monadiques

Introduction. Nous venons de voir qu'un fragment important de la logique monadique, comportant notamment les formules modélisant les syllogismes catégoriques, se ramenait au calcul des propositions. Nous considérons maintenant la logique monadique complète. La clef du traitement est la réduction des formules à la *forme simple*.

Théorème. Une formule est simple si et seulement si elle est une combinaison booléenne de SMQ et d'atomes.

Démonstration. La condition est visiblement suffisante. On établit qu'elle est nécessaire par induction sur la structure syntaxique des formules.

Corollaire. Une formule est simple et fermée si et seulement si elle est une combinaison booléenne de SMQ et d'atomes sans variable.⁷⁷

Lois de passage. Les *lois de passage* sont des schémas d'équivalence logique entre formules. Associées au théorème de l'échange,⁷⁸ elles permettent de réduire les formules à la forme simple. On a :

- $\forall x A \leftrightarrow A$ et $\exists x A \leftrightarrow A$, si A ne comporte pas d'occurrence libre de x .
- $\forall x \neg A \leftrightarrow \neg \exists x A$; $\exists x \neg A \leftrightarrow \neg \forall x A$.
- $\forall x (A \wedge B) \leftrightarrow (\forall x A \wedge \forall x B)$; $\exists x (A \vee B) \leftrightarrow (\exists x A \vee \exists x B)$.
- $\forall x (A \vee B) \leftrightarrow (\forall x A \vee B)$; $\exists x (A \wedge B) \leftrightarrow (\exists x A \wedge B)$, si B ne comporte pas d'occurrence libre de x .

Rappelons que ces règles sont valables en logique prédicative générale.

Procédure de décision pour la logique monadique. On introduit d'abord un lemme.

Lemme. Si la formule A est simple, alors il existe des formules simples A' et A'' , logiquement équivalentes à $\forall x A$ et $\exists x A$, respectivement.

Démonstration. Simple utilisation des lois de passage.

Théorème. Toute formule monadique est logiquement équivalente à une forme simple.

Démonstration. On obtient cette forme simple par application répétée du lemme.

Procédure de décision. Pour tester la validité d'une formule, on réduit sa fermeture universelle à la forme simple et on applique la technique du paragraphe 8.1.3.

Complexité. La procédure implique des réductions en forme normale (conjonctive ou disjonctive) répétées. Pour une forme prénexe comportant n quantifications alternées, n réductions peuvent être nécessaires.

Un exemple. On veut comparer les formules

$$\forall x \exists y [(Px \vee Qy) \wedge (Rx \vee Sy)] \text{ et } \exists y \forall x [(Px \vee Qy) \wedge (Rx \vee Sy)].$$

Il est clair que la première formule est conséquence logique de la seconde, mais la réciproque est fausse. Pour le voir, on réduit d'abord la première formule à la forme simple. Dans la liste ci-dessous, toutes les formules sont logiquement équivalentes.

⁷⁷Les atomes sans variable sont *true*, *false* et les propositions élémentaires. Bien que ces dernières soient assimilées à des prédicats à 0 argument, il est commode de les admettre en logique monadique. D'ailleurs, on pourrait éliminer les atomes sans variable en les remplaçant par des SMQ particuliers.

⁷⁸Le théorème de l'échange permet de remplacer une sous-formule par une sous-formule logiquement équivalente, sans changer la sémantique de départ.

$$\begin{aligned}
& \forall x \exists y [(Px \vee Qy) \wedge (Rx \vee Sy)], \\
& \forall x \exists y [(Px \wedge Rx) \vee (Px \wedge Sy) \vee (Qy \wedge Rx) \vee (Qy \wedge Sy)], \\
& \forall x [(Px \wedge Rx) \vee (Px \wedge \exists y Sy) \vee (\exists y Qy \wedge Rx) \vee \exists y (Qy \wedge Sy)], \\
& \forall x [(Px \wedge Rx) \vee (Px \wedge \exists y Sy) \vee (\exists y Qy \wedge Rx)] \vee \exists y (Qy \wedge Sy), \\
& \forall x [(Px \vee \exists y Qy) \wedge (Px \vee Rx) \wedge (Rx \vee \exists y Sy)] \vee \exists y (Qy \wedge Sy), \\
& [(\forall x Px \vee \exists y Qy) \wedge \forall x (Px \vee Rx) \wedge (\forall x Rx \vee \exists y Sy)] \vee \exists y (Qy \wedge Sy).
\end{aligned}$$

Une forme disjonctive normale équivalente est :

$$\begin{aligned}
& (\forall x Px \wedge \forall x Rx) \vee (\forall x Px \wedge \exists y Sy) \vee (\exists y Qy \wedge \forall x Rx) \vee \\
& (\exists y Qy \wedge \forall x (Px \vee Rx) \wedge \exists y Sy) \vee \exists y (Qy \wedge Sy).
\end{aligned}$$

Pour la seconde formule, on obtient une liste analogue :

$$\begin{aligned}
& \exists y \forall x [(Px \vee Qy) \wedge (Rx \vee Sy)], \\
& \exists y [\forall x (Px \vee Qy) \wedge \forall x (Rx \vee Sy)], \\
& \exists y [(\forall x Px \vee Qy) \wedge (\forall x Rx \vee Sy)], \\
& \exists y [(\forall x Px \wedge \forall x Rx) \vee (\forall x Px \wedge Sy) \vee (Qy \wedge \forall x Rx) \vee (Qy \wedge Sy)], \\
& (\forall x Px \wedge \forall x Rx) \vee (\forall x Px \wedge \exists y Sy) \vee (\exists y Qy \wedge \forall x Rx) \vee \exists y (Qy \wedge Sy).
\end{aligned}$$

Une forme disjonctive normale équivalente est la dernière ligne :

$$(\forall x Px \wedge \forall x Rx) \vee (\forall x Px \wedge \exists y Sy) \vee (\exists y Qy \wedge \forall x Rx) \vee \exists y (Qy \wedge Sy)$$

En comparant les deux formes normales, on observe que la première comporte les mêmes cubes que la seconde, plus un cube supplémentaire. Il est donc possible de rendre la première formule vraie tout en falsifiant la seconde, au moyen d'une interprétation rendant faux les quatre cubes ci-dessus et vrai le cube supplémentaire

$$(\exists y Qy \wedge \forall x (Px \vee Rx) \wedge \exists y Sy).$$

Une telle interprétation sur le domaine $\{a, b\}$ est par exemple celle qui rend vrais les atomes Pa, Qa, Rb, Sb et faux les atomes Pb, Qb, Ra, Sa .

On peut aussi appliquer la technique vue au paragraphe 8.1.3. Il faut montrer la consistance d'une conjonction de cinq formules, dont les quatre premières sont les négations des cubes communs aux deux formules étudiées et dont la cinquième est le cube supplémentaire. Les cinq membres de la conjonction sont donc

1. $\neg(\forall x Px \wedge \forall x Rx)$, soit $\exists x (\neg Px \vee \neg Rx)$;
2. $\neg(\forall x Px \wedge \exists y Sy)$, soit $\exists x \neg Px \vee \forall y \neg Sy$;
3. $\neg(\exists y Qy \wedge \forall x Rx)$, soit $\forall y \neg Qy \vee \exists x \neg Rx$;
4. $\neg \exists y (Qy \wedge Sy)$, soit $\forall y (\neg Qy \vee \neg Sy)$;
5. $\exists y Qy \wedge \forall x (Px \vee Rx) \wedge \exists y Sy$.

Tout modèle éventuel devra satisfaire $\exists y Qy$ et $\exists y Sy$ (formule 5), ce qui permet de simplifier d'emblée les formules 2 et 3 en $\exists x \neg Px$ et $\exists x \neg Rx$, respectivement. Cette simplification montre que la formule 1 est inutile et peut donc être omise. Il reste donc à trouver un modèle pour la conjonction

$$\exists x \neg Px \wedge \exists x \neg Rx \wedge \forall y (\neg Qy \vee \neg Sy) \wedge \exists y Qy \wedge \forall x (Px \vee Rx) \wedge \exists y Sy.$$

En regroupant les deux schémas universels et par renommage de y en x , cette formule se réécrit en

$$\exists x \neg Px \wedge \exists x \neg Rx \wedge \exists x Qx \wedge \exists x Sx \wedge \forall x [(\neg Qx \vee \neg Sx) \wedge (Px \vee Rx)].$$

D'après les résultats du paragraphe 8.1.3, il suffit de vérifier séparément la consistance des formules

$$\begin{aligned}
& \exists x \neg Px \wedge \forall x [(\neg Qx \vee \neg Sx) \wedge (Px \vee Rx)], \\
& \exists x \neg Rx \wedge \forall x [(\neg Qx \vee \neg Sx) \wedge (Px \vee Rx)], \\
& \exists x Qx \wedge \forall x [(\neg Qx \vee \neg Sx) \wedge (Px \vee Rx)], \\
& \exists x Sx \wedge \forall x [(\neg Qx \vee \neg Sx) \wedge (Px \vee Rx)],
\end{aligned}$$

ou encore (§ 8.1.3) des formules

$$\begin{aligned}
& \neg Px \wedge (\neg Qx \vee \neg Sx) \wedge (Px \vee Rx), \\
& \neg Rx \wedge (\neg Qx \vee \neg Sx) \wedge (Px \vee Rx), \\
& Qx \wedge (\neg Qx \vee \neg Sx) \wedge (Px \vee Rx), \\
& Sx \wedge (\neg Qx \vee \neg Sx) \wedge (Px \vee Rx),
\end{aligned}$$

ce qui est évident dans chaque cas.

Remarque. Rappelons que, dans la recherche de modèles pour ces formules, on peut se limiter aux modèles fondamentaux. On peut aussi, à partir des quatre modèles obtenus, créer un modèle commun, mais ce modèle ne sera généralement pas fondamental. Dans notre exemple, il comportera au minimum deux éléments ; le modèle à deux éléments donné plus haut rend vraies les quatre formules, en considérant l'instance $x = a$ pour les deuxième et troisième formules, et l'instance $x = b$ pour les deux autres formules.

8.2 La logique de Bernays et Schönfinkel

8.2.1 Introduction

L'introduction dans la logique prédicative des prédicats polyadiques ne pose pas de problème sémantique. En particulier, la notion d'interprétation s'étend immédiatement à ce cas. Cependant, l'introduction d'un seul prédicat à deux arguments suffit à prévenir l'existence d'une "vraie" procédure de décision. Plus précisément, la méthode des tableaux sémantiques, celle de Herbrand, celle de Hilbert et beaucoup d'autres permettent, de manière systématique, de reconnaître les formules valides et les formules inconsistantes mais, pour chacune de ces méthodes, il existe toujours certaines formules contingentes pour lesquelles aucune conclusion n'est fournie en un temps fini, et il a été démontré que cette lacune était irrémédiable.

On peut cependant, en restreignant la logique prédicative, obtenir des fragments décidables, et la logique monadique est probablement le plus connu. La limitation de l'arité des prédicats ne conduit cependant pas très loin. Une autre voie plus prometteuse est la limitation des schémas de quantification, que nous explorons ici.

8.2.2 Logique prédicative sans quantification

Si on s'interdit de quantifier les variables, celles-ci deviennent, sur le plan sémantique, indistinguables des constantes. En effet, interpréter une constante ou une variable est simplement lui associer un élément du domaine d'interprétation. On peut donc admettre qu'en l'absence de quantification, les seuls termes sont les constantes.

Une formule de la logique sans quantification est une combinaison linéaire d'atomes sans variables. Si une telle formule comporte n atomes distincts, elle admettra 2^n interprétations, exactement comme en logique propositionnelle. On notera que deux atomes sont (complètement) indépendants dès qu'ils sont syntaxiquement distincts; il n'y a pas plus de liens sémantiques entre, par exemple, $P(a, a, b)$ et $P(a, b, b)$ qu'entre $P(a, a, b)$ et $Q(c, d)$. L'étude des formules prédictives sans quantification se ramène à l'étude des formules booléennes correspondantes.

8.2.3 Logique prédicative sans alternance de quantification

L'étape suivante consiste naturellement à considérer les formules comportant une seule quantification mais, telle quelle, cette classe n'est pas bien définie du point de vue sémantique. En effet, les formules $\forall x (P(x) \wedge Q(x, a))$ et $\forall x P(x) \wedge \forall x Q(x, a)$ étant logiquement équivalentes, il n'y a pas de raison d'accepter la première et de refuser la seconde. Il est également peu indiqué d'imposer le type de quantification, puisque des formules comme $\forall x (P(x) \Rightarrow R(a))$ et $\exists x P(x) \Rightarrow R(a)$ sont logiquement équivalentes. On peut cependant imposer ce type de restriction pour les formes prénexes.

Formules existentielles pures, formules universelles pures. Une formule existentielle (universelle) pure est une formule logiquement équivalente à une forme prénex ne comportant que des quantificateurs existentiels (universels). Toutes les formules introduites au paragraphe précédent sont donc des universelles pures. Dans la mesure où toute formule se réduit aisément à la forme prénex, il n'est pas gênant de se restreindre à ces formes dans cette étude.

On sait qu'une formule est consistante si et seulement si sa fermeture existentielle est consistante; une formule est valide si et seulement si sa fermeture universelle est valide. Cependant, la fermeture existentielle de la formule $P(x) \vee \neg P(y)$ est valide sans que la matrice elle-même le soit; de même, la fermeture universelle de $P(x) \wedge \neg P(y)$ est inconsistante alors que la matrice est consistante.

Le théorème de Herbrand donne un moyen simple de tester la consistance des formes de Skolem, qui devient une procédure de décision dans le cas où on n'a pas de symboles fonctionnels. On a les résultats suivants :

Théorème. Une formule universelle pure (forme de Skolem) est consistante si et seulement si on obtient un schéma vérifonctionnellement consistant en prenant la conjonction des formules obtenus en substituant des variables libres aux variables quantifiées de la matrice.

Démonstration. Corollaire immédiat du théorème de Herbrand, les variables libres de la formule jouant le rôle de constantes de Skolem.

Théorème. Une formule existentielle pure est valide si et seulement si on obtient un schéma vérifonctionnellement valide en prenant la disjonction des formules obtenus en substituant des variables libres aux variables quantifiées de la matrice.

8.2.4 Logique prédicative avec une alternance de quantification

La technique du paragraphe précédent permet d'analyser toutes les formules monadiques et de nombreuses autres formules utiles, notamment celles dont la forme prénex comporte une seule alternance de quantificateurs.

Exemple. La formule

$$\exists x \forall y P(x, y) \Rightarrow \forall y \exists x P(x, y)$$

est logiquement équivalente à la forme prénex

$$\forall x \forall y \exists z \exists w [P(x, z) \Rightarrow P(w, y)];$$

cette formule est la fermeture universelle de

$$\exists z \exists w [P(x, z) \Rightarrow P(w, y)],$$

donc ces trois formules sont valides si et seulement si la dernière est valide, ce qui est le cas puisque la disjonction

$$[P(x, x) \Rightarrow P(x, y)] \vee [P(x, x) \Rightarrow P(y, y)] \vee [P(x, y) \Rightarrow P(x, y)] \vee [P(x, y) \Rightarrow P(y, y)]$$

est valide.

Petit théorème de Herbrand. Le théorème de Herbrand est à la base d'une procédure générale permettant de reconnaître la validité ou l'inconsistance des formules prédictives quelconques. Particularisé aux formules quantifiées pures, il permet de reconnaître aussi les formules contingentes. Dans ce paragraphe, nous nous limitons à ce cas particulier.

Nous considérons une forme prénex universelle pure S , sans variable libre mais contenant éventuellement des constantes. Le *domaine de Herbrand* H_S (ou *univers de Herbrand*) de S est l'ensemble de ces constantes s'il en existe, et le singleton $\{c\}$ sinon; ce domaine est fini. Une *interprétation de Herbrand* d'une formule en forme de Skolem S est une interprétation $\mathcal{H}$ de S dont le domaine est H_S et telle que chaque constante présente dans S soit interprétée en elle-même (si cette interprétation rend S vraie, on parle de *modèle de Herbrand*). La *base de Herbrand* B_S est l'ensemble des atomes fondamentaux. Cette base est finie car le nombre de prédicats intervenant dans S est évidemment fini. Si $|H_S| = k$, chaque prédicat d'arité n donne lieu à k^n atomes fondamentaux.

Théorème. Si $\mathcal{H}$ est une interprétation de Herbrand pour la matrice $A(x_1, \dots, x_n)$, on a $\mathcal{H}[\forall x_1 \dots \forall x_n A(x_1, \dots, x_n)] = \mathbf{V}$ si et seulement si $\mathcal{H}[A(h_1, \dots, h_n)] = \mathbf{V}$ pour tous $h_1, \dots, h_n \in H$.

Corollaire. Une forme de Skolem est vraie pour une interprétation de Herbrand si et seulement si toutes ses instances fondamentales sont vraies pour cette interprétation.

Simplification. Les interprétations de Herbrand s'identifient aux fonctions (totales) de la base de Herbrand B_H sur $\{\mathbf{V}, \mathbf{F}\}$, ou encore aux sous-ensembles de B_H , c'est-à-dire aux interprétations propositionnelles de lexique $\Pi = B_H$.

Petit théorème de Herbrand. Une formule universelle pure S est consistante si et seulement si elle admet un modèle de Herbrand.

Remarque. On voit immédiatement l'intérêt de ce théorème qui permet, lors de la recherche de modèles, de se limiter aux interprétations de Herbrand, donc à un domaine générique, unique et simple.

Démonstration. La condition est visiblement suffisante. On montre qu'elle est nécessaire en donnant une technique de transformation d'un modèle quelconque $\mathcal{I}$ (de domaine D quelconque) en un modèle de Herbrand $\mathcal{H}$ (de domaine $H = H_S$).

1. On commence par donner une fonction w qui à tout élément $h \in H$ du domaine de Herbrand H associe un élément $w(h) \in D$. Si au moins une constante apparaît dans S , toutes ces constantes sont interprétées par $\mathcal{I}$ et on pose $w(c_i) = \mathcal{I}[c_i] = I_c[c_i] \in D$; sinon, la constante arbitraire c est interprétée en un élément $d = w(a) \in D$ quelconque.

2. Pour donner une interprétation de Herbrand $\mathcal{H}$, il faut spécifier l'ensemble des atomes fondamentaux qui seront vrais dans $\mathcal{H}$.

Soient $h_1, \dots, h_n \in H$ et p un symbole prédicatif d'arité n . Pour interpréter l'atome fondamental $p(h_1, \dots, h_n)$, on pose $\mathcal{H}[p(h_1, \dots, h_n)] = I_c[p](w(h_1), \dots, w(h_n))$ ($I_c[p]$ est une fonction de D^n dans $\{\mathbf{V}, \mathbf{F}\}$).

On a donc

$$\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[p(x_1, \dots, x_n)] = \mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[p(x_1, \dots, x_n)]$$

3. Soit $\varphi(x_1, \dots, x_n)$ une matrice ne contenant aucune variable libre autre que $x_1, \dots, x_n$. On a $\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[\varphi(x_1, \dots, x_n)] = \mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[\varphi(x_1, \dots, x_n)]$
4. Toute formule de la forme $\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)$ satisfaite par $\mathcal{I}$ est aussi satisfaite par $\mathcal{H}$. On a successivement

$$\mathcal{I}[\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)] = \mathbf{V} \text{ (hypothèse),}$$

$$\mathcal{I}_{x_1/d_1, \dots, x_n/d_n}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } d_1, \dots, d_n \in D,$$

$$\mathcal{I}_{x_1/w(h_1), \dots, x_n/w(h_n)}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } h_1, \dots, h_n \in H.$$

$$\mathcal{H}_{x_1/h_1, \dots, x_n/h_n}[\varphi(x_1, \dots, x_n)] = \mathbf{V}, \text{ pour tous les } h_1, \dots, h_n \in H.$$

$$\mathcal{H}[\forall x_1 \dots \forall x_n \varphi(x_1, \dots, x_n)] = \mathbf{V}.$$

Procédures de décision Une conséquence immédiate du petit théorème de Herbrand est que, pour tester la consistance d'une forme universelle

Remarque. La théorie de Herbrand a une portée très générale; elle n'est pas restreinte au cas particulier des formes universelles pures. que ce qui vient d'être s'applique seulement aux formes de Skolem. Par exemple, la formule

$$p(a) \wedge \exists x \neg p(x)$$

est consistante, mais n'a pas de modèle de Herbrand : l'univers de Herbrand (si on le considère comme défini) serait le singleton $\{a\}$, et la formule n'admet que des modèles à deux éléments au moins. En revanche, la forme de Skolem correspondante

$$p(a) \wedge \neg p(b)$$

admet le modèle de Herbrand $\{p(a)\}$, tel que $p(a) = \mathbf{V}$ et $p(b) = \mathbf{F}$.