

Introduction à la théorie des modèles

La théorie des modèles, ou mieux, l'approche sémantique en logique, a de nombreuses applications en mathématiques, en linguistique, en philosophie, etc. Il ne peut être ici question que d'une initiation aux concepts les plus fondamentaux de ce domaine, mais on indiquera quelques implications pour les mathématiques et la philosophie de ces dernières.

Il est désormais classique, en logique mathématique, de distinguer le point de vue sémantique du point de vue syntaxique ; le premier relève de la *théorie des modèles*, tandis que le second caractérise la *théorie de la démonstration*. On peut y voir une sorte de division du travail : la théorie de la démonstration définit des théories formelles, et la théorie des modèles en donne des *interprétations*.

Pour mieux comprendre le rôle de chacune d'elles, il convient de préciser la nature des théories formelles. Il est commode, pour l'expliquer, de rappeler la manière dont Pascal, dans *De l'esprit géométrique*, justifiait la méthode axiomatique. Celle-ci était en effet présentée comme un succédané de la « méthode idéale » ; car la *méthode idéale*, dans les sciences, consisterait à ne jamais utiliser un terme sans le définir, ni avancer une proposition sans la démontrer. Or une telle méthode est inapplicable en fait, puisqu'elle conduirait à une régression à l'infini : un terme est défini à partir d'autres termes, et une proposition démontrée à partir d'autres propositions.

De là la nécessité de recourir à la méthode axiomatique ; les termes définis le sont à partir de *termes primitifs* (non définis) et les théorèmes sont démontrés à partir d'*axiomes* (propositions non démontrées). Tout ceci est bien connu. Mais pour comprendre comment on en vient à la théorie formelle, il convient de souligner qu'en passant de la méthode idéale à la méthode axiomatique nous avons dû non seulement renoncer à une forme de perfection, mais surtout nous avons changé notre conception de la *rigueur*. Dans l'esprit de la méthode idéale, en effet, être rigoureux c'était n'avoir aucun *présupposé*, tous les éléments intervenant dans la théorie ayant été justifiés. Mais la méthode axiomatique n'exclut pas les présupposés ; et elle admet que la rigueur ne consiste pas à ne pas en avoir, mais à les *explicit*er tous. C'est pourquoi on a pu reprocher à Euclide des « lacunes » dans ses démonstrations, *i.e.* le recours à des hypothèses qui n'étaient pas justifiables à partir de sa liste d'axiomes et de postulats.

Mais de ce point de vue, il manquait quelque chose d'essentiel à la méthode axiomatique. Il ne suffit pas en effet de donner une liste des termes primitifs et des axiomes pour expliciter tous les présupposés d'une théorie. Car pour dériver les théorèmes à partir des axiomes, il faut encore utiliser des règles d'inférence. La

rigueur propre à la méthode axiomatique exigerait donc que l'on explicite également la totalité des règles de déduction qui sont utilisées dans les preuves.

En tenant compte de cette remarque, on peut distinguer trois *niveaux de formalisation* d'une théorie :

Niveau 0 : il correspond à ce que l'on appelle communément la *théorie intuitive* ; pour développer la théorie, le mathématicien se base uniquement sur l'« intuition » qu'il a des objets dont elle parle. À ce niveau, il n'y a pas vraiment formalisation, même s'il peut y avoir symbolisation¹. C'est le cas de la géométrie pré-euclidienne, de l'arithmétique pré-peanienne (ou pré-dedekindienne), et de la théorie cantorienne des ensembles.

Niveau 1 : c'est la théorie axiomatique ; le mathématicien ne justifie plus ses raisonnements sur la base de ses intuitions, mais se réfère à des axiomes. C'est le cas de la géométrie euclidienne, ou de sa reconstruction hilbertienne (par exemple) ; c'est aussi celui des axiomes de Peano pour l'arithmétique, et de la théorie axiomatique « naïve » des ensembles formulée par Zermelo, Fraenkel et d'autres.

Niveau 2 : la *théorie axiomatique formelle* ; pour définir celle-ci, on explicite à la fois les axiomes « mathématiques » (ou non-logiques) au sens habituel, mais aussi les règles d'inférence, qui peuvent être considérées comme des *axiomes logiques*. Comme nous le verrons plus bas, il est facile de donner des exemples de théories axiomatiques formelles à partir des théories illustrant les niveaux inférieurs : il suffit de considérer leurs traductions dans le premier ordre. Signalons en particulier que ZFC, la théorie des ensembles de Zermelo-Fraenkel exprimée dans le premier ordre, est une théorie qui permet de construire 99 % des mathématiques usuelles. C'est ainsi la presque totalité des mathématiques qui se trouve être formalisée.

Comment les règles d'inférence peuvent-elles être explicitées ? Comment sait-on qu'on les a toutes données et qu'il n'en manque aucune à notre liste ? On doit évidemment disposer de deux notions de déductibilité. Supposons en effet donnée une liste de règles d'inférence formant nos axiomes logiques. Si l'on note,

¹ Le fait d'utiliser des notations permettant d'abrégier le discours n'implique pas en lui-même que la formalisation ait débuté, et cela en dépit du fait que de telles notations, une fois créées, peuvent largement faciliter le travail futur de formalisation.

comme on le fait classiquement, $\Phi \vdash \varphi$ le fait que l'énoncé φ (du premier ordre) est déductible à partir de l'ensemble d'énoncés Φ au moyen des règles d'inférence, nous aurons besoin d'une notion de conséquence indépendante de ces règles afin de pouvoir établir :

$$\Phi \vdash \varphi \Leftrightarrow \varphi \text{ est une conséquence de } \Phi \quad (1)$$

Dans le sens « $\Rightarrow$ », on parlera du « *théorème de fiabilité* » (les règles d'inférence sont fiables en ce sens que chaque fois qu'elles permettent de dériver un énoncé à partir d'autres énoncés, le premier est conséquence des seconds) ; dans le sens « $\Leftarrow$ », on aura le *théorème de complétude* (les règles d'inférence sont complètes en ce sens qu'elles permettent de dériver toutes les conséquences d'un ensemble d'énoncés donné).

Une telle équivalence est effectivement obtenue pour les théories du premier ordre, c'est-à-dire les théories exprimées dans des langages du premier ordre. Ceux-ci sont des langages *formels*. C'est-à-dire des langages qui ne sont d'abord définis que par leur syntaxe, sans aucune référence à la signification de leurs expressions. Pour définir un langage formel, on donne un alphabet, c'est-à-dire un ensemble de *symboles*, et des règles permettant de construire les « expressions bien formées » (EBF), celles qui sont conformes aux règles syntaxiques.

L'alphabet d'un langage du premier ordre comprend deux types de symboles : les *symboles logiques*, qui sont communs à tous les langages du premier ordre ; et les *symboles non-logiques*, qui varient d'un langage à l'autre. Les symboles logiques sont les *variables d'individu* : $x_0, x_1, x_2, \dots$ (l'ensemble $\mathcal{V}$ des variables est dénombrable) ; les *connecteurs logiques* : $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$; les *symboles de quantification* : $\forall, \exists$; le *symbole d'égalité* : $=$; et les *symboles de parenthésage* : $), ($

Les *symboles non-logiques* contiennent un ensemble possiblement vide de *constantes d'individu* ; pour tout $n \geq 1$, un ensemble possiblement vide de *symboles de prédicat n -aires* ; et pour tout $n \geq 1$, un ensemble possiblement vide de *symboles de fonction n -aires*. Les symboles logiques étant communs à tous les langages du premier ordre, on désigne un langage L par l'ensemble de ses symboles non-logiques.

Si L est un langage du premier ordre, une *expression* de L (ou une *L -expression*) est une suite finie de symboles de L . On définit deux types d'EBF : les *termes* et les *formules*. Lorsque ces expressions seront interprétées, les premiers désigneront des individus du domaine de discours, tandis que les formules seront des propositions. Les termes sont les variables et les constantes d'individu, ainsi que

toutes les expressions de la forme $f t_1 \dots t_n$, où f est un symbole de fonction et les t_i sont des termes. Les formules sont construites à partir des formules atomiques (de la forme $t_1 = t_2$ ou $R t_1 \dots t_n$, où R est un prédicat n -naire) en utilisant les connecteurs et les quantificateurs².

Pour illustrer ces notions, considérons la théorie qui nous servira de fil conducteur par la suite : l'arithmétique de Peano. Appelons structure de Dedekind un triplet $\langle X, f, x_0 \rangle$ où X est un ensemble, $f: X \rightarrow X - \{x_0\}$ est une bijection, et pour toute propriété P , si x_0 a la propriété P et si $f(x)$ a la propriété P chaque fois que x l'a, alors tous les éléments de X ont cette propriété. Évidemment, $\langle \mathbb{N}, s, 0 \rangle$ (où s est la fonction successeur) est une structure de Dedekind. En fait, celui-ci a montré que deux structures de Dedekind quelconques sont isomorphes. Autrement dit, ces conditions définissent $\mathbb{N}$ à l'isomorphisme près.

Or ces conditions sont précisément celles qui sont exprimées dans les célèbres axiomes de Peano (qui mériteraient d'ailleurs davantage le nom d'axiomes de Dedekind). En utilisant des notations symboliques, on serait tenté de les formuler ainsi :

- A1. $\forall x \neg s(x) = 0$
- A2. $\forall x \forall y (s(x) = s(y) \rightarrow x = y)$
- A3. $\forall x (x \neq 0 \rightarrow \exists y s(y) = x)$
- A4. $\forall P [(P(0) \wedge \forall x (P(x) \rightarrow P(s(x)))) \rightarrow \forall x P(x)]$

Mais A4 (le principe de récurrence) n'est pas un énoncé du premier ordre. Dans le premier ordre, en effet, on ne quantifie que sur des variables d'individu, non sur des propriétés. La quantification sur les propriétés (des individus) est du second ordre ; or la logique du second ordre n'est pas axiomatisable. Ce qui signifie que (1) n'est pas vrai pour des formules du second ordre.

Il est néanmoins possible d'exprimer le principe de récurrence dans le premier ordre en remplaçant l'axiome A4 par un schéma d'axiomes :

$$A4'. [\varphi(0) \wedge \forall x (\varphi(x) \rightarrow \varphi(s(x)))] \rightarrow \forall x \varphi(x)$$

Pour chaque formule φ du langage $L = \{0, s\}$ (du premier ordre), $A4'$ est un axiome. On donne donc par ce schéma une infinité (dénombrable) d'axiomes.

Un modèle des axiomes A1-A3 et A4' est une structure $\langle A, f, a \rangle$, où A est un ensemble, f est une fonction définie sur A et $a \in A$, qui satisfait tous ces axiomes

² En fait, nous nous intéresserons surtout à des énoncés, c'est-à-dire des formules sans variable libre.

(i.e. dans laquelle ils sont tous vrais). La notion de conséquence qui intervient dans (1) peut être définie à partir de la notion de modèle : un énoncé de L est une conséquence des axiomes s'il est vrai dans tous les modèles de ces axiomes. C'est une façon « technique » d'exprimer la notion intuitive de conséquence : une formule est conséquence des axiomes si elle est vraie toutes les fois que les axiomes sont vraies.

Comme l'indique par ailleurs (1), dans le premier ordre, il est possible de définir un ensemble fini de règles d'inférence permettant de dériver toutes les conséquences (et uniquement les conséquences) des axiomes. Comme exemple de telles règles, nous donnons celles qui concernent la négation, l'implication et le quantificateur universel. Dans ces deux derniers cas, une règle permet d'introduire l'opérateur logique, et une autre permet de l'éliminer

Règles d'introduction

Implication :

[A]³

B

$A \rightarrow B$

Quantification universelle :

A(a)

$\forall x A(x)$

si a ne figure ni dans la conclusion ni dans les prémisses non déchargées.

Négation :

[A]

$\perp$

$\neg A$

$\perp$ (EFSQ)

A

Règles d'élimination

A $A \rightarrow B$

B

$\forall x A(x)$

A(t)

A $\neg A$

$\perp$

$\neg A$

$\perp$ (RAA)

A

³ Dans ces règles, les crochets signifient que l'hypothèse en question est « déchargée » lors de l'application de la règle, ce qui veut dire que la dérivation ne dépend plus de l'hypothèse. Par exemple, lorsque l'on a démontré B en supposant A, on a démontré $A \rightarrow B$ indépendamment de l'hypothèse A.

Une *dérivation* est un arbre (finitaire⁴ et de hauteur finie) obtenu en appliquant les règles données. Elles sont les correspondants formels des démonstrations. Les feuilles de l'arbre sont les prémisses (déchargées ou non) ; la racine est la conclusion.

En exprimant l'arithmétique de Peano dans le premier ordre, on a donc gagné la complétude : toute conséquence des axiomes est démontrable à partir d'un nombre fini de règles logiques, qui peuvent être ainsi explicitement énoncées. On a de cette façon achevé l'axiomatisation jusqu'au deuxième niveau de formalisation. Mais ce n'est pas sans certains sacrifices. On peut en effet tirer du théorème de complétude le résultat suivant :

Théorème de compacité : si tout sous-ensemble fini d'un ensemble Γ d'énoncés a un modèle, alors Γ a un modèle.

Du théorème de compacité, on peut démontrer que la théorie T dont les axiomes sont $A1$ - $A3$ et $A4'$ a des modèles qui ne sont pas isomorphes au modèle standard $\langle \mathbb{N}, s, 0 \rangle$. En effet, si l'on ajoute à L une constante c et aux axiomes les énoncés :

$c \neq 0$
 $c \neq s(0)$
 $c \neq s(s(0))$
...

la théorie ainsi obtenue est telle que tous ses sous-ensembles finis ont un modèle ; d'où il suit que la théorie elle-même a un modèle. Celui-ci est également modèle de T , mais il n'est évidemment pas isomorphe au modèle « standard » $\langle \mathbb{N}, s, 0 \rangle$, puisque l'interprétation de c est un élément distinct de $0, 1, 2, 3, \dots$. Contrairement à la définition dans le second ordre, celle du premier ordre ne caractérise donc pas les entiers à l'isomorphisme près (on dit que la théorie n'est pas *catégorique*, *i.e.* ses modèles ne sont pas tous isomorphes⁵).

Or, ne l'oublions pas, un énoncé n'est démontrable dans une théorie du premier ordre que s'il est vrai dans tous les modèles de cette théorie ; on peut donc s'attendre à ce que des énoncés vrais dans le modèle standard ne soient pas démontrables, simplement parce qu'ils ne sont pas vrais dans les modèles non

⁴ Un arbre est finitaire lorsque chaque noeud n'a qu'un nombre fini de successeurs.

⁵ Ce résultat n'a rien d'étonnant si l'on a recours à un argument de cardinalité : du point de vue ensembliste, il y a autant de propriétés P des entiers qu'il y a de sous-ensembles de $\mathbb{N}$, c'est-à-dire $2^{\aleph_0}$; l'ensemble des formules est en revanche dénombrable : le principe de récurrence du premier ordre est plus faible que celui du second ordre.

standard. Cette situation est clairement mise en évidence dans le premier théorème d'incomplétude de Gödel⁶ ; d'après la preuve de ce théorème, il existe des formules $\varphi(x)$ telles que pour tout n , $\varphi(n)$ est démontrable (ici, n désigne le terme $ss\dots s(0)$ (n fois)), mais $\forall x\varphi(x)$ n'est pas démontrable (on dit que T est ω -incomplet).

Il n'est pas possible de remédier à cette situation en recourant au second ordre, puisque la logique du second ordre n'étant pas axiomatisable, la catégoricité de la théorie $A1-A4$ ne se traduirait pas par une puissance démonstrative plus grande. De plus, les propriétés P auxquelles se réfère le principe de récurrence exprimé dans le second ordre sont en fait des ensembles : n'a-t-on pas introduit dans cette quantification des hypothèses non explicitées (contrairement à notre conception de la rigueur), hypothèses constituant les axiomes de la théorie des ensembles⁷ ? On peut néanmoins recourir à de nouvelles règles d'inférence. Par exemple, on obtient la ω -logique en ajoutant aux règles d'inférence du premier ordre la ω -règle :

$$\frac{\varphi(0), \varphi(1), \dots, \varphi(n), \dots}{\forall x\varphi(x)}$$

Non seulement cette règle permet évidemment de supprimer l' ω -incomplétude, mais encore, on démontre qu'un énoncé est démontrable dans cette logique ssi il est vrai dans le modèle standard. On remarquera néanmoins que l'on a renoncé au caractère finitaire des preuves : la règle suppose une infinité de prémisses.

Les logiques plus fortes que la logique du premier ordre peuvent assurément avoir des avantages que n'a pas celle-ci, et des propriétés particulièrement attractives. Mais ne négligeons pas ce que peuvent aussi nous apporter les « limites » mêmes du premier ordre : les modèles non standard de l'arithmétique sont aussi des structures mathématiques qui sont intéressantes en elles-mêmes ; et le même argument qui a permis de le construire, appliqué à la théorie des réels, a mis en évidence l'existence du modèle qui sert aujourd'hui d'univers de base de l'analyse non standard.

⁶ Il est désormais classique de considérer dans ce contexte une théorie plus riche que celle que nous avons présentée : on ajoute au langage des symboles de fonction pour l'addition et la multiplication, et aux axiomes, les définitions récursives de ces opérations.

⁷ C'est notamment l'affirmation de Quine : la logique du second ordre est un déguisement de la théorie des ensembles.