

Jean Ladrière

Le théorème de Löwenheim-Skolem

L'une des raisons pour lesquelles on s'est efforcé de formaliser les mathématiques est que la notion d'infini, essentielle pour la pensée mathématique, est loin de présenter la même clarté et les mêmes garanties que les concepts propres aux mathématiques finies. On espérait, grâce aux ressources du langage formel, obtenir un contrôle indirect de la notion d'infini en représentant les raisonnements infinitistes dans des systèmes d'expression accessibles à des raisonnements finitistes, et ramener ainsi le domaine infini dans l'aire de sécurité de la pensée finie. L'examen attentif des langages formels a montré, de diverses manières, qu'il y a une sorte d'incapacité congénitale du formalisme à fournir une représentation adéquate des spéculations infinitistes. Le théorème de Löwenheim-Skolem fait apparaître un des aspects de cette inadéquation constitutive des systèmes formels. Pour en faire voir toute la portée, il faut le replacer dans le contexte général de l'étude des rapports entre les systèmes formels et leurs modèles.

*

Un *langage* peut être considéré comme un ensemble de règles permettant de construire, à partir de symboles élémentaires donnés, des expressions de complexité croissante et opérant une répartition de ces expressions en catégories caractérisées par leurs propriétés syntaxiques (c'est-à-dire par les modalités de leur usage dans la construction des expressions). Parmi les catégories syntaxiques d'un langage, il en est une qui occupe une situation privilégiée parce que, envisagée du point de vue sémantique, elle est faite d'expressions qui présentent un caractère de fermeture minimale : c'est celle des *propositions*. Une proposition exprime la plus petite unité de sens que l'on peut considérer comme fermée sur elle-même : elle peut être affectée par les prédicats sémantiques que l'on désigne communément par les termes « vrai » et « faux ». Un langage est dit *formalisé* lorsqu'il se présente sous forme de règles énoncées complètement, sans ambiguïté, et répondant à des critères précis d'effectivité. La notion d'effectivité peut être représentée elle-même sous forme de manipulations formelles, mais on peut dire, intuitivement, qu'elle correspond à des opérations qui se déroulent

selon un schéma canonique et peuvent s'achever en un temps fini, du type de celles que pourrait réaliser une machine. Ainsi il faut pouvoir reconnaître de façon effective si un symbole donné fait ou non partie des symboles du langage étudié, si telle expression est bien une expression de ce langage, à quelle catégorie syntaxique elle appartient, comment elle a été construite, etc.

Une *théorie* est une classe particulière de propositions d'un langage, considérées comme vraies. (Un même langage peut évidemment contenir plusieurs théories distinctes, et même éventuellement une infinité de théories.) Une théorie peut être dite *formalisée* lorsque le critère d'appartenance qui la caractérise peut être décrit au moyen de règles complètes, dépourvues d'ambiguïté et effectives. Bien entendu, il n'est possible de formaliser une théorie que dans le cadre d'un langage lui-même formalisé. Pratiquement, on a été amené, en vue de formaliser les théories que l'on s'est proposé (jusqu'ici) d'étudier de façon stricte, à construire des langages artificiels, relativement pauvres mais en principe entièrement contrôlables. On s'est d'abord préoccupé de construire des langages logiques purs, c'est-à-dire des langages permettant de représenter, en quelque sorte à vide, les procédés connus de raisonnement (tels qu'ils sont utilisés, par exemple, en mathématiques). Nous aurons à nous occuper dans ce qui suit d'un tel langage, celui de la logique des prédicats du premier ordre.

Il y a deux manières de formaliser une théorie, l'une qui est en quelque sorte extrinsèque et l'autre qui est purement intrinsèque. La méthode extrinsèque consiste à décrire la classe des propositions formant la théorie en suivant pas à pas le processus de construction de ces propositions. Pratiquement, on peut fournir une telle description en se référant à un certain domaine d'objets, supposé donné, muni d'une structure, c'est-à-dire d'une certaine collection de relations. La méthode intrinsèque consiste à isoler, au sein de la théorie, une classe particulière de propositions, appelées axiomes, que l'on peut énumérer effectivement, et à fournir des *règles* au moyen desquelles on peut obtenir toutes les propositions de la théorie au moyen de celles-là. Une théorie présentée sous cette forme est ordinairement appelée un *système formel*. La première méthode peut être baptisée *méthode sémantique*, la seconde *méthode syntaxique*. Un des problèmes importants que pose l'étude des formalismes est de comparer ces deux modes de caractérisation : une théorie étant décrite sémantiquement, on peut se demander comment l'axiomatiser, c'est-à-dire comment la décrire syntaxiquement — et inversement, une théorie étant donnée sous forme axiomatique, on peut se proposer d'en étudier la sémantique.

Illustrons ceci par le cas de la *logique des prédicats du premier ordre*, que concerne précisément le théorème de Löwenheim-Skolem. Nous désignerons dans la suite cette logique par le sigle L.P. Cette théorie a pour but de caractériser les formes de raisonnement que l'on peut effectuer au moyen des opérations élémentaires portant sur des propositions (négation, conjonction, disjonction, implication, équivalence) et des opérations de

quantification portant sur les prédicats de premier niveau (prédicats d'individus).

Le langage dans lequel est formulée la théorie LP est extrêmement simple. Il comporte une liste infinie de symboles représentant des *variables individuelles* (variables pouvant être remplacées par des noms d'individus), une autre liste infinie de symboles représentant des prédicats d'individus, les symboles désignant les opérations logiques élémentaires, les symboles de quantification, et des symboles séparatifs (par exemple des parenthèses). Les prédicats peuvent s'appliquer à un ou plusieurs arguments : ils correspondent donc soit à des propriétés d'individus soit à des relations entre deux ou plusieurs individus. En ce qui concerne les opérations logiques élémentaires, on peut se limiter, par exemple, à la négation et à l'implication, car les autres opérations peuvent être définies au moyen de celles-là. Des règles précisent quels sont les groupes de symboles qui peuvent être considérés comme des propositions. Une expression du type $P(x_1, x_2, \dots, x_n)$, où P est un symbole désignant un prédicat à n arguments et $x_1, x_2, \dots, x_n$ sont des symboles désignant des variables individuelles, est une proposition. (Cette expression signifie : *Le prédicat P est vérifié du n -uplé d'individus $x_1, x_2, \dots, x_n$. Ou encore : Les individus $x_1, x_2, \dots, x_n$ ont entre eux la relation exprimée par le prédicat P .*) Si A est une proposition, $\sim A$ (non- A) est une proposition. Si A et B sont des propositions, $A \supset B$ (A implique B) est une proposition. Si $A(x)$ est une proposition contenant la variable individuelle x , $(x)A(x)$ (pour tout x , la proposition $A(x)$ est valable, tout x a la propriété exprimée par A) et $(Ex)A(x)$ (il y a un x au moins pour lequel la proposition $A(x)$ est valable, il y a un x qui possède la propriété exprimée par A) sont des propositions. (Dans l'expression $A(x)$, A peut désigner un prédicat simple, mais peut être aussi une expression complexe, contenant d'autres variables que x et jouant le rôle d'un prédicat complexe à l'égard de x .)

La théorie LP peut être caractérisée sous forme syntaxique, c'est-à-dire présentée sous la forme d'un système axiomatique d'une manière assez simple. Les axiomes sont décrits au moyen de quelques schémas de construction¹. Ils constituent une classe infinie de propositions, contenant toutes les propositions formées selon ces schémas. Les règles de déduction sont au nombre de deux. L'une est le *modus ponens* : Si $A \supset B$ et A sont des théorèmes (c'est-à-dire soit des axiomes soit des propositions déjà dérivées des axiomes), B est aussi un théorème. L'autre est la règle de la généralisation : Si $A(x)$ est un théorème, $(x)A(x)$ est aussi un théorème². La théorie LP

1. Dans son *Introduction to mathematical logic*, par exemple, Church utilise cinq schémas d'axiomes. Trois d'entre eux correspondent à la logique des propositions et expriment certaines propriétés des opérations de négation et d'implication. Les deux autres correspondent à certaines propriétés du quantificateur universel (*Pour tout x ...*), en liaison avec l'opération d'implication. Il y a moyen de présenter les axiomes de façon directe, à condition d'utiliser une règle appropriée de substitution. Mais nous n'avons pas à nous embarrasser ici de ces détails techniques.

2. Remarquons qu'il n'est pas nécessaire de prévoir une règle pour le maniement du quantificateur existentiel (*Il y a un x tel que...*) car celui-ci peut être défini au moyen du quantificateur universel : $(Ex)A(x)$ est équivalent par définition à $\sim (x) \sim A(x)$. Les deux règles mentionnées ici sont celles de

est formée par l'ensemble des propositions que l'on peut déduire de la classe des axiomes au moyen de ces deux règles.

Mais la théorie LP peut aussi être caractérisée d'une manière sémantique, au moyen de ce qu'on appelle une *interprétation*. Soit L le langage dans lequel est formulée la théorie LP. Une interprétation est un procédé qui permet d'associer à chaque proposition de L soit la propriété « vrai » soit la propriété « faux ». Pour construire une interprétation, on se donne d'abord un certain domaine D d'objets dont les membres seront considérés comme des individus. On peut prendre par exemple pour D l'ensemble des nombres entiers. Il est possible alors d'associer aux propositions de L des énoncés relatifs au domaine D. Concrètement, une telle association sera établie au moyen de deux applications, c'est-à-dire de deux correspondances fonctionnelles (associant à un objet de l'ensemble de départ un objet et un seul de l'ensemble d'arrivée).

Appelons la première application *assignation individuelle* : une telle application associe à chaque variable individuelle de L un objet (et un seul) du domaine D. Appelons la seconde application *assignation prédicative* : une telle application associe à chaque variable prédicative à n arguments de L une relation (et une seule) définie sur D, c'est-à-dire un certain ensemble de n -uples formés d'objets de D. (Ainsi, si on prend pour D l'ensemble des entiers, on pourra associer à un prédicat P à deux arguments de L une relation binaire entre entiers, par exemple « double de ». Une telle relation peut être considérée comme l'ensemble des paires d'entiers dont le premier élément est le double du second : $(2, 1)$, $(4, 2)$, $(6, 3)$, etc.)

Une assignation individuelle et une assignation prédicative étant données, à toute proposition élémentaire de L se trouve *ipso facto* associé un énoncé relatif au domaine D. (Soit par exemple la proposition $P(x_1, x_2)$. Si l'assignation individuelle choisie fait correspondre à x_1 l'objet 4 et à x_2 l'objet 2, et si l'assignation prédicative choisie fait correspondre au prédicat P la relation « double de », notre proposition sera associée à l'énoncé « Quatre est le double de deux ».) La connaissance du domaine D nous permet de déterminer si un énoncé ainsi associé à une proposition de L est vrai ou faux.

Une interprétation du langage sera définie par les règles suivantes, qui suivent les règles de formation des propositions de L.

a) Si une proposition est de forme $P(x_1, x_2, \dots, x_n)$, elle est *vraie* ou *fausse* pour une double assignation (une assignation individuelle jointe à une assignation prédicative) donnée, suivant que cette double assignation lui associe un énoncé vrai ou faux. (Ainsi, dans notre exemple, la proposition $P(x_1, x_2)$ devait être déclarée *vraie*).

b) Une proposition de forme $\sim A$ est *vraie* ou *fausse* suivant que la proposition A est *fausse* ou *vraie*.

l'un des systèmes décrits par Church sous le titre *Calcul fonctionnel du 1^{er} ordre* dans l'ouvrage mentionné plus haut. Il existe d'autres formulations de la théorie LP, mais ces variantes n'ont pas d'intérêt pour notre propos.

c) Une proposition de forme $A \supset B$ est *vraie*, sauf si la proposition A est *vraie* et la proposition B *fausse* (auquel cas elle est *fausse*).

d) Une proposition de forme $(x)A(x)$ est *vraie* pour une assignation prédicative donnée si, pour cette assignation, la proposition $A(x)$ est *vraie* quelle que soit l'assignation individuelle choisie (c'est-à-dire quel que soit l'objet de D associé à la variable x). Elle est *fausse*, si, pour une assignation individuelle au moins, la proposition $A(x)$ est *fausse*. (Ainsi, D étant l'ensemble des entiers, si on associe à A la propriété « être un nombre premier », la proposition $(x)A(x)$ est *fausse*, car, pour une assignation individuelle qui associe à x le nombre 10 par exemple, la proposition $A(x)$, associée à l'énoncé faux « Dix est un nombre premier », doit être déclarée *fausse*.)

Une proposition est dite *valide dans un domaine donné* si elle est *vraie* pour toutes les assignations possibles (à la fois individuelles et prédictives) relativement à ce domaine, autrement dit si elle est *vraie* dans toutes les interprétations relatives à ce domaine. Elle est dite *exemplifiable dans un domaine donné* si elle est *vraie* pour une assignation individuelle et une assignation prédicative au moins relativement à ce domaine, autrement dit si elle est *vraie* dans une interprétation au moins relative à ce domaine. Elle est dite *valide* si elle est *valide* dans tous les domaines possibles et elle est dite *exemplifiable* si elle est *exemplifiable* dans un domaine au moins.

Étant donné une classe de propositions, on dit qu'elle *admet le domaine D comme modèle* (pour une interprétation donnée) si toutes les propositions de cette classe sont simultanément *exemplifiables* dans ce domaine, c'est-à-dire si, pour l'interprétation en question, elles deviennent toutes *vraies*. Comme une théorie constitue une classe de propositions, une théorie admet un modèle s'il existe un domaine dans lequel toutes les propositions de cette théorie sont simultanément *exemplifiables*. Fournir un modèle d'une théorie déterminée, c'est en somme se donner l'image de cette théorie sous la forme d'une certaine structure (le domaine muni des relations associées aux prédicats de la théorie), autrement dit c'est fournir une sorte de réalisation concrète de la théorie.

Ces notions étant fixées, nous pouvons examiner comment se présentent les relations entre la représentation sémantique et la représentation syntaxique de LP. Au point de vue sémantique, la théorie LP sera caractérisée comme l'ensemble des propositions de L qui sont *valides*, c'est-à-dire qui sont *vraies* pour toutes les interprétations possibles relativement à tous les domaines possibles. C'est là une manière de dire que les propositions de LP sont purement formelles : elles représentent des situations qui peuvent se réaliser dans tous les univers possibles. Au point de vue syntaxique, la théorie LP est caractérisée, comme on l'a vu plus haut, par un ensemble d'axiomes et deux règles.

La première question qui se pose est de savoir si la représentation syntaxique d'une théorie correspond de façon adéquate à sa représentation

sémantique. Cette question a reçu une réponse affirmative, sous la forme de deux métathéorèmes, dont le premier peut être appelé *théorème de validité* et dont le second est le célèbre *théorème d'adéquation* dû à Gödel. Le théorème de validité s'énonce comme suit : *Tout théorème de LP est une proposition valide de L*. La démonstration de ce métathéorème est fort simple. Elle consiste à montrer que les axiomes de LP sont des propositions valides, et à montrer ensuite que les règles de déduction de LP conservent la validité (c'est-à-dire fournissent des conclusions valides lorsque leurs prémisses le sont); aucune de ces tâches ne présente de difficulté. Le théorème d'adéquation s'énonce comme suit : *Toute proposition valide de L est un théorème de LP*. Ce métathéorème, dont la démonstration est beaucoup plus compliquée que celle du précédent, signifie que l'axiomatique de la logique des prédicats est formulée de telle sorte qu'elle épuise effectivement toutes les formes de raisonnement admissibles (moyennant une interprétation « naturelle » des opérations logiques élémentaires et des quantificateurs, contenue dans les règles énoncées plus haut).

Le théorème d'adéquation, tel qu'il vient d'être formulé, peut être considéré comme une conséquence d'un métathéorème qu'on pourrait appeler le *métathéorème de cohérence*, qui contient du reste l'essentiel du théorème d'adéquation. On dit qu'une classe K de propositions est *cohérente* lorsqu'il n'existe aucune proposition A telle que l'on puisse déduire de la classe K à la fois A et $\sim A$. (Une proposition est dite *déductible d'une classe K de propositions de L* s'il est possible de déduire cette proposition des axiomes de LP, enrichis des propositions de la classe K, au moyen des règles de déduction de LP.) Le métathéorème de cohérence s'énonce comme suit : *La condition nécessaire et suffisante pour qu'une classe de propositions de L soit cohérente est qu'elle admette (au moins) un modèle*. Ce théorème comprend deux parties. La première partie — *Si une classe de propositions de L admet un modèle, elle est cohérente* — peut être établie très simplement. La deuxième partie — *Si une classe K de propositions de L est cohérente, elle admet un modèle* — ne peut être établie que moyennant un raisonnement complexe qui fait apparaître toute la portée du théorème d'adéquation.

La démonstration consiste évidemment à indiquer comment on peut construire un modèle pour la classe K de propositions considérée³. La difficulté principale que l'on doit surmonter est que la classe K est en général trop restreinte pour qu'on puisse en tirer des informations suffisantes pour la construction directe d'un modèle admissible. Il faudra donc commencer par l'étendre de façon convenable. De façon précise, il faudra prévoir une extension qui réponde à deux conditions : elle devra être *complète* (c'est-à-dire que toute proposition du langage L ou sa négation devra y figurer)

3. L'esquisse de démonstration donnée ici s'inspire de l'exposé donné par Church du théorème de cohérence dans *Introduction to mathematical logic*. Mais la démonstration de Church fait intervenir directement un domaine dénombrable. On a supposé ici que le domaine D est infini sans autre précision.

et elle devra comporter au moins une exemplification de chacune des propositions existentielles qu'elle contient. (Supposons par exemple que la classe K contienne l'existentielle $(Ex)A(x)$. Notre extension devra contenir une proposition du genre $A(c)$ où c est un terme désignant un individu. La proposition $A(c)$ (*L'individu c possède la propriété A*) peut être considérée comme une exemplification de la proposition générale $(Ex)A(x)$ (*il y a un x qui possède la propriété A*).)

Pour réaliser la première condition, on élargit la classe K donnée en une *classe cohérente maximale*. Une classe cohérente maximale est une classe cohérente qui contient toutes les propositions de L compatibles avec elle (c'est-à-dire telles que leur appartenance à la classe en question ne rend pas celle-ci non-cohérente). Une classe cohérente maximale répond effectivement à la condition posée : si A est une proposition de L qui ne fait pas partie de cette classe, c'est qu'elle est incompatible avec elle — mais alors $\sim A$ est compatible avec elle (comme on peut le montrer par un raisonnement simple en se servant de certains théorèmes de LP) et donc en fait partie.

Il est possible d'ordonner les propositions de L , en rangeant d'abord les symboles de L dans un certain ordre et en déterminant, à partir de là, un ordre pour les propositions. (On peut aussi utiliser l'axiome du choix, dont il sera question plus loin.) La classe K étant donnée, il suffit de parcourir toutes les propositions de L , dans l'ordre où elles sont ainsi classées, et d'examiner, pour chacune d'elles, si elle est ou non compatible avec la classe K' déjà obtenue au moment où on procède à cet examen. Si la réponse est affirmative, on ajoute cette proposition à la classe K' et on passe à l'examen de la proposition suivante; sinon on reprend la classe K' telle quelle et on passe à l'examen de la proposition suivante. On obtient ainsi une suite infinie de classes dont chacune contient la précédente. On prend alors l'union de toutes ces classes : c'est une classe K^* définie par la condition qu'une proposition A fait partie de K^* si et seulement si elle fait partie de l'une des classes de notre suite infinie. On montre sans peine que, si K est cohérente, K^* est aussi cohérente. Et il résulte de la construction de K^* que c'est une classe cohérente maximale.

Pour réaliser la deuxième condition, on doit enrichir la langue L en lui ajoutant des constantes individuelles. Une constante individuelle est une expression qui représente un individu déterminé (à la différence d'une variable individuelle, qui représente seulement la place possible d'un individu éventuel). Considérons une classe de propositions K qui contient une proposition existentielle $(Ex)A(x)$. Nous ajoutons à L un nouveau symbole, c , qui joue le rôle d'une constante individuelle et nous étendons les règles de construction des propositions de telle sorte que c puisse servir d'argument à un prédicat. (Ainsi, si P est un prédicat à un argument, $P(c)$ sera une proposition, affirmant de l'individu c la propriété P .) Nous enrichissons la classe K de la proposition $A(c)$. En procédant de la même manière pour toutes les propositions existentielles de K , nous obtenon-

une nouvelle classe K' , qui contient K et qui répond à la condition posée : K' contient, pour chaque existentielle, une exemplification correspondante. On montre facilement que, si la classe K est cohérente, la classe K' est aussi cohérente.

Pour démontrer le théorème de cohérence, il faut réaliser ces deux conditions simultanément. Mais on ne peut arriver à ce résultat d'un seul coup ; on doit passer par une suite infinie d'étapes, en suivant une sorte de parcours en zig-zags (où l'on saute alternativement d'une condition à l'autre). Soit K_0 une classe cohérente de propositions de L . Première étape : nous immergeons K_0 dans une classe cohérente maximale K^*_0 . Ainsi nous avons réalisé la première condition. Deuxième étape : nous enrichissons le langage L de façon à associer à toutes les propositions existentielles de K^*_0 une exemplification. Nous obtenons ainsi un langage L_1 et une classe K_1 qui répond à notre seconde condition. Cette classe K_1 est cohérente. Mais, en général, elle ne répondra plus à la première condition relativement au langage L_1 . Nous devons donc agir avec K_1 , relativement à L_1 , comme nous avons agi avec K_0 relativement à L . Nous allons donc immerger K_1 dans une classe cohérente maximale K^*_1 . Puis nous allons enrichir le langage L_1 de façon à associer à toutes les propositions de K^*_1 une exemplification. Nous obtenons ainsi un langage L_2 et une classe K_2 qui répond à notre seconde condition. Et ainsi de suite, indéfiniment. Il nous suffit maintenant de construire un langage L' en prenant l'union des langages L , L_1 , L_2 , etc., et de construire une classe K' en prenant l'union des classes K^*_0 , K^*_1 , etc. La classe K' est cohérente relativement à L' et elle répond à nos deux conditions : elle est complète et elle contient au moins une exemplification de toutes les propositions existentielles qui en font partie.

Nous pouvons maintenant construire un modèle pour la classe K' . Il suffit pour cela de faire choix d'un domaine infini D et de décrire une interprétation pour laquelle toutes les propositions de K' sont *vraies* relativement à ce domaine. Nous pouvons prendre comme domaine D un ensemble infini quelconque. Pour définir une interprétation, nous devons faire choix d'une assignation individuelle et d'une assignation prédicative. Nous choisirons comme assignation individuelle une application qui fait correspondre à tout terme individuel de L' (variable individuelle ou constante individuelle) un objet et un seul du domaine D , deux termes distincts étant associés à des objets distincts. (Ainsi un terme individuel de L' a pour image un objet bien déterminé de D et tout objet de D ne peut être l'image que d'un terme individuel de L' au plus.) Nous définirons d'autre part une assignation prédicative comme suit. Soit P une variable prédicative de L' à n arguments. Nous lui associons une relation R à n objets, définie sur D , comme suit : la relation R existe ou n'existe pas entre les objets a_1 , a_2 , ..., a_n de D suivant que la proposition $P(t_1, t_2, \dots, t_n)$, où $t_1, t_2, \dots, t_n$ sont les termes individuels de L' dont les images respectives (par l'assignation individuelle) sont $a_1, a_2, \dots, a_n$, fait partie ou ne fait pas partie de la

classe K' . Ainsi, si la proposition $P(t_1, t_2, \dots, t_n)$ fait partie de K' , l'énoncé $R(a_1, a_2, \dots, a_n)$ est vrai, et si cette proposition ne fait pas partie de K' , cet énoncé est faux. Autrement dit : à une variable prédicative P est associé l'ensemble des n -uples d'objets de D , images, par l'assignation individuelle, des n -uples de termes individuels de L' qui sont arguments de P dans une proposition au moins de K' . Les règles qui déterminent la valeur de vérité d'une proposition de K' , sur la base des deux assignations qui viennent d'être définies, sont celles qui ont été énoncées plus haut dans la définition de la notion d'interprétation.

Il reste à montrer que pour l'interprétation ainsi décrite, le domaine D constitue bien un modèle pour K' . Pour cela examinons les différents types possibles de propositions de K' .

a) Si une proposition de K' est du type $P(t_1, t_2, \dots, t_n)$ où P est une variable prédicative et $t_1, t_2, \dots, t_n$ des termes individuels de L' , il résulte de nos définitions que cette proposition est vraie si elle fait partie de K' et fausse si elle n'en fait pas partie.

b) Soit une proposition complexe, du type $\sim A$, $A \supset B$ ou $(x)A(x)$. Nous supposons que, pour les propositions A , B et toutes les propositions du type $A(t)$, où t est un terme individuel de L' , on a déjà établi la propriété à démontrer, à savoir : une proposition est vraie si elle fait partie de K' , fausse si elle n'en fait pas partie.

b 1) Examinons le cas d'une proposition de type $\sim A$. Si $\sim A$ fait partie de K' , A ne peut en faire partie, car la classe K' est cohérente. Mais alors, en vertu de l'hypothèse d'induction, A est faux. Et en vertu de notre règle d'interprétation, $\sim A$ est vrai. Et si $\sim A$ ne fait pas partie de K' , A doit en faire partie, car K' est complète (comme classe maximale). Mais alors, en vertu de l'hypothèse d'induction, A est vrai. Et en vertu de notre règle d'interprétation, $\sim A$ est faux.

b 2) Examinons ensuite le cas d'une proposition de type $A \supset B$. Si $A \supset B$ fait partie de K' , ou bien $\sim A$ doit faire partie de K' , ou bien B doit en faire partie (les deux conditions pouvant du reste être réunies). Car si ni $\sim A$ ni B n'appartiennent à K' , comme A et $\sim B$ doivent alors appartenir à K' (qui est une classe complète), $A \supset B$ ne peut appartenir à K' . (Si $A \supset B$ faisait partie de K' dans ces hypothèses, on pourrait dériver de K' , par la règle du *modus ponens*, $\sim A$, et K' serait non-cohérente.) Si $\sim A$ appartient à K' , en vertu de ce qu'on vient de démontrer, $\sim A$ est vrai et donc A faux. Et si B appartient à K' , en vertu de l'hypothèse d'induction, B est vrai. Donc $A \supset B$ ne peut faire partie de K' que si A est faux ou B vrai. Mais alors, en vertu de la règle d'interprétation, $A \supset B$ est vrai (car $A \supset B$ ne peut être faux que si A est vrai et B faux en même temps). Si $A \supset B$ ne fait pas partie de K' , comme on vient de le voir, A est vrai et B faux. Et dans ce cas $A \supset B$ est faux.

b 3) Reste à examiner le cas d'une proposition de type $(x)A(x)$. Si $(x)A(x)$ fait partie de K' , en utilisant l'un des axiomes de LP et le *modus ponens*, on peut dériver de K' n'importe quelle proposition de type $A(t)$ où t est

un terme individuel quelconque de L' ⁴. Une telle proposition, étant compatible avec K' , doit faire partie de K' puisque K' est une classe maximale. En vertu de l'hypothèse d'induction, toutes ces propositions de type $A(t)$ sont donc vraies. Mais alors, en vertu de la règle d'interprétation, $(x)A(x)$ est vrai. (Cette règle nous dit en effet que $(x)A(x)$ est vrai si, quelle que soit l'assignation individuelle choisie, $A(x)$ est vrai. Nous obtenons ici : $A(t)$ est vrai, quel que soit le terme t . L'assignation individuelle que nous avons choisie fait correspondre à tout terme t un objet de D . Dire que $A(t)$ est vrai quel que soit t est équivalent à dire que $A(x)$ est vrai quelle que soit l'assignation individuelle choisie, c'est-à-dire quel que soit l'objet de D associé à x .) Si $(x)A(x)$ ne fait pas partie de K' , $\sim (x)A(x)$ doit en faire partie, car K' est une classe complète. Mais $\sim (x)A(x)$ est équivalent par définition à $(\exists x) \sim A(x)$. Et en vertu de la seconde condition réalisée par la classe K' , cette classe doit contenir la proposition $\sim A(c)$, pour une certaine constante c . En vertu de notre hypothèse d'induction, cette proposition doit donc être vraie et donc la proposition $A(c)$ doit être fausse. Mais alors, en vertu de notre règle d'interprétation, $(x)A(x)$ doit être faux. (Il existe en effet au moins une constante c pour laquelle $A(c)$ est faux. Mais dire cela équivaut à dire qu'il existe au moins une assignation individuelle pour laquelle $A(x)$ est faux. Or c'est là la condition de fausseté prévue pour $(x)A(x)$ par la règle d'interprétation.)

En raisonnant par induction sur la construction des propositions (c'est-à-dire en partant de propositions élémentaires et en passant de proche en proche aux propositions complexes construites au moyen des opérations logiques élémentaires et des quantificateurs), nous pouvons conclure, sur la base des résultats précédents, que toute proposition de K' est vraie dans l'interprétation proposée, autrement dit que le domaine D (pour cette interprétation), constitue un modèle de K' . Mais comme la classe K dont nous sommes partis fait partie elle-même de la classe K' , toutes les propositions de K sont vraies dans notre interprétation et donc le domaine D constitue un modèle pour K . Le théorème de cohérence est donc établi. On peut montrer, à titre de corollaire, que si une proposition A n'est pas dérivable de la classe K , il existe un modèle de K dans lequel la proposition A est fausse. (Cela signifie que, dans l'interprétation qui définit le modèle, la proposition A est fausse.)

Comme on le voit, la démonstration du théorème de cohérence revient à exhiber une structure (le domaine D muni des relations qui ont été définies sur lui par l'assignation prédicative) qui fournit une sorte de réalisation concrète des propriétés formelles exprimées par les propositions de K . Le théorème d'adéquation est une conséquence immédiate du théorème de cohérence.

4. Bien entendu, lorsqu'on élargit le langage L en un langage L' , on admet que les axiomes et règles s'étendent à toutes les expressions de L' . En particulier, l'un des axiomes de L prend la forme : $(x)A(x) \subset A(t)$, où t est une variable individuelle ou une constante individuelle quelconque de L' .

Il n'a été question jusqu'ici que de l'existence d'un modèle. On peut s'interroger sur la cardinalité du modèle. (Sans entrer ici dans des définitions précises, on pourra dire que la cardinalité d'un ensemble est la propriété commune à cet ensemble et à tous ceux qui lui sont équipotents. On dit que deux ensembles sont équipotents s'il existe entre eux une correspondance biunivoque, qui associe à tout élément de l'un un élément et un seul de l'autre et réciproquement. La cardinalité d'un ensemble fini est donnée par le nombre entier qui correspond au nombre d'éléments de cet ensemble. La cardinalité de l'ensemble des entiers est appelée *puissance du dénombrable*. Tout ensemble qui est équipotent (au sens susdit) à l'ensemble des entiers a cette même cardinalité et on dit qu'il est *dénombrable*. Un ensemble infini qui n'est pas équipotent à l'ensemble des entiers est dit *non-dénombrable*. (C'est le cas par exemple de l'ensemble des nombres réels, ou de l'ensemble des fonctions d'entiers, ou de l'ensemble des suites quelconques d'entiers.) Cantor a montré qu'il est possible de construire une hiérarchie de types d'infinité, commençant par la puissance du dénombrable, et dont les termes successifs peuvent être définis de proche en proche, indéfiniment.) En particulier on peut se demander s'il est toujours possible, quelle que soit la classe K de propositions (de L) considérées, de trouver un modèle qui a le même degré d'infinité que l'ensemble des entiers.

C'est précisément à cette question que répond le théorème de Löwenheim-Skolem. Ce théorème est la généralisation d'une propriété qui résulte d'un théorème dû à Löwenheim. Le théorème de Löwenheim affirme que si une proposition de L est valide dans un domaine infini dénombrable, elle est valide dans n'importe quel domaine (non-vide). On en tire le corollaire suivant : *Si une proposition de L est exemplifiable dans un domaine non-vide quelconque (de quelque degré d'infinité qu'il soit), elle est exemplifiable dans un domaine dénombrable*. Skolem a généralisé cette propriété en l'étendant à une classe quelconque de propositions : *Si toutes les propositions d'une classe quelconque de propositions de L sont simultanément exemplifiables dans un domaine non-vide quelconque, elles sont toutes simultanément exemplifiables dans un domaine dénombrable*. On pourra formuler de façon plus compacte le théorème de Löwenheim-Skolem comme suit : *Si une classe de propositions de L admet un modèle quelconque, elle admet un modèle dénombrable*⁵.

Il y a moyen de démontrer ce théorème en incorporant en quelque sorte sa démonstration dans celle du théorème de cohérence⁶. Dans l'esquisse

5. Löwenheim a présenté son théorème dans *Über Möglichkeiten im Relativkalkül*. Skolem en a donné une démonstration simplifiée dans *Logisch-kombinatorische Untersuchungen* etc. C'est dans ce même article qu'il a donné sa généralisation du théorème. Dans *Einige Bemerkungen zur axiomatischen Begründung der Mengenlehre*, il a donné une démonstration du théorème de Löwenheim sans l'axiome du choix. Quelques années plus tard, dans *Über einige Grundlagenfragen der Mathematik*, il a donné une nouvelle version de cette démonstration. Il a développé son interprétation du théorème, dans son application à la théorie des ensembles, dans *Sur la portée du théorème de Löwenheim-Skolem*.

6. C'est précisément ainsi que procède Church dans *Introduction to mathematical logic*.

qui a été donnée ci-dessus de cette dernière démonstration, il a été question simplement d'un domaine D infini. On peut montrer qu'il est possible de prendre pour domaine D l'ensemble des entiers ou, ce qui revient au même, un ensemble ayant même puissance que l'ensemble des entiers. (On doit pour cela utiliser une énumération des termes individuels de L' et des variables prédicatives de L , c'est-à-dire associer un numéro d'ordre — sous la forme d'un nombre entier — à chaque terme individuel de L' et à chaque variable prédicative de L . Mais cela ne présente aucune difficulté de principe.) On obtient alors la seconde partie du théorème de cohérence sous la forme suivante : *Si une classe de propositions de L est cohérente, elle admet un modèle dénombrable.* En joignant cette propriété à la première partie du théorème de cohérence, on obtient le théorème de Löwenheim-Skolem : *Si une classe de proposition de L admet un modèle (quelconque), elle admet un modèle dénombrable.*

Mais on peut aussi donner du théorème de Löwenheim-Skolem une démonstration indépendante de celle du théorème de cohérence, qui fait mieux apparaître sa véritable nature. Cette démonstration consiste à extraire du modèle dont l'existence est affirmée par hypothèse un modèle dénombrable⁷. Le théorème de Löwenheim-Skolem peut facilement être étendu au cas où le langage L dans lequel est formulée la théorie LP contient non seulement des variables individuelles mais aussi un certain nombre (éventuellement une collection infinie dénombrable) de constantes individuelles. Nous nous placerons dans ce cas. Soit donc une classe cohérente K de propositions de L et un domaine D qui, pour une interprétation donnée, constitue un modèle pour K . Nous allons exhiber un sous-ensemble dénombrable de D qui constituera également un modèle pour K .

L'interprétation qui fait de D un modèle pour K se base sur une certaine assignation individuelle et une certaine assignation prédicative. L'assignation individuelle fait correspondre à chaque terme individuel de L (variable ou constante) un objet et un seul du domaine D . Dans le raisonnement qui va suivre, nous allons considérer la partie de cette assignation qui concerne les constantes comme fixée et nous allons restreindre de façon appropriée la partie de cette assignation qui concerne les variables. Par ailleurs nous ne modifierons en rien l'assignation prédicative.

Une proposition de K contiendra, en général, des quantificateurs. Mais elle peut aussi contenir des variables qui ne sont pas liées par des quantificateurs. (Elle peut naturellement ne contenir que des variables de ce genre.) Mais nous pouvons transformer les propositions de K de telle sorte qu'elles ne contiendront plus que des variables quantifiées et que, de plus, les quantificateurs universels et les quantificateurs existentiels apparaîtront sous forme de deux groupes distincts. La métathéorie syntaxique de L

7. Les explications qui suivent ne constituent nullement une démonstration au sens strict. Elles sont destinées simplement à faire apparaître l'idée essentielle du mécanisme de démonstration. Cet exposé s'inspire en partie de la présentation donnée du théorème par Paul Cohen dans *Set theory and the continuum hypothesis*.

nous apprend en effet qu'il est toujours possible de donner à une proposition de L une forme canonique, que nous appellerons *forme normale d'exemplification*, dans laquelle toutes les variables sont quantifiées et dans laquelle, de plus, tous les quantificateurs universels précèdent les quantificateurs existentiels. Nous disposons d'un métathéorème sémantique affirmant qu'une proposition de L est exemplifiable dans un domaine donné si et seulement si sa forme normale d'exemplification est exemplifiable dans ce domaine. Nous pouvons donc, dans notre démonstration, supposer que toutes les propositions de K ont été remplacées par leur forme normale d'exemplification.

Une proposition quelconque de K se présentera donc sous la forme : $(x_1)(x_2)...(x_n)(E\gamma_1)(E\gamma_2)...(E\gamma_m)A$ (où l'expression A contient les variables mentionnées dans les quantificateurs et aucune autre). Dire qu'une proposition de ce genre est exemplifiable dans D , c'est dire que, quels que soient les objets de D associés à $x_1, x_2, \dots, x_n$, il existe au moins une assignation individuelle qui associe à $\gamma_1, \gamma_2, \dots, \gamma_m$ des objets tels que, pour une telle assignation, la proposition en question est vraie relativement à D .

Nous avons supposé que le langage L contient des constantes. L'assignation individuelle qui rend les propositions de K vraies dans D associe à ces constantes des objets de D . Soit C un sous-ensemble arbitraire de D contenant tous ces objets. Comme les constantes de L forment au plus un ensemble infini dénombrable, nous pouvons toujours nous arranger pour que C soit au plus dénombrable. (Si le nombre de constantes de L est fini, C pourra naturellement être fini.) Nous allons nous servir de C comme noyau de la construction de notre modèle dénombrable. (S'il n'y avait pas de constantes dans L , on pourrait prendre pour noyau un sous-ensemble fini ou dénombrable quelconque de D .) Toutes les assignations individuelles dont il va être question ci-après associeront aux constantes de L (qui figurent dans les propositions de K) des objets de C , exactement comme l'assignation de départ. Comme l'ensemble C sera inclus dans notre modèle, le cas des constantes individuelles est ainsi réglé une fois pour toutes. Passons à l'examen du sort réservé *aux variables individuelles*.

Considérons pour commencer le cas d'une proposition de K qui ne contiendrait pas de quantificateurs existentiels. Soit par exemple $(x_1)...(x_n)A$ une proposition de ce genre. Toute assignation individuelle qui associe aux variables $x_1, \dots, x_n$ des objets de C rend automatiquement la proposition en question vraie relativement à C (et donc à tout domaine contenant C), puisque cette proposition est vraie relativement à D et donc est vraie quels que soient les objets de D associés à ces variables. (Nous admettrons que plusieurs variables x_i peuvent avoir la même image dans C .)

Passons maintenant au cas d'une proposition qui contiendrait des quantificateurs existentiels. Soit par exemple $(x_1)...(x_n)(E\gamma_1)...(E\gamma_m)A$ une proposition de ce genre. Appelons S cette proposition. Considérons toutes les assignations qui associent aux variables $x_1, \dots, x_n$ respectivement les objets $a_1, \dots, a_n$ de C . Parmi ces assignations, il y en a au moins une qui associe

à $y_1, y_2, \dots, y_m$ respectivement les objets, $b_1, b_2, \dots, b_m$ de D de telle sorte que, pour cette assignation, la proposition S est vraie relativement à D . Il peut y avoir une infinité d'assignations de ce genre. Choisissons-en une, arbitrairement, et ajoutons à l'ensemble C les objets $b_1, b_2, \dots, b_m$ qu'elle associe respectivement à $y_1, y_2, \dots, y_m$ ⁸. Répétons ce choix pour toutes les assignations, du type décrit ci-dessus, qui associent à $x_1, x_2, \dots, x_n$ un n -uplet d'objets de C . Comme l'ensemble C est au plus dénombrable, il ne peut y avoir qu'une infinité dénombrable de n -uplets de C et donc de choix. Pour chaque choix, nous ajoutons à C un m -uplet d'objets de D . Comme la réunion d'un ensemble dénombrable de m -uplets d'objets est dénombrable, nous aurons donc, au terme de nos choix, ajouté à C un ensemble dénombrable d'objets de D . Remarquons que nous utilisons, dans cette construction, le principe du choix. Ce principe est un axiome célèbre de la théorie des ensembles qui affirme ce qui suit : étant donné une collection infinie d'ensembles, il existe un ensemble formé en prélevant un élément et un seul dans chacun des ensembles de la collection⁹.

En procédant de même pour toutes les propositions de K (du type considéré), nous obtenons un ensemble C_1 , qui contient C , qui est un sous-ensemble de D , et qui est dénombrable. Comme les propositions de K sont des suites finies de symboles de L et comme l'ensemble des symboles de L est dénombrable, il ne peut y avoir au plus qu'une infinité dénombrable de propositions dans K , *a fortiori* de propositions de K du type considéré. Comme une collection dénombrable d'ensemble dénombrables est dénombrable, nous n'aurons finalement ajouté à C qu'un ensemble dénombrable et l'ensemble résultant C_1 est donc lui-même dénombrable.

Nous n'avons cependant pas le droit de dire que C_1 est un modèle pour K . Considérons en effet une proposition telle que $(x)(\exists y)A(x, y)$. Si nous associons à x un objet pris dans C , alors, d'après notre construction, nous pouvons trouver une assignation qui associe à y un objet de C_1 tel que, pour cette assignation, notre proposition est vraie relativement à C_1 . Mais pour avoir le droit de dire que notre proposition est exemplifiable dans C_1 , nous devrions pouvoir affirmer cela quel que soit l'objet associé à x . Or C_1 est plus riche que C . Si nous associons à x un objet de C_1 qui ne se trouve pas dans C , nous ne sommes pas du tout assurés de pouvoir trouver dans C_1 lui-même un objet qui, associé à y , rendra notre proposition vraie dans C_1 .

Nous devons donc reprendre notre construction en faisant cette fois jouer à C_1 le rôle que jouait C tout à l'heure. Les mêmes considérations vont pouvoir se répéter, et nous allons aboutir ainsi à un ensemble C_2 , qui contient

8. Dans sa démonstration, Skolem fait intervenir, pour représenter ce choix, des fonctions qui dépendent des variables $x_1, x_2, \dots, x_n$. Ainsi la variable y_1 serait remplacée par une fonction $f_1(x_1, x_2, \dots, x_n)$, qui, à tout n -uplet d'objets associés aux variables $x_1, x_2, \dots, x_n$, fait correspondre un objet b_1 choisi parmi ceux pour lesquels la proposition S est vraie dans D .

9. Skolem lui-même a indiqué comment modifier la démonstration de façon à éviter le recours à l'axiome du choix. La démonstration donnée par Church ne fait pas intervenir l'axiome du choix; mais elle se base sur une énumération des expressions de L .

C_i , qui est un sous-ensemble de D , et qui est dénombrable. Mais à propos de cet ensemble C_2 , nous devons faire la même remarque que pour C_1 . Aussi devons-nous de nouveau recommencer notre construction, en partant cette fois de C_2 . Et ainsi de suite. Formons l'ensemble-union de tous les ensembles C_i ainsi obtenus, c'est-à-dire l'ensemble D^* défini comme suit : la condition nécessaire est suffisante pour qu'un objet appartienne à D^* est qu'il appartienne à un certain ensemble C_i . Comme les ensembles C_i forment une suite dénombrable et sont eux-mêmes tous dénombrables, l'ensemble D^* est lui-même dénombrable. En effet, l'union d'une collection dénombrable d'ensembles dénombrables est dénombrable.

Cet ensemble D^* constitue un modèle dénombrable pour notre classe K de propositions. Autrement dit, toutes les propositions de K sont simultanément exemplifiables dans D^* . Ceci est évident pour le cas d'une proposition qui ne contiendrait que des quantificateurs universels. Prenons le cas d'une proposition qui contient de plus des quantificateurs existentiels, par exemple $(x_1) \dots (x_n) (E y_1) \dots (E y_m) A$. Quels que soient les objets associés aux x_i , en vertu de notre construction, nous pourrions toujours trouver une assignation qui rendra notre proposition vraie dans D^* . En effet, les objets associés aux x_i doivent appartenir à un certain ensemble C_k . Nous sommes certains de pouvoir trouver dans C_{k+1} un m -uplet d'objets tels que, en les associant aux y_i , on obtienne une assignation qui rend notre proposition vraie relativement à D^* .

La plus remarquable conséquence du théorème de Löwenheim-Skolem concerne la théorie des ensembles, et se présente sous la forme d'un paradoxe dont la signification est considérable. Il est possible de formaliser la théorie des ensembles dans le cadre de la théorie des prédicats du premier ordre. Plus précisément, il est possible de représenter cette théorie sous forme d'un système axiomatique et d'exprimer ces axiomes sous forme de propositions du langage L , à condition d'inclure dans celui-ci deux constantes prédictives, la relation d'appartenance (exprimant l'appartenance d'un élément à un ensemble) et la relation d'égalité. Les axiomes de la théorie des ensembles forment une classe infinie (dénombrable) de propositions. Moyennant les axiomes et les règles de la théorie LP, on peut en dériver les théorèmes (connus) de la théorie des ensembles. Comme le théorème de Löwenheim-Skolem peut être étendu sans difficulté au cas d'un langage L muni de constantes prédictives, il s'applique à l'axiomatique de la théorie des ensembles. Et nous avons le résultat paradoxal : si la théorie des ensembles est exemplifiable, elle est exemplifiable dans un domaine dénombrable. Cette propriété est paradoxale, parce qu'il est possible de représenter, dans le cadre de la théorie axiomatique des ensembles, le célèbre raisonnement de Cantor qui prouve l'existence d'ensembles non-dénombrables. Ce raisonnement montre que, pour obtenir un ensemble non-dénombrable, il suffit par exemple de prendre l'ensemble de tous les sous-ensembles (finis ou infinis) de l'ensemble

des nombres entiers. Comme on peut donner un modèle dénombrable pour la théorie des ensembles, il doit donc exister dans ce modèle un objet *ND* qui a les propriétés d'un ensemble non-dénombrable.

Les ensembles de la théorie axiomatisée ont pour images des objets du modèle. Et la relation d'appartenance de la théorie axiomatisée a pour image une relation du même type, *R*, entre objets du modèle. L'ensemble non-dénombrable, qui a pour image l'objet *ND* du modèle, est un ensemble d'ensembles. Aux ensembles qui appartiennent à l'ensemble non-dénombrable de la théorie correspondent des objets qui ont avec l'objet *ND* la relation *R*. Nous pouvons donc dire que l'objet *ND* « contient » (au sens, bien entendu, de la relation *R*, définie sur le modèle) une infinité non-dénombrable d'objets du modèle. Autrement dit, si l'objet *ND* existe (et il doit exister pour que le modèle offre en effet une exemplification à la théorie), il existe un sous-ensemble du modèle qui est non-dénombrable. Mais le modèle lui-même constitue un ensemble dénombrable. Comment un ensemble dénombrable peut-il contenir un ensemble non-dénombrable ?

Pour résoudre ce paradoxe, il faut se rappeler ce que signifie « existence d'un ensemble non-dénombrable » : un ensemble est non-dénombrable s'il n'est pas possible d'établir une correspondance biunivoque (un à un) entre cet ensemble et l'ensemble des entiers. Dans le cadre d'une théorie des ensembles, une correspondance biunivoque est représentée par un ensemble (l'ensemble représentatif de la relation de correspondance biunivoque). Il existe effectivement dans le modèle au moins un ensemble non-dénombrable. Cela signifie qu'il n'existe pas, dans le modèle, d'ensemble effectuant la mise en correspondance biunivoque entre cet ensemble et l'ensemble des entiers (lui-même représenté dans le modèle). Mais dès qu'on se place en dehors du modèle, on peut effectivement établir une correspondance biunivoque entre cet ensemble et l'ensemble des entiers. Autrement dit, les ensembles qui représentent dans le modèle les ensembles non-dénombrables de la théorie sont eux-mêmes des ensembles dénombrables.

On peut faire apparaître d'une autre manière cette limitation du modèle. Appelons *Ee* l'ensemble de tous les sous-ensembles de l'ensemble des entiers, ou, plus simplement, l'ensemble des ensembles d'entiers. Il existe en particulier, dans le modèle dénombrable de la théorie des ensembles, un ensemble qui représente l'ensemble *Ee*. Mais il est dénombrable, alors que *Ee* ne l'est pas. Cela signifie qu'il ne contient pas tous les sous-ensembles de l'ensemble des entiers. Ceci est vrai du reste non seulement pour l'ensemble *Ee* mais pour tous les ensembles infinis : tout ensemble infini a son représentant dans le modèle et on peut former dans le modèle l'ensemble *E* de ses sous-ensembles mais cet ensemble *E* ne contient pas tous les sous-ensembles de l'ensemble en question.

Il y a donc un décalage systématique entre la théorie formalisée et la métathéorie. Il existe des ensembles qui, du point de vue de la théorie, sont non-dénombrables et qui apparaissent comme tels dans tous les modèles de la théorie, mais qui apparaissent comme dénombrables lorsqu'on les regarde

de l'extérieur pour ainsi dire, c'est-à-dire du point de vue de la métathéorie. Le concept de cardinalité de la théorie n'est donc pas le même que celui de la métathéorie. Celle-ci se place dans la perspective de ce qu'on pourrait appeler la théorie naïve des ensembles, qui utilise la notion d'ensemble sans lui donner une représentation formelle. La théorie axiomatisée n'utilise la notion d'ensemble que selon le sens précis que lui donnent les axiomes : seuls existent pour elle les ensembles dont l'existence est affirmée par les axiomes ou par les théorèmes que l'on peut en déduire. Skolem interprète cette situation en disant que, par suite de l'axiomatisation, les concepts de la théorie des ensembles et par conséquent, par leur intermédiaire, tous les concepts mathématiques, se trouvent relativisés. Ils n'ont plus un sens absolu, donné en soi, indépendamment de toute représentation, mais ils prennent des sens différents selon le domaine dans lequel ils se réalisent. Le théorème de Löwenheim-Skolem a incontestablement la portée que lui attribue Skolem. Mais il n'en reste pas moins que les notions mathématiques intuitives continuent à jouer, à l'égard des concepts axiomatisés, un rôle régulateur. C'est parce qu'il en est ainsi que l'on parle de l'inadéquation des systèmes axiomatiques.

Pour préciser en quoi consiste au juste cette inadéquation, dans le cas qui nous intéresse, il convient de rapprocher la situation de la théorie axiomatique des ensembles de celle de la théorie axiomatique des nombres entiers. Il existe, comme on sait, un système d'axiomes pour l'arithmétique, dont l'essentiel est dû à Peano. Ce système peut être exprimé dans le cadre du langage L (enrichi de certaines constantes). Comme le principe d'induction n'est qu'un schéma d'axiomes et correspond en réalité à une infinité dénombrable d'axiomes¹⁰, le système d'axiomes de l'arithmétique constitue un ensemble infini de propositions de L . Or on peut démontrer assez aisément que si un ensemble fini ou dénombrable de propositions de L admet un modèle, il admet n'importe quel modèle de cardinalité plus grand. On devrait s'attendre à ce que la théorie axiomatique des nombres n'admette pour modèle que le domaine (dénombrable) formé par l'ensemble des entiers. Or, en vertu de la propriété qu'on vient d'énoncer, si elle admet ce modèle (ce qui est effectivement le cas), elle doit admettre aussi n'importe quel modèle non-dénombrable. Cela indique qu'elle ne réussit pas à caractériser l'ensemble des entiers.

Skolem a du reste établi un résultat qui est d'une certaine manière plus impressionnant encore¹¹ : il est possible de construire un modèle pour la théorie axiomatique de l'arithmétique qui est dénombrable mais dont le

10. Le principe d'induction affirme, de chaque propriété d'entiers P : si la propriété P est vraie pour 0 et si, lorsqu'elle est vraie pour n , elle est vraie pour $(n + 1)$, elle est vraie pour tout entier. Nous avons donc une proposition distincte pour chaque propriété P .

11. Dans *Über die Unmöglichkeit etc.* et *Über die Nicht-Charakterisierbarkeit etc.* Il a donné une autre présentation de cette question dans *Peano's axioms and models of arithmetic*.

type d'ordre est donné par un ordinal plus grand que celui qui caractérise la suite des entiers. (On pourrait traduire intuitivement cette propriété en disant : dans ce modèle, après les nombres entiers, il y a encore d'autres objets.) Ainsi, même sans dépasser la cardinalité du dénombrable, on peut trouver pour l'axiomatique des entiers des modèles différents de l'ensemble des entiers. Appelons un modèle de ce genre, dont le type d'ordre ou la cardinalité n'est pas le même que celui des entiers, un modèle *non-régulier*.

L'existence des modèles non-réguliers pour l'arithmétique formalisée traduit d'une autre manière l'inadéquation de l'axiomatique des ensembles. Le principe d'induction fait en effet apparaître la notion *propriété d'entiers*. Or une propriété d'entiers peut être considérée (extensionnellement) comme un ensemble d'entiers : c'est l'ensemble des entiers qui vérifient la propriété considérée. Comme il est question, dans le principe d'induction, d'une propriété « quelconque » et donc, par le fait même, d'un sous-ensemble « quelconque » de l'ensemble des entiers, on doit s'attendre en quelque sorte *a priori*, sur la base du théorème de Löwenheim-Skolem, à ce que l'axiomatique des entiers soit inadéquate (puisque le concept d'ensemble est relatif au modèle dans lequel il se réalise). Plus précisément, le principe d'induction, tel qu'il est formulé dans le cadre d'une théorie axiomatisée, ne peut correspondre au principe d'induction de la théorie intuitive des nombres. Selon celle-ci, le principe d'induction s'applique à tous les prédicats d'entiers. Dans le cadre d'une théorie axiomatisée, le principe d'induction s'applique aussi à « tous » les prédicats d'entiers : mais la théorie ne connaît que les prédicats d'entiers qui peuvent être représentés dans le cadre formel qu'elle constitue. Ou, pour employer le vocabulaire ensembliste, la théorie ne connaît que les ensembles d'entiers dont elle peut établir l'existence par ses moyens propres. Or il n'est pas possible de représenter dans un système axiomatique tous les prédicats d'entiers.

On peut montrer cela au moyen du célèbre argument de la diagonale, qui a servi à Cantor à démontrer que l'ensemble des parties d'un ensemble a une cardinalité plus grande que celle de l'ensemble donné¹². Soit un système axiomatique pour la théorie des nombres entiers. On peut énumérer tous les ensembles d'entiers dont l'existence peut être prouvée dans ce système. En effet, il suffit d'appliquer à ce système le théorème de Löwenheim-Skolem pour obtenir un modèle dénombrable du système : or ce modèle devra contenir les objets qui correspondent aux ensembles d'entiers du système. Du reste l'existence d'un ensemble d'entiers est affirmée sous forme d'une proposition de longueur finie. Or il ne peut y avoir qu'une infinité dénombrable de propositions de ce genre dans un système axiomatique (puisque la liste des symboles du système est dénombrable). Nous pouvons, en dehors du système, définir un ensemble d'entiers *E* comme suit : un nombre entier *n* fait partie de *E* si et seulement si il n'appartient pas au *n*^e ensemble de notre énumération. L'ensemble *E* ne peut certainement pas appartenir à

12. C'est ce que fait remarquer Wang dans *On formalization*.

notre énumération. Il est en effet distinct de chaque ensemble de cette énumération au moins par le nombre n (correspondant au numéro d'ordre de cet ensemble dans l'énumération). Ainsi il n'est pas possible de représenter dans un système axiomatique tous les prédicats d'entiers (de l'arithmétique intuitive) pas plus qu'il n'est possible de représenter dans un système axiomatique tous les ensembles de la théorie intuitive des ensembles. Et cela explique qu'une théorie axiomatique des nombres admette d'autres modèles que l'ensemble des entiers : elle est en quelque sorte trop pauvre pour caractériser entièrement le domaine des entiers, elle comporte comme une frange d'indétermination qui la rend apte à représenter d'autres structures que celle des entiers.

Le théorème de Löwenheim-Skolem a été généralisé par Henkin ¹³ pour une théorie logique incluant une théorie de types (c'est-à-dire admettant une hiérarchie infinie de niveaux de prédicats et, à chaque niveau, la quantification des prédicats du niveau précédent), sous la forme suivante : un ensemble K de propositions de cette théorie admet un modèle dénombrable si et seulement si chaque sous-ensemble fini de K admet un modèle (quelconque). Henkin a montré comment, à partir de cette généralisation, on peut construire des modèles non-réguliers pour toute théorie axiomatisée de l'arithmétique, soulignant ainsi d'une nouvelle manière la parenté entre le théorème de Löwenheim-Skolem, et donc l'inadéquation de la théorie axiomatique des ensembles, et l'inadéquation de la théorie axiomatique des nombres.

La notion du *modèle non-régulier* a été généralisée par Rosser et Wang ¹⁴ pour une très vaste classe de systèmes formels. Ces auteurs ont montré qu'il existe des systèmes qui admettent des modèles réguliers et des modèles non réguliers (en ce sens généralisé) et des systèmes qui n'admettent que des modèles non réguliers. Cela fait apparaître d'une nouvelle manière la relativité des concepts mathématiques, mise en évidence par Skolem. Des notions de base telles que les notions d'égalité, de nombre entier, de nombre ordinal, d'ensemble, ont un sens « classique », quand elles sont représentées dans des systèmes à modèles réguliers. Mais il existe des systèmes d'une grande puissance logique qui n'admettent pas de modèles réguliers et dans lesquels ces notions prennent un sens nouveau, généralisé. Mais ces systèmes sont en quelque sorte moins sûrs que les systèmes classiques : ils occupent une situation intermédiaire entre la sécurité logique des systèmes classiques et la contradiction.

Mais, sous ces généralisations, nous devons retrouver l'essentiel de la signification de notre théorème. Comme nous l'avons vu, il entraîne l'inadéquation de la théorie axiomatisée des ensembles et de la théorie axioma-

13. Dans *Completeness in the theory of types*.

14. Dans *Non-standard models for formal logics*.

tisée des nombres. Or, de part et d'autre, ce qui permet de mettre en évidence cette inadéquation, c'est l'argument de la diagonale. On l'a vu clairement dans le cas de l'arithmétique. En ce qui concerne les ensembles, on a vu que la théorie axiomatisée des ensembles admet un modèle dénombrable parce qu'elle ne permet pas de représenter tous les ensembles de la théorie intuitive. Or lorsque celle-ci démontre l'existence d'ensembles non-dénombrables, elle fait intervenir (comme le principe d'induction de l'arithmétique intuitive) tous les ensembles d'entiers. C'est bien cette totalité des ensembles d'entiers qui est invoquée dans l'argument de la diagonale. Certes, il est possible de représenter le raisonnement de la diagonale dans le cadre d'un système axiomatique, mais il y perd en quelque sorte sa vertu : il ne peut plus produire effectivement le non-dénombrable, parce que la totalité qui intervient dans le système n'a plus un sens absolu, mais est relative aux possibilités de représentation du système.

Rappelons l'essentiel de l'argument de la diagonale. Soit E un ensemble quelconque et PE l'ensemble des parties (ou du sous-ensemble) de E . On va montrer qu'il n'existe pas de correspondance biunivoque entre E et PE . Supposons qu'il existe une telle correspondance, et désignons par Pa l'objet de PE associé à l'objet a de E . Étant donné un sous-ensemble Pa de E , a (qui lui est associé) peut lui appartenir ou ne pas lui appartenir. Formons un ensemble D (ensemble diagonal) au moyen de la condition suivante : la condition nécessaire et suffisante pour qu'un objet a de E fasse partie de D est qu'il ne fasse pas partie du Pa correspondant. Cet ensemble D sera nécessairement distinct de tous les sous-ensembles Pa : si un Pa contient l'objet a , alors D ne peut le contenir, et si un Pa ne contient pas l'objet a , alors D doit le contenir. Or l'ensemble D est bien un sous-ensemble de E . On conclut facilement à partir de là que la cardinalité de PE est plus élevée que celle de E .

Comme on le voit, la définition de l'ensemble D fait intervenir la totalité des sous-ensembles Pa et présuppose que tous ces sous-ensembles sont connus (puisque l'on doit pouvoir déterminer pour chacun s'il contient ou non l'objet a de E qui lui est associé par hypothèse). En ce sens D est une entité du second degré : la connaissance de D présuppose la connaissance complète de tous les Pa , c'est-à-dire de tous les sous-ensembles de E . Comme D est lui-même un sous-ensemble de E , d'une certaine manière il se présuppose lui-même. Cet ensemble D a donc un caractère paradoxal, qui provient de ce qu'il est une sorte de réflexion de l'ensemble PE en lui-même. La théorie intuitive réussit à dépasser le dénombrable et à introduire des niveaux de plus en plus élevés d'infinité, mais c'est parce qu'elle considère comme réalisé un ensemble d'opérations qui n'est pas réalisé et qui n'est même pas réalisable effectivement : le parcours de tous les sous-ensembles de l'ensemble de départ. Si elle peut se détacher du dénombrable, c'est donc parce qu'elle n'est pas effective et opère sous le signe du « comme si ».

La théorie intuitive admet qu'elle a à sa disposition tous les sous-ensembles Pa sans se préoccuper de la manière dont ces ensembles lui sont accessibles.

En réalité, elle ne fait que thématiser une visée qui en général enveloppe déjà des thématisations préalables. Dans un premier moment, elle constitue la visée de la suite des entiers, en se représentant la possibilité d'une itération indéfinie de l'opération qui consiste à passer d'un entier au suivant. Dans un deuxième moment, elle prend cette visée comme objet et traite l'ensemble des entiers comme s'il était effectivement disponible, au même titre que les nombres entiers eux-mêmes. Elle se sert de cet ensemble comme d'un modèle pour se représenter de la même manière n'importe quel sous-ensemble infini de l'ensemble des entiers, ou, pour le dire plus simplement, n'importe quel ensemble d'entiers. Dans un troisième moment, elle constitue la visée de la suite des ensembles d'entiers, en se servant du modèle de la suite des entiers et en feignant de croire que l'on peut en effet traiter les ensembles d'entiers comme les entiers. Dans un quatrième moment elle thématise cette visée, sous la forme d'une définition (celle de l'ensemble diagonal) qui inclut la donnée de la totalité des ensembles d'entiers, qui considère donc cette totalité comme présente à la manière d'un objet entièrement disponible. Ultérieurement elle répète cette suite de démarches, en se servant du modèle de la construction de l'ensemble des sous-ensembles de l'ensemble des entiers pour construire l'ensemble des sous-ensembles d'un ensemble quelconque. Lorsque la théorie intuitive parle de « tous » les ensembles d'entiers, ou de « tous » les prédicats d'entiers, elle exprime par là la thématisation de sa visée et ce « tous » a ainsi un sens réflexif. Il y a réflexion de la théorie intuitive en elle-même en ce sens qu'il y a en elle passage du plan de l'expression au plan de l'acte et réciproquement : la visée est un acte, non une expression, mais elle présuppose des expressions, et la thématisation, objectivant la visée, reprend l'acte que constitue celle-ci dans une expression d'un niveau supérieur.

Le système formel ne connaît pas ces changements de plan : il se meut tout entier à un même plan, qui est celui des expressions. Les ensembles dont il parle doivent être effectivement présents pour lui, c'est-à-dire représentables au moyen d'expressions du système. Si le système répond aux conditions d'effectivité auxquelles répondent les systèmes classiques, du type de ceux dont il a été question dans ce qui précède, il n'admettra qu'une infinité au plus dénombrable de symboles et des expressions de longueur finie, et ne pourra donc contenir qu'une infinité dénombrable d'expressions¹⁵. On ne pourra donc trouver dans de tels systèmes qu'une

15. Nous devons faire remarquer ici que l'on a entrepris depuis un certain temps déjà des recherches sur des systèmes formels qui ne répondent plus aux conditions d'effectivité des systèmes classiques. Il s'agit de systèmes qui contiennent une infinité non dénombrable de symboles, ou (et) des règles de caractère infinitiste (comportant une infinité de prémisses), ou (et) des expressions d'une longueur infinie. L'introduction de tels systèmes correspond évidemment à un élargissement de la notion de système formel et les propriétés qu'ils présentent diffèrent de celles des systèmes classiques. On a pu retrouver cependant pour certains d'entre eux un théorème assez analogue au théorème de Löwenheim-Skolem, qui assigne une cardinalité minimale aux modèles du système (la cardinalité de l'ensemble des parties de l'ensemble des symboles du système). Il reste en tout cas que, pour de tels systèmes comme pour les systèmes classiques, n'existent que les ensembles représentables par des expressions du système.

infinité dénombrable d'ensembles (en particulier, s'il s'agit d'un système axiomatique pour l'arithmétique, une infinité dénombrable d'ensembles d'entiers). Certes il est possible d'exprimer dans un système formel le raisonnement de la diagonale et de se donner ainsi une image, dans le système, des opérations qui permettent à la théorie intuitive de passer au non-dénombrable. Mais cette possibilité même, on l'a vu, traduit la limitation du système : c'est précisément parce qu'il ne contient pas assez d'ensembles qu'il peut démontrer l'existence, en son sein, d'ensembles non-dénombrables. Le regard extérieur qui analyse le système du point de vue de la théorie intuitive le restitue pour ainsi dire à sa limitation et fait rentrer dans le dénombrable ce qui, pour le système, était non-dénombrable.

Lorsque le système, imitant la pensée intuitive, parle de la totalité des ensembles d'entiers, il ne se réfère, en fait, qu'à la totalité des ensembles d'entiers qu'il contient effectivement, non à la totalité virtuelle d'une visée. Les démarches formelles reproduisent bien les démarches de la pensée intuitive, mais d'une manière en quelque sorte tout extérieure; il n'y a pas, dans le système, de véritable réflexion, car il n'y a pas en lui d'acte au sens propre. L'acte de visée, par lequel la pensée intuitive anticipe ses opérations possibles sans s'assurer qu'elle pourra effectivement leur faire correspondre des expressions, est représenté dans le système par un opérateur de totalisation qui ne fait, lui, que récapituler des expressions que l'on peut effectivement retrouver une à une. (Lorsqu'un système formel fait intervenir l'expression *pour tout x* , cela signifie que la variable x peut être remplacée par n'importe quel objet d'un certain domaine préassigné.) Ainsi lorsque le système représente la thématization réflexive de la pensée intuitive, il ne fait que reprendre des expressions dans une expression.

Mais cela indique que les expressions d'un système formel ne sont pas vraiment expressives; elles représentent, elles fournissent un équivalent tangible des opérations de la pensée intuitive, elles ne signifient pas comme celle-ci. La vertu représentative du système, qui fonde d'ailleurs sa fécondité en tant qu'instrument de clarification, est précisément liée à cette abstraction qui détache, dans le système, les expressions des significations. Le système formel exhibe pour ainsi dire à l'état pur, dans une sorte d'en-soi autosuffisant, le schème de structures possibles. La méthode des modèles est précisément destinée à mettre en évidence les structures compatibles avec le système, c'est-à-dire répondant à la forme générale qu'il prescrit. Mais la forme devient objet. C'est pour cela qu'il est possible de l'analyser au moyen de modèles; c'est pour cela aussi que le système peut offrir une image entièrement intelligible des démarches de la pensée intuitive. L'intelligibilité dont il s'agit ici est celle de l'opération : c'est dans ses opérations propres de construction et de dérivation que le système nous fait comprendre, à la manière d'une machine analogique, comment procède la pensée intuitive. La forme opératoire objectivée travaille en quelque sorte par elle-même. En tant qu'elle est opératoire, elle est effectivement productive, elle exhibe des schèmes, elle fait voir des enchaînements. En tant qu'elle

est objectivée, elle reste enfermée dans sa propre effectivité et lorsqu'elle représente l'infini, elle le matérialise dans ses propres limites.

La pensée intuitive, par contre, n'isole pas ses expressions de l'intention significative qui les porte. C'est pourquoi elle peut incorporer à ses produits ses propres visées totalisatrices, c'est-à-dire son mouvement même. La réflexion qu'elle comporte n'est pas simple récapitulation, mais transgression : la totalisation, supposée accomplie, nous transporte réellement au-delà de ce qui était déjà présent. Aussi, lorsque la pensée intuitive se représente l'infini, le saisit-elle comme l'horizon toujours ouvert de démarches indéfiniment réitérables et superposables, comme le champ toujours disponible d'une incessante transgression. C'est son rapport à cet horizon qui la fait significative; en se réfléchissant, elle s'annonce toujours à elle-même. Elle a besoin, certes, du formalisme pour se rendre claires pour elle-même les démarches qu'elle a déjà accomplies. Elle ne semble pas pouvoir s'en remettre au formalisme de décider de ses propres possibilités. La fécondité du formalisme est rétrospective, celle de la pensée intuitive anticipatrice. La force du formel est de rendre présent; mais la présentification de l'infini le finitise. La force de l'intuition est d'annoncer; mais la pure annonce risque toujours de se perdre dans l'indicible.

(Septembre 1967.)

BIBLIOGRAPHIE

CHURCH (Alonzo), *Introduction to mathematical logic*, vol. I, Princeton University Press, 1956, X + 376 p.

COHEN (Paul), *Set theory and the continuum hypothesis*, New York, Amsterdam, W. A. Benjamin, 1966 154 p.

HENKIN (Leon), « Completeness in the theory of types », *Journal of Symbolic Logic*, vol. 15, 1950, p. 81-91.

LÖWENHEIM (Leopold), « Über Möglichkeiten im Relativkalkül », *Mathematische Annalen*, vol. 76, 1915, p. 447-470.

ROSSER (J. Barkley) and WANG (Hao), « Non-standard models for formal logics », *Journal of Symbolic Logic*, vol. 15, 1950, p. 113-129.

SKOLEM (Thoralf), « Logisch-kombinatorische Untersuchungen über die Erfüllbarkeit oder Beweisbarkeit mathematischer Sätze nebst einem Theoreme über dichte Mengen », *Skrifter utgitt av Videnskapsselskapet i Kristiana, I. Matematisk-naturvidenskabelig Klasse*, 1920, n° 4, 1920, 36 p. — « Einige Bemerkungen zur axiomatischen Begründung der Mengenlehre », in *Wissenschaftliche Vorträge gehalten auf den fünften Kongress der skandinavischen Mathematiker in Helsingfors von 4. bis 7. Juli 1923*, Helsingfors, 1923, p. 217-232. — « Über einige Grundlagenfragen der Mathematik », *Skrifter utgitt av Det Norske Videnskaps-Akademi i Oslo, I. Matematisk-naturvidenskabelig Klasse*, 1929, n° 4, 1929, 49 p. — « Über die Unmöglichkeit einer vollständigen Charakterisierung der Zahlenreihe mittels eines endlichen Axiomensystems », *Norsk Matematisk Forenings Skrifter, séries II*, n° 10, 1933, p. 73-82. — « Über die Nicht-charakterisierbarkeit der Zahlenreihe mittels endlich oder abzählbar unendlich vieler Aussagen mit ausschliesslich Zahlenvariablen. *Fundamenta Mathematicae*, vol. 23, 1934, p. 150-161. — « Sur la portée du théorème de Löwenheim-Skolem, in *Les entretiens de Zürich sur les fondements et la méthode des sciences mathématiques*, 6-9 décembre 1938, p. 25-47. — « Peano's axioms and models of arithmetic », in *Mathematical interpretations of formal systems*, Amsterdam, North-Holland Publishing Company, 1955, p. 1-14.

WANG (Hao), « On formalization », *Mind*, vol. 64, 1955, p. 226-238.