

Université de Nancy I
IREM

INTRODUCTION

A LA

THEORIE DES ENSEMBLES

par Jean-Pierre FERRIER

Octobre 1975

PREFACE

Ce traité de théorie des ensembles fait partie du cours de deuxième année de la maîtrise ès-Sciences Mathématiques, tel qu'il est enseigné à Nancy. Il est donc plus particulièrement destiné aux étudiants qui se préparent au métier d'enseignant du second degré. On s'est refusé pour cette raison à une présentation purement dogmatique de la théorie. L'enseignement universitaire est suffisamment rempli de développements de cette sorte qui risquent de donner à la longue une fausse idée de la science. Pour faire apparaître au contraire les motivations qui ont présidé à l'introduction des notions nouvelles, on a choisi de s'appuyer sur l'histoire des mathématiques au moment de la création de la théorie des ensembles.

On n'a cependant pas voulu écrire un ouvrage d'histoire. L'excellent livre de Jean Cavaillès, heureusement complété par la correspondance entre Cantor et Dedekind, remplit parfaitement ce rôle. De plus, le présent exposé s'adresse à des étudiants qui ne possèdent que des fragments de la théorie, alors qu'un traité historique n'est profitable qu'à celui qui en a déjà une bonne connaissance. Aussi, a-t-on évité de s'appesantir sur tous les essais et erreurs qui ont précédé au dégagement des concepts fondamentaux. Dans un souci de clarté, on s'est même écarté de l'ordre chronologique, en feignant d'ignorer dans un premier temps l'importance de l'étude des dérivés d'ensembles de points de la droite dans la création de la théorie; de la même façon ont été mises à part les considérations relatives à la construction des nombres entiers. L'articulation retenue a été finalement la suivante : l'étude de problèmes d'équipotence conduit aux antinomies, puis à la formalisation de la théorie des ensembles; cette dernière permet la construction des nombres entiers; ainsi se trouve faite la transition avec la notion d'ordre et les théorèmes fondamentaux sur les ensembles bien ordonnés.

Si l'introduction des nombres transfinis se trouve ainsi repoussée en fin de traité, c'est aussi pour assurer à ce dernier le caractère le plus élémentaire possible. On a cependant cherché à éviter l'erreur, faussement justifiée par des motifs pédagogiques, qui consiste à admettre au passage des résultats non triviaux. Il est en effet essentiel que de futurs enseignants aient compris que la chaîne du raisonnement mathématique ne supporte pas

l'absence d'un maillon. Aussi n'a-t-on pas, au moment de la construction des nombres entiers, supposé connue la notion de cardinal. Celle-là n'est donnée que beaucoup plus tard, après celle de type d'ordre. Il est certes possible d'introduire les nombres cardinaux autrement, à l'aide du symbole $\aleph$ de Hilbert par exemple; cette façon de procéder, très peu intuitive, est plutôt réservée à des ouvrages beaucoup plus ambitieux. D'autre part, l'étude directe des cardinaux finis est difficile, sans préparer à celle des ordinaux quelconques qui la rend pourtant a posteriori inutile.

Pour toutes ces raisons, les nombres entiers sont introduits suivant la méthode employée par Dedekind lui-même, à partir de l'axiome de l'infini. On sait bien que cet axiome n'est pas utile pour définir le cardinal d'un ensemble fini si l'on dispose du lemme de Zermelo. La présentation choisie a néanmoins le mérite d'être très élémentaire et de faire intervenir rapidement la notion de bon ordre, préparant ainsi les généralisations ultérieures du concept de nombre. On n'a pas en revanche insisté sur certaines axiomatisations, comme celle de Peano, qui n'ont pas de répercussions sur les mathématiques actuelles.

Le lecteur que rebuterait l'étude des ensembles bien ordonnés quelconques, pourra donc se contenter de la première partie de l'exposé, qui se suffit à elle-même. Celui qui chercherait au contraire la présentation la plus rapide, pourra aisément la reconstituer en modifiant quelque peu l'ordre des paragraphes et en négligeant certains développements. Il est ainsi souhaité que chacun puisse trouver la manière lui convenant le mieux pour aborder la théorie.

Faut-il ajouter qu'on n'a pas voulu présenter un modèle pour l'enseignement secondaire ? A ce niveau-là, bien peu de définitions devraient être explicitement données, contrairement à une certaine mode. Cependant savoir plus pour pouvoir enseigner moins impose aux futurs maîtres l'acquisition d'une vision cohérente sur le sujet.

Bibliographie

On trouvera un exposé très élémentaire de la théorie des ensembles, avec démonstration en exercices de certains théorèmes dans

Seymour Lipschutz.- Set Theory and Related Topics; Shame's outline series, Mc Graw-Hill Book Company, New-York.

D'une ambition voisine du présent traité, citons

J.-M. Exbrayat et P. Mazet.- Algèbre I, notions fondamentales de la théorie des ensembles; notions modernes de mathématiques, Hatier, Paris,

et d'une présentation nettement plus difficile, centrée sur la partie axiomatique de la théorie,

Nicolas Bourbaki.- Eléments de mathématique I, livre I, théorie des ensembles; Actualités scientifiques et industrielles, Hermann, Paris.

Jean-Louis Krivine.- Théorie axiomatique des ensembles; Collection Sup, Presses Universitaires de France, Paris.

Pour ce qui est de la partie historique, le traité est inspiré de

Jean Cavailles.- Philosophie mathématique; Collection Histoire de la Pensée, Hermann, Paris.

Nicolas Bourbaki.- Histoire des Mathématiques; Hermann, Paris.

René Baire.- Leçons sur les fonctions discontinues; Collection de monographies sur la théorie des fonctions; Gauthier-Villars, Paris.

et de quelques conférences de Jean-Louis Ovaert à l'I.R.E.M. de Nancy. Indiquons enfin, à l'intention de ceux qui voudraient une information plus complète sur les systèmes d'axiomes, les premiers chapitres de

A.A. Fraenkel, Y. Bar-Hillel, A. Levy.- Foundations of Set Theory; Studies in logic and the foundations of Mathematics; North-Holland Publishing Company, Amsterdam.

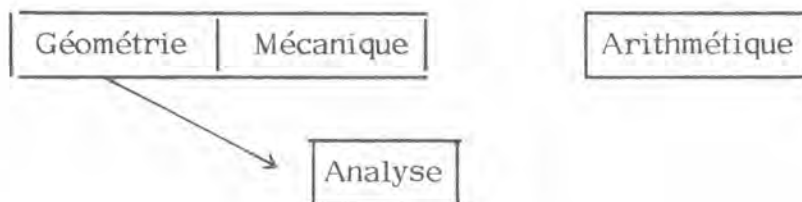
TABLE DES MATIERES

PREFACE	1
BIBLIOGRAPHIE	3
TABLE DES MATIERES	4
APERCU GENERAL	5
CHAPITRE I.- PUISSANCE DES ENSEMBLES.....	8
1. La Crise de l'infini	8
2. Ecueils et promesses	11
3. Le dénombrable et le continu	14
4. Exploration du continu	14
5. Développements	17
CHAPITRE II.- AXIOMES DE ZERMELO-FRAENKEL	19
1. Découverte d'antonimies	19
2. Axiome de Zermelo	20
3. Remarques	24
4. Constructions complémentaires	25
CHAPITRE III.- CORRESPONDANCES , APPLICATIONS	28
1. Correspondances	28
2. Applications	29
3. Injections, surjections, bijections	31
4. Notions complémentaires	36
CHAPITRE IV.- FAMILLES	39
1. Familles	39
2. Intersection et réunion d'une famille d'ensembles	42
3. Produit d'une famille d'ensembles	43
4. Somme disjointe d'une famille d'ensembles	44
CHAPITRE V.- ENSEMBLES FINIS	46
1. Définition des ensembles finis	46
2. Propriétés des ensembles finis	47
CHAPITRE VI.- NOMBRES ENTIERS	49
1. N-ensembles	49
2. Axiome de l'infini	53
3. Opérations sur les nombres entiers	56
4. Retour sur les ensembles finis	57
CHAPITRE VII.- ENSEMBLES DERIVES	59
CHAPITRE VIII.- ENSEMBLES BIEN ORDONNES	63
1. Propriétés des ensembles bien ordonnés	63
2. Comparaison des ensembles bien ordonnés	65
3. Théorèmes fondamentaux	66
CHAPITRE IX.- ORDINAUX, CARDINAUX	71
1. Nombres ordinaux	71
2. Nombres cardinaux	73

APERÇU GENERAL

Les mathématiques, telles que les avaient conçues les Grecs, étaient restées, jusqu'au dix-septième siècle, divisées en deux branches, l'Arithmétique et la Géométrie, qui s'opposaient à plusieurs égards. La première donnait souvent lieu à des calculs formels mais rigoureux; la seconde faisait au contraire appel en grande partie à l'intuition.

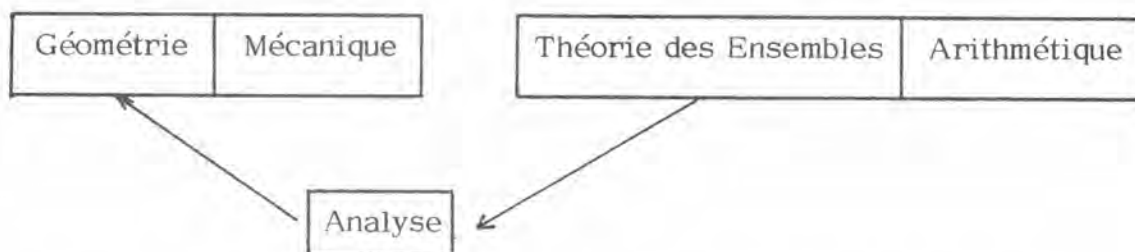
Au moment de la création du calcul infinitésimal (Leibniz, Newton), la Mécanique s'ajoute à la Géométrie. L'Analyse naît plus tard (Euler) pour unifier certains calculs apparaissant sous des formes différentes dans des problèmes de Géométrie ou de Mécanique (aires, moments d'inertie). On utilise alors les nombres réels, mais ceux-là ne sont que les points de la droite géométrique. Ainsi l'Analyse se trouve directement fondée sur la Géométrie, ce que résume le schéma suivant.



Le dix-huitième siècle voit se développer l'Analyse. Le langage (fonction à la place de mouvement, fonction continue, intégrale) et les techniques (construction des fonctions spéciales) se perfectionnent, mais dans une absence totale de rigueur. On aboutit à des erreurs (confusion faite par Cauchy entre continuité et continuité uniforme, utilisation abusive de séries divergentes) ou à la découverte de phénomènes qui heurtent l'intuition (celle par Bolzano de courbes continues n'admettant de tangente en aucun point).

On débouche alors sur une crise de l'Analyse. La nécessité apparaît de fonder cette dernière sur des bases solides, que l'on cherche naturellement

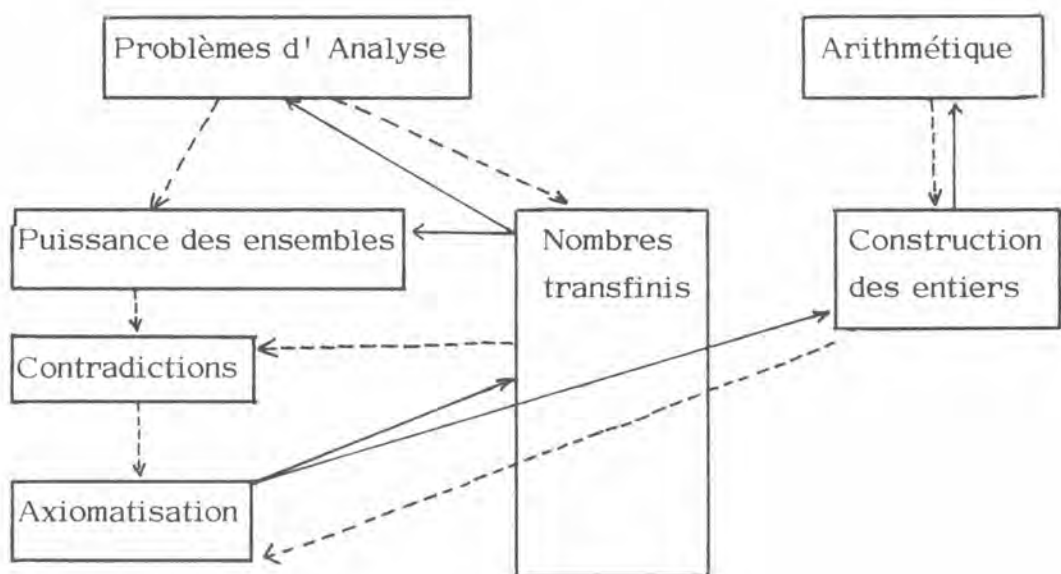
du côté de l'Arithmétique. Ce rôle est rempli par la Théorie des Ensembles qui se construit à la fin du dix-neuvième siècle et au début du vingtième. C'est sur elle que se fonde maintenant l'Analyse, qui sert elle-même de fondement à la Géométrie.



Dans la réalité la genèse de la Théorie des Ensembles n'est pas simple. Le langage de la théorie est en fait créé (Bolzano, puis surtout Cantor et Dedekind) dès la fin du dix-neuvième siècle.

Ce sont encore des problèmes d'Analyse, liés en particulier aux séries trigonométriques, qui y ont contribué. La classification de certains ensembles de points de la droite conduit ainsi à l'introduction des nombres transfinis. C'est à propos de ces recherches, et des essais de construction des nombres réels, que l'on est amené à comparer entre eux les différents infinis, à étudier ce que l'on appelle la puissance des ensembles (Cantor). Peu à peu, tous ces travaux permettent de dégager les notions essentielles restées jusque là enchevêtrées.

Une fois les concepts fondamentaux établis, il est facile de reconstruire l'Arithmétique, puis l'Analyse. Malheureusement le changement de siècle voit surgir des contradictions qui semblent aboutir à une condamnation sans appel de la théorie (Poincaré). Une axiomatisation doit alors être entreprise, qui était en fait préparée par les travaux autour de la notion de nombre entier (Dedekind). Les axiomes prennent des formes diverses (Zermelo, Fraenkel, Von Neumann). A partir de là, il devient enfin possible de construire les entiers, les nombres transfinis, et de définir correctement la puissance des ensembles.



PUISSANCE DES ENSEMBLES

1. La Crise de l'infini.

Les mathématiques ont utilisé depuis le début des raisonnements de théorie des ensembles. La Géométrie, par exemple, à propos des intersections de courbes et surfaces, fait largement intervenir les notions d'appartenance et d'inclusion. C'est seulement le mélange entre la notion d'ensemble et celle de grandeur, et en particulier la considération d'ensembles infinis comme tels, qui se heurte au refus des mathématiciens.

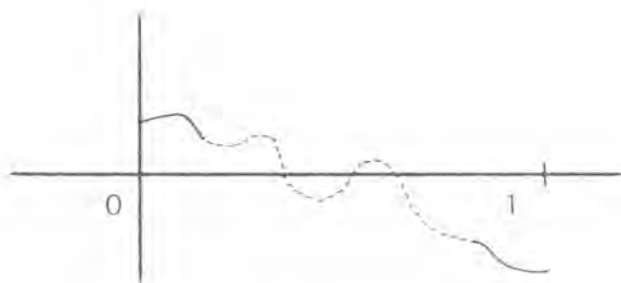
De tels ensembles, qui nécessitent de concevoir dans une même pensée une infinité d'objets, relèvent de ce que l'on appelle l'"infini actuel", lequel est soigneusement évité. Si on a, par exemple, le droit de dire qu'un point appartient à une droite, il est en revanche interdit de considérer que la droite est composée de points. Le langage des "lieux géométriques" est caractéristique de cette attitude. Cela ne pose pas de problème particulier tant que les seuls ensembles rencontrés sont d'un type simple et peuvent être complètement décrits par un nombre fini de paramètres.

En Analyse, on oppose à l'"infini actuel" interdit, l'"infini potentiel" qui est la faculté pour une grandeur toujours finie d'augmenter indéfiniment. C'est ainsi qu'il est permis de parler de limite d'une suite, de série convergente. Cette attitude prudente permet le développement de l'Analyse Classique, et en particulier celui du Calcul Infinitésimal; elle a le mérite de mettre les mathématiques à l'abri des controverses interminables, comme celle qui résulte de la considération d'infiniment grands ou petits.

Une autre façon de refuser l'"infini actuel" consiste à raisonner "en compréhension", et à remplacer chaque fois l'ensemble des éléments qui possèdent une propriété par la propriété elle-même. C'est ainsi qu'on parlera, en Arithmétique, de la "propriété commune des nombres entiers qui sont multiples de 6". Il faut reconnaître qu'il y a là une grande part d'hypocrisie.

Ce sont les nécessités de l'Analyse qui vont imposer aux mathématiques de sortir de leur réserve. Nous allons choisir un exemple, tiré de l'étude de fonctions d'une variable réelle.

Comme le remarque Bolzano dès 1799, la décomposition en éléments simples des fractions rationnelles à coefficients réels repose sur le théorème suivant : toute fonction numérique continue f définie sur l'intervalle $[0, 1]$ telle que $f(0) \geq 0$, $f(1) \geq 0$, s'annule au moins une fois sur l'intervalle.

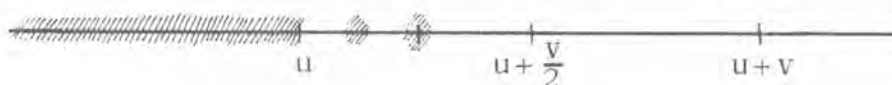


Cet énoncé prend un caractère intuitif pour celui qui s' imagine la courbe représentative; on pourrait s'en tenir là, mais il est bien difficile pour un mathématicien de résister à la tentation d'en donner une démonstration, c'est à dire de construire effectivement un nombre de $[0, 1]$ où la fonction donnée s'annulera. Désignant par M l'ensemble des nombres u de $[0, 1]$ tels que $f(u) \geq 0$ (*), on se ramène aisément à démontrer le nouveau théorème : tout ensemble majoré non vide de nombres réels admet une borne supérieure.

Si en effet U est la borne supérieure de M , on ne peut avoir $f(U) > 0$ sinon il existerait un nombre $\eta > 0$ tel que $f(U + \eta) > 0$, ni $f(U) < 0$ sinon il existerait un nombre $\eta > 0$ tel que $f(U - x) < 0$ pour $0 \leq x \leq \eta$ et $U - \eta$ serait un majorant de M .

Le caractère abstrait du nouvel énoncé le rend moins intuitif que le précédent. Essayons d'en donner une démonstration et pour cela considérons de façon générale un ensemble M de nombres réels contenant un nombre u et majoré par un nombre $u + v$. Une méthode consiste à cerner la borne supérieure cherchée par approximations successives.

(*) Si, comme le fait Bolzano, on raisonne en compréhension, l'introduction de cet ensemble exige une périphrase appropriée. A ce propos, il convient de signaler que nous évitons, d'une manière générale, d'employer la terminologie de l'époque pour ne pas introduire de confusion, de parler par exemple de "grandeur" au lieu de nombre, de "continuité" au lieu de connexité ...



On examine si $u + \frac{v}{2}$ majore M , puis, s'il en est ainsi, si $u + \frac{v}{2^2}$ majore $M, \dots$. Si l'on obtient une infinité de réponses affirmatives, le nombre u est le plus grand élément de M et par suite la borne supérieure cherchée. Dans le cas contraire on arrive à une réponse négative pour $u + \frac{v}{2^{m_1}}$. On s'intéresse alors à $u + \frac{v}{2^{m_1}} + \frac{v}{2^{m_1+1}}$, $u + \frac{v}{2^{m_1}} + \frac{v}{2^{m_1+2}}, \dots$. On ne peut pas avoir une infinité de réponses affirmatives sinon $u + \frac{v}{2^{m_1}}$ majorerait M . On arrive donc nécessairement à une réponse négative pour $u + \frac{v}{2^{m_1}} + \frac{v}{2^{m_2}}$ avec $m_2 > m_1$. En poursuivant indéfiniment, on met ainsi en évidence une suite

$$u, \quad u + \frac{v}{2^{m_1}}, \quad u + \frac{v}{2^{m_1}} + \frac{v}{2^{m_2}}, \dots$$

de majorants de M . La borne supérieure cherchée sera alors la limite de cette suite.

Il reste cependant à établir que la suite en question est bien convergente. Cela est encore intuitif: on sent bien que, dans un développement dyadique par exemple, la limite peut être calculée avec une précision arbitraire. Mais tout cela ne démontre cependant pas que la limite existe.

Essayons une autre méthode et introduisons l'ensemble M' des nombres réels qui minorent un nombre M et l'ensemble M'' des nombres réels qui majorent strictement tout nombre de M . Il est facile de voir que tout nombre de M' minore tout nombre de M'' et que l'ensemble des nombres réels est réunion disjointe des ensembles non vides M' et M'' . On donne à une telle partition le nom de "coupure". Le second théorème est alors une conséquence facile de la propriété suivante:

(C) Pour toute coupure dans les nombres réels, il existe un nombre réel U tel que les ensembles de la coupure soient définis par l'un ou l'autre des systèmes d'inégalités $u' \leq U < u''$ ou $u' < U \leq u''$.

Cette propriété exprime la connexité des nombres réels; elle aussi découle de la connaissance intuitive que l'on croit avoir de la droite géométrique, mais on n'est pas plus avancé pour cela.

Cependant le travail qui vient d'être fait n'est pas complètement inutile. Nous avons dégagé un certain nombre de propriétés des nombres réels et les avons reliées entre elles. Une attitude consiste alors à prendre les propriétés

comme axiomes définissant implicitement les nombres réels. Cependant le mathématicien ne sera pleinement satisfait que lorsqu'il aura donné des nombres réels une définition complète.

Une première idée, exploitée par Dedekind - Tannery, consiste à remarquer que les nombres rationnels ne possèdent pas la propriété (C). Les nombres réels sont alors introduits à partir de l'ensemble Λ (*) des coupures dans les nombres rationnels. Dans cette définition, la notion de nombre réel utilise la donnée de deux ensembles infinis de nombres rationnels. Il n'est pas ici possible d'esquiver leur considération par un raisonnement en compréhension. On se trouve en présence d'ensembles infinis qu'il faut manipuler comme des êtres mathématiques ordinaires.

Une autre façon de procéder, employée par Heine - Cantor, part de la notion de suite convergente. On peut en effet remarquer que la suite mise en évidence dans la recherche de la borne supérieure est une suite de Cauchy. On peut encore construire les nombres réels à partir des suites de Cauchy de nombres rationnels. Dans ce cas le nombre réel apparaîtra comme un ensemble infini d'une complexité encore plus grande, à savoir une classe d'équivalence de suites de nombres rationnels.

2. Ecueils et promesses.

C'est Bolzano qui étudie le premier les ensembles infinis en tant que tels dans son ouvrage sur les "Paradoxes de l'Infini". Le problème est de comparer entre eux divers ensembles; il introduit à ce sujet la notion d'équipotence qui est définie comme suit:

Définition 1.- On dit que deux ensembles E, F sont équipotents ou encore ont même puissance s'il existe une bijection entre E et F.

(*) La définition des nombres réels à partir des coupures dans les nombres rationnels a l'inconvénient de laisser une ambiguïté pour les coupures définissant un nombre rationnel. Une amélioration consiste à utiliser les sections commençantes ouvertes de nombres rationnels. Ces questions sont développées en détail dans une autre partie du cours de maîtrise.

Cette définition n'est pas sans susciter des paradoxes. Galilée déjà remarquait que l'application $n \mapsto n^2$ est une bijection des entiers naturels sur leurs carrés, ce qui allait à l'encontre de la règle du "tout plus grand que la partie", et ne faisait qu'alimenter la méfiance envers l'infini actuel.

La notion d'équipotence semble donc beaucoup trop grossière pour être digne d'intérêt. C'est du moins l'impression qu'on peut avoir lorsqu'en 1873 Cantor se pose le problème de savoir si l'ensemble $\mathbb{N}$ des nombres entiers est équipotent à l'ensemble $\mathbb{R}$ des nombres réels. L'intuition paraît forcer la conviction, car $\mathbb{N}$ est formé de points isolés alors que $\mathbb{R}$ constitue un continu. Cependant cet argument est sans valeur; on démontre d'ailleurs que l'ensemble $\mathbb{Q}$ des nombres rationnels est équipotent à $\mathbb{N}$. Autrement dit en posant la

Définition 2. – On dit qu'un ensemble E est dénombrable s'il est équipotent à $\mathbb{N}$,

il est possible de démontrer la

Proposition 1. – L'ensemble $\mathbb{Q}$ des nombres rationnels est dénombrable.

Dire qu'un ensemble est dénombrable signifie encore qu'il est possible de ranger ses éléments en une suite infinie $x_0, x_1, \dots, x_n, \dots$. C'est ce qu'il s'agit de faire pour l'ensemble des nombres rationnels. Posons $x_0 = 0$; le complémentaire de 0 dans $\mathbb{Q}$ est en bijection avec l'ensemble des couples (p, q) où $p \in \mathbb{Z}^*$, $q \in \mathbb{N}^*$ et p, q sont premiers entre eux.

Il n'existe aucun couple tel que $|p| + q = 0$ ni $|p| + q = 1$; on s'intéresse aux couples tels que $|p| + q = 2$ et on pose

$$x_1 = -1, \quad x_2 = 1,$$

puis aux couples tels que $|p| + q = 3$ et on définit

$$x_3 = -2, \quad x_4 = 2,$$

$$x_5 = -\frac{1}{2}, \quad x_6 = \frac{1}{2}.$$

Pour chaque entier naturel n , il n'existe qu'un nombre fini de couples (p, q) admissibles tels que $|p| + q \leq n$; on les classe en commençant par ceux pour lesquels l'entier q est le plus petit. De cette façon on arrive de proche en proche à ranger tous les nombres rationnels en une suite infinie et la proposition est démontrée.

On objectera que si $\mathbb{Q}$ n'est pas constitué de points isolés, il n'a pas la propriété suivante que possède en revanche $\mathbb{R}$: quels que soient le point a et le nombre réel r , il existe un point b dont la distance à a soit exactement r .

Il faut cependant rejeter ce nouvel argument : le premier faisait appel à une notion topologique alors que le second utilise une notion métrique. Or une bijection arbitraire n'est astreinte à conserver aucune de ces notions.

Le problème de l'équipotence qui vient d'être posé prend un intérêt tout particulier à la lumière d'un énoncé que Dedekind démontre en 1873 également.

Proposition 2.- L'ensemble des nombres réels qui sont algébriques est dénombrable.

On dit qu'un nombre réel est algébrique s'il est solution d'une équation algébrique

$$a_m X^m + a_{m-1} X^{m-1} + \dots + a_0 = 0,$$

dans laquelle les coefficients $a_0, \dots, a_m$ sont des entiers relatifs, avec $a_m \neq 0$, le degré m de l'équation étant lui-même arbitraire.

On s'intéresse à l'ensemble des nombres algébriques tels que

$$(m+1) |a_m| + \dots + 2 |a_1| + |a_0| = n,$$

où n est un entier naturel fixé et où $a_0, \dots, a_m$ sont des entiers naturels, coefficients d'une équation de degré m dont le nombre considéré est solution. Puisque $a_m \neq 0$, on a $m+1 \leq n$, de sorte que pour n fixé, il y a un nombre fini de possibilités pour m . Pour n, m fixés, il n'y a encore qu'un nombre fini de possibilités pour $a_0, \dots, a_m$. Enfin pour $a_0, \dots, a_m$ fixés, il y a au plus m solutions à l'équation algébrique de degré m correspondante.

A chaque valeur de l'entier n correspond donc un nombre fini de nombres algébriques. En donnant à n des valeurs successives, on peut donc ranger les nombres algébriques en une suite infinie.

Si l'on sait alors démontrer que $\mathbb{N}$ et $\mathbb{R}$ ne sont pas équipotents, on sait en déduire que l'ensemble des nombres algébriques n'est pas équipotent à $\mathbb{R}$. En particulier ces deux ensembles diffèrent et ainsi se trouve démontrée par la théorie des ensembles l'existence de nombres réels qui ne sont pas algébriques.

3. Le dénombrable et le continu.

C'est toujours en 1873 que Cantor apporte lui-même la solution du problème posé :

Proposition 3.- Les ensembles $\mathbb{N}$ et $\mathbb{R}$ ne sont pas équipotents.

Supposons en effet, par l'absurde, que tous les nombres réels puissent être rangés en une suite infinie $x_0, x_1, \dots, x_n, \dots$. Nous construisons par récurrence une suite d'intervalles $I_n = [a_n, b_n]$ de longueur non nulle de la façon suivante : $I_0 = [a_0, b_0]$ est un intervalle ne contenant pas x_0 ; supposant l'intervalle $I_{n-1} = [a_{n-1}, b_{n-1}]$ construit, nous choisissons un intervalle $I_n = [a_n, b_n]$ contenu dans I_{n-1} et ne contenant pas x_n .

Les intervalles I_n ainsi construits sont évidemment emboîtés ; un résultat fondamental de topologie de $\mathbb{R}$ montre qu'il s'agit d'une suite de compacts emboîtés non vides. Par conséquent les intervalles I_n ont un point commun x . Mais puisque $x_n \notin I_n$, on a nécessairement $x \neq x_n$ pour tout n , ce qui contredit l'hypothèse.

En fait Cantor avait commencé par donner du résultat une démonstration un peu plus compliquée. Il faut surtout noter l'intervention d'une propriété de topologie dans la démonstration, à savoir le fait que des intervalles $[a_n, b_n]$ emboîtés non vides ont toujours un point commun. Cette propriété était connue par Dedekind sous le nom de "principe de continuité".

4. Exploration du continu.

Une fois acquise l'existence d'au moins deux "puissances infinies" distinctes, l'idée vient naturellement d'étudier systématiquement les ensembles infinis. Ainsi dès 1874, Cantor se pose un nouveau problème : est-ce que $\mathbb{R}$ et $\mathbb{R}^2$, et plus généralement $\mathbb{R}^n$ sont équipotents ? Une telle éventualité semble absurde car cela reviendrait à dire que plusieurs variables indépendantes peuvent être ramenées à une seule. Cependant

Proposition 4.- Les ensembles $\mathbb{R}$ et $\mathbb{R}^n$ sont équipotents pour tout entier $n > 1$.

Ce n'est qu'en 1877 que Cantor aboutit à la solution. Il s'intéresse d'abord au problème sous une forme un peu différente et démontre que le segment $[0, 1]$ et le carré $[0, 1] \times [0, 1]$ sont équipotents.

L'idée de la démonstration est simple; elle est exposée dans une lettre à Dedekind du 20 juin 1977. Soient x et y des nombres réels de $[0, 1]$; on les écrit sous forme d'un développement décimal

$$x = \frac{\alpha_1}{10} + \frac{\alpha_2}{10^2} + \dots + \frac{\alpha_n}{10^n} + \dots$$

$$y = \frac{\beta_1}{10} + \frac{\beta_2}{10^2} + \dots + \frac{\beta_n}{10^n} + \dots,$$

où les α_n, β_n sont des entiers compris entre 1 et 9. De ces nombres x, y , on déduit un autre nombre z de $[0, 1]$ donné par

$$z = \frac{\alpha_1}{10} + \frac{\beta_1}{10^2} + \frac{\alpha_2}{10^3} + \frac{\beta_2}{10^4} + \dots + \frac{\alpha_n}{10^{2n-1}} + \frac{\beta_n}{10^{2n}} + \dots$$

Autrement dit, on obtient le développement décimal de z en intercalant les chiffres du développement de y dans ceux du développement de x . Inversement, partant d'un nombre z de $[0, 1]$ on peut obtenir x ou y par extraction des chiffres d'ordre impair ou pair du développement décimal de z . Ainsi on pense obtenir une bijection entre $[0, 1] \times [0, 1]$ et $[0, 1]$.

Il y a, hélas, une difficulté que ne manque pas de relever Dedekind. En effet le développement décimal n'est pas unique; il ne l'est pas pour les nombres décimaux de $]0, 1[$: un tel nombre peut s'écrire à l'aide d'une suite constamment égale à 0 à partir d'un certain rang ou d'une suite constamment égale à 9 à partir d'un certain rang. Si l'on veut qu'il y ait unicité, il faut exclure l'une des deux écritures, par exemple celle qui utilise, pour un nombre décimal de $]0, 1[$, une suite constamment égale à 0 à partir d'un certain rang.

Il n'y a alors plus de difficulté pour construire suivant le procédé indiqué le nombre z à partir des nombres x et y ; ainsi peut-on définir sans ambiguïté une application de $[0, 1] \times [0, 1]$ dans $[0, 1]$ dont on vérifie aussitôt qu'elle est injective. Malheureusement cette application n'est pas surjective. Par exemple le développement décimal autorisé

0, 1840905070.....

provient des développements décimaux

0, 14957..... et 0, 80000.....,

dont le second n n'est pas autorisé. Le nombre réel correspondant n n'est donc pas dans l'image de l'application considérée.

Nous avons seulement pu mettre en bijection le carré $[0, 1] \times [0, 1]$ avec une partie de $[0, 1]$ à savoir le complémentaire des nombres réels de $[0, 1]$ ayant un développement décimal $0, \alpha_1 \alpha_2 \alpha_3 \dots$ dans lequel tous les termes de rang impair ou bien tous les termes de rang pair sont nuls à partir d'un certain rang autre que le premier. Ce résultat surprend autant que celui annoncé : le carré est équipotent à une partie de l'intervalle.

Cantor corrige sa démonstration en utilisant le résultat suivant : tout nombre irrationnel de $]0, 1[$ peut être représenté par un développement en fraction continue

$$x = \frac{1}{\alpha_1 + \frac{1}{\alpha_2 + \frac{1}{\alpha_3 + \dots + \frac{1}{\alpha_n + \dots}}}}$$

et l'ensemble E des nombres irrationnels de $]0, 1[$ est de cette manière en correspondance bijective avec l'ensemble des suites $\alpha_1, \alpha_2, \dots, \alpha_n, \dots$ de nombres entiers > 0 . En utilisant le même procédé que précédemment, on construit alors une bijection de $E \times E$ sur E . La démonstration s'achève alors aisément si l'on sait construire une bijection de $[0, 1]$ sur E .

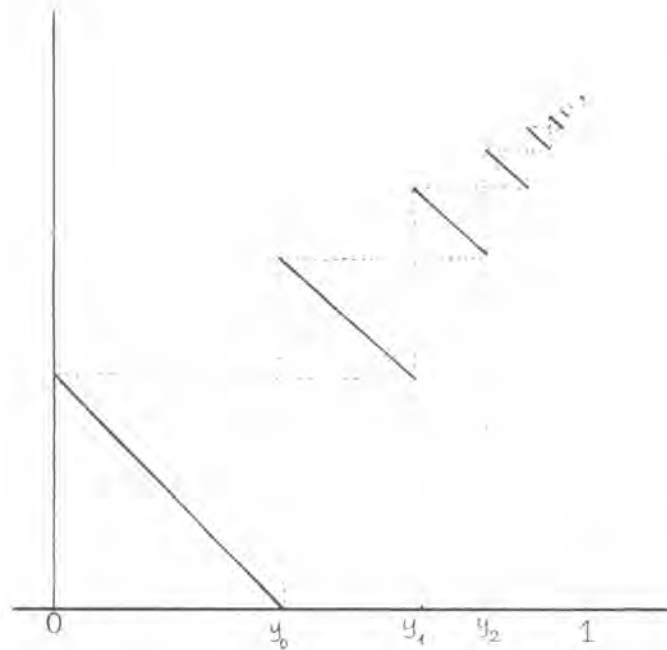
D'abord $[0, 1]$ est réunion disjointe de E et de l'ensemble $[0, 1] \cap \mathbb{Q}$ qui, d'après la proposition 1, peut être rangé en suite infinie $x_0, x_1, \dots, x_n, \dots$. Choisissons dans E une suite $y_0, y_1, \dots, y_n, \dots$ strictement croissante et tendant vers 1, et désignons par F le complémentaire dans $[0, 1]$ de l'ensemble des points de cette suite. Nous définissons une bijection de F sur E en posant

$$f(x_n) = y_n \quad \text{pour } n = 0, 1, \dots$$

et

$$f(x) = x \quad \text{pour } x \in E \cap F.$$

Il ne reste plus qu'à construire une bijection de $[0, 1]$ sur F . On se ramène pour cela à mettre en bijection les intervalles $[y_n, y_{n+1}[$ et $]y_n, y_{n+1}[$, ou encore les intervalles $[0, 1]$ et $]0, 1]$, ce qui peut être fait explicitement par la fonction dont le graphe est donné par



(L'extrémité droite de chaque segment représenté est supposée ne pas appartenir au graphe).

Nous laissons enfin au lecteur le soin de construire à partir de là une bijection entre $\mathbb{R}$ et $\mathbb{R}^2$ et d'étendre la propriété à $\mathbb{R}^n$.

5. Développements.

Nous n'allons pas entreprendre dès maintenant une étude systématique de la puissance des ensembles infinis car nous forgerons plus tard des outils qui rendront cette dernière beaucoup plus aisée. Signalons cependant que, presque vingt ans plus tard, en 1892, Cantor donne une nouvelle démonstration de la proposition 3, fondée sur le procédé diagonal, et qui est la suivante : on suppose par l'absurde que $\mathbb{R}$ est dénombrable ; il en est alors de même de l'ensemble $[0, 1[$ dont on peut ranger les éléments en une suite infinie $x_0, x_1, \dots, x_n, \dots$. Ecrivons pour tout n le développement décimal

$$0, \alpha_{n1} \alpha_{n2} \dots \alpha_{np} \dots$$

de x_n , qui est bien déterminé si l'on interdit les suites égales à 9 à partir d'un certain rang. Pour tout $n = 0, 1, \dots$, choisissons un nombre entier β_n compris entre 0 et 8 et distinct de α_{nn} . Le nombre réel x ayant pour développement décimal

$$0, \beta_1 \beta_2 \dots \beta_n \dots$$

appartient à $]0, 1[$ et est distinct de tous les nombres x_n , ce qui est absurde.

La même méthode permet d'établir que l'ensemble $\mathcal{P}(\mathbb{N})$ de parties de $\mathbb{N}$ n'est pas dénombrable. En effet une partie de $\mathbb{N}$ peut être représentée par une fonction caractéristique, c'est à dire une suite formée de 0 et de 1. On pourrait d'ailleurs montrer, en utilisant encore les développements en fraction continue des nombres irrationnels de $]0, 1[$, que $]0, 1[$ est équipotent à $\mathcal{P}(\mathbb{N})$; nous ne le ferons pas car les théorèmes généraux permettront d'obtenir ce résultat de façon plus aisée. Une généralisation est donnée par le

Théorème (Cantor). - Soit E un ensemble quelconque. Il n'existe pas d'application surjective de l'ensemble E sur l'ensemble $\mathcal{P}(E)$ des parties de E.

Soit en effet f une application de E dans $\mathcal{P}(E)$ et désignons par A l'ensemble des éléments x de E qui n'appartiennent pas à la partie $f(x)$. Clairement A est une partie de E et cependant A n'est pas l'image de f . Si l'on a en effet $A = f(a)$, l'hypothèse $a \in A$ est absurde puisqu'elle implique $a \in f(a)$ et donc $a \notin A$ d'après la définition de A , ainsi que l'hypothèse contraire qui implique $a \notin f(a)$ et donc $a \in A$.

Le théorème de Cantor permet de construire, à côté des deux puissances infinies déjà étudiées, le dénombrable et le continu, d'autres puissances infinies. Ainsi $\mathcal{P}(\mathbb{R})$ n'est pas équipotent à $\mathbb{R}$, ni à $\mathbb{N}$. Il est même possible de construire une infinité de puissances infinies distinctes à savoir

$$\mathbb{N}, \mathcal{P}(\mathbb{N}), \mathcal{P}(\mathcal{P}(\mathbb{N})), \mathcal{P}(\mathcal{P}(\mathcal{P}(\mathbb{N}))) \dots$$

Note : La proposition 4 conduit à s'interroger sur le sens à donner à l'expression "coordonnées indépendantes". Cet énoncé montre en effet que $\mathbb{R}^2$ peut être décrit à l'aide d'une seule variable réelle. Il ne faut cependant pas en tirer des conclusions trop hâtives : les changements de coordonnées utilisés par les géomètres dans l'étude des courbes et surfaces sont toujours supposés au moins continus dans les deux sens; or l'application bijective de $[0, 1] \times [0, 1]$ qui est décrite non seulement n'est pas bicontinue, mais suivant les termes de Dedekind lui-même est "telle que toute partie connexe, si petite soit-elle, a une image complètement déchirée".

II

AXIOMES DE ZERMELO-FRAENKEL

1. Découverte d'antinomies.

L'étude entreprise par Cantor allait déboucher, à la fin du dix-neuvième siècle, sur la découverte de contradictions, ou antinomies, qui devaient provisoirement ébranler la construction exposée au chapitre qui précède.

La première contradiction, connue sous le nom d'antonomie Burali-Forti, apparaît dans la théorie des nombres transfinis vers les années 1895-1897. Nous en parlerons plus loin, nous contentant de noter qu'elle n'a pas eu un retentissement considérable; elle se situe en effet dans un domaine très particulier et la réaction a été de mettre en doute la construction des nombres transfinis plutôt que les bases de la théorie des ensembles.

Une autre contradiction est découverte en 1899 par Cantor lui-même: Considérons l'ensemble Ω de tous les objets mathématiques; les parties de Ω sont elles-mêmes des éléments de Ω de sorte que Ω contient $\mathcal{P}(\Omega)$; on obtient facilement une surjection de Ω sur $\mathcal{P}(\Omega)$ en associant à toute partie de Ω cette partie elle-même et à tout objet qui n'est pas une partie de Ω une partie choisie une fois pour toutes; cela contredit alors directement le théorème de Cantor énoncé au chapitre précédent.

En 1903, Russel, après Zermelo, donne à cette contradiction une forme extrêmement simple, qui porte son nom: considérons l'ensemble E des objets mathématiques x tels que $x \notin x$; l'hypothèse $E \in E$ est contradictoire puisqu'elle implique $E \notin E$ à cause de la définition E et l'hypothèse $E \notin E$ également. Cette dernière contradiction se situe maintenant au fondement même de la théorie.

Vers la même époque surgissent un certain nombre d'antinomies sémantiques comme l'antonomie Richard, celle du menteur... qui ont trait à ce que l'on pourrait appeler la "logique naïve". Nous ne nous étendrons pas sur ce

sujet.

Aux yeux de certains mathématiciens, comme par exemple Poincaré qui était pourtant au début un défenseur de la théorie des ensembles, le discrédit est alors jeté vers tout ce qui touche à cette théorie. Il semble au contraire à d'autres, comme Hilbert, qu'il serait dommage de renoncer à ce qui commence à apparaître comme un puissant outil de démonstration : "Du paradis mathématique que Cantor a créé pour nous, dit-il, nul ne doit nous chasser". Le problème consiste alors à essayer d'éliminer les contradictions de la théorie. Après un certain nombre de tentatives philosophiques vouées à l'échec pour distinguer les propriétés pathologiques des autres, on s'est tourné vers une formalisation de la théorie qui pouvait seule résoudre la difficulté.

2. Axiomes de Zermelo.

La théorie des ensembles utilise en principe, pour éviter les embûches de la langue courante, un langage formalisé, à savoir un langage doté du calcul des prédicats du premier ordre et du signe d'égalité $=$. Ce calcul contient des connecteurs exprimant la négation, la disjonction, la conjonction, l'implication et l'équivalence, pour lesquels nous utiliserons les signes "non", "ou", "et", $\Rightarrow$ et $\Leftrightarrow$, ainsi que des quantificateurs existentiel et universel, $\exists$ et $\forall$. La théorie des ensembles introduit un signe spécifique, à savoir le signe d'appartenance $\in$, qui est un signe relationnel de poids 2.

Si l'on ne veut pas se référer à un traité de logique, on considère simplement une collection d'objets entre lesquels la théorie des ensembles introduit une relation binaire

$$x \in y$$

Cette relation, dite relation d'appartenance, se lit "x appartient à y" ou "x est un élément de y" (*).

Il faut noter qu'on ne définit pas le mot ensemble ou le mot élément; ces mots sont en général considérés comme synonymes d'objets de la collection considérée: on emploie simplement l'un ou l'autre pour attirer l'attention du lecteur sur la place de l'objet en question dans la relation d'appartenance $x \in y$.

Nous allons maintenant énumérer un certain nombre d'axiomes de la théorie, en adoptant l'exposition originale de Zermelo, qui date de 1908. Dans ce qui suit, la relation non $(x \in y)$ sera notée $x \notin y$.

(*) On trouve dans certains ouvrages l'expression barbare "x est élément de y"; elle provient probablement d'un excès de pédantisme qui voudrait que la relation $x \in y$ exprime seulement une propriété de x, sans référence au fait fondamental qu'un ensemble est "constitué d'éléments".

Axiome I de détermination ou d'extensionnalité.

Cet axiome signifie qu'un ensemble est entièrement déterminé par ses éléments; ou encore qu'un ensemble peut être défini en extension, c'est à dire par une énumération complète de ses éléments. Il s'énonce :

Soient x, y des ensembles; si tout élément de x est un élément de y et inversement, alors $x = y$,

ce qui, en langage formalisé peut s'écrire

$$(\forall x)(\forall y) \left[(\forall z)((z \in x) \iff (z \in y)) \implies (x = y) \right] .$$

Il est possible de donner une autre forme à l'axiome I en introduisant la

Définition 1.- Soient x, y des ensembles; on dit que x est inclus dans y , ou que x est une partie ou un sous-ensemble de y , et on note $x \subset y$, si tout élément de x est un élément de y .

Il convient de distinguer soigneusement les relations d'appartenance $x \in y$ et d'inclusion $x \subset y$. Cette distinction a constitué un pas important dans l'élaboration progressive de la théorie des ensembles. Malheureusement le langage mathématique, qui s'est formé avant le théorie, multiplie les exemples de confusions. Par exemple on dira pour $x \in y$, en plus des expressions citées "y contient x", et, bien que plus rarement, "x est contenu dans y". Or ces deux tournures s'emploient, et très couramment, pour exprimer l'inclusion $x \subset y$. On voit d'ailleurs que pour interpréter ces relations comme une propriété de y on ne dispose pratiquement que d'une seule et même expression.

On vérifie maintenant aussitôt que l'axiome I s'énonce :

Soient x, y des ensembles; si $x \subset y$ et $y \subset x$, alors $x = y$.

Axiome II des ensembles élémentaires.

Cet axiome se décompose lui-même en plusieurs parties.

a) Axiome de l'ensemble vide

On pose l'existence d'un ensemble n'ayant aucun élément; en langage formalisé cela s'écrit

$$(\exists x)(\forall y)(y \notin x) .$$

D'après l'axiome d'extensionnalité, un tel ensemble est nécessairement unique. On lui donne le nom d'ensemble vide et on le note $\emptyset$.

b) Axiome de l'ensemble à un élément

On pose, pour tout ensemble x , l'existence d'un ensemble ayant pour seul élément x ; cela peut s'écrire

$$(\forall x)(\exists y)(\forall z) \left[(z \in y) \iff (z = x) \right] .$$

Comme précédemment un tel ensemble est unique; il est noté $\{x\}$.

c) Axiome de la paire

On pose, pour des ensembles x, y , l'existence d'un ensemble dont les éléments sont x et y .

Cet ensemble est encore unique, il est noté $\{x, y\}$. Remarquons aussitôt que $\{x, x\} = \{x\}$.

Ce dernier axiome nous permet d'introduire une notion nouvelle.

Définition 2.- Soient a, b , des ensembles. On appelle couple des ensembles a, b , et on note (a, b) l'ensemble

$$\{\{a\}, \{a, b\}\}.$$

On emploie parfois le terme de "paire ordonnée" pour désigner le couple. L'intérêt de la notion de couple réside dans la

Proposition 1.- Soient a, b, a', b' des ensembles; alors $(a, b) = (a', b')$ si et seulement si $a = a', b = b'$.

La condition est évidemment suffisante à cause des propriétés de l'égalité. Supposons inversement que $(a, b) = (a', b')$, c'est à dire

$$\{\{a\}, \{a, b\}\} = \{\{a'\}, \{a', b'\}\}.$$

Si $a = b$, alors l'ensemble de gauche se réduit à $\{\{a\}\}$ et n'a qu'un élément; il doit alors en être de même de l'ensemble de droite, ce qui exige $a' = b'$; dans ce cas, $\{\{a\}\} = \{\{a'\}\}$, d'où $\{a\} = \{a'\}$ et $a = a' = b = b'$.

Si $a \neq b$, l'ensemble de gauche a deux éléments, donc aussi celui de droite, ce qui exige $a' \neq b'$. De plus, ou bien

$$\{a\} = \{a', b'\} \quad \text{et} \quad \{a, b\} = \{a'\}$$

ou bien

$$\{a\} = \{a'\} \quad \text{et} \quad \{a, b\} = \{a', b'\}.$$

La première hypothèse est impossible car $\{a', b'\}$ a deux éléments; par suite $a = a'$ puis $b = b'$.

Axiome III de séparation ou de compréhension.

Il s'agit d'un schéma d'axiomes qui exprime qu'un ensemble peut être défini en séparant les éléments d'un ensemble donné qui satisfont une certaine propriété. Il s'énonce

Etant donnés un ensemble x et une relation $R(u)$ à un argument, il existe un ensemble y dont les éléments sont exactement les éléments de x qui satisfont la relation donnée.

On peut écrire :

$$(\forall x)(\exists y)(\forall z) \left[(z \in y) \Leftrightarrow ((z \in x) \text{ et } R(z)) \right] .$$

De façon générale on dit qu'une relation $R(u)$ à un argument est collectivisante s'il existe un ensemble dont les éléments sont exactement les ensembles qui satisfont cette relation. L'axiome de compréhension signifie simplement que toute relation de la forme

$$(u \in a) \text{ et } R(u)$$

est collectivisante. Il faut noter qu'il existe des relations non collectivisantes; par exemple la relation $u \notin u$ ne l'est pas : supposons qu'il existe un ensemble a dont les éléments soient exactement les ensembles x tels que $x \notin x$; on vérifie alors aussitôt que $a \in a$ équivaut à $a \notin a$, ce qui est contradictoire.

C'est la restriction imposée aux relations par l'axiome III qui permet d'éviter la considération d'ensembles "trop gros" et par conséquent les contradictions citées au début du chapitre.

Axiome IV de l'ensemble des parties.

On pose, pour tout ensemble x , l'existence d'un ensemble dont les éléments sont les parties de x .

Cela peut s'écrire

$$(\forall x)(\exists y)(\forall z) \left[(z \in y) \Leftrightarrow (z \subset x) \right] .$$

Un tel ensemble est bien sûr unique; on l'appelle ensemble des parties de x et on le note $\mathcal{P}x$, $\mathcal{P}(x)$ ou $\mathcal{P}(x)$. On notera que cet axiome signifie simplement que la relation $u \subset a$ est collectivisante.

Axiome V de l'ensemble réunion

Cet axiome permet de réunir les ensembles d'un ensemble donné; il s'énonce : Etant donné un ensemble x , il existe un ensemble dont les éléments sont les éléments des éléments de x .

En langage formalisé:

$$(\forall x)(\exists y)(\forall z) \left[(z \in y) \Leftrightarrow (\exists t)((z \in t) \text{ et } (t \in x)) \right] ,$$

ou encore :

la relation $(\exists t)((u \in t) \text{ et } (t \in a))$ est collectivisante.

Nous arrêtons provisoirement l'énumération des axiomes de Zermelo: nous aurons l'occasion d'introduire quelques autres axiomes au cours des prochains chapitres.

3. Remarques.

Nous avons donné un groupe de cinq axiomes, en restant à peu près fidèles à la présentation adoptée par Zermelo. Il est bien sûr possible de développer une axiomatique semblable en introduisant des variantes dans l'agencement des axiomes.

Une première remarque concerne l'axiome d'extensionnalité. Cet axiome donne en réalité une caractérisation de l'égalité entre ensembles. Nous avons jusqu'ici considéré la notion d'égalité comme préexistant à la théorie. Un autre point de vue, qui est celui de Fraenkel par exemple, consiste au contraire à se servir de l'axiome pour définir l'égalité entre ensembles. Autrement dit on pose, compte tenu de la définition 1, la

Définition .- Soient x, y des ensembles; on dit que x et y sont égaux et on pose $x = y$ si à la fois $x \subset y$ et $y \subset x$.

Naturellement il faut que la relation $x = y$ ainsi définie possède les propriétés de l'égalité, à savoir qu'étant donnée une relation $R(x)$, si $R(a)$ est vraie et $a = b$, alors $R(b)$ est vraie. Nous avons à effectuer cette vérification en particulier pour les relations $x \in a$ et $a \in x$. Il n'y a pas de difficulté pour la seconde; pour la première en revanche, il faut introduire un axiome :

Soient x, y, z des ensembles; si $x \in z$ et $x = z$, alors $y \in z$.

Nous ne vérifierons pas que cet axiome suffit à conférer à la relation $x = y$ les propriétés souhaitées, car nous ne voulons pas insister sur la manière dont sont construites les relations.

A cet égard les deux attitudes que nous venons de distinguer semblent s'opposer: dans un cas on place l'égalité dans la logique (*) et dans le

(*) D'un point de vue plus naïf cela revient à interpréter l'égalité $a=b$ comme le fait que a et b représentent le même objet mathématique, la différence éventuelle d'apparence n'étant due qu'à l'utilisation de symboles abrégiateurs.

second cas on la définit. Cependant, tous les objets mathématiques étant des ensembles, au moins ceux que l'on rencontre dans les cours de maîtrise, l'égalité n'aura plus à être définie. De cette façon, le choix d'une attitude ou d'une autre n'a pas d'incidence sur le reste des mathématiques.

On se gardera ainsi dans tous les cas de se croire obligé de définir l'égalité entre certains objets mathématiques, de prétendre par exemple que deux groupes sont dits égaux s'ils ont même ensemble sous-jacent et même loi : un groupe est la donnée d'un couple (G, τ) où G est un ensemble et τ une application de $G \times G$ dans G ayant certaines propriétés; comme on le sait $(G, \tau) = (G', \tau')$ si et seulement si $G = G'$ et $\tau = \tau'$. Il n'est en revanche pas nuisible de rappeler ou de caractériser cette égalité.

Une seconde remarque concerne l'indépendance des axiomes. Par exemple, l'axiome de l'ensemble à un élément est une conséquence immédiate de l'axiome de la paire.

Compte tenu de l'axiome de séparation, l'axiome de l'ensemble vide peut d'ailleurs être remplacé par l'axiome :

Il existe un ensemble.

En effet, si a est un ensemble, l'ensemble des éléments x de a tels que $x \notin x$ est un ensemble qui n'a aucun élément.

L'axiome de la paire peut être de la même façon remplacé par l'axiome

Si x, y sont des ensembles, il existe un ensemble dont x, y soient des éléments.

Si a est un tel ensemble, on vérifie aussitôt que l'ensemble des éléments z de a tels que $z = x$ ou $z = y$ est l'ensemble paire cherché.

4. Constructions complémentaires.

A partir de maintenant, nous ne donnerons plus la traduction en langage formalisé des énoncés qui interviendront. Nous utiliserons le plus souvent des majuscules $E, F, G, X, Y, Z \dots$ pour désigner des ensembles, notant par des minuscules $a, b, c, x, y, z \dots$ leurs éléments (qui sont bien sûr eux aussi des ensembles).

Intersection de deux ensembles : Soient E, F des ensembles; les éléments communs à E et F , qui sont par exemple les éléments x de E vérifiant $x \in F$, constituent, d'après l'axiome de séparation, un ensemble. Cet ensemble est appelé intersection des ensembles E et F et noté $E \cap F$.

Complémentaire d'une partie : Soient E un ensemble et A une partie de E ; les éléments de E qui n'appartiennent pas à A constituent, toujours d'après l'axiome de séparation, un ensemble. Cet ensemble est appelé complémentaire de A dans E et noté $\complement_E A$ ou encore $E \setminus A$; c'est une nouvelle partie de E .

Réunion de deux ensembles : Soient E, F deux ensembles; nous allons montrer que les ensembles x qui appartiennent à E ou à F constituent un ensemble. Cet ensemble sera appelé réunion des ensembles E et F et noté $E \cup F$.

On considère pour cela la paire $\{E, F\}$ et l'ensemble somme de cette paire. Les éléments de ce dernier ensemble sont exactement les x pour lesquels il existe un élément X de $\{E, F\}$ tel que $x \in X$, c'est à dire les x pour lesquels $x \in E$ ou $x \in F$; l'ensemble somme de $\{E, F\}$ répond à la question.

Nous ne donnerons pas les propriétés des opérations de réunion, intersection, passage au complémentaire sur l'ensemble des parties d'un ensemble; ces propriétés sont bien connues et le lecteur est renvoyé aux ouvrages élémentaires qui en traitent.

Produit de deux ensembles : Soient E, F des ensembles; nous allons montrer que les couples (x, y) où $x \in E, y \in F$ constituent un ensemble. Cet ensemble sera appelé produit des ensembles E et F et noté $E \times F$.

Il s'agit en somme de voir que la relation

$$(\exists t)(\exists u) [(x = (t, u)) \text{ et } (t \in E) \text{ et } (u \in F)]$$

est collectivisante en x . Pour pouvoir appliquer l'axiome de séparation, il suffit de montrer que cette relation implique l'appartenance à un ensemble fixe. Or, si $x = (t, u)$ avec $t \in E$ et $u \in F$, par définition

$$x = \{\{t\}, \{u\}\}.$$

Chaque ensemble $\{t\}, \{u\}$ est respectivement une partie de E, F ; chacun est encore une partie de $E \cup F$, c'est à dire un élément de $\mathcal{P}(E \cup F)$. Par suite x est une paire de $\mathcal{P}(E \cup F)$, et donc un élément de $\mathcal{P}(\mathcal{P}(E \cup F))$.

Triplet : Soient a, b, c des ensembles; on définit le triplet (a, b, c) de ces ensembles comme l'ensemble

$$((a, b), c).$$

On vérifie aussitôt que $((a, b), c) = ((a', b'), c')$ si et seulement si $a = a', b = b', c = c'$.

Nous laissons au lecteur le soin de définir les ensembles $E \cap F \cap G$, $E \cup F \cup G$, $E \times F \times G \dots$ en suivant les méthodes qui précèdent.

III

CORRESPONDANCES, APPLICATIONS

1. Correspondances.

La notion de correspondance est une généralisation de celle d'application dans laquelle on n'attache pas à un élément de l'ensemble de départ une valeur unique mais un ensemble de valeurs possibles.

Définition 1.- On appelle correspondance un triplet (E, F, Γ) où E, F sont des ensembles et Γ est une partie de l'ensemble $E \times F$.

L'ensemble E est appelé ensemble de départ ou source de la correspondance; l'ensemble F est appelé ensemble d'arrivée ou but; l'ensemble Γ est appelé graphe.

La relation $(x, y) \in \Gamma$ se lit " y est associé à x dans la correspondance". L'ensemble des $x \in E$ auxquels est associé au moins un $y \in F$ est l'ensemble de définition ou le domaine de la correspondance; l'ensemble des $y \in F$ qui sont associés à au moins un $x \in E$ est l'ensemble des valeurs ou image.

Donnons tout de suite quelques exemples :

a) Etant donné un ensemble E , on appelle diagonale de $E \times E$ l'ensemble des couples (x, y) de $E \times E$ tels que $x = y$. La correspondance (E, E, Δ) est la correspondance identique de E ; à chaque $x \in E$, cette correspondance associe le seul élément x .

b) Soit (E, F, Γ) une correspondance; on note $\bar{\Gamma}^1$ l'ensemble des couples (y, x) de $F \times E$ tels que $(x, y) \in \Gamma$. La correspondance $(F, E, \bar{\Gamma}^1)$ est la correspondance réciproque de la correspondance considérée.

c) Soient (E, F, Γ) et (F, G, Γ') deux correspondances, le but de la première étant la source de la seconde. Désignons par $\bar{\Gamma}'^1 \circ \bar{\Gamma}^1$ l'ensemble

des couples (x, z) de $E \times G$ tels qu'il existe $y \in F$ vérifiant $(x, y) \in \Gamma$ et $(y, z) \in \Gamma'$. La correspondance $(E, G, \Gamma' \circ \Gamma)$ est la correspondance composée des correspondances données.

On notera que la composée de (E, F, Γ) et de (F, E, Γ) n'est pas nécessairement la correspondance identique de E .

Remarque : on se gardera bien de confondre une correspondance (E, F, Γ) avec son graphe Γ . La seule donnée du graphe ne permet pas en effet de déterminer l'ensemble de départ ou celui d'arrivée. En revanche l'ensemble de définition et celui des valeurs ne dépendent que du graphe.

On peut introduire la notion de graphe en dehors de celle de correspondance. Un graphe est par définition un ensemble dont tous les éléments sont des couples. Etant donné un graphe Γ , on démontre qu'il existe toujours des ensembles E, F tels que $\Gamma \subset E \times F$; ou encore, tout graphe est le graphe d'au moins une correspondance. En effet la relation $(x, y) \in \Gamma$ n'est autre que

$$\{\{x\}, \{x, y\}\} \in \Gamma.$$

Elle implique que $\{x, y\}$ appartient à l'ensemble somme de Γ , puis que x, y appartiennent à l'ensemble somme G de ce dernier ensemble; il en résulte que $\Gamma \subset G \times G$, ce qu'il fallait démontrer.

2. Applications.

On dit qu'un graphe Γ est fonctionnel si les relations $(x, y) \in \Gamma$ et $(x, y') \in \Gamma$ impliquent $y = y'$.

Définition 2.- Une application est une correspondance (E, F, Γ) telle que Γ soit un graphe fonctionnel et que E soit l'ensemble de définition.

Dire que E est l'ensemble de définition signifie qu'à tout $x \in E$ est associé au plus un élément de F . Une correspondance (E, F, Γ) est donc une application si à tout élément de E est associé un élément et un seul de F . Si f désigne l'application (E, F, Γ) , on convient de noter $f(x)$ l'unique élément de F qui est associé par f à un élément x de E .

On emploie souvent la notation $f : E \rightarrow F$ pour désigner une application de E dans F et une notation comme $f : x \mapsto \sqrt[3]{x}$ pour désigner l'application qui fait correspondre l'élément $\sqrt[3]{x}$ à l'élément x , les ensembles de définition de valeurs étant supposés précisés par ailleurs. On pourra ainsi écrire :

"Etant données des applications injectives $f : E \rightarrow F$ et $g : F \rightarrow G$, l'application $g \circ f : E \rightarrow G$ est injective".

"La fonction réelle de variable réelle $f : x \mapsto x^4 \sin x^3$ est dérivable".

On évitera en revanche de mêler ces deux symbolismes et d'employer une écriture comme

$$\text{"Soit } f : \mathbb{R}_+ \rightarrow \mathbb{R} \\ x \mapsto f(x) = x + \sqrt{x(x^2+1)} \text{"}$$

à la fois redondante, inesthétique et encombrante. Dans ce dernier cas il vaut mieux écrire : "soit f l'application de $\mathbb{R}_+$ dans $\mathbb{R}$ définie par

$$f(x) = x + \sqrt{x(x^2+1)} \text{"}$$

En principe le mot de fonction est synonyme de celui d'application. Son usage est plutôt réservé à certains cas particuliers : si l'espace d'arrivée est $\mathbb{R}$ (resp. $\mathbb{C}$) on parlera de fonction numérique réelle (resp. complexe); si c'est un espace vectoriel on parlera de fonction vectorielle... Dans d'autres cas, comme pour les fonctions rationnelles, l'usage a consacré un abus de langage : ainsi, lorsqu'un corps a été choisi, celui des nombres réels par exemple, par la fonction rationnelle f définie par

$$f(x) = \frac{x+1}{x^2+4x+4},$$

on attend l'application de $\mathbb{R} \setminus \{-2\}$ dans $\mathbb{R}$ qui est définie par la formule qui précède. En fait la véritable donnée est la fraction rationnelle

$$\frac{X+1}{X^2+4X+4};$$

il lui est associé sans ambiguïté une application dont le domaine de définition est l'ensemble des éléments du corps qui s'annulent par le dénominateur, dans une écriture sous forme irréductible. Si l'on doit effectuer des compositions, on le fait toujours sur les fractions (*).

Ainsi l'idée de fonction peut parfois se rattacher à celle d'application d'un sous-ensemble d'un ensemble donné dans un autre et c'est peut-être la raison pour laquelle le mot "fonction" désigne, dans l'enseignement du second degré, une correspondance dont le graphe est fonctionnel. Cette attitude attire plus de complications que d'avantages : la composition est une opération délicate et, surtout l'ensemble de définition n'apparaît pas toujours clairement (penser à $\frac{x}{x}$, $(\sqrt{x})^4$, $x^{4/2}$, 2^x ...), ce qui nécessite de le rappeler chaque fois.

(*) Une autre convention consiste à ajouter un élément à l'infini au corps, mais nous n'insisterons pas sur ce point.

Exemples : a) La correspondance identique (E, E, Δ) est une application; elle est encore appelée application identique de E et notée Id_E .

b) Soit E un ensemble; la correspondance $(\emptyset, E, \emptyset)$ est une application qui prend nom d'application vide.

c) Soient E, F des ensembles et a un élément de F . La correspondance $(E, F, E \times \{a\})$ est une application; l'ensemble des valeurs de cette application a au plus un élément: on dit que c'est une application constante.

d) Si $f = (E, F, \Gamma)$ et $g = (F, G, \Gamma')$ sont des applications, la correspondance composée de f et g est une application appelée encore application composée de f et g ; on la note $g \circ f$.

On vérifie aussitôt la propriété $(g \circ f)(x) = g(f(x))$.

On notera en revanche que la correspondance réciproque d'une application n'est en général pas une application; nous reviendrons plus loin sur ce point.

e) Soient E, F des ensembles. Considérons l'ensemble Γ des couples $((x, y), x)$ de $(E \times F) \times E$. La correspondance $(E \times F, E, \Gamma)$ est une application, appelée première projection de $E \times F$ et notée pr_1 ; elle vérifie $\text{pr}_1((x, y)) = x$. On définit de façon semblable la seconde projection pr_2 .

Etant donnés deux ensembles E, F , les applications de source E et de but F constituent un ensemble appelé ensemble des applications de E dans F et noté $\mathcal{F}(E, F)$; cet ensemble se définit simplement comme une partie de l'ensemble $\{E\} \times \{F\} \times \mathcal{G}(E \times F)$ des correspondances de source E et de but F .

3. Injections, surjections, bijections.

Définition 3.- Soient E, F des ensembles et f une application de E dans F .

On dit que f est injective, ou que f est une injection de E dans F , si la relation $f(x) = f(x')$ implique la relation $x = x'$.

On dit que f est surjective, ou que f est une surjection de E sur F si l'image de E par f est égale à F tout entier; on dit encore que f est une application de E sur F .

On dit que f est bijjective, ou que f est une bijection de E sur F si f est à la fois injective et surjective.

Une bijection d'un ensemble E sur lui-même s'appelle encore une permutation de E .

Dire que f est injective signifie que tout élément y de F est associé à au plus un élément de E par f . Si Γ est le graphe de f , c'est encore dire que Γ^{-1} est un graphe fonctionnel.

Dire que f est surjective signifie que tout élément y de F est associé à un élément de E par f . Par suite f est bijective si et seulement si tout élément de F est associé à un élément et un seul de E ; c'est encore dire que la correspondance Γ^{-1} est une application.

Il est classique de paraphraser ces définitions d'une autre manière en introduisant la notion d'équation. Etant donnés deux ensembles E, F , et deux applications f, g , de E dans F , la relation

$$(1) \quad f(x) = g(x)$$

où x est un élément de E , s'appelle équation définie par f et g . De façon non abrégée cette relation s'écrit

$$(\exists y) [(x, y) \in \Gamma \text{ et } (x, y) \in \Gamma']$$

où Γ, Γ' sont les graphes respectifs de f, g .

Tout élément x de E qui satisfait (1) s'appelle solution de l'équation.

Un cas particulier fréquent est celui où g est l'application constante égale à y_0 ; l'équation (1) s'écrit dans ce cas sous la forme

$$f(x) = y_0.$$

Dire que f est injective (resp. surjective, bijective) signifie que pour tout élément y de E , l'équation $f(x) = y$ admet au plus une solution (resp. au moins une solution, une solution et une seule).

a) Soient E un ensemble et A une partie de E . L'application de A dans E , qui à chaque élément x de A associe l'élément x lui-même, est une injection; on l'appelle injection canonique de A dans E .

b) Soit E un ensemble; l'application de E dans $E \times E$ qui à l'élément x de E associe le couple (x, x) est une injection; on l'appelle application diagonale de E .

c) Si f est une bijection de E sur F , l'application réciproque f^{-1} est une bijection réciproque de f . La relation $x = f^{-1}(y)$ équivaut à la relation $y = f(x)$; on a

$$\begin{array}{c} \xrightarrow{-1} \\ \xleftarrow{-1} \\ f = f \end{array}$$

et $f^{-1} \circ f$ (resp. $f \circ f^{-1}$) est l'application identique de E (resp. de F).

Exemples : a) La correspondance identique (E, E, Δ) est une application; elle est encore appelée application identique de E et notée Id_E .

b) Soit E un ensemble; la correspondance $(\emptyset, E, \emptyset)$ est une application qui prend nom d'application vide.

c) Soient E, F des ensembles et a un élément de F . La correspondance $(E, F, E \times \{a\})$ est une application; l'ensemble des valeurs de cette application a au plus un élément: on dit que c'est une application constante.

d) Si $f = (E, F, \Gamma)$ et $g = (F, G, \Gamma')$ sont des applications, la correspondance composée de f et g est une application appelée encore application composée de f et g ; on la note $g \circ f$.

On vérifie aussitôt la propriété $(g \circ f)(x) = g(f(x))$.

On notera en revanche que la correspondance réciproque d'une application n'est en général pas une application; nous reviendrons plus loin sur ce point.

e) Soient E, F des ensembles. Considérons l'ensemble Γ' des couples $((x, y), x)$ de $(E \times F) \times E$. La correspondance $(E \times F, E, \Gamma')$ est une application, appelée première projection de $E \times F$ et notée pr_1 ; elle vérifie $\text{pr}_1((x, y)) = x$. On définit de façon semblable la seconde projection pr_2 .

Etant donnés deux ensembles E, F , les applications de source E et de but F constituent un ensemble appelé ensemble des applications de E dans F et noté $\mathcal{F}(E, F)$; cet ensemble se définit simplement comme une partie de l'ensemble $\{E\} \times \{F\} \times \mathcal{G}(E \times F)$ des correspondances de source E et de but F .

3. Injections, surjections, bijections.

Définition 3.- Soient E, F des ensembles et f une application de E dans F .

On dit que f est injective, ou que f est une injection de E dans F , si la relation $f(x) = f(x')$ implique la relation $x = x'$.

On dit que f est surjective, ou que f est une surjection de E sur F si l'image de E par f est égale à F tout entier; on dit encore que f est une application de E sur F .

On dit que f est bijjective, ou que f est une bijection de E sur F si f est à la fois injective et bijective.

Une bijection d'un ensemble E sur lui-même s'appelle encore une permutation de E .

Dire que f est injective signifie que tout élément y de F est associé à au plus un élément de E par f . Si Γ est le graphe de f , c'est encore dire que Γ^{-1} est un graphe fonctionnel.

Dire que f est surjective signifie que tout élément y de F est associé à un élément de E par f . Par suite f est bijective si et seulement si tout élément de F est associé à un élément et un seul de E ; c'est encore dire que la correspondance Γ^{-1} est une application.

Il est classique de paraphraser ces définitions d'une autre manière en introduisant la notion d'équation. Etant donnés deux ensembles E, F , et deux applications f, g , de E dans F , la relation

$$(1) \quad f(x) = g(x)$$

où x est un élément de E , s'appelle équation définie par f et g . De façon non abrégée cette relation s'écrit

$$(\exists y) [(x, y) \in \Gamma \text{ et } (x, y) \in \Gamma']$$

où Γ, Γ' sont les graphes respectifs de f, g .

Tout élément x de E qui satisfait (1) s'appelle solution de l'équation.

Un cas particulier fréquent est celui où g est l'application constante égale à y_0 ; l'équation (1) s'écrit dans ce cas sous la forme

$$f(x) = y_0.$$

Dire que f est injective (resp. surjective, bijective) signifie que pour tout élément y de E , l'équation $f(x) = y$ admet au plus une solution (resp. au moins une solution, une solution et une seule).

a) Soient E un ensemble et A une partie de E . L'application de A dans E , qui à chaque élément x de A associe l'élément x lui-même, est une injection; on l'appelle injection canonique de A dans E .

b) Soit E un ensemble; l'application de E dans $E \times E$ qui à l'élément x de E associe le couple (x, x) est une injection; on l'appelle application diagonale de E .

c) Si f est une bijection de E sur F , l'application réciproque f^{-1} est une bijection réciproque de f . La relation $x = f^{-1}(y)$ équivaut à la relation

$y = f(x)$; on a

$$\begin{matrix} \xrightarrow{-1} \\ \xrightarrow{-1} \\ f^{-1} = f \end{matrix}$$

et $f^{-1} \circ f$ (resp. $f \circ f^{-1}$) est l'application identique de E (resp. de F).

d) Comme on le vérifie facilement, la composée de deux injections (resp. de deux surjections, de deux bijections) est une injection (resp. une surjection, une bijection).

Soient E, F des ensembles et f une bijection de E sur F . Pour toute application g de E dans E , on appelle transmuée de g par f l'application

$$f \circ g \circ f^{-1}$$

de F dans F . On vérifie que l'application qui à g associe $f \circ g \circ f^{-1}$ est une bijection de $\mathcal{F}(E, E)$ sur $\mathcal{F}(F, F)$.

Soit f une application d'un ensemble E dans un ensemble F . On dit que f est inversible à gauche s'il existe une application g de F dans E telle que $g \circ f$ soit l'application identique de E . Une telle application g s'appelle un inverse à gauche ou une rétraction de f .

Proposition 1.- On suppose que l'ensemble E est non vide. Pour que l'application f soit inversible à gauche, il faut et il suffit qu'elle soit injective.

La condition est évidemment nécessaire car si g est un inverse à gauche de f , la relation $f(x) = f(x')$ implique $g(f(x)) = g(f(x'))$ c'est à dire $x = x'$. Inversement, si f est injective, choisissons arbitrairement un élément a de E et considérons l'application g de F dans E obtenue en définissant $g(x)$ pour $x \in F$ comme

- l'unique élément y de E tel que $x = f(y)$ si x appartient à l'image de E par f ,

- l'élément a si x n'appartient pas à l'image de E par f .

On vérifie aussitôt que g est un inverse à gauche de f .

On dit que f est inversible à droite s'il existe une application h de F dans E telle que $f \circ h$ soit l'application identique de F . Une telle application h s'appelle un inverse à droite ou section de f .

Proposition 2.- Pour que l'application f soit inversible à droite, il faut et il suffit qu'elle soit surjective.

Proposition 3.- Pour qu'une application soit inversible à gauche et à droite, il faut et il suffit qu'elle soit bijective; dans ce cas les inverses à gauche et à droite coïncident avec l'application réciproque.

Nous avons vu que si f est bijective, l'application réciproque f^{-1} est à

la fois un inverse à gauche et à droite de f .

D'autre part si f possède un inverse à droite h , alors f est surjective puisque l'on a pour tout élément y de F

$$y = f(x)$$

avec $x = g(y)$.

Supposons maintenant que f possède aussi un inverse à gauche g ; ou bien $E = \emptyset$ et alors $g = \emptyset$, puisque h est une application de F dans E , de sorte que f, g, h sont des applications vides; ou bien E est non vide et il résulte de la proposition 1 que f est injective. Par suite f est bijective et on vérifie aussitôt que $g = h = f^{-1}$.

Ainsi la proposition 3 est démontrée; pour achever de prouver la proposition 2, il reste à établir que si l'application f est surjective, elle est inversible à droite. On recherche une application h de F dans E telle que

$$(2) \quad f(h(x)) = x$$

pour tout élément x de E .

Désignons par $\varphi(x)$ l'ensemble des éléments y de F tels que $f(y) = x$; cette notation fonctionnelle est justifiée: $\varphi(x)$ est l'image de x par l'application φ de F dans $\mathcal{P}(E)$ dont le graphe est l'ensemble des couples (x, y) tels que $(z \in y) \iff ((z, x) \in \Gamma)$. La relation (2) peut encore s'écrire

$$h(x) \in \varphi(x).$$

Pour chaque x , l'ensemble $\varphi(x)$ est non vide puisque f est surjective. On est ainsi ramené, après un changement de notations, à la propriété suivante qui fera l'objet d'un nouvel axiome :

Axiome VII (du choix)

Quels que soient les ensembles E, F et l'application f de E dans l'ensemble des parties non vides de F , il existe une application g de E dans F telle que

$$g(x) \in f(x)$$

pour tout élément x de E .

Puisque $f(x)$ est non vide, il est toujours possible de choisir un élément de cet ensemble. Si on se réfère à la notion naïve d'application, il suffit de faire un tel choix pour tout x et de considérer l'application g qui associe à x l'élément de $f(x)$ choisi; l'existence de g ne semble ainsi pas soulever

difficulté particulière. Malheureusement cette démarche est insuffisante : nous avons donné d'une application une définition précise; la construction d'une application d'un ensemble dans un autre repose par essence sur celle d'un ensemble qui est son graphe.

Désignons par Γ le graphe d'une application g de E dans F répondant à la question. La donnée de l'application f nous permet de considérer l'ensemble G des parties $\{x\} \times f(x)$ de $E \times F$ où x est un élément de E ; les éléments de G sont des ensembles deux à deux disjoints. Nécessairement Γ rencontre chacun de ces ensembles suivant un élément et un seul; en effet

$$\Gamma \cap (\{x\} \times f(x)) = (x, g(x)).$$

Inversement il est facile de voir que si un ensemble Γ possède cette propriété, alors l'ensemble $\Gamma \cap \subseteq G$ est le graphe d'une application g de E dans F qui répond à la question posée. De cette manière, l'axiome VII découle de l'axiome suivant dont on peut montrer qu'il lui est équivalent modulo les axiomes I à V :

Axiome VII'

Si E est un ensemble dont les éléments sont des ensembles non vides deux à deux disjoints, il existe un ensemble F qui rencontre chaque ensemble de E suivant un élément et un seul.

Nous verrons au chapitre suivant une autre forme équivalente de cet axiome dont le nom vient de ce qu'il correspond à la possibilité d'effectuer en une seule opération un nombre quelconque, fini ou infini, de choix.

Pour revenir à l'aspect historique, il faut noter que l'axiome du choix a suscité une longue controverse parmi les mathématiciens. Son introduction date de 1902 à propos de la comparaison par Beppo-Levi de la puissance d'un ensemble d'ensembles disjoints non vides avec celle de la réunion de ces derniers. Cependant Cantor l'avait déjà utilisé sans le dire à maintes reprises. Sa formulation explicite date en réalité de 1904; elle est due à Zermelo à la suite d'une suggestion de Erhard Schmidt et précède la première démonstration que Zermelo a donnée du théorème de bon ordre.

L'utilisation inconsciente de l'axiome, comme celle qui en était faite par Cantor, ne soulevait pas de critiques particulières. Cependant les conséquences du théorème de bon ordre, qui sera établi plus loin, sont si surprenantes que la formulation explicite de l'axiome qui en était à la base fut rejetée par de nombreux mathématiciens. Seul Poincaré contesta la démonstration, dans son aspect non prédicatif.

Actuellement, l'axiome du choix est généralement admis aussi bien en algèbre qu'en analyse; certains mathématiciens cependant s'attachent à discerner les énoncés qui en sont indépendants des autres.

Voici un exemple de propriété que l'on peut déduire de l'axiome du choix :

Il existe des ensembles de la droite qui ne sont pas mesurables pour la mesure de Lebesgue.

Il revient au même de construire des ensembles du cercle unité U du plan complexe qui ne sont pas mesurables pour la mesure superficielle $d\theta$ du cercle. Soit α un nombre irrationnel; dans le groupe multiplicatif U , considérons le sous-groupe G des $\exp(2\pi i n \alpha)$ où x parcourt l'ensemble $\mathbb{Z}$ des entiers relatifs. D'après l'axiome du choix, il existe un ensemble E de représentants pour la relation d'équivalence définie par G dans U . Alors U est réunion disjointe et dénombrable des parties $\exp(2\pi i n \alpha) \cdot E$; puisque la mesure $d\theta$ est invariante par rotation, si l'une était mesurable, toutes le seraient et auraient même mesure; suivant que cette mesure serait nulle ou strictement positive, la mesure totale de U serait nulle ou infinie, ce qui est dans tous les cas absurde.

4. Notions complémentaires.

Soit f une application d'un ensemble E dans un ensemble F . Pour toute partie A de E , on note $f(A)$ et on appelle image de A par f l'ensemble des $y \in F$ tels qu'il existe $x \in A$ pour lequel $y = f(x)$.

Pour toute partie B de F , on note $f^{-1}(B)$ et on appelle image réciproque de B par f l'ensemble des $x \in E$ tels que $f(x) \in B$.

Nous n'insisterons pas sur les formules qui mêlent les opérations d'image ou d'image réciproque et de réunion, intersection ou complémentaire; rappelons seulement que s'il n'y a aucune difficulté pour l'image réciproque et que si $f(A \cup B) = f(A) \cup f(B)$, on a en revanche seulement

$$f(A \cap B) \subset f(A) \cap f(B),$$

et il n'existe dans le cas général aucune relation entre $f(A)$ et $f^{-1}(f(A))$. De plus $f^{-1}(f^{-1}(f(A))) = A$ alors que $f^{-1}(f(A)) \supset A$ et que l'égalité n'a lieu, en général, que pour une application injective.

Remarques

- a) on a $f\langle\emptyset\rangle = \emptyset$; l'image de E par f est l'image de f .
- b) on a évidemment $f\langle\{x\}\rangle = \{f(x)\}$. Par abus de langage on dit encore que $f(x)$ est l'image de x par f .
- c) il est classique d'utiliser les parenthèses ordinaires pour noter images et images réciproques; on notera ainsi $f(A)$ au lieu de $f\langle A\rangle$ et $f^{-1}(A)$ au lieu de $f^{-1}\langle A\rangle$. La notation fonctionnelle est parfaitement licite: f détermine des applications de $\mathcal{P}(E)$ dans $\mathcal{P}(F)$ et de $\mathcal{P}(F)$ dans $\mathcal{P}(E)$; l'abus de langage consiste à noter f et f^{-1} les applications en question. Cela n'est pas sans danger; on observera à ce sujet que $f(x)$ est un élément de F alors que $f(A)$ en est une partie.
- d) si y est un élément de F , on note encore $f^{-1}(y)$ et on appelle image réciproque de y l'ensemble $f^{-1}\langle\{y\}\rangle$. Cette notation prête aussi à confusion; lorsque f est bijective, l'application réciproque est notée f^{-1} et $f^{-1}\langle\{y\}\rangle = \{f^{-1}(y)\}$ diffère de $f^{-1}(y)$.

Soient f une application d'un ensemble E dans lui-même et A une partie de E . On dit que A est une partie stable par f si elle vérifie l'inclusion $f\langle A\rangle \subset A$. On dit qu'un élément x de E est fixe ou invariant par f si $f(x) = x$. Une partie fixe est une partie dont tous les éléments sont fixes. Toute partie fixe est stable mais la réciproque est fausse; on se gardera de confondre les deux propriétés.

On considère deux ensembles E, F et un sous-ensemble E' de E . Si f, g sont des applications de E dans F , on dit que f et g coïncident sur E' si $f(x) = g(x)$ pour tout élément x de E' ; cela signifie que les restrictions à E' de f et g sont égales.

Soit $f = (E, F, \Gamma)$ une application de E dans F . Si E' est un sous-ensemble de E , la correspondance $(E', F, \Gamma \cap E' \times F)$ est une application appelée restriction de f à l'ensemble E' .

La restriction de f à E' s'obtient encore en composant à droite f avec l'injection canonique de E' dans E .

Inversement si g est une application de E' dans F , on appelle prolongement de g à E toute application f de E dans F dont la restriction à E' soit g .

On utilise moins souvent la restriction de l'espace d'arrivée. Soit F' un sous-ensemble de F , pour que E soit l'ensemble de définition de la correspondance $(E, F', \Gamma \cap E \times F')$, il faut et il suffit que F' contienne

l'image de f . Dans ce cas $(E, F', \Gamma \cap E \times F')$ est une application appelée application déduite de f par restriction de l'espace d'arrivée. Pour la désigner, on parlera de l'application f considérée comme à valeurs dans F' .

Il peut arriver que l'on ait besoin de restreindre à la fois la source et le but. Il convient alors de restreindre d'abord la source, puisque cela peut se faire dans tous les cas, puis de vérifier ensuite que la restriction du but est possible.

Il est au contraire toujours possible d'étendre l'espace d'arrivée; si on se donne un ensemble F' contenant F , il suffit de composer f à gauche avec l'injection canonique de F dans F' . Pour désigner l'application ainsi définie, on parlera de l'application f , considérée comme une application à valeurs dans F' .

IV

FAMILLES

1. Familles

Soit I un ensemble. On appelle famille indexée par I le couple (I, Γ) de I et d'un graphe fonctionnel Γ dont l'ensemble de définition est I . L'ensemble I est appelé ensemble d'indices de la famille. Si $x = (I, \Gamma)$ est une famille, on emploie usuellement la notation

$$(x_i)_{i \in I}$$

pour mettre en évidence son ensemble d'indices.

Comme on le voit aussitôt, la donnée de la famille (I, Γ) équivaut à celle du graphe Γ (*). La notion de famille est d'autre part très voisine de celle d'application. De façon précise, si E est un ensemble, on dit que la famille (I, Γ) est une famille d'éléments de E si $\Gamma \subset I \times E$; cela signifie que l'image de Γ , qui est appelée ensemble des valeurs de la famille, est contenue dans E ; dans ce cas (I, E, Γ) est une application. Toute famille est une famille d'éléments d'au moins un ensemble E ; par opposition avec les applications, on ne met pas dans la donnée d'une famille celle d'un tel ensemble E .

La différence concerne non seulement la terminologie mais encore les notations. Soit $x = (I, \Gamma)$ une famille; si j est un élément de I , il existe un y unique tel que $(j, y) \in \Gamma$. Pour le désigner, au lieu de la notation fonctionnelle $x(j)$, on préfère la notation x_j (ou x_j^j , j_x , j_x suivant les cas).

(*) Il reviendrait donc au même de définir une famille comme un graphe fonctionnel, ce qui est fait dans de nombreux ouvrages. Si nous avons préféré mettre explicitement l'ensemble d'indices dans la donnée, c'est pour se mettre en accord avec l'écriture

$$(x_i)_{i \in I}$$

et l'expression

"famille indexée par I ",

qui veulent que l'on explicite toujours l'ensemble d'indices lorsqu'on introduit une famille.

Donnons quelques exemples de familles :

a) La famille vide correspond à $I = \Gamma = \emptyset$; on la note par exemple $(x_i)_{i \in \emptyset}$.

b) Lorsque I est la paire $\{1, 2\}$, une famille $(x_i)_{i \in \{1, 2\}}$ est déterminée par son graphe qui est l'ensemble $\{(1, x_1), (2, x_2)\}$. Quand nous aurons défini 1 et 2 , nous verrons qu'il s'agit d'ensembles distincts, de sorte que l'ensemble précédent est entièrement déterminé par la connaissance du couple. Pour cette raison, on confondra souvent la famille en question avec ce couple. Inversement à tout couple (a, b) est associée une famille $(x_i)_{i \in \{1, 2\}}$ telle que $x_1 = a$ et $x_2 = b$.

c) Soit E un ensemble et Δ_E la diagonale de $E \times E$. On dit que (E, Δ_E) est la famille associée à l'ensemble E , ou encore la famille obtenue en indexant les éléments de E par eux-mêmes. Cette famille se note $(x)_{x \in E}$.

De cette manière la notion de famille peut être considérée comme une généralisation de la notion d'ensemble. Il faut bien se garder de confondre la famille $(x_i)_{i \in I}$ avec l'ensemble de ses valeurs, qui est parfois noté $\{x_i\}_{i \in I}$. On distinguera ainsi la famille $(x_n)_{n \in \mathbb{N}}$ telle que $x_0 = 1$, $x_1 = -1$, $x_2 = 1$, $x_4 = -1 \dots$ et l'ensemble $\{-1, 1\}$ de ses valeurs.

Certaines définitions mathématiques concernent des familles ou des ensembles. Ainsi, dans un espace vectoriel, on définit des familles libres et des parties libres : par exemple la famille finie $(x_i)_{i \in I}$ est libre si pour toute combinaison linéaire $\sum \lambda_i x_i$ la relation $\sum \lambda_i x_i = 0$ implique $\lambda_i = 0$ pour tout i ; la partie finie A est libre si la famille obtenue en indexant l'ensemble A par lui-même est libre, ce qui signifie que pour toute combinaison linéaire $\sum \lambda_x x$, la relation $\sum \lambda_x x = 0$ implique $\lambda_x = 0$ pour tout x ; on peut encore dire que pour toute combinaison linéaire finie $\sum_{i=1}^n \lambda_i x_i$ où les x_i sont dans A et tous distincts, alors $\sum \lambda_i x_i = 0$ implique $\lambda_i = 0$ pour tout i .

d) Lorsqu'une famille est indexée par l'ensemble $\mathbb{N}$ des entiers naturels, elle prend le nom de suite.

e) A la notion d'application surjective, rien ne correspond pour les familles. En revanche, à celle d'application injective, correspond celle de famille sans répétition : une famille $(x_i)_{i \in I}$ est dite sans répétition si $x_i = x_j$ implique $i = j$.

f) Si $(x_i)_{i \in I}$ est une famille indexée par I et J une partie de I , la famille indexée par J dont le graphe est l'ensemble des couples (i, x_i) où i

parcourt J , est appelée sous-famille obtenue par restriction à J de l'ensemble des indices.

Par exemple une famille quelconque d'un espace vectoriel est libre si toutes ses sous-familles finies sont libres.

On se gardera de confondre la notion de sous-famille avec celle de suite extraite. Si $(x_n)_{n \in \mathbb{N}}$ est une suite, on appelle suite extraite de $(x_n)_{n \in \mathbb{N}}$ une suite du type $(y_n)_{n \in \mathbb{N}}$ où $y_n = x_{p_n}$ et où $(p_n)_{n \in \mathbb{N}}$ est une suite strictement croissante de nombres entiers. Cette notion correspond à la notion de composition pour les applications.

Notons à ce sujet que la supériorité de la notation fonctionnelle apparaît lorsqu'on considère des composées. Dans l'exemple qui précède, la donnée de la suite $(p_n)_{n \in \mathbb{N}}$ correspond à celle d'une application φ de $\mathbb{N}$ dans $\mathbb{N}$ et on préfère souvent la notation mixte

$$(x_{\varphi(n)})_{n \in \mathbb{N}}.$$

Quant à la notion de réciproque, elle n'a pas de traduction commode pour les familles.

famille	application
$(x_i)_{i \in I}$	$i \mapsto x(i)$
x_j	$x(j)$
ensemble d'indices	ensemble de définition
famille vide	application vide
famille associée à E	application identique de E
famille sans répétition	application injective
sous-famille	application restreinte

Il est aisé de vérifier que les familles indexées par I et à valeurs dans E constituent un ensemble $\Phi(I, E)$. Il est en revanche contradictoire de parler de l'ensemble des familles indexées par I , sauf si I est vide.

Il est tout aussi contradictoire de parler de l'ensemble de familles d'éléments d'un ensemble donné, ce que l'on aurait une tendance assez naturelle à faire par analogie avec l'ensemble des parties.

2. Intersection et réunion d'une famille d'ensembles

Soit $(E_i)_{i \in I}$ une famille d'ensembles. On cherche à construire un ensemble E tel que $u \in E$ si et seulement si $u \in E_i$ pour tout $i \in I$; si I est non vide, choisissons un élément i_0 de I ; la relation que u doit satisfaire s'écrit encore

$$(3) \quad (\forall i) [(i \in I \setminus \{i_0\}) \Rightarrow (u \in E_i)] \quad \text{et} \quad (u \in E_{i_0}),$$

où $(u \in E_i)$ est la relation $(\exists x) [(u \in x) \text{ et } ((i, x) \in \Gamma)]$ si Γ est le graphe de la famille. D'après l'axiome de séparation, elle définit un ensemble qui est appelé intersection de la famille $(E_i)_{i \in I}$ et noté $\bigcap_{i \in I} E_i$.

Si $I = \emptyset$, la relation que u doit satisfaire s'écrit

$$(\forall i) [(i \in \emptyset) \Rightarrow (u \in E_i)].$$

Elle est toujours satisfaite et ne peut être collectivisante.

Si on se donne un ensemble E , pour toute famille non vide $(A_i)_{i \in I}$ de parties de E , on a $\bigcap_{i \in I} A_i \subset E$; la relation (3) implique en effet $u \in E$. On convient de définir l'intersection de la famille, pour I éventuellement vide, en ajoutant la condition $u \in E$, de sorte que si I est vide $\bigcap_{i \in I} A_i = E$. Il faut être prudent car cette notion dépend, comme on le voit, de l'ensemble E choisi.

On considère toujours une famille $(E_i)_{i \in I}$ d'ensembles; on cherche maintenant à construire un ensemble E tel que $u \in E$ si et seulement s'il existe $i \in I$ tel que $u \in E_i$. Si $\{E_i\}_{i \in I}$ désigne l'ensemble des valeurs de la famille $(E_i)_{i \in I}$, l'ensemble $\mathfrak{S}\{E_i\}_{i \in I}$ somme du précédent a la propriété recherchée. Cet ensemble est appelé réunion de la famille $(E_i)_{i \in I}$ et noté

$$\bigcup_{i \in I} E_i.$$

Un cas particulier important est celui de la famille $(E)_{E \in \mathfrak{E}}$ associée à un ensemble $\mathfrak{E}$ d'ensembles. Dans ce cas, $\mathfrak{E}$ est l'ensemble des valeurs de la famille et la réunion de la famille coïncide avec l'ensemble somme $\mathfrak{S}\mathfrak{E}$. Ainsi est justifiée la notation

$$\bigcup_{E \in \mathfrak{E}} E$$

pour l'ensemble somme qui apparait comme réunion d'une famille particulière d'ensembles.

Soit $(A_i)_{i \in I}$ une famille de parties d'un ensemble E ; on dit que $(A_i)_{i \in I}$ est un recouvrement de E si $\bigcup_{i \in I} A_i = E$; on dit que $(A_i)_{i \in I}$ est une partition de E si l'on a en outre $A_i \cap A_j = \emptyset$ pour $i \neq j$.

3. Produit d'une famille d'ensembles

Soit $(E_i)_{i \in I}$ une famille d'ensembles; on définit un nouvel ensemble, appelé produit de la famille $(E_i)_{i \in I}$ et noté

$$\prod_{i \in I} E_i$$

dont les éléments sont exactement les familles $(x_i)_{i \in I}$ telles que $x_i \in E_i$ pour tout $i \in I$. Il suffit pour cela de trouver un ensemble fixe contenant chaque famille $(x_i)_{i \in I}$ ayant la propriété ci-dessus; on peut choisir l'ensemble $\phi(I, \bigcup_{i \in I} E_i)$ des familles indexées par I d'éléments de $\bigcup_{i \in I} E_i$.

Soit j un élément de I ; l'application de $\prod_{i \in I} E_i$ dans E_j qui associe à un élément $x = (x_i)_{i \in I}$ du produit l'élément x_j de E_j est appelée projection d'ordre j et notée pr_j . On a donc:

$$x = (pr_i(x))_{i \in I}.$$

Proposition 4. - Soit $(E_i)_{i \in I}$ une famille d'ensembles; pour tout ensemble F et toute famille $(f_i)_{i \in I}$ dans laquelle f_i est une application de F dans E_i pour tout $i \in I$, il existe une application f et une seule de F dans $\prod_{i \in I} E_i$ telle que

$$(4) \quad f_i = pr_i \circ f$$

quel que soit $i \in I$.

L'unicité est évidente car la relation (4) implique $f(x) = (f_i(x))_{i \in I}$ pour tout $x \in F$. Inversement cette dernière relation permet de définir f .

On peut appliquer cet énoncé pour ramener un système d'équations à une équation unique. Considérons un ensemble E , une famille $(F_i)_{i \in I}$ d'ensembles et des familles $(f_i)_{i \in I}, (g_i)_{i \in I}$ telles que f_i, g_i soient des applications de E dans F_i . On appelle système d'équations associé à ces familles la relation

$$(\forall i) (f_i(x) = g_i(x))$$

dans laquelle $x \in E$. Si f, g sont des applications de E dans $\prod_{i \in I} F_i$

qui sont données pour la proposition 4, on voit que ce système est équivalent à l'équation

$$f(x) = g(x).$$

Si $(E_i)_{i \in I}$ est une famille vide d'ensembles, alors sa réunion est vide et $\prod_{i \in I} E_i = \{\emptyset\}$. Si l'un des ensembles E_i est vide, on voit facilement que le produit $\prod_{i \in I} E_i$ est encore vide. Peut-on en revanche affirmer que $\prod_{i \in I} E_i$ est non vide dans chaque cas où aucun E_i n'est vide? La démonstration de cet énoncé exige l'axiome du choix. Plus précisément l'axiome

(VII'') Tout produit d'une famille $(E_i)_{i \in I}$ d'ensembles non vides est non vide

est équivalent à l'axiome du choix.

Posons l'axiome (VII) et considérons une famille $(E_i)_{i \in I}$ non vide d'ensembles non vides. Si f désigne l'application de I dans $\mathcal{P}(\bigcup_{i \in I} E_i)$ qui à chaque $i \in I$ associe la partie non vide E_i de $\bigcup_{i \in I} E_i$, l'axiome VII assure l'existence d'une application g de I dans $\bigcup_{i \in I} E_i$ telle que $g(i) \in f(i) = E_i$ pour tout $i \in I$. La famille $(g(i))_{i \in I}$ est alors un élément du produit.

Posons réciproquement l'axiome (VII'') et considérons une application f d'un ensemble E dans l'ensemble des parties non vides d'un ensemble F . Si alors G est le produit de la famille $(f(x))_{x \in E}$, nous savons que G est non vide et contient une famille $y = (y_x)_{x \in E}$. Le couple $g = (y, F)$ est alors une application de E dans F telle que $g(x) \in f(x)$ pour tout $x \in E$.

Il résulte de ce qui précède que si chaque E_i est non vide, chaque projection pr_i est surjective; la vérification de ce point est laissée au lecteur.

De même que nous avons confondu la famille $(x_i)_{i \in \{1, 2\}}$ avec le couple (x_1, x_2) , nous identifions le produit $\prod_{i \in \{1, 2\}} E_i$ avec le produit $E_1 \times E_2$.

4. Somme disjointe d'une famille d'ensembles

Soit $(E_i)_{i \in I}$ une famille d'ensembles; on appelle somme disjointe de cette famille et on note $\bigsqcup_{i \in I} E_i$ l'ensemble des couples (i, x) tels que $i \in I$ et $x \in E_i$. Cet ensemble est défini comme une partie du produit $I \times \bigcup_{i \in I} E_i$.

Soit j un indice de I . L'application φ_j qui à $x \in E_j$ associe le couple

(j, x) est une bijection de E_j sur le sous-ensemble $\{j\} \times E_j$ de la somme disjointe. On vérifie que les $\{i\} \times E_i$ sont deux à deux disjoints et que $\bigsqcup_{i \in I} E_i$ est leur réunion.

Proposition 5.- Soit $(E_i)_{i \in I}$ une famille d'ensembles; pour tout ensemble F et toute famille $(f_i)_{i \in I}$, où f_i est une application de E_i dans F pour tout $i \in I$, il existe une application f et une seule de $\bigsqcup_{i \in I} E_i$ dans F telle que

$$(5) \quad f_i = f \circ \varphi_i$$

quel que soit $i \in I$.

L'unicité résulte de ce que la relation (5) implique $f((i, x)) = f_i(x)$ pour chaque $x \in E_i$; cette dernière relation permet inversement de définir f .

Si on considère la famille des injections canoniques des E_i dans leur réunion, la proposition 5 met en évidence une application de la somme disjointe $\bigsqcup_{i \in I} E_i$ sur la réunion $\bigcup_{i \in I} E_i$; c'est la seconde projection.

ENSEMBLES FINIS

1. Définition des ensembles finis

Nous avons vu que les ensembles infinis contredisaient la règle du "tout plus grand que la partie". Cette idée va nous servir maintenant pour définir les ensembles finis et infinis. De façon précise, si Y est une partie d'un ensemble X , on dit que Y est une partie stricte de X si de plus $Y \neq X$. Nous posons alors, suivant Dedekind, la

Définition 1.- On dit qu'un ensemble est fini s'il n'est équipotent à aucune de ses parties strictes. Un ensemble est dit infini s'il n'est pas fini.

On vérifie aussitôt que si un ensemble X est équipotent à un ensemble fini (resp. infini), il est lui-même fini (resp. infini). De plus

Proposition 1.- Soit X un ensemble; les propriétés qui suivent sont équivalentes :

- (i) X est fini
- (ii) toute injection de X dans lui-même est surjective
- (iii) toute surjection de X sur lui-même est injective.

Démonstration.- Tout d'abord les assertions (ii) et (iii) sont équivalentes; supposons (ii) : si f est une surjection de X sur lui-même, alors f admet une section g qui est, par nature, injective; par suite g est bijective et $f = g^{-1}$ est elle-même bijective. Supposons inversement (iii): on peut supposer $X \neq \emptyset$; si f est une injection de X dans lui-même, f admet une rétraction g qui est surjective; alors g est bijective et $f = g^{-1}$ l'est aussi.

Il nous reste à montrer que les assertions (i) et (ii) sont équivalentes; si X n'est pas fini, il existe, par définition, une bijection f de X sur une de ses parties strictes; cette bijection définit une injection non surjective de X dans lui-même. Inversement si X est fini, toute injection f de X dans lui-même définit une bijection de X sur $f(X)$; par suite $f(X) = X$ et f est surjective.

Corollaire.- Soient X, Y des ensembles finis équipotents et f une application de X dans Y ; il y a alors équivalence pour f entre les propriétés suivantes :

- (i) f est injective
- (ii) f est surjective
- (iii) f est bijective.

Il suffit de considérer une bijection φ de Y sur X et de considérer l'application $f \circ \varphi$ de Y dans lui-même (ou bien l'application $\varphi \circ f$ de X dans lui-même).

2. Propriétés des ensembles finis

Nous indiquons d'abord un certain nombre de propriétés qui découlent facilement de la définition.

Proposition 2.- Si X est un ensemble fini, toute partie de X est finie.

Soit Y une partie de X ; supposons par l'absurde que Y admette une partie stricte Z telle qu'il existe une bijection φ de Y sur Z . Nous définissons une application ψ de X dans $Z \cup \bigcap_X Y$ en posant

$$\begin{aligned}\psi(x) &= \varphi(x) & \text{si } x \in Y, \\ \psi(x) &= x & \text{si } x \in \bigcap_X Y.\end{aligned}$$

On vérifie aussitôt que ψ est bijective; d'autre part $Z \cup \bigcap_X Y$ est évidemment une partie stricte de X .

Corollaire.- Soient X un ensemble fini et Y un ensemble. L'ensemble Y est fini dans les deux cas suivants:

- (i) il existe une injection de Y dans X
- (ii) il existe une surjection de X sur Y .

En effet l'assertion (i) implique que Y est équipotent à une partie de X ; d'autre part l'assertion (ii) implique l'assertion (i) comme on le voit en considérant une section.

Proposition 3.- Si X est un ensemble fini et a un ensemble, alors l'ensemble $X \cup \{a\}$ est fini.

Raisonnons en effet par l'absurde en supposant l'existence d'une bijection f de $X \cup \{a\}$ sur une partie stricte Y de ce dernier ensemble. Ou bien $f(X) \subset X$ et f induit une bijection de X sur $f(X)$; il est facile de voir que $f(X)$ est une partie stricte de X car $f(X) = X$ impliquerait $f(a) = a$ et

et $f(X \cup \{a\}) = X \cup \{a\}$; ce qui contredit l'hypothèse. Ou bien il existe $b \in X$ tel que $f(b) = a$, ce qui implique $f(a) \in X$. Si on considère alors l'application g de X dans lui-même telle que $g(b) = f(a)$ et $g(x) = f(x)$ pour $x \in X, x \neq b$, alors g est injective et $g(X)$ est une partie stricte de X car $g(X) = X$ impliquerait encore $f(X \cup \{a\}) = X \cup \{a\}$; l'hypothèse est encore contredite.

A partir de là, nous pouvons construire quelques exemples d'ensembles finis. Tout d'abord l'ensemble vide $\emptyset$ est fini. D'autre part, si X est fini, alors $X \cup \{X\}$ est fini. Sont par suite finis les ensembles

$$\{\emptyset\}$$

$$\{\emptyset, \{\emptyset\}\}$$

$$\{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}$$

etc.

Ce procédé de génération d'ensembles finis qui consiste à ajouter à un ensemble ce même ensemble comme élément est particulièrement important ; nous y reviendrons à propos des nombres entiers.

Notons en particulier que des ensembles à un ou deux éléments sont toujours finis.

Nous citons maintenant d'autres propriétés que nous ne sommes pas en mesure de démontrer tout de suite.

Proposition 4.- Soit X un ensemble fini ; alors $\wp X$ est fini. Si de plus chaque élément de X est fini, $\mathfrak{G} X$ est fini.

On en déduit que si X, Y sont des ensembles finis, alors $X \cup Y, X \times Y, \mathcal{F}(X, Y)$ sont des ensembles finis. Si $(X_i)_{i \in I}$ est une famille finie (c'est à dire indexée par un ensemble fini) d'ensembles finis, alors $\bigcup_{i \in I} X_i$, $\bigcap_{i \in I} X_i$ sont des ensembles finis.

VI

NOMBRES ENTIERS

L'examen des ensembles finis $\emptyset, \{\emptyset\}, \{\emptyset, \emptyset\} \dots$ que l'on construit de proche en proche en prenant chaque fois la réunion de X avec $\{X\}$, laisse présager qu'une véritable étude des ensembles finis ne peut pas se faire sans la considération de relations d'ordre.

C'est d'ensembles ordonnés qu'il va donc être question dans ce chapitre. Nous ne reviendrons pas sur la définition d'une relation d'ordre ou d'un ensemble ordonné; nous nous contenterons de construire explicitement un ensemble ordonné particulièrement important, qui est celui des nombres entiers.

1. $\mathbb{N}$ -ensembles

Nous partons de la connaissance intuitive que chacun a des nombres entiers et nous cherchons à caractériser ces derniers par les propriétés de leur ordre. Ils constituent un ensemble ordonné E qui est bien sûr non vide et qui possède la propriété fondamentale

- (i) Toute partie non vide de E possède un plus petit élément.

Cependant cette propriété ne suffit pas; elle n'empêche pas E d'être un ensemble très petit, un ensemble ordonné à un élément par exemple. Aussi doit-on ajouter la propriété

- (ii) E n'a pas de plus grand élément.

Malheureusement cela n'empêche pas E d'être trop grand; si nous considérons par exemple l'ensemble



des nombres réels de la forme $1 - 1/n$ ou $2 - 1/n$, où n est un entier ≥ 2 , il est facile de vérifier que les propriétés (i) et (ii) sont satisfaites. On constate que la partie de cet ensemble constituée des nombres $1 - 1/n$ est isomorphe à l'ensemble des entiers; elle n'a pas de plus grand élément mais est majorée par des éléments de l'ensemble, comme $2 - 1/2$. C'est cette éventualité que nous écartons en ajoutant la propriété

(iii) Toute partie majorée non vide de E possède un plus grand élément.

Nous nous arrêtons provisoirement dans l'énumération des propriétés et introduisons pour les résumer la

Définition 1.- On appelle N -ensemble tout ensemble ordonné non vide qui possède les propriétés (i), (ii) et (iii).

Soit maintenant E un N -ensemble; nous allons étudier quelles propriétés nous pouvons déduire de la définition. Tout d'abord il résulte de la propriété (i) que E est totalement ordonné. De plus :

a) Puisque E est non vide, il résulte de (i) que E possède un plus petit élément e . D'autre part si x est un élément de E , d'après (ii) ce n'est pas le plus grand élément; par conséquent l'intervalle $]x, \rightarrow$ est non vide et possède un plus petit élément x^* . Il n'y a alors aucun élément entre x et x^* ; pour cette raison, x^* est appelé successeur de x .

b) Soit maintenant un élément x de E distinct de e ; l'intervalle $S_x = (\leftarrow, x[$ est majoré et possède d'après (iii) un plus grand élément $*x$ qui est appelé prédécesseur de x ; on a clairement $(*x)^* = x$.

Nous avons de plus :

Proposition 1.- (principe de récurrence simple).- Soit P une partie du N -ensemble E telle que

a) $e \in P$

b) la relation $x \in P$ implique $x^* \in P$.

Alors $P = E$.

Démonstration.- Si P diffère de E , le complémentaire de P dans E possède d'après (i) un plus petit élément x . A cause de l'hypothèse a), on a $x \neq e$. Par suite, x admet un prédécesseur y qui appartient alors à P . L'hypothèse b) implique $y^* \in P$, soit $x \in P$, ce qui est absurde.

Proposition 2 (principe de récurrence sur tous les prédécesseurs). -

Soit P une partie du $\mathbb{N}$ -ensemble E telle que la relation $(\leftarrow, x] \subset P$ implique $x \in P$.

Alors $P = E$.

Démonstration .- Si P diffère de E , $\bigcap_E P$ possède un plus petit élément x et $(\leftarrow, x] \subset P$; on obtient ainsi une contradiction avec l'hypothèse.

Il faut noter que cette dernière forme du principe de récurrence n'utilise que la propriété (i); nous y reviendons dans les chapitres ultérieurs.

Le principe de récurrence s'étend aux relations $R(x)$ dans lesquelles x est astreint à rester dans E . Si par exemple

- a) $R(e)$ est vraie,
- b) $R(x)$ implique $R(x^*)$,

alors $R(x)$ est vraie pour tout $x \in E$. Il suffit de considérer la partie P de E des éléments vérifiant la relation.

Il en est de même du principe de récurrence sur tous les prédécesseurs.

La donnée d'un $\mathbb{N}$ -ensemble E permet également d'effectuer ce que l'on appelle une construction par récurrence.

Nous allons nous placer d'abord dans le cas de la récurrence simple. Soit F un ensemble. On cherche à construire une famille $(x_p)_{p \in E}$ de F telle que x_e soit un élément a donné de F et sachant que l'on dispose d'un "procédé" pour construire l'élément x_{p^*} lorsque x_p est connu. La donnée du "procédé" correspond simplement à celle d'une application ϕ de F dans lui-même et il s'agit de construire une famille $(x_p)_{p \in E}$ de F telle que

$$x_e = a$$

et

$$x_{p^*} = \phi(x_p).$$

Nous allons voir que cela est toujours possible et de façon unique. Pour montrer l'existence d'une solution, considérons l'ensemble des éléments p de E tels qu'il existe une application f_p unique de $(\leftarrow, p]$ dans E vérifiant pour $q < p$ les relations

$$f_p(e) = a$$

et

$$f_p(q^*) = \phi(f_p(q)).$$

En appliquant le principe de récurrence simple, il est facile de montrer que P est l'ensemble E tout entier. L'unicité montre de plus que pour $p \leq m$, la restriction de f_m à $(\leftarrow, p]$ coïncide avec f_p ; autrement dit, pour $q \leq p$ on a $f_m(q) = f_p(q)$. Si nous posons alors

$$x_p = f_p(p),$$

on vérifie aisément que la famille $(x_p)_{p \in E}$ possède les propriétés cherchées. La démonstration de l'unicité ne présente pas de difficulté.

Dans d'autres cas, le "procédé" de construction suppose non seulement la connaissance de l'élément précédemment construit, mais aussi celle de l'étape p à laquelle on se trouve ou encore celle de tous les éléments précédemment construits. Dans ce dernier cas, cela se traduit par la donnée d'une famille $(\phi_p)_{p \in E}$, où ϕ_p applique $\mathcal{F}((\leftarrow, p[, F)$ dans F , et on cherche une famille $(x_p)_{p \in E}$ telle que $x_p = \phi_p(f_p)$, où f_p est l'application associée à la suite $(x_q)_{q < p}$. Pour $p = e$, l'ensemble $\mathcal{F}((\leftarrow, e, F)$ est réduit à l'application vide et nécessairement $x_e = \phi_e(\{\phi\})$; ainsi la donnée de la famille $(\phi_p)_{p \in E}$ contient-elle celle du premier élément de la suite cherchée. Pour montrer l'existence d'une solution on considère l'ensemble P des éléments p de E tels qu'il existe une fonction $f_p \in \mathcal{F}((\leftarrow, p), F)$ unique vérifiant pour chaque $q < p$ la relation $f_p(q) = \phi_p(g_{p,q})$, dans laquelle $g_{p,q}$ est la restriction de f_p à l'intervalle $(\leftarrow, p[$. En appliquant le principe de récurrence sur tous les prédécesseurs, on montre que $p = E$ et on définit alors la famille $(x_p)_{p \in E}$ par $x_p = f_p^*(p)$. L'unicité s'établit sans grande difficulté.

Il peut arriver, aussi bien dans la construction par récurrence simple que dans la récurrence sur l'ensemble des prédécesseurs que l'on n'ait pas de "procédé" imposé, mais que l'on dispose au contraire d'une certaine liberté. Dans le premier cas, cela correspond à la donnée d'une application de F dans l'ensemble des parties non vides de F : l'existence d'un "procédé" résulte aussitôt de l'axiome du choix. Dans le second cas, la liberté est traduite par la donnée d'une famille $(\psi_p)_{p \in E}$ où ψ_p applique $\mathcal{F}((\leftarrow, p[, F)$ dans l'ensemble des parties non vides de F ; il faut appliquer l'axiome du choix à l'application $(p, f) \mapsto \psi_p(f)$ de l'ensemble des couples (p, f) dans l'ensemble des parties non vides de F pour en déduire l'existence d'un "procédé" $(\phi_p)_{p \in E}$.

Nous sommes maintenant en mesure d'établir que les propriétés (i),

(ii), (iii) suffisent à caractériser l'ordre de l'ensemble des entiers parmi ceux des ensembles ordonnés non vides. De façon précise

Théorème 1 (d'isomorphisme). - Si E, E' sont deux $\mathbb{N}$ -ensembles, il existe une bijection strictement croissante et une seule de E dans E' .

On remarque en effet qu'une bijection strictement croissante f de E dans F associe au plus petit élément e de E le plus petit élément e' de E' et vérifie pour tout élément p de E la relation

$$f(p^*) = (f(p))^* .$$

Si nous posons $f(p) = x_p$ et si ϕ est l'application de F dans lui-même qui, à l'élément p' de F associe son successeur p'^* , on a alors

$$x_e = e'$$

et

$$x_{p^*} = \phi(x_p) .$$

On est ainsi ramené à la construction par récurrence d'une suite (x_p) . Il reste simplement à voir que si on pose inversement $f(p) = x_p$, on définit une application strictement croissante de E sur E' . Le fait que f est strictement croissante, i.e. que $f(m) > f(p)$ pour $m > p$ se voit par récurrence sur m dans le $\mathbb{N}$ -ensemble $]p, \rightarrow)$. Enfin f est surjective : dans le cas contraire, il existerait un plus petit $p' \in F$ qui ne soit pas dans l'image ; en considérant son prédécesseur ${}^*p'$ on aboutit à une absurdité.

2. Axiome de l'infini

Nous avons obtenu l'unicité, à isomorphisme près, d'un $\mathbb{N}$ -ensemble. Le problème se pose de savoir si l'existence d'un $\mathbb{N}$ -ensemble découle des axiomes jusqu'ici introduits. Remarquons qu'un $\mathbb{N}$ -ensemble E est infini ; en effet l'application

$$x \longrightarrow x^*$$

est une bijection de E sur le complémentaire dans E du plus petit élément e . Par conséquent, l'existence d'un $\mathbb{N}$ -ensemble implique déjà celle d'un ensemble infini.

Il est assez facile d'imaginer que cela n'est pas une conséquence des axiomes I à V et VII ; ces axiomes permettent en effet de construire de nouveaux ensembles à partir d'ensembles déjà connus, mais si l'on part d'ensembles finis, ils n'introduisent que de nouveaux ensembles finis.

Aussi pose-t-on un

Axiome VI de l'infini : il existe un ensemble infini.

La question initialement posée est alors résolue par le

Théorème 2 (Dedekind). - L'axiome de l'infini équivaut à l'existence d'un $\mathbb{N}$ -ensemble.

Soit X un ensemble infini; il existe donc une bijection f de X sur une de ses parties strictes. Pour chaque élément x de E , nous désignons par C_x la plus petite partie de X stable par f contenant x (c'est à dire l'intersection des parties stables par f contenant x).

Nous choisissons un élément e de X qui n'appartient pas à $f(X)$ et désignons par E l'ensemble C_e . Si p, q sont des éléments de E , nous posons

$$p \leq q$$

si $C_q \subset C_p$. Nous allons montrer que la relation $p \leq q$ est une relation d'ordre sur E qui en fait un $\mathbb{N}$ -ensemble.

Lemme. - $f(C_p) = C_{f(p)} = C_p \setminus \{p\}$.

Tout d'abord $p \in C_p$ et $f(C_p) \subset C_p$, de sorte que $C_p \supset \{p\} \cup f(C_p)$; inversement $\{p\} \cup f(C_p)$ est stable par f et contient p , donc contient C_p , de sorte que $C_p = \{p\} \cup f(C_p)$. Comme $f(C_p) \neq C_p$ (c'est vrai si $p = e$ et pour $p \neq e$ cela impliquerait la stabilité par f de $E \setminus C_p$ qui contient e en étant distinct de E), on a $f(C_p) = C_p \setminus \{p\}$.

Sachant que $f(p) \in f(C_p)$, on a $C_{f(p)} \subset f(C_p)$; inversement $\{p\} \cup C_{f(p)}$ est stable et contient p , donc C_p , ce qui achève la démonstration.

La relation $p \leq q$ considérée est évidemment réflexive et transitive; l'antisymétrie résulte de ce que, par le lemme, p est l'unique élément de $C_p \setminus f(C_p)$. Le lemme montre encore que $f(p)$ est le successeur de p dans cette relation d'ordre.

Démontrons la propriété (i) : soit pour cela une partie F non vide de E sans plus petit élément. Considérons l'ensemble G des minorants de F , lequel ne rencontre pas F . Cet ensemble contient e et est stable par f ; en effet, pour tout $p \in G$ et tout $q \in F$ on a $q < p$ c'est à dire $f(q) \leq p$. Finalement $G = E$, ce qui est absurde.

La propriété (ii) est immédiate puisque $f(p) > p$. Enfin la propriété (iii) résulte de ce que si une partie majorée non vide F de E n'a pas de

plus grand élément, l'ensemble G des éléments de E qui minore un élément de F est stable par f . Comme G contient e , alors $G = E$, ce qui absurde.

Nous posons une fois pour toutes l'axiome de l'infini. L'ensemble des nombres entiers est introduit par le

Théorème 2. - Il existe un $\mathbb{N}$ -ensemble et un seul tel que pour tout élément p de E , l'on ait $p = (\leftarrow, p[$.

L'unicité est facile. Si E, E' sont deux tels $\mathbb{N}$ -ensembles, on montre par récurrence sur p , que l'unique bijection croissante f de E sur E' vérifie $f(p) = p$.

Pour l'existence, considérons d'abord un $\mathbb{N}$ -ensemble E . Nous nous intéressons aux éléments p de E tels que $(\leftarrow, p[$ soit isomorphe de façon unique à un ensemble ordonné α_p tel que pour tout $q \in \alpha_p$ on ait $q = (\leftarrow, q[$. Clairement le plus petit élément e de E possède la propriété avec $\alpha_e = \{\emptyset\}$. Si un élément p de E possède la propriété, alors p^* la possède aussi car, une fois remarqué que $\alpha_p \notin \alpha_p$, l'unique solution est donnée par l'ensemble

$$\alpha_{p^*} = \alpha_p \cup \{\alpha_p\},$$

muni de la relation d'ordre prolongeant celle de α_p et pour laquelle l'élément α_p est plus grand que tous les autres et l'isomorphisme de $(\leftarrow, p^*[$ sur α_{p^*} prolongeant le précédent et envoyant p sur α_p . Il résulte alors du principe de récurrence simple que tous les éléments de E possèdent la propriété considérée. A cause de l'unicité, si $p \leq q$, l'ensemble α_p est l'intervalle $(\leftarrow, p[$ de α_q , avec la relation d'ordre induite et les isomorphismes s'induisent. Pour achever la démonstration, il suffirait de pouvoir considérer la réunion des ensembles α_p , qui est encore l'ensemble des α_p .

Cela nécessite une version plus forte de l'axiome de compréhension que celle que nous avons donnée. Sachant qu'une relation $R(u, v)$ est dite fonctionnelle en u si les relations $R(x, y)$ et $R(x, y')$ impliquent $y = y'$, l'axiome en question s'énonce :

Pour tout ensemble x et toute relation $R(u, v)$ fonctionnelle en u , il existe un ensemble y dont les éléments sont les ensembles z vérifiant

$$(\exists t) [(t \in x) \text{ et } R(t, z)] .$$

En appliquant ce schéma d'axiomes à la relation $(u \in E) \text{ et } (v = \alpha_u)$, qui est fonctionnelle en u , d'après l'unicité de α_u , on obtient l'existence

d'un ensemble dont les éléments sont les α_p .

L'unique ensemble qui possède les propriétés du théorème 2 est noté $\mathbb{N}$. Ses éléments sont appelés nombre entiers naturels. Les premiers nombres entiers sont

$$\begin{aligned}0 &= \emptyset \\1 &= \{\emptyset\} \\2 &= \{\emptyset, \{\emptyset\}\}\end{aligned}$$

etc.

3. Opérations sur les nombres entiers.

L'addition dans $\mathbb{N}$ est donnée par la

Proposition 3.- Pour tout élément n de $\mathbb{N}$, il existe une bijection croissante et une seule f_n de $\mathbb{N}$ sur l'intervalle $[n, \rightarrow)$. On pose alors

$$n + m = f_n(m).$$

Cela résulte de ce que l'intervalle $[n, \rightarrow)$ est un $\mathbb{N}$ -ensemble. Nous laissons au lecteur le soin d'établir par récurrence les propriétés de l'addition dans $\mathbb{N}$. Notons seulement que $n^* = n + 1$.

La multiplication dans $\mathbb{N}$ se fait par la

Proposition 4.- Pour tout élément n de $\mathbb{N}$, il existe une application g_n de $\mathbb{N}$ dans lui-même et une seule telle que $g_n(0) = 0$ et

$$g_n(m+1) = g_n(m) + n.$$

On pose maintenant

$$nm = g_n(m).$$

La construction de g_n se fait par récurrence; les propriétés de la multiplication et les propriétés de distributivité par rapport à l'addition se démontrent sans difficulté.

On construit de même l'exponentiation par la

Proposition 5.- Pour tout élément n de $\mathbb{N}$, il existe une application h_n de $\mathbb{N}$ dans lui-même et une seule telle que $h_n(0) = 1$ et

$$h_n(m+1) = n h_n(m).$$

On pose ici

$$n^m = h_n(m).$$

4. Retour sur les ensembles finis.

Tout nombre entier naturel est un ensemble fini; cela se démontre par récurrence, sachant que $n^* = n \cup \{n\}$, à l'aide de la proposition 3 du chapitre V. Une propriété fondamentale des ensembles finis est donnée par le

Théorème 3.- Tout ensemble fini est équipotent à un nombre entier naturel et un seul.

Ce nombre entier naturel s'appelle le nombre d'éléments ou cardinal de l'ensemble X et se note $\text{Card}(X)$.

Démonstration. L'unicité est évidente, car, étant donnés deux nombres entiers différents, il y a en a toujours un qui est une partie stricte de l'autre.

Pour établir l'existence, donnons-nous un ensemble X qui ne soit, par l'absurde, équipotent à aucun nombre entier naturel; nous allons voir que X est infini et construire, pour cela, une injection de $\mathbb{N}$ dans X .

Il s'agit simplement de construire une suite $(x_n)_{n \in \mathbb{N}}$ sans répétition d'éléments de X . Nous allons procéder pour cela par récurrence sur l'ensemble des prédécesseurs. Supposons donc les x_p construits pour $p < n$ de façon que la famille $(x_p)_{p < n}$ soit sans répétition. Si l'ensemble des valeurs de cette famille était X tout entier, il en résulterait une bijection de $n = (\leftarrow, n[$ sur X , ce qui est contraire à l'hypothèse; par conséquent il existe au moins un élément de X qui n'est pas égal à l'un des x_p ; choisissons en un qui sera x_n ; la famille $(x_p)_{p < n+1}$ est encore sans répétition.

Si l'on veut mettre en forme ce qui précède, on peut procéder comme suit : pour tout entier n et pour toute application f de n dans X , désignons par $\psi_n(f)$

- l'ensemble X lui-même si f n'est pas injective,
- le complémentaire de l'image de f si f est injective;

cet ensemble est non vide, sinon X serait équipotent à n . Il s'agit alors d'appliquer le procédé de récurrence sur l'ensemble des prédécesseurs décrit au paragraphe 1.

Nous laissons maintenant au lecteur le soin de démontrer par récurrence sur $\text{Card}(X)$ ou $\text{Card}(Y)$ que $X \cup Y$ est fini lorsque X et Y le sont et que

$$\text{Card}(X \cup Y) \leq \text{Card}(X) + \text{Card}(Y).$$

De même $X \times Y$, $\mathcal{F}(X, Y)$ sont finis lorsque X, Y le sont et

$$\text{Card}(X \times Y) = \text{Card}(X) \times \text{Card}(Y),$$

$$\text{Card } \mathcal{F}(X, Y) = \text{Card}(X)^{\text{Card}(Y)};$$

en particulier

$$\text{Card } \mathcal{P}(X) = 2^{\text{Card}(X)}.$$

VII

ENSEMBLES DERIVES

Nous revenons dans ce chapitre sur quelques problèmes d'analyse qui ont contribué à la genèse de la théorie des ensembles.

Certains d'entre eux proviennent de l'étude du mouvement des cordes vibrantes, abordée depuis Bernouilli et Euler et dans laquelle on recherche des solutions (fonctions représentant le mouvement de la corde à un moment donné) sous forme de séries de fonctions trigonométriques

$$(1) \quad f(x) = \sum_{n=0}^{\infty} (a_n \sin nx + b_n \cos nx) \quad .$$

En 1807, Fourier donne un procédé pour construire le développement en série (1) qui porte son nom d'une fonction périodique f quelconque; comme il est bien connu, les coefficients a_n et b_n sont donnés par

$$b_0 = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) dx$$

et pour $n \geq 1$

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin nx dx \quad , \quad b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos nx dx \quad ,$$

lorsque la période est 2π . Ce n'est que bien plus tard, avec Heine, Dirichlet et Riemann, que se pose le problème de l'intégrabilité des fonctions considérées, de la convergence de la série, puis celui de savoir si la série représente bien la fonction de départ.

Nous allons concentrer notre attention sur le premier problème cité :

Pour quelles fonctions f définies sur un intervalle fermé borné $[a, b]$, est-il possible de définir l'intégrale

$$\int_a^b f(x) dx \quad ?$$

L'intégrale telle que la définissait Cauchy s'appliquait à des fonctions continues, monotones. Afin de simplifier l'exposition nous allons supposer connue l'intégrale d'une fonction continue sur un intervalle $[a, b]$ et voir comment il est possible d'étendre la définition à des fonctions ayant un certain nombre de discontinuités constituant un ensemble P de points de $[a, b]$.

Soit donc f une telle fonction; dans un souci de simplicité, on supposera que f est bornée.

1) Supposons d'abord que P ne rencontre pas l'intervalle $]a, b[$; en coupant l'intervalle en deux au besoin, on se ramène à étudier le cas où $P = \{b\}$. Sur chaque intervalle $[a, b - \varepsilon]$, avec $\varepsilon > 0$, la fonction f est continue de sorte que l'intégrale

$$\int_a^{b-\varepsilon} f(x) dx$$

est définie. Il est facile de voir que cette expression a une limite lorsque ε tend vers zéro puisque f est bornée. On posera alors par définition

$$\int_a^b f(x) = \lim_{\substack{\varepsilon \rightarrow 0 \\ \varepsilon > 0}} \int_a^{b-\varepsilon} f(x) dx.$$

2) Supposons que l'ensemble P soit fini; on se ramène au cas précédent en décomposant l'intervalle $[a, b]$ en un nombre fini d'intervalles.

Introduisons maintenant l'ensemble P' des points d'accumulation de P , encore appelé ensemble dérivé de P . Dire que P est fini signifie simplement que $P' = \emptyset$; nous allons voir qu'il est encore possible de définir l'intégrale dans des cas où P' est non vide.

3) Supposons que P' ne rencontre pas $]a, b[$; on se ramène au cas où $P' = \{b\}$, c'est à dire au cas où $P \cap [a, b - \varepsilon]$ est fini pour tout $\varepsilon > 0$. Dans ce cas l'intégrale a été construite sur l'intervalle $[a, b - \varepsilon]$ en 2); on la définit sur $[a, b]$ par un passage à la limite comme dans 1).

4) Si maintenant P' est fini, en raisonnant comme dans 2) on se ramène au cas 3).

Il est possible de considérer l'ensemble dérivé de P' , soit P'' , puis de proche en proche pour tout entier n , l'ensemble dérivé $P^{(n)}$ d'ordre n de P . Dire que P' est fini signifie que $P'' = \emptyset$. Plus généralement

5) Si $P^{(n)}$ est vide pour un entier n , alors l'intégrale de f se définit par récurrence à l'aide des procédés décrits en 1) et 2).

La suite $P^{(n)}$ est, pour $n \geq 1$, une suite décroissante d'ensembles fermés de $[a, b]$. On peut poser

$$P^{(\omega)} = \bigcap_{n \geq 1} P^{(n)}.$$

Dire que $P^{(\omega)} = \emptyset$ équivaut, d'après le théorème sur les ensembles compacts emboîtés, à $P^{(n)} = \emptyset$ pour un entier n ; cette condition ne fait donc que résumer l'ensemble des conditions 5). Cependant il n'y a pas de raison que le procédé s'arrête; en posant

$$P^{(\omega+1)} = (P^{(\omega)})',$$

on saura encore définir l'intégrale si

$$6) \quad P^{(\omega+1)} = \emptyset,$$

ou encore, en posant la définition convenable, si

$$7) \quad \text{il existe un entier } n \text{ tel que } P^{(\omega+n)} = \emptyset.$$

On voit ainsi apparaître de nouveaux nombres au delà des nombres entiers, que nous avons notés $\omega, \omega+1, \dots, \omega+n, \dots$, auxquels on pourrait encore ajouter $2\omega, 2\omega+1, \dots, 3\omega, \dots, \omega^2, \dots, \omega^n, \dots, \omega^\omega \dots$. Nous les introduirons de façon rigoureuse au chapitre IX sous le nom de nombres ordinaux ou nombres transfinis.

Heureusement, il existe une définition plus simple que l'intégrale qui généralise tous les cas que nous avons mentionnés, à savoir celle qu'a donné Riemann et sur laquelle nous ne reviendrons pas. En effet

Proposition 1 (du Bois Reymond).— Soit f une fonction bornée sur l'intervalle $[a, b]$. Pour que f soit intégrable au sens de Riemann, il suffit que l'ensemble P de ses points de discontinuité puisse être inclus dans une réunion finie d'intervalles dont la somme des longueurs est arbitrairement petite.

et

Théorème 1 (Cantor).— Avec les notations et hypothèses qui précèdent, si P' est dénombrable f est intégrable au sens de Riemann.

La démonstration de la proposition 1 étant évidente, nous allons seulement donner celle du théorème 1.

On peut écrire le complémentaire de $\bar{P} = P \cup P'$ dans $[a, b]$ comme réunion d'une suite d'intervalles disjoints (I_n) telle que la somme 1 des longueurs des I_n soit $b-a$: supposons par exemple que $a=0, b=1$ et posons

$$f(x) = x - s(x)$$

où $s(x)$ est la somme des longueurs des intervalles $I_n \cap [0, x]$. Clairement f est continue et varie entre 0 et 1. Comme f est constante sur chaque I_n l'ensemble de ses valeurs est dénombrable car $P \cup P'$ l'est. Comme f doit prendre toutes les valeurs entre 0 et 1 cela impose que $1 = 1$.

Le cas où P' est dénombrable généralise tous les cas 1), 2), ..., 7)... A chaque étape de la récurrence on réunit en effet un ensemble dénombrable d'ensembles du type précédent et une réunion dénombrable d'ensembles dénombrables est dénombrable.

VIII

ENSEMBLES BIEN ORDONNES

1. Propriétés des ensembles bien ordonnés

Nous avons introduit au chapitre VI les nombres entiers à partir des propriétés (i), (ii) et (iii) de leur ordre. La première de ces propriétés joue un rôle particulièrement important dans le principe de récurrence. Nous allons étudier maintenant de façon systématique les ensembles ordonnés qui possèdent la propriété en question. De façon précise nous posons la

Définition 1.- On appelle ensemble bien ordonné un ensemble ordonné E tel que toute partie non vide de E possède un plus petit élément.

Tout ensemble bien ordonné est totalement ordonné; la réciproque n'est pas vraie: par exemple, les ensembles totalement ordonnés $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ ne sont pas bien ordonnés.

Soit E un ensemble bien ordonné; si E est non vide, E possède un plus petit élément, de plus:

- a) tout élément x de E qui n'est pas le plus grand élément de E possède un successeur x^* .
- b) toute partie P de E qui est majorée admet une borne supérieure.

On peut vérifier que si un ensemble totalement ordonné vérifie les deux propriétés a) et b), alors E est bien ordonné; si en effet P est une partie non vide de E , considérons l'ensemble Q des minorants de P . Clairement la borne supérieure $\sup Q$ de Q est un minorant de P . De plus, $\sup Q \in P$ sinon $(\sup Q)^*$ est encore un minorant de P .

En revanche un élément d'un ensemble bien ordonné distinct du plus petit élément n'admet pas nécessairement de prédécesseur.

Si E est un ensemble totalement ordonné, et si x est un élément de E , on note S_x , et on appelle section commençante ouverte définie par x ,

l'intervalle $(\leftarrow, x[$ des éléments y de E tels que $y < x$. On appelle de façon générale section commençante une partie S de E telle que si $x \in S$ et $y \leq x$, alors $y \in S$; autrement dit, si $x \in S$, alors $S_x \subset S$. On dit qu'une section commençante S est stricte si $S \neq E$. Toute partie S_x est évidemment une section commençante stricte. On a pour un ensemble bien ordonné une réciproque :

Théorème 1.- Soit E un ensemble bien ordonné; l'application $x \mapsto S_x$ est une bijection croissante de E sur l'ensemble des sections commençantes strictes de E ordonné par l'inclusion.

Le problème est de montrer que l'application en question est surjective. Soit donc S une section commençante stricte de E . Alors $\bigcup_E S$ est non vide et admet un plus petit élément x . On vérifie alors que $S = S_x$. On aurait pu définir x comme la borne supérieure de E .

Théorème 2 (principe de récurrence transfinie).- Soit E un ensemble bien ordonné; si une partie P de E est telle que $S_x \subset P$ implique $x \in P$, alors $P = E$.

Il s'agit d'une généralisation du principe de récurrence sur tous les prédécesseurs. La démonstration est la même.

Il faut noter que les hypothèses " $e \in P$ " et " $x \in P$ implique $x^* \in P$ " n'impliquent pas nécessairement $P = E$. En revanche, on pourrait remplacer l'hypothèse du théorème 2 par les hypothèses " $x \in P$ implique $x^* \in P$ " et "pour toute partie majorée Q de P , la borne supérieure de Q appartient à P ".

On laisse au lecteur le soin de traduire le principe de récurrence transfinie en termes de relations.

Nous allons donner pour terminer quelques procédés permettant de construire des ensembles bien ordonnés :

- 1) Si E est un ensemble ordonné, toute partie de E est bien ordonnée pour l'ordre induit.
- 2) Si E est un ensemble ordonné, l'ensemble ordonné obtenu en adjoignant à E un élément qui majore tous les éléments de E (sans changer la relation d'ordre sur E) est bien ordonné.
- 3) Si $\mathcal{E}$ est un ensemble d'ensembles bien ordonnés tel que si E, E' sont des ensembles quelconques de $\mathcal{E}$, l'un est une section commençante de l'autre, alors $\bigcup_{E \in \mathcal{E}} E$ est bien ordonné pour l'ordre défini sur la réunion en "recollant" les ordres des ensembles de $\mathcal{E}$.

2. Comparaison des ensembles bien ordonnés

Nous introduisons une notion nouvelle.

Définition 2.- Soient E, E' des ensembles bien ordonnés; on dit qu'une application f de E dans E' est un bon morphisme si f est strictement croissante et si $f(E)$ est une section commençante de E' .

Il revient au même de dire que f est un isomorphisme de E sur une section commençante de E' . On vérifie encore que cela revient à imposer

$$f(S_x) = S_{f(x)}$$

pour tout élément x de E .

L'intérêt des bons morphismes réside dans la

Proposition 1.- Soient E, E' des ensembles bien ordonnés et f, g des applications de E dans E' . On suppose que f est un bon morphisme et que g est strictement croissante, alors $f \leq g$.

S'il n'en est pas ainsi, il existe un plus petit élément x dans E tel que $f(x) > g(x)$. Alors $g(x)$ appartient à $S_{f(x)}$, soit à $f(S_x)$ et s'écrit $f(y)$ pour un certain $y < x$. On a alors $f(y) \leq g(y)$ et $g(y) < g(x)$.

Finalement,

$$g(x) = f(y) \leq g(y) < g(x),$$

ce qui est absurde.

Corollaire.- Soient E, E' des ensembles bien ordonnés,

- 1) Il existe alors au plus un bon morphisme de E dans E' .
- 2) S'il existe un bon morphisme f de E dans E' et un bon morphisme g de E' dans E , alors E, E' sont isomorphes.

La première partie est une conséquence immédiate de la proposition 1 : si f_1, f_2 sont deux bons morphismes, on peut leur faire jouer le rôle de f ou g . Pour établir la seconde partie, il suffit de voir que $f \circ g = \text{Id}_{E'}$ et $g \circ f = \text{Id}_E$. Par exemple, $f \circ g$ est un bon morphisme de E' dans E' comme composé de deux bons morphismes; comme $\text{Id}_{E'}$ est aussi un bon morphisme de E' dans E' , il y a égalité.

Nous sommes maintenant en mesure de prouver le
Théorème 3 (de comparaison).- Soient E, E' des ensembles bien ordonnés. On a l'une des propriétés suivantes qui s'excluent mutuellement :

- a) E est isomorphe à une section commençante stricte de E' .
- b) E' est isomorphe à une section commençante stricte de E .
- c) E et E' sont isomorphes.

Le fait que a), b) et c) s'excluent mutuellement résulte du corollaire qui précède. Il reste à prouver que soit E est isomorphe à une section commençante de E' , soit l'inverse. Nous introduisons pour cela l'ensemble $\mathcal{S}$ des sections commençantes S de E telles qu'il existe un bon morphisme de S dans F ; le bon morphisme est unique et sera noté f_S . Si nous posons alors

$$M = \bigcup_{S \in \mathcal{S}} S,$$

il est clair que M est une section commençante de E . On définit une application f de M dans E' en posant $f(x) = f_S(x)$ pour $x \in S$ et en vérifiant que la définition est indépendante de la partie S contenant x , à cause de l'unicité des bons morphismes. De plus f est un bon morphisme : en effet, f est strictement croissante et

$$f(M) = \bigcup_{S \in \mathcal{S}} f(S)$$

est une section commençante de E' .

La démonstration est alors achevée si $M = E$ ou $f(M) = E'$; il reste à écarter le cas où l'on aurait $M \neq E$ et $f(M) \neq E'$. En désignant par x, y les plus petits éléments de $\bigcup_E M$ et $\bigcup_{E'} f(M)$, on vérifie que $S = M \cup \{x\}$ est une section commençante de E et que f , prolongée à x par $f(x) = y$, est un bon morphisme : f est strictement croissante et $f(M \cup \{x\}) = f(M) \cup \{f(y)\}$ est une section commençante de E' . Tout cela est absurde car S doit être incluse dans M .

Corollaire.- Soit P une partie d'un ensemble bien ordonné E ; alors P est isomorphe à une section commençante de E .

Il s'agit d'exclure le cas où il existerait un isomorphisme f de E sur une section commençante stricte de P ; en appliquant la proposition 1 au bon morphisme Id et à l'application strictement croissante f de E dans lui-même, on obtiendrait $f(x) \geq x$ pour tout $x \in E$. Alors $f(P)$ ne peut être une section commençante stricte de P car $f(P)$ n'admet pas de majorant strict.

3. Théorèmes fondamentaux

Nous démontrons en premier le théorème de Zorn, contrairement à l'ordre historique.

Théorème 4 (Zorn).- Tout ensemble ordonné E contient une partie bien ordonnée sans majorant strict.

Lorsque E est vide, il n'y a rien à démontrer, de sorte que E sera supposé non vide. Considérons alors l'application φ de $\mathcal{P}(E)$ dans $\mathcal{P}(E)$ qui à une partie A de E associe l'ensemble des majorants stricts de A lorsqu'il en existe et la partie A elle-même dans le cas contraire. Il est facile de voir que $\varphi(A)$ n'est jamais vide. Il résulte alors de l'axiome du choix que l'on peut trouver une application f de $\mathcal{P}(E)$ dans E telle que $f(A) \in \varphi(A)$ pour toute partie A de E . Autrement dit $f(A)$ est un majorant strict de A lorsqu'il en existe et un élément de A dans le cas contraire.

Nous introduisons l'ensemble $\mathcal{E}$ des parties A de E qui sont bien ordonnées pour l'ordre induit et qui vérifient

$$f(A \cap S_x) = x$$

pour tout élément x de A . Le noeud de la démonstration repose dans ce que si A, B sont des ensembles de $\mathcal{E}$, alors l'une de ces parties est une section commençante de l'autre.

Pour établir ce fait, nous introduisons l'ensemble C des éléments x de $A \cap B$ tels que

$$A \cap S_x \subset B \quad \text{et} \quad B \cap S_x \subset A.$$

D'abord C est une section commençante de A : en effet $C \subset A$; si $x \in C$ et $y \in A$, $y < x$ alors $y \in A \cap S_x \subset B$ de sorte que $y \in A \cap B$; de plus, $S_y \subset S_x$ de sorte que $A \cap S_y \subset B$ et $B \cap S_y \subset A$. De même C est une section commençante de B . Il suffit maintenant de montrer que $C = A$ ou $C = B$. Dans le cas contraire, soient a, b les plus petits éléments de $A \setminus C, B \setminus C$. Puisque $C = A \cap S_a$, on a $f(C) = a$; de même, $f(C) = b$, de sorte que $a = b$ et que cet élément appartient à $A \cap B$. De plus, $A \cap S_a = C \subset B$ et $B \cap S_b = C \subset A$. Par conséquent $a \in C$, ce qui est absurde.

Considérons maintenant l'ensemble

$$M = \bigcup_{A \in \mathcal{E}} A.$$

D'après ce que nous avons vu, la partie M est bien ordonnée. On vérifie aisément par ailleurs que $M \in \mathcal{E}$. Si M n'a pas de majorant strict, le théorème est démontré. Dans le cas contraire, $m = f(M)$ est un majorant strict de M . La partie $A = M \cup \{m\}$ est alors bien ordonnée et si $x \in A$, on a $f(A \cap S_x) = x$: c'est clair pour $x \in M$ car $f(A \cap S_x) = f(M \cap S_x)$ et si $x = m$ aussi puisque $f(A \cap S_m) = f(M)$. Par conséquent $A \in \mathcal{E}$, ce qui est absurde car A devrait être incluse dans M .

La version que nous avons donnée du théorème de Zorn n'est pas celle qui est utilisée en pratique. Pour énoncer cette dernière, nous avons besoin de la

Définition 3. - On dit qu'un ensemble ordonné E est inductif si toute partie totalement ordonnée de E est majorée.

Alors

Théorème 4 bis. - Tout ensemble ordonné inductif E possède un élément maximal, c'est à dire un élément x n'admettant pas de majorant strict.

En effet, d'après le théorème 4, l'ensemble E contient une partie bien ordonnée A sans majorant strict; puisque A est en particulier totalement ordonnée, A est majorée par un élément x . Il est clair que x ne peut admettre lui-même de majorant strict.

Remarque. - Si on avait remplacé les termes "totalement ordonné" par "bien ordonné" dans la définition d'un ensemble ordonné inductif, le théorème 4 bis resterait valable et serait même en principe plus fort. On ne l'a pas fait pour une raison simple : tel qu'il est énoncé, le théorème 4 bis ne se réfère pas à la notion d'ensemble bien ordonné; cette dernière n'a servi que pour la démonstration.

Exemples

1) Tout espace vectoriel E sur un corps K admet une base. Considérons en effet l'ensemble $\mathcal{L}$ des parties libres de E , ordonné par inclusion ; si $\mathcal{P} \in \mathcal{L}$ est un ensemble totalement ordonné de parties libres de E , il est facile de montrer que

$$M = \bigcup_{L \in \mathcal{P}} L$$

est encore une partie libre; or M majore $\mathcal{P}$. Cela montre que $\mathcal{L}$ est un ensemble inductif. D'après le théorème 4 bis, il existe dans $\mathcal{L}$ un élément maximal L_0 . Il reste à voir que L_0 est une base; autrement dit que L_0 est génératrice. Dans le cas contraire, il existerait un vecteur v de E n'appartenant pas au sous-espace vectoriel engendré par L_0 . Alors $L_0 \cup \{v\}$ serait encore une partie libre, ce qui est absurde car elle majorerait strictement L_0 .

2) En particulier il existe une base de $\mathbb{R}$ comme espace vectoriel sur $\mathbb{Q}$; une telle base s'appelle une base de Hamel. On peut démontrer qu'une telle

base a nécessairement la puissance du continu. En effet, chaque nombre réel peut être représenté par une famille $(\lambda_i)_{i \in F}$ de nombres rationnels indexée par une partie finie F de la base.

La donnée d'une base de Hamel permet de construire des applications $\mathbb{Q}$ -linéaires de $\mathbb{R}$ dans $\mathbb{R}$. Remarquons d'ailleurs que pour qu'une application f de $\mathbb{R}$ dans $\mathbb{R}$ soit $\mathbb{Q}$ -linéaire, il suffit qu'elle soit additive, c'est à dire vérifie

$$f(x+y) = f(x) + f(y)$$

pour x, y dans $\mathbb{R}$. En effet, cela implique pour tout entier relatif n , la relation

$$f(nx) = nf(x),$$

d'où, pour $n \neq 0$,

$$f(x) = \frac{1}{n} f(nx),$$

puis pour tout rationnel p/q la relation

$$f\left(\frac{p}{q}x\right) = \frac{p}{q} f(x).$$

Les applications $\mathbb{Q}$ -linéaires de $\mathbb{R}$ dans $\mathbb{R}$ sont donc les homomorphismes du groupe additif $\mathbb{R}$.

Si $(e_i)_{i \in I}$ est une base de Hamel, pour toute famille $(x_i)_{i \in I}$ de nombres réels, il existe une application $\mathbb{Q}$ -linéaire et une seule f de $\mathbb{R}$ dans $\mathbb{R}$ telle que $f(e_i) = x_i$ quel que soit $i \in I$. Il en résulte que l'ensemble des homomorphismes du groupe additif $\mathbb{R}$ a la puissance de $\mathcal{F}(I, \mathbb{R})$, et donc une puissance strictement supérieure à celle du continu.

On vérifie par ailleurs que l'ensemble des homomorphismes continus de $\mathbb{R}$ dans $\mathbb{R}$ a la puissance du continu. Cela démontre l'existence d'homomorphismes non continus.

Une autre application du théorème 4 bis est le théorème suivant dont on pourrait déduire inversement le théorème de Zorn.

Théorème 5 (Zermelo). - Tout ensemble peut être muni d'une relation d'ordre pour laquelle il est bien ordonné.

Soit E un ensemble. On considère l'ensemble $\mathcal{E}$ des couples (A, R) où A est une partie de E et R une relation de bon ordre sur A , muni de la relation d'ordre pour laquelle $(A, R) \leq (A', R')$ si (A, R) est une section commençante de (A', R') . L'ensemble ordonné $\mathcal{E}$ est inductif: si une partie $\mathcal{P}$ de $\mathcal{E}$ est totalement ordonnée, en posant

$$M = \bigcup_{A \in \text{pr}_1(\mathcal{P})} A$$

et en considérant sur M la relation d'ordre obtenue en "recollant" les relations $R \in \text{pr}_2(\mathcal{P})$, on obtient un ensemble bien ordonné et un majorant de $\mathcal{P}$. Il existe par conséquent dans $\mathcal{E}$ un élément maximal (A, R) . Alors $A = E$, sinon en considérant un élément $x \notin A$, la partie $A \cup \{x\}$ et la relation d'ordre sur $A \cup \{x\}$ qui coïncide avec R sur A et pour laquelle x majore les éléments de A , on obtiendrait un majorant strict de (A, R) .

IX

ORDINAUX, CARDINAUX

Nous utiliserons dans ce chapitre les théorèmes du chapitre VIII pour introduire les nombres transfinis et préciser la notion de puissance d'un ensemble.

1. Nombres ordinaux

Les nombres ordinaux sont caractérisés parmi les ensembles bien ordonnés de la même façon que l'ensemble $\mathbb{N}$ l'est parmi les $\mathbb{N}$ -ensembles.

Définition 1.- On dit qu'un ensemble bien ordonné E est un ordinal si l'on a $S_x = x$ pour tout élément x de E .

Dans un ordinal E , la relation $x \leq y$ équivaut donc à $x \subset y$; par conséquent il existe au plus une relation d'ordre sur un ensemble donné pour laquelle cet ensemble soit un ordinal. De la même façon, la relation $x < y$ équivaut à $x \in y$.

Proposition 1.- Si deux ordinaux α, β sont isomorphes, alors $\alpha = \beta$.

Supposons en effet qu'il existe un isomorphisme f de α sur β ; on démontre alors que $f(x) = x$ pour tout $x \in \alpha$ par récurrence transfinie. Supposons en effet que l'on ait $f(y) = y$ pour tout $y < x$. Il vient

$$f \langle S_x \rangle = S_x.$$

Or le premier membre est encore $S_{f(x)}$, de sorte que $f(x) = x$.

Remarque.- La notation $f(S_x)$ prête à confusion ici; son utilisation conduit à une simplification erronée de la démonstration précédente.

Corollaire.- Soient α, β des ordinaux; on a l'une des propriétés suivantes qui s'excluent:

- a) α est une section commençante stricte de β .
- b) β est une section commençante stricte de α .
- c) $\alpha = \beta$.

C'est une conséquence de la proposition 1 et du théorème 3 du chapitre VIII. Notons que a) s'écrit encore $\alpha \in \beta$ et b) s'écrit de même $\beta \in \alpha$.

Si α, β sont des ordinaux, nous notons $\alpha \leq \beta$ la relation d'inclusion $\alpha \subset \beta$. Le corollaire qui précède exprime le fait qu'il s'agit d'une relation d'ordre total et que la relation $\alpha < \beta$ n'est autre que la relation $\alpha \in \beta$.

Ces notations sont en accord avec celles utilisées lorsque α, β sont des éléments d'un autre ordinal γ . Notons également que si α est un ordinal, on ne peut avoir $\alpha \in \alpha$ car cela équivaut à $\alpha < \alpha$.

L'ensemble $\mathbb{N}$ est un ordinal. De plus :

- 1) Toute section commençante d'un ordinal est un ordinal; en particulier, tout élément d'un ordinal est un ordinal; tout entier naturel est un ordinal; $\mathbb{N}$ est donc l'ensemble des ordinaux strictement inférieurs à $\mathbb{N}$.
- 2) Si α est un ordinal, alors $\alpha^* = \alpha \cup \{\alpha\}$ est un ordinal.
- 3) Si $\mathcal{E}$ est un ensemble d'ordinaux, alors $\bigcup_{\alpha \in \mathcal{E}} \alpha$ est un ordinal : en effet, si α, β sont des éléments de $\mathcal{E}$, l'un est une section commençante de l'autre.

Nous sommes maintenant en mesure d'expliquer l'antinomie Burali-Forti; elle consiste à constater que parler de l'ensemble des ordinaux est contradictoire. Si $\mathcal{E}$ désigne en effet cet ensemble, l'ordinal $\omega = \bigcup_{\alpha \in \mathcal{E}} \alpha$ majore évidemment les ordinaux de $\mathcal{E}$; or $\omega \cup \{\omega\}$ majore strictement ω puisque $\omega \notin \omega$.

Les ordinaux ne constituent donc pas un ensemble. Quand on parle de relation d'ordre entre ordinaux, il ne s'agit pas d'une relation d'ordre sur un ensemble. Cette relation est une relation de bon ordre dans le sens suivant :

Tout ensemble non vide d'ordinaux possède un plus petit élément.

Soit en effet $\mathcal{E}$ un ensemble non vide d'ordinaux. Alors $\beta = \bigcap_{\alpha \in \mathcal{E}} \alpha$ est un ordinal qui minore évidemment les ordinaux de $\mathcal{E}$. De plus, $\beta \in \mathcal{E}$ car dans le cas contraire on aurait $\beta \in \alpha$ pour tout $\alpha \in \mathcal{E}$ et par suite $\beta \in \beta$, ce qui est impossible.

La proposition 1 est complétée par le

Théorème 1. - Soit E un ensemble bien ordonné; il existe un ordinal et un seul qui soit isomorphe à E ; cet ordinal est noté $\text{Ord}(E)$ et appelé type d'ordre de E .

L'unicité résulte de la proposition 1. On démontre par récurrence transfinie qu'il existe pour tout élément x de E un ordinal α_x et un seul qui soit isomorphe à S_x . On vérifie enfin que les α_x constituent un ensemble, qui est un ordinal isomorphe à E .

2. Nombres cardinaux

Il s'agit de donner un sens précis à la notion de puissance. En 1884, Frege adopte la définition très simple suivante : la puissance d'un ensemble E est l'ensemble de tous les ensembles équipotents à E . Malheureusement, il est facile de vérifier que cette définition est contradictoire : on ne peut pas parler de l'ensemble des ensembles équipotents à E .

Nous allons faire intervenir les nombres ordinaux. On sait, par le théorème de Zermelo, que tout ensemble E peut être muni d'une relation de bon ordre, puis par le théorème 1, que l'ensemble ordonné obtenu est isomorphe à un ordinal. Par suite

Proposition 2.- Tout ensemble est équipotent à un ordinal.

Cet ordinal n'est cependant pas unique ; par exemple $\mathbb{N}$ et $\mathbb{N} \cup \{\mathbb{N}\}$ sont équipotents. Aussi adopte-t-on la

Définition 2.- Soit E un ensemble ; on appelle cardinal de E et on note $\text{Card}(E)$ le plus petit ordinal équipotent à E .

Les propriétés du cardinal de E sont données par la

Proposition 3.- Soient E, F des ensembles.

- 1) E et $\text{Card}(E)$ sont équipotents.
- 2) $\text{Card}(E) = \text{Card}(F)$ si et seulement si E, F sont équipotents.
- 3) $\text{Card}(E) \leq \text{Card}(F)$ si et seulement si il existe une injection de E dans F (resp. une surjection de F sur E ou bien si $E = \emptyset$); en particulier si $E \subset F$, alors $\text{Card}(E) \leq \text{Card}(F)$.
- 4) $\text{Card}(E) < \text{Card}(\mathcal{P}(E))$.

La propriété 1) résulte de la définition ; par suite, si $\text{Card}(E) = \text{Card}(F)$, alors E, F sont équipotents. Inversement, si E, F sont équipotents, E et $\text{Card}(F)$ le sont, ce qui implique $\text{Card}(E) \leq \text{Card}(F)$. De la même façon $\text{Card}(F) \leq \text{Card}(E)$ et la propriété 2) est complètement démontrée.

Si $\text{Card}(E) \leq \text{Card}(F)$, alors $\text{Card}(E) < \text{Card}(F)$ et on en déduit une injection de E dans F . Inversement, si une telle injection existe, on en

déduit une injection f de E dans $\text{Card}(F)$; alors $f(E)$ est isomorphe à une section commençante de E de sorte que $\text{Ord}(f(E)) \leq \text{Card}(F)$. Par suite, $\text{Card}(f(E)) \leq \text{Card}(F)$ puis $\text{Card}(E) \leq \text{Card}(F)$.

La propriété 4) n'est enfin que la traduction du théorème de Cantor établi au chapitre I.

Le point 3) et la propriété d'antisymétrie de la relation d'ordre entre cardinaux se traduisent par le

Corollaire (Théorème de Cantor-Bernstein).- Soient E, F des ensembles. S'il existe une injection de E dans F et une injection de F dans E , il existe une bijection de E sur F .

On dit qu'un ordinal α est un cardinal si $\alpha = \text{Card}(\alpha)$; cela signifie que α n'est équipotent à aucun ordinal $\beta < \alpha$.

Lorsque E est un ensemble fini, la définition 2 coïncide avec celle déjà donnée et les nombres entiers naturels sont des cardinaux.

Proposition 4.- Les propriétés qui suivent sont équivalentes :

- 1) n est un entier naturel.
- 2) n est un cardinal fini.
- 3) n est un ordinal fini.
- 4) n est un ordinal non équipotent à n^* .

On sait que 1) implique 2) et 2) implique évidemment 3). De plus, 3) implique 4) car si un ordinal n est équipotent à $n^* = n \cup \{n\}$, alors l'ensemble n^* est infini, ce qui est impossible lorsque l'ensemble n est fini. Enfin 4) implique 1) car si un ordinal n n'est pas un entier naturel, alors $n \geq \mathbb{N}$ de sorte que l'application qui à 0 associe $\{n\}$, à $x \in \mathbb{N} \setminus \{0\}$ associe *x et à $x \in n \setminus \mathbb{N}$ associe x lui-même est une bijection de n sur n^* .

L'ensemble $\mathbb{N}$ est un cardinal; c'est le plus petit ordinal infini.

Nous allons, pour terminer, décrire rapidement quelques opérations sur les cardinaux qui étendent celles données sur les nombres entiers.

1) Soit $(\alpha_i)_{i \in I}$ une famille de cardinaux; on pose

$$\sum_{i \in I} \alpha_i = \text{Card} \left(\bigsqcup_{i \in I} \alpha_i \right).$$

En particulier $\alpha + \beta = \text{Card}(\alpha \sqcup \beta)$.

2) On pose

$$\prod_{i \in I} \alpha_i = \text{Card} \left(\prod_{i \in I} \alpha_i \right).$$

En particulier $\alpha \cdot \beta = \text{Card} (\alpha \times \beta)$.

3) On pose

$$\alpha^\beta = \text{Card } \mathcal{G}(\beta, \alpha).$$

On laisse au lecteur le soin de démontrer un certain nombre de propriétés concernant ces opérations; par exemple :

$$\text{Card} \left(\bigcup_{i \in I} E_i \right) \leq \sum_{i \in I} \text{Card}(E_i),$$

l'égalité ayant lieu si les ensembles sont disjoints,

$$\text{Card} \left(\prod_{i \in I} E_i \right) = \prod_{i \in I} \text{Card}(E_i),$$

$$\text{Card } \mathcal{G}(E, F) = (\text{Card } F)^{\text{Card } E}.$$

Enfin un cardinal α est fini si et seulement si $\alpha \neq \alpha + 1$.

