



7. LES DÉCORS PLAFONNANTS, tel l'Oculus de la chambre des époux d'Andrea Mantegna (ci-dessus) ou les trompe-l'œil sur voûtes sont des représentations en perspective où l'artiste adopte un «point de vue» où il s'affranchit de la verticalité.

appelle une involution et auquel on préfère aujourd'hui le terme de birapport de quatre points) entre les segments définis par six points d'une droite reste inchangé par projection centrale ou par projection parallèlement à une direction donnée.

Les découvertes de Desargues ont rendue possible une théorie générale des projections, à laquelle s'emploieront les géomètres de la première moitié du XIX^e siècle, Gaspard Monge, Jean-Victor Poncelet et Michel Chasles, pour ne citer qu'eux. À son tour, la géométrie projective construite par ces mathématiciens conduira à entrevoir la possibilité des géométries non euclidiennes et à en concevoir des modèles euclidiens. On y rencontre aussi des représentations de l'infini. Par exemple, dans les modèles d'Henri Poincaré ou d'Eugenio Beltrami pour le «plan hyperbolique», l'infini est ramené à distance finie en assimilant les points à l'infini à un cercle ; l'intérieur de ce cercle constitue le «plan», et les «droites» de ce plan sont par définition des arcs circulaires orthogonaux au contour. Il s'agit d'un modèle euclidien d'une géométrie non euclidienne, puisque, par un point extérieur à une «droite», passent une infinité de «droites» parallèles à la droite en question. On notera que les extrémités des droites sont situées sur le contour, c'est-à-dire à l'infini : avec un tel modèle, le monde clos des Anciens fait en quelque sorte son retour...

La perspective cavalière

Les projections parallèles, projections où les rayons qui partent de l'objet à représenter sont parallèles à une direction donnée, semblent connues de longue date, au moins depuis l'architecte romain Vitruve (I^{er} siècle de notre ère) : beaucoup de dessins, qu'ils soient à vocation artistique ou technique, semblent relever de ce que l'on nomme aujourd'hui la perspective cavalière ou d'une projection orthogonale (projection dans une direction orthogonale au plan du dessin). En fait, il ne s'agit que de tentatives empiriques de représenter localement la profondeur pour des objets isolés. Ils ne peuvent en aucun cas relever d'une volonté de «représenter l'espace» de façon globale, cohérente et homogène. On ne peut véritablement parler de pers-

pective cavalière qu'à partir du moment où ses règles ont été codifiées géométriquement. Or celles-ci, bien que nettement plus simples que celles de la perspective centrale, n'ont été édictées qu'après celles de la perspective centrale du Quattrocento.

Les premiers dessins en perspective parallèle systématique apparaissent dans les théâtres de machines (des plans de machines) ou dans les ouvrages d'ingénieurs du XVI^e siècle, et les premiers traités de perspective cavalière ou militaire sont le fait de théoriciens français du début du XVII^e siècle. Quant aux architectes, qui ont donné la perspective centrale aux peintres, ils ne se sont saisis que tardivement, à partir de la fin du XIX^e siècle, de la perspective cavalière et de l'axonométrie (la représentation des trois dimensions par projection sur trois axes tracés dans le plan du dessin).

Dans les perspectives parallèles, on considère les rayons visuels parallèles. Comme l'a montré Desargues, on peut aussi dire que ces rayons visuels se rencontrent à l'infini. En d'autres termes, la perspective cavalière équivaut à une perspective centrale où l'œil du dessinateur est «rejeté à l'infini». En 1925, le peintre russe Lissitzky s'est exclamé : «Le suprématisme a fait reculer l'extrémité de la pointe de la pyramide visuelle à l'infini... [Ce faisant] il a inventé la dernière illusion : l'extensibilité infinie vers l'arrière-plan ou l'avant-plan.» Il lui semblait qu'on en avait fini avec la peinture du sujet renaissant, que le point de vue n'existait plus. Le peintre s'assimile alors au Créateur par le regard venu de l'infini. Il nous paraît pourtant que, dans une perspective parallèle, le point de vue est toujours présent, tout rejeté à l'infini qu'il soit et justement du fait qu'il est rejeté : une exclusion n'est pas une absence.

Jean-Pierre LE GOFF est professeur de mathématiques et d'histoire des mathématiques à l'IUFM (Institut universitaire de formation des maîtres) de Caen et à l'IREM (Institut de recherches sur l'enseignement des mathématiques) de Basse-Normandie (Université de Caen).

René TATON, *L'œuvre mathématique de Girard Desargues*, PUF, 1951 (réédition Vrin, 1988).

Philippe COMAR, *La perspective en jeu : Les dessous de l'image*, Éditions Gallimard/La Découverte, n° 138, 1992.

Jean-Pierre LE GOFF, *Desargues et la naissance de la géométrie projective*, in

Desargues en son temps, sous la direction de Jean Dhombres et Joël Sakarovich, Blanchard, 1994.

Jean-Pierre LE GOFF, *Piero della Francesca. De la perspective en peinture*, préface de Hubert Damisch et postface de Daniel Arasse, in *Medias Res*, 1998.

Les Cahiers de la perspective, sous la direction de Didier BESSOT et Jean-Pierre LE GOFF, IREM de Basse-Normandie, Université de Caen, 1981-1996.

D. BESSOT et J.-P. LE GOFF, *Mais où est donc passée la troisième dimension ?* in *Histoires de problèmes, histoire des mathématiques*, Éditions Ellipses, 1993.

L'ensemble triadique de Cantor

ANDRÉ DELEDICQ

Les préposés au nettoyage de la planète Ter découvrent un ensemble de Cantor qui remplit un segment de largeur 1, mais son complémentaire aussi...



Dans les écoles de Gullicantor, une province de Ter, l'éducation était stricte : le plus gros travail dévolu aux garçons était le nettoyage des barres de cuivre que les Gullicantoriens utilisaient comme décoration.

Les règles de nettoyage étaient strictes : il fallait, dans l'empire de Ter, d'abord nettoyer le tiers central de chaque barre, puis les tiers centraux de chaque tiers situés de part et d'autre du tiers que l'on venait de nettoyer, et ainsi de suite. Indéfiniment. Les mathématiciennes de Ter avaient montré que ce nettoyage de chaque barre ne prend qu'un temps fini si l'on suppose que la durée de nettoyage d'un segment ne dépend que de sa longueur.

Comment savoir si toute la barre, que l'on suppose de longueur 1, est bien nettoyée par ce processus, s'interrogeaient les jeunes Gullicantoriens ? À chaque promotion de jeunes garçons, les mathématiciennes de l'empire expliquaient que, le tiers central étant nettoyé ($1/3$), ils nettoyaient ensuite un tiers des deux tiers restants, soit $(1/3) \times (2/3)$, et ainsi de suite, ce qui faisait : $1/3 + 2/3 \times (1/3 + 2/3 \times (1/3 + 2/3 \times (1/3 + \dots)))$

La formule traduit effectivement le processus de nettoyage, s'extasiaient les jeunes garçons : un tiers, plus un tiers sur les deux tiers restants, plus encore un tiers sur les deux tiers restants, et ainsi de suite, soit, après simplification : $(1 + 2/3 + 4/9 + 8/27 + \dots) / 3$,

ce qui est la série géométrique de raison $2/3$ divisée par 3.

Nous savons, continuaient les garçons, que la somme : $1 + a + a^2 + \dots + a^n \dots$ vaut $1/(1 - a)$. Ici, a est égal à $2/3$ et la somme est égale à $1/(1 - 2/3)$, soit exactement 3. Comme cette somme est multipliée par $1/3$, cela fait 1 : à la fin, la barre tout entière serait nettoyée.

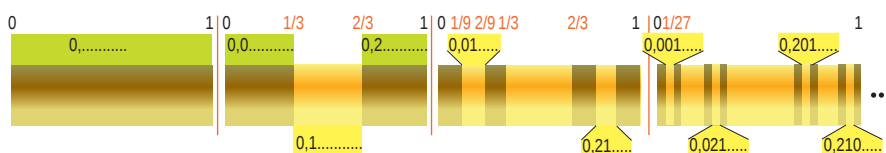
Hélas, des garçons de Ter n'avaient pas la bosse des maths, ce qui ennuyait les autorités. Celles-ci pensaient instituer un système de quotas pour les postes de mathématiciens à l'Université. Les mathématiciennes expliquaient le phénomène de nettoyage en tirant parti du fait que les habitants de Ter, leurs noms l'indique, n'ont que trois doigts. Les Terriens comptent ainsi en base trois et écrivent la suite des entiers : 0, 1, 2, 00, 01, 02, 10, 11, 12, 20, 21, 22... ; de même, les nombres fractionnaires, exprimés en base 3, ne comportent que des 0, des 1 et des 2.

Les points de la barre qui seront nettoyés sont ceux qui, à un moment quelconque du découpage en trois, ont un chiffre «1», comme on le voit sur la figure ci-dessous. Et les nombres dont l'écriture ne contient que des 0 et des 2 sont des abscisses de points de la barre qui resteront sales, par exemple 0,202020, 0,200200200 et bien d'autres. Même à la fin du nettoyage, tous ces points ne seront pas nettoyés : l'ensemble de ces

points, que l'on dénomme une poussière de Cantor, restera sale.

Pour écrire les abscisses des points non nettoyés, on n'a besoin que de deux doigts, celui qui marque 0 et celui qui marque 2, remarqua la mathématicienne. Imaginons maintenant, continua-t-elle, que, comme les habitants de la planète Bis, nous n'ayons que deux doigts et que nous ne puissions calculer qu'en base deux. En comptant sur nos doigts nous pourrions encore écrire les abscisses des nombres du segment. Quand j'écris un nombre avec deux doigts, poursuivit-elle, nous pouvons reconnaître l'abscisse d'un point ou d'un autre selon la manière dont on interprète l'écriture : – si nous pensons qu'il s'agit d'une écriture des Bissiens, en base 2, nous obtenons tous les points du segment, – si nous pensons que c'est une écriture de Terriens, en base 3, nous obtenons tous les points non nettoyés.

L'ensemble de ces points non nettoyés est donc en bijection avec l'ensemble des points du segment : il semblerait que l'on n'ait rien nettoyé du tout ! Quel travail, s'exclamèrent en chœur les jeunes garçons, et nous ne savons même pas si on a nettoyé tout le segment ou si il est encore entièrement non nettoyé...



André DELEDICQ est maître de conférence. Il travaille à l'Institut de recherche pour l'enseignement des mathématiques (IREM) de l'Université de Paris 7.

André DELEDICQ et Francis CASIRO, *Apprivoiser l'infini*, ACL-Les Éditions du Kangourou, 1997.

L'infini, pierre de touche du constructivisme

ALLAN CALDER

Selon les mathématiciens constructivistes, un objet mathématique n'existe que si on peut le construire. Cette définition de l'existence mathématique oppose, depuis plus d'un siècle, constructivistes et formalistes.



La plus haute perfection de Dieu est la possibilité de créer un ensemble infini, et son immense bonté le conduit à le créer.

Georg Cantor

Si des hommes communiquaient avec une autre forme de vie dans l'Univers, les deux civilisations auraient, pensent les mathématiciens réalistes, en commun une forme élémentaire de mathématiques qu'elles utiliseraient pour dialoguer. Depuis l'époque de Platon, de nombreux mathématiciens et philosophes pensent que les mathématiques existent extrinsèquement, en dehors de la connaissance que nous en avons, ce qui leur confère une sorte de vérité absolue. La tâche du mathématicien consiste à découvrir cette vérité.

Cette croyance en une mathématique «préexistante» n'est pas partagée par tous : au XIX^e siècle, le mathématicien allemand Leopold Kronecker était persuadé que Dieu avait créé les entiers et que «tout le reste est l'œuvre de l'homme». Selon ce point de vue, le mathématicien ne découvre pas les mathématiques, il les invente.

Mathématiques découvertes ou mathématiques inventées, la question n'est pas futile : selon leur option philosophique, les mathématiciens orientent différemment leurs travaux. Ainsi, l'idée des mathématiques constructives, selon laquelle l'existence

d'un objet mathématique n'est établie que lorsqu'on a montré comment le construire, est fondée sur la conviction que les mathématiques sont inventées. Les conséquences de ce choix sur des sujets aussi importants que la validité de démonstrations mathématiques ou la signification de l'existence mathématique sont mises en lumière par les méthodes et les idées des mathématiques grecques, qui ont guidé l'évolution des mathématiques pendant les 25 derniers siècles.

Pour les Grecs, les mathématiques étaient presque réduites à la géométrie, et les *Éléments* d'Euclide en étaient le fleuron. Les méthodes déductives de l'école éléatique de philosophie (Parménide, Zénon, etc.), la logique aristotélicienne et l'idéalisme platonicien s'y fondaient en un système axiomatique. Euclide partit d'objets non définis, tels que le point, la droite et le plan, ainsi que d'un certain nombre d'axiomes, c'est-à-dire d'intuitions concernant ces objets ; il leur appliqua les règles de la déduction aristotélicienne pour obtenir de nouvelles propositions, ou théorèmes.

Bien que des objets comme le point (qui n'a pas de dimension) et la droite (qui a une longueur, mais pas de largeur) n'existent pas dans le monde réel, nous avons généralement le sentiment de savoir ce qu'ils sont. De surcroît, les théorèmes d'Euclide étaient confirmés

empiriquement : à l'aide d'un papier et d'un crayon, on vérifie, par exemple, que deux triangles dont les côtés sont égaux ont des angles égaux.

À l'époque moderne, les mathématiques ont continué à dépendre de la méthode axiomatique ; dans nombre de cas, cette méthode a été appliquée à des objets auxquels manquaient les vertus intuitives ou «constructibles» de la géométrie euclidienne. Pour les constructivistes, de telles applications ont faussé le sens des mathématiques. C'est pourquoi un certain nombre de propositions ont été avancées, ces 100 dernières années, afin de renouer avec la tradition intuitive. Nous examinerons ici cette longue controverse, dont les aspects les plus importants apparaissent au cours de l'évolution mathématique du concept d'infini.

Les constructions de l'infini

Les paradoxes de Zénon sur la divisibilité à l'infini et la continuité échaudèrent les mathématiciens grecs qui devinrent extrêmement prudents dans l'utilisation des infinis. L'objet fondamental des mathématiques grecques était la longueur, et Euclide, par exemple, se référait, non pas à des lignes sans fin, mais à des segments de droite extensibles à des longueurs arbitraires. Cette notion d'infini «potentiel» – on considère des lignes finies en

se réservant la possibilité de les étendre à l'infini – est une extrapolation raisonnable de l'expérience. Le concept opposé, l'infini «actuel», selon lequel il existe réellement des lignes de longueur infinie, est métaphysiquement différent ; il demande un effort d'imagination vertigineux. Aussi les mathématiciens ont-ils évité les infinis actuels pour se limiter aux infinis potentiels.

Au XVII^e siècle, René Descartes et Pierre de Fermat, en introduisant les coordonnées cartésiennes, transformèrent les problèmes sur les formes géométriques en problèmes sur les nombres. Cette numérisation de la géométrie conféra au nombre, plutôt qu'à la longueur, la place prépondérante dans les mathématiques. Les bases des travaux de Newton et de Leibniz sur le calcul infinitésimal étaient posées sous la forme d'un ensemble de lois pour manipuler des nombres qui sont infiniment petits sans être nuls.

Ce concept de quantité infiniment petite était troublant, à une époque où les objets mathématiques étaient porteurs de la même réalité que, par exemple, les électrons aujourd'hui. Il fallut un certain temps avant que l'indéniable utilité du calcul infinitésimal oblige les mathématiciens à accepter de plus ou moins bon gré l'idée des infinis actuels qui en résultait. Pendant les deux siècles qui suivirent, les mathématiciens forgèrent avec succès des techniques de calcul tout en contestant leurs bases philosophiques. À la fin du XIX^e siècle, cet irritant grain de sable au fond de l'huître mathématique produisit une perle : la théorie des ensembles.

Depuis Georg Cantor, à qui l'on doit les premières ébauches de la théorie des ensembles, un ensemble est défini comme une collection d'objets, réels ou abstraits, munie d'une loi qui détermine leur appartenance à l'ensemble. Cette idée simple eut de profondes

conséquences sur le concept d'infini. Des ensembles comme les nombres réels sont, par leur nature même, des infinis actuels. L'étude de tels objets conduisit Cantor à tirer la conclusion déconcertante que, tout comme la cardinalité des ensembles finis (ou leur nombre d'éléments), la cardinalité des ensembles infinis pouvait être mesurée : certains ensembles étaient plus infinis que d'autres!

Cantor définit alors toute une hiérarchie d'ensembles infinis : l'ensemble des nombres rationnels, des nombres irrationnels (des nombres réels qui ne sont pas des fractions), des nombres transcendants (des nombres réels qui ne sont racine d'aucune équation de type $a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = 0$, où les nombres a_n sont des entiers relatifs), etc. Selon Cantor, deux ensembles ont la même taille, ou cardinalité, quand tous les éléments de l'un peuvent être, un à un, associés à tous les



1. DANS L'ÉCOLE D'ATHÈNES, qu'il exécuta au début du XVI^e siècle, Raphaël dépeint les grands philosophes et mathématiciens de la Grèce antique : au centre se tiennent Platon (la main levée) et Aristote ; en bas à gauche, Pythagore médite sur son système de proportions, alors qu'Euclide trace un cercle sur une ardoise, en bas à droite. Pour les Grecs, les mathématiques se limitaient presque exclusivement à la géométrie, et c'est dans son traité de géométrie, les *Éléments*, qu'Eu-

clide réunit pour la première fois les bases d'un système axiomatique. Dans ce système, certaines affirmations évidentes, initiatives ou axiomes, permettent d'obtenir, grâce aux règles aristotéliennes de la déduction, de nouveaux résultats nommés théorèmes. Les mathématiciens constructivistes proposent de reconstruire les mathématiques sans l'aide des concepts non intuitifs, tels que l'infini actuel, introduits en mathématiques depuis quelques siècles.

éléments de l'autre. Cette manière de faire correspondre deux ensembles est une bijection.

Ainsi un ensemble infini est dénombrable s'il a la même cardinalité que l'ensemble des naturels positifs 1, 2, 3... (voir *L'infini est-il paradoxal en mathématiques?*, pages 30 à 38). Beaucoup de mathématiciens contemporains de Cantor trouvèrent que l'idée d'ensemble infini était absurde. Cantor lui-même admit qu'il avait été acculé «presque contre son gré» à considérer des ensembles infinis comme des entités uniques. Cette suppression de la distinction entre infini potentiel et infini actuel était «en contradiction avec des traditions appréciées». Et, plus choquant encore que la notion d'ensemble,

était l'usage que Cantor en faisait ; les mathématiciens furent scandalisés par sa démonstration que l'ensemble des nombres transcendants est infini non dénombrable.

Preuves non constructives

Un nombre transcendant est un nombre qui n'est pas algébrique, et un nombre x est algébrique s'il est solution d'une équation de la forme indiquée précédemment. Par exemple, les nombres rationnels sont des nombres réels qui peuvent se mettre sous la forme p/q , avec p et q entiers, q non nul ; ainsi, tout nombre rationnel x est algébrique, car il vérifie l'équation $qx - p = 0$. Cantor démontra qu'il

existe des nombres, les nombres transcendants, qui ne remplissent pas ces conditions ; il montra surtout que les nombres transcendants sont en nombre infini, en commençant par préciser que l'ensemble des nombres algébriques est infini dénombrable. Autrement dit, il mit en évidence une bijection entre l'ensemble des entiers naturels et l'ensemble des nombres algébriques.

L'ensemble des nombres réels n'étant pas dénombrable, il doit donc exister des nombres réels «laissés pour compte», c'est-à-dire qui ne sont pas algébriques. Supposons maintenant que l'ensemble de ces nombres transcendants soit dénombrable. Il n'est pas difficile de démontrer que l'ensemble formé de l'union de deux ensembles dénom-

2. Les ordinaux, des infinis toujours plus grands

Les premiers nombres connus furent les nombres entiers naturels 0, 1, 2, 3..., qui permettent de dénombrer les ensembles d'objets discontinus : des moutons, des sacs de blé, etc. On représente géométriquement ces nombres entiers naturels par des points régulièrement répartis sur une demi-droite, la distance entre l'origine et un point de la demi-droite représentant un nombre entier naturel. Les Égyptiens, les Grecs, les Arabes et les Européens de la Renaissance ont généralisé le concept de nombre entier naturel et défini les nombres rationnels, ou fractions, grâce aux opérations «algébriques» que sont la somme, la différence, le produit et le quotient de deux nombres entiers naturels.

Après les nombres rationnels, les Grecs introduisirent les nombres réels pour des raisons géométriques : $\sqrt{2}$, par exemple, est la longueur de la diagonale d'un carré de côté 1. Avec ces généralisations, tous les points d'une droite correspondent à un nombre réel, et inversement. Les opérations algébriques appliquées à ces nombres conduisirent à diverses généralisations, tels les nombres complexes.

Vers la fin du XIX^e siècle, le mathématicien allemand Georg Cantor voulut généraliser le concept de nombre, non pas de façon algébrique, comme l'avaient fait ses prédécesseurs, mais du point de vue de leur ordre : autrement dit, il voulut, à partir du principe de dénombrement utilisé dans la construction des nombres entiers, généraliser la relation «est supérieur à» (par exemple, 7 est supérieur à 4) et conserver l'acte de compter par lequel on place un nouveau nombre après des nombres déjà définis. Les nombres infinis, ou «transfinités», inventés par Cantor permettent de compter plus loin qu'avec les nombres entiers naturels. Ainsi les ordinaux sont, soit les nombres entiers naturels, soit les nombres transfinités.

Comment Cantor a-t-il proposé de dépasser l'infini des nombres entiers naturels ? On l'imagine en reprenant l'interprétation géométrique des nombres par des points d'une demi-droite : pour dépasser l'infini que représente cet ensemble de points, on juxtapose seulement une deuxième demi-droite à côté de la première, puis une troisième, et ainsi de suite (voir la figure ci-contre). On juxtapose ainsi 10, ou 100 demi-droites, et comme le nombre de demi-droites n'est pas limité, on obtient finalement un quart de plan tout entier plein de demi-droites. Pourquoi s'arrêter là et ne pas chercher ensuite à empiler de tels plans ? Les possibilités de juxtaposition sont illimitées, et les objets obtenus sont les ordinaux.

Revenons à l'image intuitive des demi-droites et examinons les ordinaux successifs que nous créons par le procédé précédent. Après avoir compté tous les nombres entiers naturels (en nombre infini), on introduit le nombre transfini ω , que l'on représente comme l'origine de la deuxième demi-droite. En se déplaçant ensuite sur cette demi-droite, on compte $\omega + 1$, $\omega + 2$, $\omega + 3$... Puis, quand on a ainsi compté tous les nombres de la forme $\omega + k$, où k est un nombre entier naturel, on introduit le nombre $\omega + \omega$, que l'on note aussi $\omega \times 2$, qui constitue l'origine de la troisième demi-droite. Enfin, quand on a compté tous les nombres de la forme $(\omega \times p) + k$ (où p et k sont des nombres entiers naturels quelconques), on introduit le nombre transfini ω^2 : à ce stade, on a formé le quart de plan évoqué précédemment. Ainsi, le nombre ordinal ω^2 qui vient après tous les nombres de ce réseau régulier plan est placé au-dessus de 0 dans un plan superposé au premier. On empilera alors d'autres plans pour faire $\omega^2 \times 2$, $\omega^2 \times 3$...

Dans un espace de dimension n supérieure à 3, l'ordinal ω^n est encore visualisable, mais la représentation des ordinaux a des limites : l'ordinal ω^n qui vient après tous les ω^n n'est pas représentable aussi simplement. Le leitmotiv de la construction des nombres transfinités est «après» : étant donné un ensemble X d'ordinaux, on définit un nouvel ordinal x par la condition que x soit plus grand (ou «après») que tous les éléments de l'ensemble X .

À ce stade, nous n'avons pas encore défini rigoureusement les ordinaux. La définition la plus simple est fondée sur un concept de «bon ordre», c'est-à-dire une façon de classer (d'ordonner) les nombres entiers naturels. Le bon ordre le plus simple correspond à celui que nous utilisons quand nous énumérons les nombres entiers naturels et que nous avons mis en œuvre précédemment : son principe tient dans «est après». Toutefois d'autres bons ordres sont possibles. Par exemple, choisissons d'ordonner deux à deux les nombres entiers non nuls en utilisant leur décomposition en facteurs premiers : tout nombre entier naturel non nul est le produit des nombres premiers 2, 3, 5, 7, 11, 13... avec des exposants convenables. Par exemple, 84 est égal au produit $2^2 \times 3^1 \times 5^0 \times 7^1$. On obtient un bon ordre si, pour comparer deux nombres entiers naturels, on considère leur décomposition en facteurs premiers et l'on cherche le plus grand facteur qui n'a pas le même exposant dans les deux nombres ; on décide de placer

brables est dénombrable (cela ne pose aucun problème si un ensemble, ou les deux, sont finis ; si tous les deux sont infinis, pour montrer que leur union est dénombrable, il suffit de mettre l'un d'eux en bijection avec l'ensemble des nombres impairs, et l'autre avec les nombres pairs). Alors, si l'hypothèse était vraie, l'ensemble des nombres réels (qui est constitué de l'union des nombres algébriques et des nombres transcendants) serait dénombrable, ce qui est une contradiction. Par conséquent, l'hypothèse selon laquelle l'ensemble des nombres transcendants est dénombrable est fautive : cet ensemble est donc infini non dénombrable.

La démonstration de Cantor utilise la méthode classique du raison-

nement par l'absurde, où l'on démontre une proposition en montrant que sa négation conduit à une contradiction. Des mathématiciens comme Kronecker furent scandalisés par une telle démonstration, parce que, dans la tradition mathématique, il n'y a pas d'existence sans construction : pour établir l'existence d'un objet mathématique, on doit spécifier par quelle méthode on peut, au moins en principe, construire l'objet. Que cette construction soit véritablement effectuée, ou même faisable, en un temps «raisonnablement» court, n'est pas une nécessité, du moment que la construction comporte un nombre fini d'étapes et que le passage d'une étape à la suivante est explicite. La démon-

stration de Cantor ne remplissait pas ces conditions : elle ne permettait pas de construire des nombres transcendants et certainement pas une infinité d'entre eux!

La démonstration de Cantor fut un des premiers exemples d'une preuve d'existence pure. Elle n'indique pas comment obtenir un nombre transcendant, mais établit l'existence de tels nombres en montrant la contradiction qu'entraînerait leur non-existence.

La preuve d'existence pure

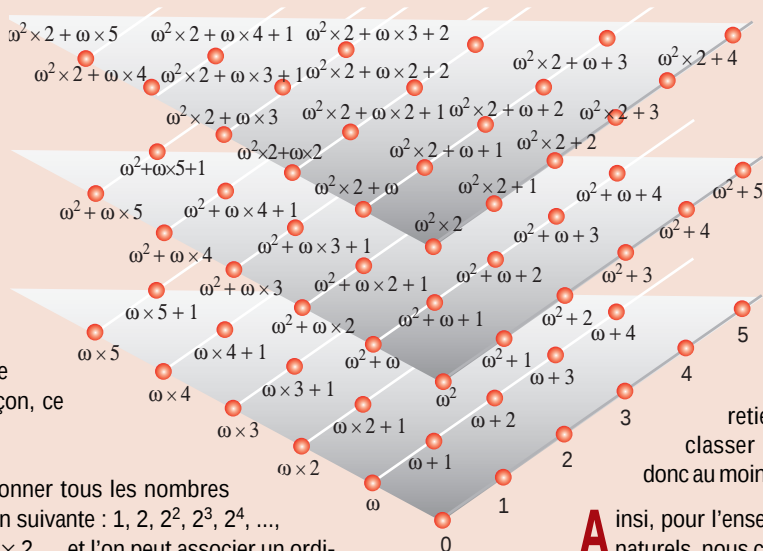
Le spectre de l'infini se dessine. Un raisonnement par l'absurde qui établit l'existence d'un objet dans un ensemble fini est accepté par tous les

après le nombre pour lequel cet exposant est le plus grand. Par exemple, pour les nombres 11 191 et 32 984, on cherche d'abord les décompositions $19^2 \times 31$ et $23 \times 7 \times 19 \times 31$. Le facteur 31 ayant le même exposant dans les deux décompositions, on compare les exposants de 19 : il est supérieur dans 11 191, de sorte qu'on classe ce nombre après 32 984. De la même façon, ce bon ordre place 2^{100} avant 5.

Ce bon ordre permet d'ordonner tous les nombres entiers non nuls de la façon suivante : 1, 2, 2^2 , 2^3 , 2^4 , ..., 3×2 , 3×2^2 , 3×2^3 , ..., 3^2 , $3^2 \times 2$, ... et l'on peut associer un ordinal à chaque nombre de cette suite : on associe l'ordinal 0 à 1, l'ordinal 1 à 2, l'ordinal 2 à 2^2 , l'ordinal 3 à 2^3 , ..., l'ordinal ω à 3 , l'ordinal $\omega + 1$ à 3×2 , l'ordinal $\omega \times 2$ à 3^2 , etc. Cette relation de bon ordre permet de représenter tous les nombres ordinaux jusqu'à ω^ω (non compris).

Les ordinaux sont à l'origine de la grande peur qu'ont ressentie les mathématiciens au début du siècle : le mathématicien italien Cesare Burali Forti (1861-1931) mit en évidence le paradoxe suivant : par définition, tout élément d'un ordinal est un ordinal, et la relation d'appartenance entre ordinaux est un bon ordre. L'ensemble des ordinaux est donc un ordinal x , et x est un élément de x (parce que x est formé de tous les ordinaux). Or le point x de l'ensemble x des ordinaux vérifie la proposition « x appartient à x », ce qui contredit l'irréflexivité de la relation d'appartenance sur x . Pour lever ce paradoxe, le logicien britannique Bertrand Russell (1872-1970) et d'autres mathématiciens imposèrent une définition *ad hoc* : une «classe» d'éléments définis par une propriété commune n'est un ensemble que dans certaines conditions, et, notamment, la classe de tous les ordinaux n'est pas un ensemble. De ce fait, on ne peut plus conclure « x appartient à x » dans la démonstration du paradoxe.

Nous avons vu que les bons ordres et les ordinaux sont deux manières de décrire la même réalité : à tout bon ordre correspond



un ordinal unique. Le mathématicien Ernst Zermelo (1871-1953) a montré (en utilisant l'axiome du choix, controversé) que, pour tout ensemble, on peut trouver une relation de bon ordre qui ordonne les éléments de l'ensemble. Si l'on oublie les noms des éléments et que l'on retient seulement la façon de classer ces éléments, on associe donc au moins un ordinal à un ensemble.

Ainsi, pour l'ensemble des nombres entiers naturels, nous connaissons au moins deux bons ordres : celui que nous utilisons quotidiennement pour classer les nombres (0, 1, 2, 3, ...) et celui que nous avons construit à partir de la décomposition en facteurs premiers (0, 1, 2, 2^2 , 2^3 , ..., 3×2 , 3×2^2 , ...). Le premier bon ordre est associé à l'ordinal ω , et le second à l'ordinal ω^ω .

Le plus petit des ordinaux associés à un ensemble est le «cardinal» (ou nombre d'éléments) de l'ensemble. Cette définition généralise aux ensembles infinis une notion familière dans le cas des ensembles finis. On montre facilement que les nombres entiers naturels sont des cardinaux, et l'on énumère les cardinaux infinis dans ce que l'on nomme la suite des alephs. Ainsi aleph zéro est le cardinal de l'ensemble des nombres entiers naturels ; c'est aussi le plus petit des cardinaux infinis, et il est confondu avec l'ordinal ω . Cantor montra d'autre part que le cardinal de l'ensemble des nombres réels est supérieur ou égal à aleph un (c'est ce que l'on nomme le paradoxe de Cantor). Il ruina sa santé à essayer de démontrer que le cardinal de l'ensemble des nombres réels était exactement égal à aleph un (hypothèse du continu). En fait, l'hypothèse du continu n'est pas décidable sur la seule base des axiomes de la théorie des ensembles.

Jean-Yves GIRARD
Institut de mathématiques (CNRS), Marseille-Luminy



Jan Brouwer (1881-1966)

mathématiciens, car on peut toujours (en principe) exhiber l'objet cherché en dressant la liste exhaustive des éléments de l'ensemble. Ce raisonnement ne s'applique pas aux nombres transcendants, eux-mêmes éléments de l'ensemble infini des nombres réels :

pour cette raison, de nombreux mathématiciens rejetèrent la démonstration de Cantor, arguant qu'une contradiction ne remplaçait pas un exemple tangible.

Cette controverse était loin d'être apaisée quand David Hilbert publia, en 1889, une preuve d'existence pure encore plus litigieuse. La démonstration concernait un problème posé par le mathématicien allemand Paul Gordon. Ce problème, l'une des questions mathématiques les plus brûlantes du temps, conjecturait l'existence d'ensembles (finis) d'objets présentant des propriétés particulières dans un certain nombre de cas. Hilbert, en raisonnant par l'absurde, montra que les ensembles en question existaient toujours. Cantor avait prouvé l'existence d'un ensemble d'objets dont on connaissait au moins quelques repré-

sentants : on avait découvert certains nombres transcendants quelque 30 ans auparavant. Hilbert soutenait à présent qu'il avait établi l'existence d'objets que nul n'avait jamais vus et que personne ne savait construire. «Ce ne sont plus des mathématiques, s'écria Gordon, c'est de la théologie!»

Néanmoins, la preuve d'existence pure faisait merveille dans les mains de Hilbert, et comme la théorie des ensembles commençait à apporter d'importants résultats mathématiques, les objections s'estompèrent. Gordon en vint même à reconnaître que «la théologie, elle aussi, avait ses mérites». Après la mort de Kronecker, en 1891, l'opposition aux méthodes non constructives diminua d'intensité, jusqu'à ce que le mathématicien hollandais Luitzen Brouwer lui insufflât un sang neuf, en 1907. Selon Brouwer, les mathématiques perdraient tout sens si on ne faisait pas une distinction nette entre l'existence réelle, ou constructive, et l'existence pure. Ces deux notions avaient été confondues dans le prolongement de la logique classique du fini à l'infini, notamment à propos de la loi du tiers exclu.

Le tiers exclu... exclu

La loi aristotélicienne du tiers exclu postule que toute proposition est soit vraie, soit fausse. Sans cette loi, il n'est évidemment pas de preuve par l'absurde. Par exemple, dans la démonstration de l'existence des nombres transcendants par Cantor, l'hypothèse que tout nombre est algébrique mène à une contradiction ; ainsi, selon le principe du tiers exclu, il doit exister des nombres transcendants. La logique aristotélicienne fait si bien partie de la vie quotidienne que la remettre en cause paraît extravagant. La majorité des gens sont persuadés du bien-fondé de cette loi : si un sac contient un nombre fini de pommes, personne ne doutera qu'il y a, soit une pomme rouge dans le sac, soit qu'il n'y en a pas. Personne n'a besoin de sortir les pommes du sac pour s'assurer que l'une de ces deux assertions est exacte (si l'on veut savoir laquelle est vraie, on répondra à la question après un temps fini).

La situation est différente quand on applique la loi du tiers exclu à un ensemble infini, par exemple un sac contenant une infinité de pommes, car on ne pourra jamais les examiner

L'INFINI ACTUEL ET LEIBNIZ

«Je suis tellement pour l'infini actuel qu'au lieu d'admettre que la nature l'abhorre je tiens qu'elle l'affecte partout, pour mieux marquer la perfection de son auteur. Ainsi, je crois qu'il n'y a aucune partie de la matière qui ne soit, je ne dis pas divisible, mais actuellement divisée, et, par conséquent, la moindre particule doit être considérée comme un monde plein d'une infinité de créatures différentes.»

Leibniz

toutes. Si l'on peut affirmer que le sac contient au moins une pomme rouge quand on en a tiré une, on ne peut pas conclure que le sac ne contient pas de pomme rouge tant qu'on n'a pas retiré toutes les pommes, ce qui est impossible puisqu'elles sont en nombre infini.

L'infini actuel n'étant pas présent dans la nature, les mathématiciens n'en ont donc pas d'idée intuitive. Pour appliquer la loi du tiers exclu à des ensembles infinis, ils doivent extrapoler à partir de leur connaissance des ensembles finis. Cette extrapolation semble d'autant plus délicate à justifier que les Grecs, à qui l'on doit cette loi du tiers exclu, avaient la plus grande méfiance envers les infinis actuels. Brouwer a même observé que les résultats de telles extrapolations pourraient ne pas être «constructivement» justes, car il serait possible de prouver l'existence d'objets «idéaux» impossibles à construire. Brouwer avançait qu'en donnant à de tels objets idéaux la même «réalité» qu'aux objets constructibles on portait atteinte à la certitude mathématique. Les objets idéaux en appellent d'autres et, bientôt, on ne sait plus si les mathématiques ont un fondement réel ou non.

Les critiques de Brouwer arrivèrent à un moment où la théorie des ensembles était vulnérable. Après 1900, la réorganisation ensembliste des mathématiques avait abouti à la formulation de paradoxes. Par exemple, comme, pour tout ensemble, on peut trouver un ensemble de cardinalité supérieure, Cantor avait indiqué qu'il devait exister un ensemble V de cardinalité supérieure à celle de l'ensemble universel U , qui contient tous les ensembles ; d'autre part, comme U contient tout, tous les éléments de V doivent appartenir à U , et la cardinalité de V doit être inférieure ou égale



3. L'UTILISATION DU TIERS EXCLU est délicate dans le cas d'ensembles infinis. Ainsi comment décider si un ensemble dénombrable de pommes contient ou non une pomme rouge sans les examiner toutes, ce qui est impossible? Les démonstrations par contradiction reposent sur la loi du tiers exclu, qui interdit que toute proposition soit autre que vraie ou fausse. Les constructivistes rejettent l'application d'une telle loi à des ensembles infinis, parce qu'il peut être impossible de déterminer laquelle des deux propositions est vraie. En médaillon, le père du constructionnisme, Luitzen Egbertus Jan Brouwer (1881-1966).

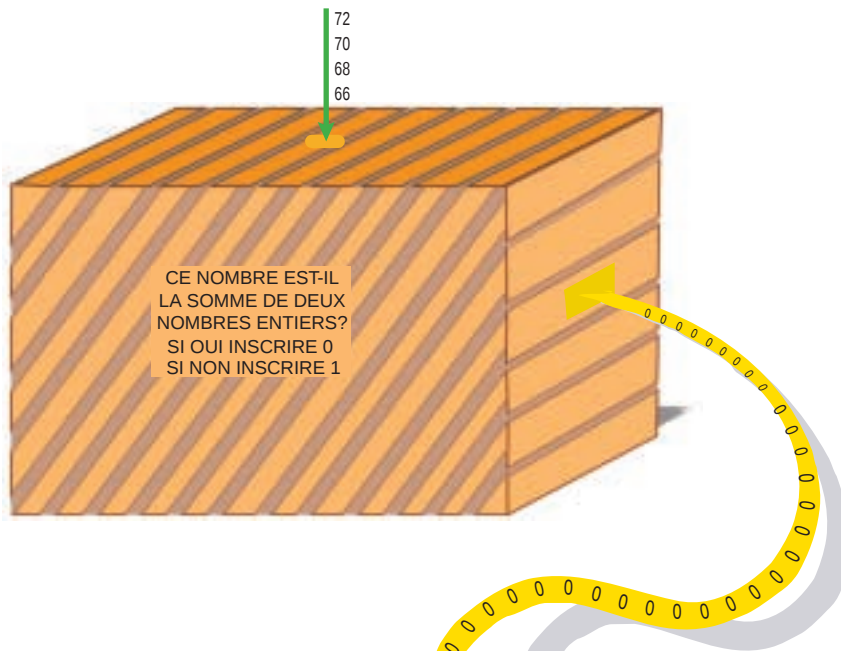
à celle de U . Des paradoxes de ce genre sapèrent les fondements des mathématiques ensemblistes et, selon Brouwer, les erreurs provenaient de l'introduction des objets idéaux en mathématiques.

Brouwer était convaincu que l'on résoudrait la crise des paradoxes si l'on rejetait objets et arguments idéaux, c'est-à-dire si l'on rebâtissait les mathématiques avec des méthodes exclusivement constructives. Comme ce programme, nommé intuitionnisme, aurait fait s'écrouler la presque totalité de l'édifice, il ne fut pas bien reçu, même de la part des mathématiciens ralliés à la cause de Brouwer. Ils étaient bien plus disposés à suivre le plan de Hilbert, qui ne cherchait pas à purifier les mathématiques du dedans, mais à les justifier du dehors. Ainsi les mathématiciens ne seraient pas, selon les propres mots de Hilbert, chassés du paradis que Cantor leur avait ouvert.

Hilbert croyait que, même si les objets idéaux n'avaient pas de signification propre, on pourrait en déduire des objets et des théorèmes importants et utiles. Selon lui, on pouvait espérer redonner un sens à la notion de certitude en transformant les mathématiques en un système axiomatique purement formel : une sorte de jeu dont l'intérêt résiderait non pas dans les objets qu'il manipule, mais dans les règles et dans les relations que ces dernières créent entre ces objets. Hilbert se hasardait même à dire que les mathématiques auraient un sens le jour où le système se révélerait cohérent, c'est-à-dire dépourvu de contradiction.

La théorie de la preuve

Le programme de Hilbert - le formalisme - avait pour premier objectif la représentation des mathématiques sous forme d'un système axiomatique formel. Il espérait qu'en transformant ainsi les données mathématiques en chaînes de symboles abstraits, combinées selon les lois de la logique, tout raisonnement porteur de germe de paradoxe serait démasqué. Mieux encore, une fois les mathématiques formalisées, on pourrait entreprendre de démontrer leur cohérence. À cette fin, Hilbert introduisit une nouvelle théorie, la théorie de la preuve, où il était possible de formuler des énoncés signifiants à partir des signes et



4. LA CONJECTURE DE GOLDBACH énonce que tout entier naturel pair est la somme de deux nombres premiers. Bien que cette proposition reste non démontrée, il existe un procédé fiable pour savoir si un nombre pair donné est la somme de deux nombres premiers. On peut donc programmer un ordinateur pour qu'il examine chaque nombre pair séparément et qu'il inscrive un 0 si ce nombre est la somme de deux nombres premiers, et un 1 dans le cas contraire. L'ensemble des nombres pairs étant infini, la chaîne des 0 (avec peut-être des 1) l'est aussi. Est-il correct d'appliquer la loi aristotélicienne du tiers exclu à cette chaîne infinie, et de dire qu'elle renferme un 1 ou qu'elle n'en renferme pas? Le problème est qu'on ne saura peut-être jamais, faute de démonstration, laquelle de deux hypothèses est la bonne.

des configurations dénués de sens du système axiomatique. Cette métamathématique devait être entièrement constructive, voire «finitiste», par essence. Il devenait, par conséquent, impensable de mettre en doute la preuve métamathématique de la cohérence des mathématiques.

Brouwer refusa de croire que montrer la cohérence des mathématiques suffirait à leur redonner un sens et passa la majeure partie de sa vie à essayer de montrer que l'on pouvait «construire» les mathématiques. Son programme ne

fut guère populaire : les tentatives intuitionnistes étaient lourdes, artificielles et difficiles à utiliser. Quand Brouwer entreprit de débarrasser les mathématiques des notions idéales, tout en conservant certains aspects éprouvés de longue date, il fit appel à des notions qui parurent tout aussi idéalistes que celles qu'il éliminait.

Au total, ce n'est pas à cause d'une foi intense dans les mathématiques que les mathématiciens rejetèrent l'intuitionnisme, mais plutôt à cause de l'absence d'une telle foi. Beaucoup



5. L'AXIOME DU CHOIX, l'un des axiomes le plus utilisés et le plus controversés de la théorie des ensembles, admet que l'on peut toujours, partant d'une collection quelconque d'ensembles non vides, choisir un représentant dans chaque ensemble. Autrement dit, pour chacune des deux collections représentées ci-dessus, il devrait être possible d'indiquer une règle pour choisir des éléments, de sorte que deux personnes différentes opèrent les mêmes choix. Par exemple, pour la collection infinie de paires de chaussures, la loi «Choisissez la chaussure droite de chaque paire» est suffisante, car les deux chaussures sont dissemblables. Il n'existe pas de telle loi pour une collection infinie de paires de chaussettes toutes semblables.

6. LA DÉFINITION DES NOMBRES RÉELS CONSTRUCTIFS.

La question est : pourquoi ne définit-on pas les nombres constructifs à partir des développements décimaux ?

L'idée la plus naturelle pour définir un nombre réel constructif consiste à dire que c'est un nombre réel pour lequel on connaît un moyen numérique (c'est-à-dire algorithmique) d'en calculer les décimales aussi loin qu'on le veut. Appelons un tel nombre «décimalement constructif».

En voici deux exemples x et y :

x est égal à 1, défini comme le nombre dont la suite des décimales est 1,0000 ...

y est égal à 0 si tout nombre pair est somme de deux nombres premiers (conjecture de Goldbach, non démontrée aujourd'hui) et est égal à $1/10^n$ sinon, en prenant pour n le plus petit entier tel que $2n$ ne s'écrit pas comme somme de deux nombres premiers.

Le nombre y est «décimalement constructif», car, si vous voulez connaître ses dix premières décimales, il suffit de vérifier que 2, 4, 6, 8, 10, 12, 14, 16, 18, 20 s'écrivent chacun comme somme de deux nombres premiers (cinq minutes suffisent pour s'en assurer). Donc x est inférieur à $1/10^{10}$ et s'écrit 0,0000000000... : on connaît avec certitude les dix premiers chiffres décimaux de y .

Considérons maintenant $z = x - y$. Puis-je connaître ses 10 premières décimales ? Non. Sans savoir si la conjecture de Goldbach est vraie ou fausse, on ne peut même pas connaître la première décimale de z . En effet, si la conjecture de Goldbach est vraie, alors $z = 1$, alors que, si elle est fausse, $z = 0,99999...99990000...$, les 9 se prolongeant jusqu'au chiffre n correspondant à un contre-exemple, $2n$, de la conjecture.

Cela montre que la différence de deux nombres «décimalement constructifs» n'est pas nécessairement un nombre «décimalement constructif». C'est pourquoi on adopte, pour définir les nombres constructifs, une méthode ne se référant pas aux chiffres décimaux. Cette définition consiste à dire que x est un nombre réel constructif si l'on connaît un moyen numérique de l'approcher d'aussi près qu'on le veut par des nombres rationnels. Avec cette définition, les nombres x et y proposés ci-dessus sont des nombres réels constructifs, ainsi que $z = x - y$ (en vérifiant la conjecture de Goldbach jusqu'à 20, on en déduit que $0 \leq 1 - z < 1/10^{10}$). Plus généralement, avec la définition adoptée, la somme de deux nombres constructifs l'est, ainsi que leur différence ou leur produit, etc.

Il est amusant de noter qu'Alan Turing, dans son article fondateur de l'informatique théorique de 1936, introduisait la notion de nombre réel calculable en se référant aux décimales (comme nous l'avons fait pour les nombres «décimalement constructifs»). Quelques mois après, prenant conscience de ce que sa notion était malcommode (la différence de deux nombres calculables ne l'est pas nécessairement), il suggérait une modification de sa définition.

étaient inquiets de voir leur discipline perdre son sens véritable ; la seule manière de garder le bateau à flots leur parut être d'accepter certaines notions idéales. Citons encore Hilbert : «À côté de l'immense essor des mathématiques modernes, quel sens auraient ces pitoyables vestiges, maigres résultats isolés, incomplets et sans lien, obtenus par les intuitionnistes ? » Le formalisme de Hilbert offrait à beaucoup un compromis attirant. Si l'on pouvait montrer que les mathématiques ne contenaient aucune contradiction, on serait sûr que, si deux personnes (ou la même en deux occasions différentes) jouaient aux mathématiques, elles n'aboutiraient pas à deux réponses exactes opposées.

En 1931, Kurt Gödel montra que le programme formaliste de Hilbert était

voué à l'échec. Il prouva qu'aucune métamathématique ne permet de démontrer qu'un système formel d'axiomes est cohérent si ce système est suffisamment riche pour inclure une arithmétique élémentaire ; dans le meilleur des cas, on peut déterminer quand le système est incohérent. Entre-temps, la théorie des ensembles et le formalisme s'étaient tant imposés en mathématiques qu'il n'était plus question de les rejeter sous prétexte que la quête fondamentale du formalisme était une cause perdue.

Aussi le formalisme est-il devenu l'un des moteurs de la pensée mathématique, et l'intuitionnisme, une curiosité mathématique désuète. Les mathématiciens optent pour l'élégance des démonstrations d'existence pure (même quand ils ont à leur disposition des démonstrations constructives des

mêmes résultats), et le bien-fondé de telles mathématiques est rarement remis en cause. Quelques querelles subsistent au sujet de la théorie des ensembles (notamment l'emploi de l'axiome du choix, dont il sera question plus loin), mais la controverse sur les méthodes de la théorie ensembliste de Cantor et Hilbert est terminée.

Le constructivisme revisité

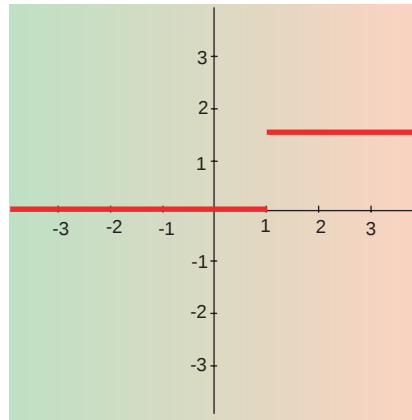
Le mathématicien Eric Bishop, en publiant en 1967 son livre, *Les fondements des mathématiques constructives*, raviva l'intérêt pour les mathématiques constructives. Contrairement à ses prédécesseurs, Bishop s'était appliqué à montrer que les résultats obtenus par des méthodes constructives peuvent être aussi élégants et puissants que ceux qui sont obtenus par des méthodes formalistes. Alors que Brouwer voulait construire les certitudes en mathématiques en partant d'un système d'axiomes intuitifs pour l'arithmétique et travailler en remontant, Bishop aborde la «constructivisation» des mathématiques selon une optique toute différente. Son but n'est pas de refonder les mathématiques, mais de les pratiquer de façon constructive ; ainsi son point de départ est l'ensemble des nombres rationnels, muni des opérations classiques, et il progresse de façon «réaliste» à partir de là. Selon Bishop, on peut atteindre la certitude en mathématique en rejetant toute notion «idéaliste» et en vérifiant en permanence que les objets ou théorèmes obtenus conservent un sens intuitif. Un grand nombre de résultats importants peuvent être obtenus par cette méthode directe.

Dans les mathématiques constructives de Bishop, la définition des nombres réels ou celle des fonctions doivent être repensées radicalement. Comme dans les mathématiques non constructives, un nombre réel est défini s'il peut être approché par des nombres rationnels arbitrairement proches. D'un point de vue constructif, on doit, pour que la définition de notre nombre réel soit considérée comme valide, donner un moyen effectif (fini) pour l'approcher par des nombres rationnels à n'importe quel degré de précision voulu. Autrement dit, il doit être possible, après un nombre fini d'étapes, d'obtenir un nombre rationnel qui approche le nombre réel en question avec la précision demandée. Cela n'est pas tou-

jours aussi facile qu'il peut le paraître : par exemple, se référer simplement aux développements décimaux poserait des problèmes insurmontables.

Une autre conséquence de la définition constructive que donne Bishop d'un ensemble est surprenante: il est impossible de décomposer l'ensemble des nombres réels en deux ensembles non vides n'ayant aucun élément en commun. Cette restriction sur le recouvrement de l'ensemble des nombres réels se répercute sur le type de fonctions à valeurs réelles que l'on peut définir constructivement. La construction proposée par Bishop repose sur l'utilisation des fonctions en escalier (voir la figure 7).

Les critères d'acceptabilité des mathématiques constructives étant plus rigoureux que ceux des mathématiques non constructives, tout théorème vérifié constructivement l'est aussi non constructivement. Néanmoins, deux théorèmes distincts constructivement peuvent avoir la même interprétation non constructive. D'un autre côté, on peut transformer n'importe quel théorème non constructif en un théorème constructif, à la condition d'ajouter «en supposant le principe du tiers exclu». De plus, des démonstrations non constructives peuvent être réinterprétées constructivement. Bishop affirmait qu'une fois les enseignements positifs de son programme tirés, les mathématiques actuelles, telles qu'elles sont pratiquées, sortiraient de leur isolement pour devenir une branche des mathématiques constructives. Néanmoins, la plupart des mathématiciens n'ont pas encore trouvé de raison suffisante pour changer radicalement la conception de leur recherche. Les mathématiques ont tant progressé que personne ne peut affirmer que le manque de réalité que leur reprochent les constructivistes soit bien fondé. Après tout, l'utilité des mathématiques dans le monde réel n'est plus à démontrer. Cette «bataille entre chats et chiens mathématiciens», comme disait Albert Einstein, apparaît dénuée de tout intérêt à nombre de scientifiques. En outre, les problèmes majeurs qui font évoluer les mathématiques sont d'ordre presque purement intuitif, interdisant ainsi de trop grands écarts par rapport à la réalité : la conviction que les mathématiques ne sont pas un jeu dénué de sens semble donc se confirmer.



7. LES FONCTIONS EN ESCALIER sont des fonctions très souvent appliquées aux nombres réels, mais qui ne peuvent être définies constructivement. En mathématiques constructives, une fonction sur les nombres réels est une loi qui associe à tout nombre réel r un autre nombre réel $f(r)$. Par exemple, la loi qui définit la fonction en escalier ci-dessus est : si r est inférieur à 1, $f(r)$ vaut 0, et si r est supérieur ou égal à 1, $f(r)$ vaut 1. Dans certains cas, il est possible de définir un nombre réel, sans pour autant pouvoir dire, au bout d'un nombre fini d'étapes, s'il est inférieur, supérieur ou égal à 1. Par conséquent, il est impossible d'attribuer la valeur 0 ou 1 à l'image de ce nombre.

Même si la révolution constructive n'est pas vraiment une révolution, il n'y a pas de doute possible sur le rôle déterminant que joueront à l'avenir les mathématiques constructives dans l'évolution de la discipline. Le récent engouement pour la «constructivisation» de certains résultats montre que, dans certains cas, il y a de gros avantages à prendre les problèmes sous cet angle nouveau. On y gagne non seulement une meilleure analyse, mais aussi des théorèmes plus puissants : ainsi, lorsque l'existence constructive d'un nombre a été mise en évidence, il peut être calculé, en pratique comme en théorie.

La majorité des mathématiciens d'aujourd'hui, formés au moule du formalisme depuis plusieurs générations, ne doutent du bien-fondé des mathématiques que lorsque l'on met à jour d'insurmontables contradictions inhérentes au système axiomatique formel des mathématiques. Même si les mathématiciens pensent que de telles contradictions peuvent éventuellement apparaître, ils restent certains de pouvoir remettre les choses au point par quelques modifications mineures des axiomes. Dans ces conditions, la question de savoir si les mathématiques abstraites ont un sens réel ou non ne peut même pas être formulée.

Aux mathématiciens qui en viendraient à douter de la validité de leur science, on répète le conseil que le philosophe français Jean d'Alembert, plagiant Pascal, donnait à ceux qui doutaient du calcul infinitésimal : «Continuez, et la foi viendra.»

Toutefois, Jean-Paul Delahaye a remarqué qu'aujourd'hui la physique utilise sans aucune retenue les mathématiques non constructives et en particulier les nombres réels qui sont sans réalité physique : en physique, aucune mesure n'est infiniment précise!

Les mathématiques du continu sont donc un instrument formel mal ajusté aux besoins véritables du physicien. Les mathématiques constructives lui conviendraient mieux, puisqu'elles ne manipulent que des objets finiment définis et que l'infini y est, en quelque sorte, toujours potentiel, comme dans le monde physique. Il est donc à prévoir que, dans un avenir plus ou moins lointain, les physiciens en viendront à utiliser les mathématiques constructives (ou certaines de leurs variantes) pour éviter d'avoir à utiliser les formalismes surdéterminants (qui en disent plus qu'on ne le souhaiterait) des mathématiques usuelles non constructives.

Préoccupés par bien d'autres choses, et effrayés par la difficulté des mathématiques constructives, très peu de physiciens ont étudié à ce jour cette possibilité d'affinement de leurs outils théoriques.

Allan CALDER est professeur à l'Université d'État du Nouveau-Mexique.

E. BISHOP, *The Foundations of Constructive Analysis*, McGraw-Hill, 1967.

E. BISHOP et D. BRIDGES, *Constructive Analysis*, Springer-Verlag, 1985

A.S. TROELSTRA et D. VAN DALEN, *Constructivism in Mathematics: an Introduction*, North-Holland, 1988.

M.J. BEESON, *Foundations of Constructive Mathematics: Metamathematical Studies*, Springer-Verlag Berlin/Heidelberg/New York, 1985.

M.B. POUR-EL et J.I. RICHARDS, *Computability in Analysis and Physics*, Springer-Verlag, 1989.

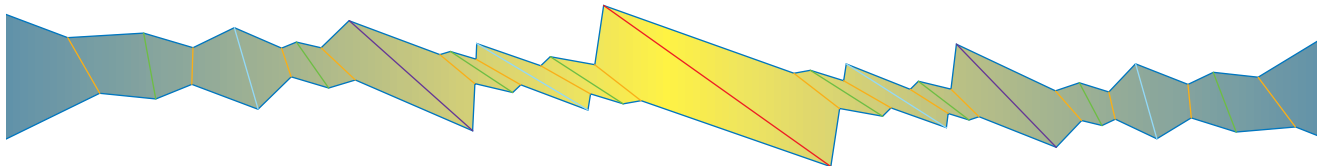
John MYHILL, *What is a Real Number?*, in *The American Mathematical Monthly*, vol. 79, n° 7, pp. 748-754.

I. LAKATOS, *A Renaissance of Empiricism in the Recent Philosophy of Mathematics*, in *Mathematics, Science and Epistemology*, Cambridge University Press, 1978.

L'infini en géométrie

MARCEL BERGER

Lors d'un dîner, un géomètre recourt aux éléments du décor pour illustrer le rôle important de l'infini en géométrie : l'analyste et l'algébriste seront convaincus...



Au sortir de la conférence, l'algébriste, l'analyste et le géomètre, dignes représentants de la trinité des mathématiques, méditent la phrase d'Hermann Weyl : «Les mathématiques sont la science de l'infini.» L'algébriste, sûr de lui, dit qu'en algèbre l'infini est omniprésent. L'analyste répond qu'il en est de même dans sa discipline. Puis tous deux se tournent, perplexes, vers le géomètre, silencieux. Observé avec insistance, celui-ci affirme également : «Oh ! c'est également vrai en géométrie.» Devant les regards interrogateurs de ses deux amis, il pour-

suit : «Bien sûr, la géométrie semble être bien installée à distance finie, les cas d'égalité des triangles ou la longueur d'un cercle égale à 2π fois le rayon en sont une illustration... Où donc se loge l'infini en géométrie ? Partout ! Puis-je vous proposer de vous le montrer devant un bon dîner ?» Les deux autres, intrigués, acceptent. Sur le chemin, le géomètre énumère en silence les différents domaines de sa discipline dans lesquels l'infini joue un rôle important. Dans chacun d'eux, l'infini est présent sous trois aspects : lorsque l'on considère des objets en nombre infini ; quand on s'intéresse

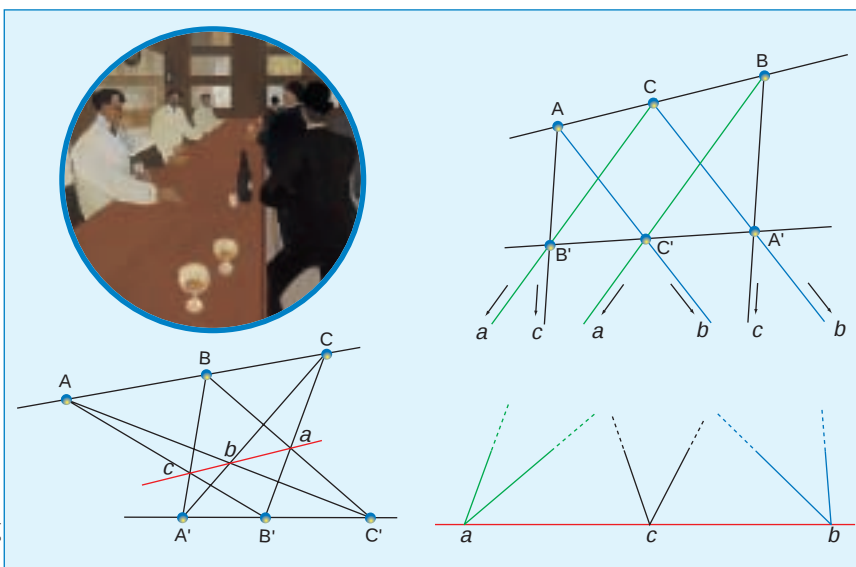
à des espaces de dimension infinie ; enfin, quand une opération est répétée un nombre infini de fois.

Nos trois mathématiciens franchissent le seuil de l'hôtel des Trois Faisans et patientent, alignés dans le corridor, en attendant l'hôtesse des lieux. Là, au bar, trois notaires passent le temps. Le géomètre ne peut manquer une telle occasion. «En guise d'introduction, vous voyez, messieurs, nous trois alignés ici, et ces trois hommes là-bas au bar, m'évoquent le théorème de Pappus (290 - 350), dont une démonstration aisée recourt à l'infini.»

L'infini selon Pappus

L'algébriste fouille sa mémoire : «Ce théorème n'affirme-t-il pas que, pour deux triplets de points alignés (A, B et C) et (A', B' et C') sur deux droites distinctes, les points d'intersection a de (BC') et $(B'C)$, b de (AC') et $(A'C)$ et c de (AB') et $(A'B)$ sont alignés (voir la figure 1) ?»

Le géomètre confirme : «Exact, et la démonstration requiert de voir ce résultat dans un plan projectif, c'est-à-dire un plan affine (dans lequel les distances ne sont pas prises en compte), auquel on a adjoint une droite, dite à l'infini. Les points de cette droite sont les directions de toutes les droites du plan (voir *De la perspective à l'infini géométrique*, pages 66 à 72). Il n'y a plus de droites parallèles : celles qui l'étaient dans le plan seul sont maintenant sécantes en un point de la droite de l'infini. On décide, pour démontrer Pappus, que cette droite est la droite (ab) , ce que l'on peut toujours faire par une projection convenable. Ainsi, dans sa nouvelle



1. LE THÉORÈME DE PAPPUS affirme que les points d'intersection (a, b, c) des droites joignant convenablement les points de deux triplets alignés, (A, B, C) et (A', B', C') sont eux aussi alignés (à gauche). Ce théorème se démontre facilement avec l'infini : dans un plan dit projectif, on décide que les points de la droite (ab) , dite droite de l'infini, sont les directions de toutes les droites du plan ; deux droites parallèles sont alors sécantes en un point de cette droite. Puisque a et b sont sur cette droite, les droites (CB') et $(C'B)$ d'une part, et (AC') et $(A'C)$ d'autre part, sont parallèles. Démontrer le théorème de Pappus revient à montrer que le point c est aussi sur la droite de l'infini, donc que les droites (AB') et $(A'B)$ sont parallèles (à droite) ; le théorème de Thalès le montre aisément.

configuration, les droites (AC') et $(A'C)$ d'une part, et (BC') et $(B'C)$ d'autre part sont parallèles. Montrer que les points a , b et c sont alignés revient à montrer que le point c est aussi sur la droite de l'infini, c'est-à-dire que les droites de la troisième paire, (AB') et $(A'B)$, sont parallèles : le théorème de Thalès le montre aisément».

Après une petite réflexion, l'algébriste dit : «Oui, évidemment, cette démonstration est beaucoup plus simple que de se placer dans un repère et de calculer les coordonnées de tous les points.»

Enfin, tout trois assis devant leur table, la rondeur des assiettes et la circularité des verres inspirent le géomètre.

Des surfaces et des volumes

«Vous connaissez tous les deux cette méthode qui requiert l'infini pour calculer la longueur d'un cercle. On y inscrit un carré, puis, bâti sur ce carré, on construit un octogone régulier, dont on divise ensuite encore en deux chacun des huit côtés, afin d'avoir un polygone régulier à 16 côtés, etc. On calcule les longueurs successives de ces polygones et, à la limite, quand le nombre de côtés tend vers l'infini, on obtient la longueur du cercle. Grâce à une telle itération, on calcule aussi la longueur d'une courbe quelconque : on y inscrit des

polygones de côtés de plus en plus petits, et la longueur de la courbe est la limite à l'infini, quand elle existe, de la suite des longueurs de ces polygones.»

Je crois pourtant me souvenir, interrompt l'algébriste, que l'on sait fabriquer des courbes de longueur infinie, même entre deux points quelconques de cette courbe.

Le géomètre reprend : «Oui, en revanche, on sait qu'il y a toujours une longueur lorsque la courbe est donnée par une équation suffisamment douce, c'est-à-dire que la courbe n'a pas de pic, ou, en termes plus mathématiques, qu'elle est dérivable en chaque point.

Dans un plan, pour calculer l'aire d'un triangle, on ajoute un second triangle égal, donc de même aire : du parallélogramme obtenu, on ôte d'un côté, et on ajoute de l'autre un triangle égal afin de former un rectangle dont l'aire est égale au produit de la longueur par la largeur. Ainsi, l'aire d'un triangle est égale au demi produit de la base par la hauteur.

Toutefois, le calcul du volume d'un polyèdre est plus subtil. Dans l'espace, on peut toujours décomposer un parallélépipède en tétraèdres ; cependant, on ne sait pas de façon évidente s'ils ont même volume, même s'ils ont la même hauteur et la même aire de base. Ils ne sont pas déduits l'un de l'autre par un déplacement. Pour montrer

qu'ils ont même volume, on recourt à l'infini : on découpe les tétraèdres en tranches horizontales de plus en plus fines ; chaque tranche contient un prisme et un morceau supplémentaire. Dans les deux tétraèdres, les prismes de deux tranches de même niveau ont le même volume, puisque l'aire de leur base est identique (voir la figure 5). Lorsque le nombre de prismes tend vers l'infini, la somme des volumes des morceaux restants devient négligeable : les tétraèdres ont donc même volume, car ils sont tous construits d'un nombre infini de prismes de même volume.

Ainsi, on ne procède pas dans l'espace de la même façon que dans le plan, par des raisonnements de découpage en un nombre fini de morceaux de volume égaux déduits les uns des autres par des transformations de l'espace. Inutile de rêver qu'un jour un génie trouvera un tel découpage, en 1903, Max Dehn (1878 - 1952) a démontré que ce rêve est impossible. Le cas du plan reste donc exceptionnel, l'espace nécessite l'infini.»

Et les cercles ?

La serveuse apporte alors les apéritifs et intervertit deux verres. Encouragé par l'abondance des cercles des verres des assiettes, des soucoupes... disposées sur la table, le géomètre saisit l'occasion

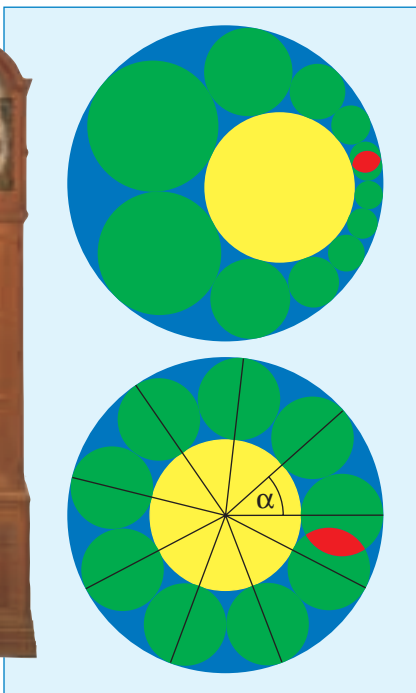


Musée des Beaux-Arts de Nantes

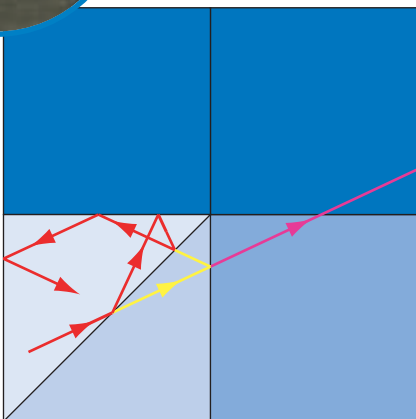
2. À L'HÔTEL DES TROIS FAISANS, on passe le temps, mais on parle aussi de l'infini en géométrie.



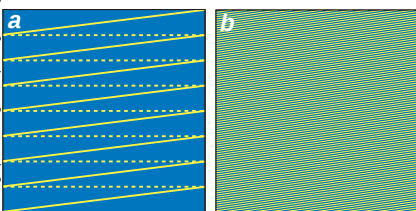
Cheltenham Art Gallery



3. L'ÉTUDE DES CERCLES (en vert) construits tangents à deux autres donnés (en haut) dont l'un est inclus dans l'autre (en bleu et en jaune) est évidente après une inversion : cette transformation géométrique rend concentriques les deux premiers cercles. Ainsi, les cercles construits sont égaux et il suffit d'étudier les rotations d'un angle α .



Photographie - Wingfield Sporting Gallery, Londres



4. LE COMPORTEMENT de deux particules qui rebondissent l'une contre l'autre sur un segment (à gauche, en haut) ou les rotations d'un angle dans un cercle (à droite) sont équivalents à des itérations d'une même opération. Dans le premier cas, la position des deux particules est représentée par un point dans un repère (un axe par particule). Puisqu'une particule ne peut pas passer par-dessus l'autre, la trajectoire de ce point est confinée dans un

et poursuit son exposé : «Avec la démonstration du théorème de Pappus, je vous ai montré que l'adjonction d'une droite à l'infini est parfaitement adaptée pour des problèmes de la géométrie des points et des droites. En revanche, elle n'est pas appropriée quand on s'intéresse à une géométrie de l'ensemble des cercles du plan euclidien, c'est-à-dire que l'on prend en compte la distance. La transformation du plan nommée inversion se révèle parfois pertinente. Elle consiste à se donner un point du plan, le centre d'inversion O . Tout point m autre que le centre est envoyé sur le point m' aligné avec O et m , et tel que le produit des longueurs $Om.Om'$ est constant : cette inversion transforme les cercles en cercles, sauf ceux passant par le centre O , qui sont transformés en droites.

À nouveau des itérations

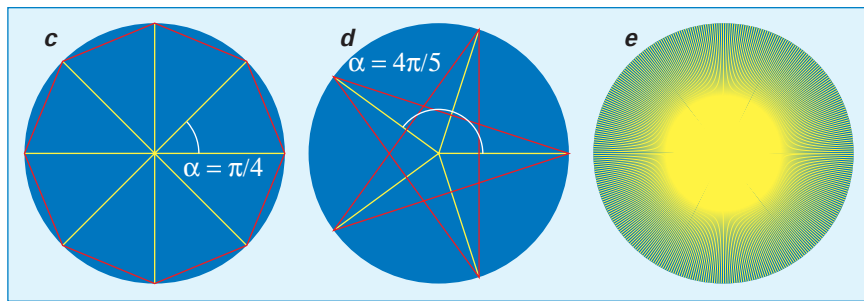
Grâce à cette inversion, l'étude de problèmes sur les cercles devient plus facile : par exemple, Jacob Steiner (1796-1863) a étudié le comportement de cercles tangents à deux autres donnés, le premier étant inclus dans le second (voir la figure 3). Un cercle est construit tangent aux deux premiers, le suivant est tangent aux trois précédents, et ainsi jusqu'à avoir parcouru le périmètre du cercle intérieur : le dernier cercle construit est-il tangent au premier construit, ou le chevauche-t-il ? Une judicieuse inversion rend concentriques les deux cercles donnés au départ. Les cercles ajoutés sont donc égaux, et il s'agit maintenant d'étudier le comportement de la droite joignant le centre des cercles concentriques (voir la figure 3) à celui de ceux construits : à chaque nouveau cercle, la droite pivote d'un angle constant.»

Le son d'un carillon ponctue le discours du géomètre, dont le regard est attiré par les aiguilles de l'horloge : «De la même façon que l'aiguille de cette horloge tourne à chaque minute...

...nous pouvons poursuivre avec l'étude de telles itérations. Quel est l'effet sur un point du cercle de l'itération d'une rotation d'angle donné α ? Quand $\alpha = 2\pi/k$ (k est un entier), les images successives d'un point de départ donné constituent les sommets d'un polygone régulier à k côtés inscrit dans le cercle, c'est le cas de la pendule pour laquelle $k = 60$; lorsque $\alpha = 2\pi p/q$, où p et q sont deux entiers premiers entre eux, le polygone obtenu est étoilé, à q côtés et tournant p fois (voir la figure 4). Mais que se passe-t-il quand k est un nombre irrationnel ? La trajectoire du point ne se refermera jamais, quel que soit le nombre d'itérations, même infini.

Toutefois, on peut en savoir plus. On montre que la trajectoire est partout et également dense dans le cercle : il y a une égale répartition des points «visités», aucune zone du cercle n'est négligée. En termes mathématiques, pour un ensemble infini $\{C_i\}_{i=1,2,3,\dots}$ de points du cercle, on dit qu'il y a égale répartition sur un arc m de longueur L du cercle lorsque, k tendant vers l'infini, la limite du nombre de points $C_{i(i=1,2,3,\dots,k)}$ qui sont sur l'arc m est toujours égale à $L/2\pi$ (si le cercle est de rayon unité), cela quel que soit l'arc m .)

L'analyste coupe ici le géomètre et lui demande, entre deux bouchées : «C'est bien joli, les itérations, mais elles restent des opérations discrètes. Or, dans la nature, un grand nombre de systèmes sont continus dans le temps. Les géomètres étudient-ils aussi ces phénomènes lorsque le temps continue très longtemps, c'est-à-dire tend vers l'infini ?»



triangle rectangle isocèle. On développe ce triangle dans un carré, puis dans le plan entier : la trajectoire est alors une droite infinie que l'on ramène dans un seul carré. Dans le second, on s'intéresse à la trajectoire d'un rayon du cercle. Certaines trajectoires sont périodiques (a, c et d), elles repassent par leur point de départ, d'autres, apériodiques (b et e), approchent d'autant plus que l'on veut toutes les positions possibles.

«Bien sûr! Le billard que j'aperçois là-bas m'inspire. En mécanique, on étudie le mouvement de deux particules de même masse sur un intervalle. Ces particules repartent en sens inverse avec la même vitesse quand elles se rencontrent. Quand on représente la position des deux particules par un point du plan, (les deux axes représentent chacun les positions d'une particule), la trajectoire de celui-ci, confinée dans un triangle rectangle isocèle, est celle d'une boule de billard qui se réfléchit sur les côtés avec un angle d'incidence égal à l'angle de réflexion par rapport à la normale, la droite perpendiculaire au côté. Nous simplifions le problème en développant la trajectoire par symétrie : le billard devient alors un carré (voir la figure 4). Puis, quand on développe la symétrie dans le plan tout entier, la trajectoire devient une droite infinie.

Vous remarquerez qu'il suffit de savoir ce qui se passe là où la trajectoire rencontre un bord du carré : le problème est alors de nouveau celui de l'itération d'un même procédé. Ici donc, deux cas se distinguent selon que la pente, qui est le rapport des vitesses initiales des deux particules, est en rapport rationnel ou irrationnel avec la longueur du carré (la longueur de l'intervalle).

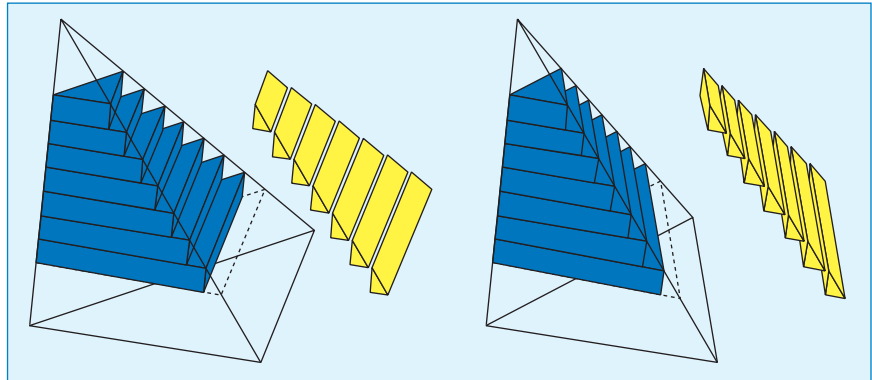
Dans le cas irrationnel, la trajectoire est partout dense et également répartie, c'est-à-dire que les deux particules occuperont toutes les positions possibles, et ce de façon équitablement répartie dans le temps. Sur un ordinateur, la trajectoire noircira l'écran de plus en plus, et de façon partout identique. Seule varie la vitesse de remplissage de l'écran. Dans le cas rationnel, la trajectoire est périodique ; après un temps, le mouvement reprend comme au départ : l'écran est seulement biffé de quelques lignes.»

L'algébriste et l'analyste, songeurs, demandent alors : «Que se passe-t-il pour deux particules de masses différentes?»

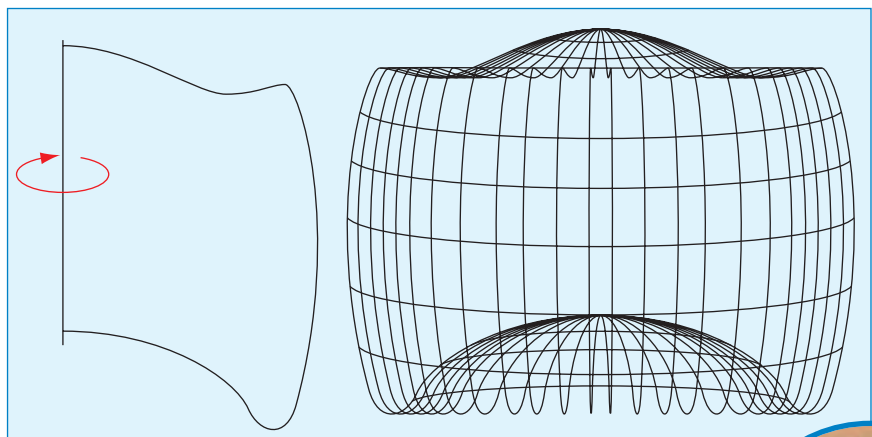
La réponse est immédiate : «Un calcul du choc élastique de particules de masses différentes, en tenant compte de la conservation de l'énergie et de la quantité de mouvement, montre que le problème devient alors celui du billard dans un triangle rectangle, mais non isocèle. Ici, les géomètres n'ont que peu de résultats, car ils ne peuvent plus développer le triangle en un pavage de carrés. Le résultat dépend essentiellement du rapport des masses ; plus précisément, il s'agit de savoir quand l'angle α (un des angles du triangle), dont la tangente

est égale au produit des racines carrées des deux masses, est en rapport rationnel ou non avec π . Dans le cas rationnel, les trajectoires sont de trois types : dans des cas rares, elles sont périodiques ;

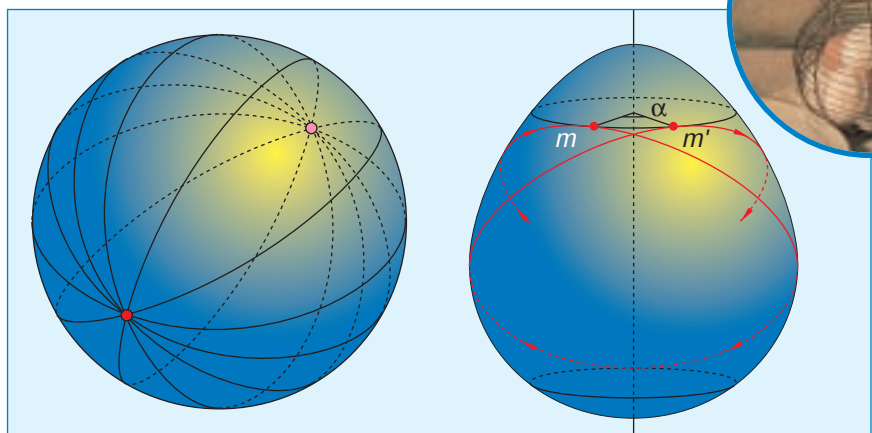
d'autres sont partout denses et bien réparties, c'est-à-dire qu'une zone du billard ne sera pas remplie avant une autre ; enfin, le plus souvent, elles sont partout denses, mais mal réparties. Dans



5. EN DÉCOUPANT DEUX TÉTRAÈDRES de même hauteur et de même aire de base en tranches, on montre qu'ils ont le même volume. Les tranches sont composées d'un prisme (en bleu) et d'un bout restant (en jaune). Dans chaque tétraèdre, les prismes du même niveau, ont un volume égal. Lorsque le nombre de prismes est infini, le volume des bouts restants est égal à zéro : les tétraèdres ont donc même volume.



6. LES SURFACES DE ZOLL sont des surfaces de révolution dont, à l'instar de la sphère, toutes les trajectoires d'une particule lancée à la surface sont des géodésiques périodiques.



7. SUR UNE SPHÈRE (à gauche), toutes les trajectoires d'une particule lancée sur la surface sans frottements et sans gravité sont des géodésiques périodiques : ce sont les grands cercles. Sur une surface de révolution (à droite), telle celle d'un œuf, les trajectoires (en rouge) sont des géodésiques périodiques pourvu qu'à partir d'un point m d'un parallèle quelconque (en noir), elles repassent sur ce même parallèle par un point m' décalé de m d'un angle α en rapport rationnel avec 2π .

tous les cas, les directions prises par une trajectoire en chaque point du billard sont en nombre fini : le système n'est donc pas ergodique en phase. Cependant, il est ergodique en position, car la trajectoire passe par tous les points. Par ailleurs, on sait que l'ordre de grandeur du nombre de trajectoires périodiques de longueur inférieure ou égale à une distance donnée croît avec le carré de la longueur de l'intervalle.»

L'analyste coupe le géomètre : «Et dans le cas irrationnel?»

«J'y venais. À partir de l'étude du cas rationnel et en approximant les triangles à angle irrationnel par de nombreux triangles à angle rationnel, on montre que l'ergodicité de phase est obtenue à l'aide de la seule ergodicité d'espace du cas rationnel. Cependant, c'est une existence complètement abstraite ! On ne connaît aucune valeur de m et de m' qui permette de vérifier cette assertion. D'autre part, on ne connaît pas l'ordre de grandeur du nombre de trajectoires périodiques de longueur inférieure ou égale à une distance donnée.»

rieure ou égale à une distance donnée.»

Le géomètre, intarissable, se remémorait les différents domaines de sa discipline où les itérations sont indispensables : «Puisque nous parlons d'itérations, une autre apparaît également dans la démonstration de l'inégalité isopérimétrique qui affirme que, parmi toutes les courbes fermées de longueur donnée, le cercle est celle qui entoure la plus grande aire. La démonstration de Wilhelm Blaschke (1885 - 1952), dans les années 1920, utilise une opération appelée symétrisation de Steiner (voir la figure 8). À partir d'un domaine dont le périmètre est une courbe quelconque, cette opération consiste, pour une direction donnée du plan, à découper le domaine en tranches infiniment fines, puis à enfiler les tranches symétriquement autour d'un axe perpendiculaire à la direction donnée. Le domaine obtenu est de même aire, mais de périmètre plus petit que celui du domaine initial. Cette réduction de périmètre est due à une propriété des trapèzes, auxquels les tranches sont assimilées : le périmètre

d'un trapèze est minimal quand les deux côtés opposés non parallèles ont même longueur.»

L'algébriste confirme qu'un rapide calcul le montre.

Encore et toujours des itérations

Le géomètre poursuit, imperturbable : «L'itération de ces symétrisations avec infiniment de directions différentes a une limite : c'est un cercle, car lui seul est symétrique dans toutes les directions. De surcroît, la symétrisation fonctionne en toute dimension.»

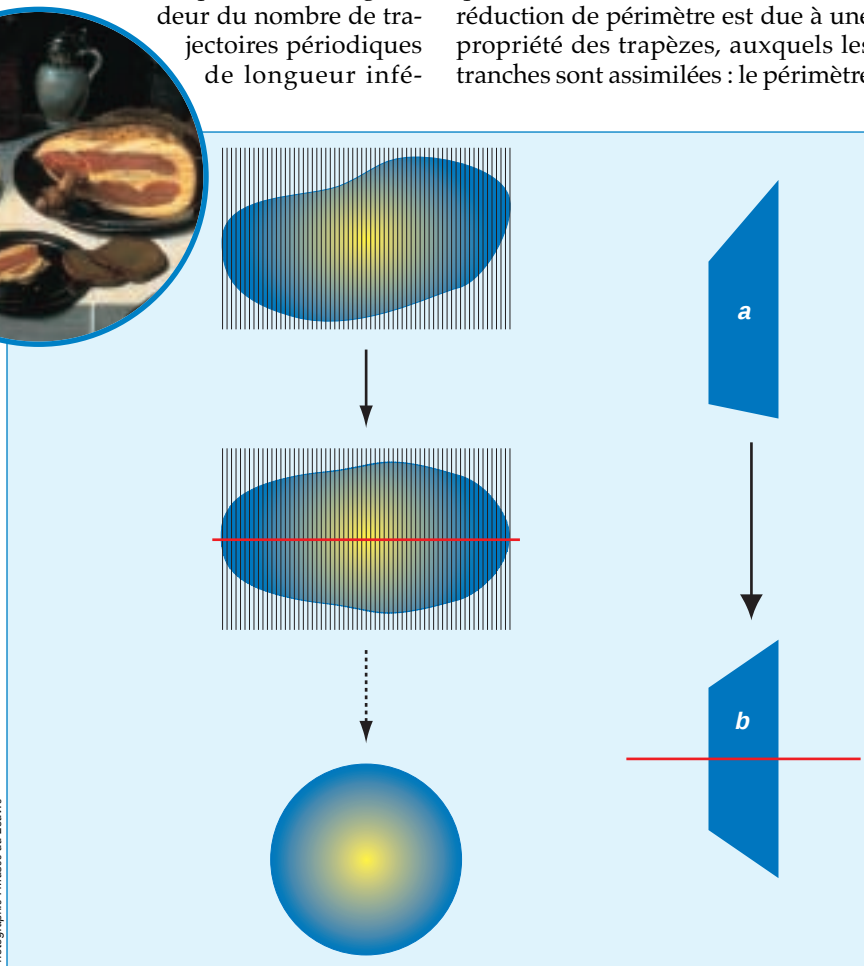
L'analyste et l'algébriste demandent alors : «Existerait-il une méthode plus naturelle pour minimiser le rapport de la longueur de la courbe frontière sur l'aire intérieure d'un domaine D ?»

Le géomètre répond : «Dans la réalité, on étudie plutôt le rapport du carré de la longueur sur l'aire pour manipuler des nombres sans dimensions. Il faut d'abord prouver l'existence d'une limite – une courbe – à ce rapport. Dans le plan, cette existence est démontrée. Toutefois, ce problème d'existence est à considérer ici dans l'espace infini de toutes les courbes : jusque dans les années 1960, dans des espaces de dimensions infinies, aucun théorème d'existence d'objets limites au sein de familles d'objets géométriques, telles les courbes, et de limites de variables elles aussi géométriques, telle l'aire d'un domaine, n'était disponible. Ce n'est qu'à la suite des travaux d'Herbert Federer, alors à l'Université Brown, qui constituent la théorie géométrique de la mesure, que l'existence de telles limites a été prouvée. On conclut alors, et ce en dimension quelconque, que la limite est, en deux dimensions, un disque, et, au-delà, une boule.

Puisque nous évoquons la démonstration de l'inégalité isopérimétrique, nous pouvons aussi penser l'augmentation du rapport du carré de la longueur sur l'aire, non plus par l'opération de symétrisation, mais par une transformation continue, c'est-à-dire la déformation d'une courbe C en une courbe de plus en plus proche d'un cercle. Pour ce faire, on se souvient que les cercles sont des courbes dont la courbure est constante.»

L'algébriste interrompt ici le géomètre : «Tu peux nous rappeler ce qu'un géomètre entend par courbure?»

«Bien sûr ! Pour déterminer la cour-



Photographie : Musée du Louvre

8. LA SYMÉTRISATION d'un domaine réduit son périmètre en maintenant son aire constante. On découpe le domaine initial en tranches infiniment fines (à gauche), puis on rassemble ces tranches symétriquement selon un axe perpendiculaire à celui des tranches. Lorsque cette opération est itérée un nombre infini de fois, le domaine tend vers un cercle qui minimise le périmètre. Celui-ci diminue en vertu d'une propriété des trapèzes de même aire (à droite) : le périmètre du trapèze b , symétrique, est inférieur à celui du trapèze a .

bure, le géomètre cherche le cercle osculateur à la courbe en un point. Pour ce faire, on considère le cercle tangent à la courbe en ce point et passant par un point voisin. Quand ce point voisin se rapproche du point étudié, ce cercle a une limite, c'est le cercle osculateur en ce point et son rayon est le rayon de courbure de la courbe. La courbure, elle, est par définition l'inverse de ce rayon. Ensuite, en chacun de ses points, la courbe se déplace le long de sa normale (la droite perpendiculaire à la tangente en ce point) d'une quantité inversement proportionnelle à la courbure en ce point (voir la figure 9). De manière plus précise, on cherche une famille de courbes $C(t)$, t est un paramètre continu, par exemple le temps, tel que la dérivée le long des normales soit inversement proportionnelle à la courbure. L'équation aux dérivées partielles qui décrit cette famille de courbes admet des solutions pour des temps assez petits. En revanche, l'existence pour un temps infini est une étude difficile.»

L'analyste interrompt de nouveau le géomètre : «Lorsque t tend vers l'infini, les courbes $C(t)$ ne convergent-elles pas vers un point, plutôt que vers un cercle?»

«Oui, cependant, convenablement renormalisées, c'est-à-dire que l'échelle de chaque courbe est modifiée de façon à conserver l'aire du domaine, les courbes $C(t)$ convergent vraiment vers un cercle. Enfin, en 1986, Gage a montré que le rapport isopérimétrique décroît lorsque t augmente. Cette technique d'évolution géométrique peut être utilisée dans des contextes plus généraux, par exemple pour obtenir des géodésiques périodiques sur les surfaces. Sujet que j'aimerais maintenant aborder avec vous, si vous le permettez.»

Les géodésiques du dessert

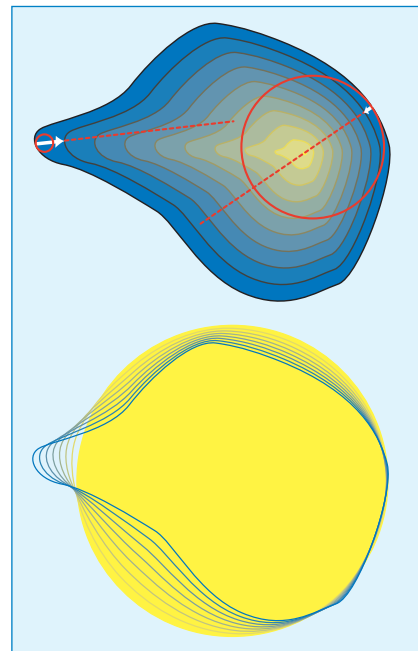
«Nous étudions un système dynamique continu dans lequel on décrit le mouvement d'une particule qui se meut sur une surface S donnée de l'espace, sans gravité. Depuis Johann Bernouilli (1667 - 1748), qui l'a enseigné à son non moins illustre élève Leonhard Euler (1707 - 1783), on sait qu'un tel mouvement suit une géodésique de la surface, c'est-à-dire une courbe dont l'accélération dans l'espace est normale à la surface, ou encore des courbes qui minimisent le chemin parcouru entre deux points voisins. Que devient cette

trajectoire, cette géodésique, lorsqu'on la prolonge indéfiniment? Est-elle périodique? À l'inverse, recouvre-t-elle toute la surface? Plus encore, prend-elle toutes les directions en chaque point? Et s'il existe des trajectoires périodiques, peut-on évaluer leur nombre en fonction de leur longueur? En 1905, Henri Poincaré a été le premier à s'intéresser à ce problème, car il s'agit du modèle mathématique le plus simple se rapprochant de celui qui décrit le Système solaire.

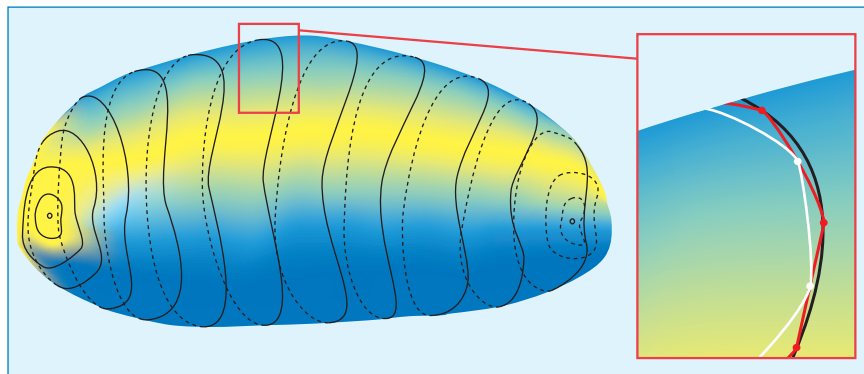
Dans le cas de la sphère, toutes les trajectoires, en nombre infini, sont périodiques et de même période : ce sont les grands cercles (les équateurs) de la sphère. Par ailleurs, d'autres surfaces, telles les surfaces de Zoll, qui ne sont pas des sphères, ont cette propriété (voir la figure 6). Pour toute surface de révolution, il y a une infinité de géodésiques périodiques. Sur une telle surface, la trajectoire d'une particule oscille entre deux parallèles. De par la révolution de la surface, une géodésique qui part d'une parallèle y reviendra après un certain temps, mais en un point qui a tourné d'un angle α par rapport au point initial. Nous sommes alors de nouveau dans le cas du billard que nous avons déjà évoqué : les géodésiques périodiques correspondent aux angles α proportionnels à π ; elles sont en nombre infini.

Pour les surfaces générales, on cherche d'une part à montrer l'existence de géodésiques périodiques, et à les évaluer asymptotiquement ; d'autre part, à déterminer l'ergodicité (en espace ou en phase) des trajectoires non périodiques. Posé simplement, le problème n'en demeure pas moins ardu : Poincaré échoua à prouver l'existence d'une

seule géodésique périodique. La raison essentielle de cette difficulté résultait, comme pour l'inégalité isopérimétrique, de l'absence d'un théorème d'existence de limites dans un espace géométrique de dimension infinie, ici celui des courbes fermées d'une surface. Ce n'est qu'en 1917 que George Birkhoff (1884-1944) a démontré, sans recours à un théorème d'existence d'une forme



9. LE PÉRIMÈTRE LE PLUS COURT pour une surface d'aire donnée (en bleu) est un cercle. L'une des méthodes pour le montrer consiste à déplacer les points du périmètre initial le long de leur normale d'une distance (en blanc) inversement proportionnelle au rayon du cercle osculateur (en rouge). En haut, le domaine tend vers un point, cependant, lorsque les domaines sont renormalisées, c'est-à-dire maintenues à aire constante, le domaine tend vers un cercle (en bas).



10. UNE TAPISSERIE est une famille de courbes qui «recouvrent» entièrement une surface. À partir d'une tapisserie quelconque (à gauche), on identifie une géodésique périodique en découpant les courbes en fragments d'une longueur inférieure à une distance donnée. Puis, entre deux points consécutifs, on prend le chemin le plus court (en rouge, à droite) ; sur une sphère, on prendrait les morceaux des grands cercles entre chaque paire de points. Enfin, des milieux des côtés du polygone géodésique, on trace un autre polygone strictement plus court (en blanc). La plus longue courbe d'une tapisserie économique est une géodésique périodique.

limite, l'existence d'une géodésique périodique pour une surface quelconque. Tout d'abord, il considère toutes les tapisseries de la surface, c'est-à-dire les familles de courbes allant d'une courbe réduite à un point m à une autre courbe réduite à un point m' : cette

famille de courbes recouvre entièrement la surface étudiée (voir la figure 10). Dans chacune des tapisseries, il y a une courbe de plus grande longueur, et l'on recherche une tapisserie pour laquelle la courbe de plus grande longueur est la plus petite possible.»

L'algébriste se permet une remarque : «En quelque sorte, il s'agit de rechercher une tapisserie économique.»

«À quelques subtilités près, c'est exact ! La plus grande courbe de cette tapisserie optimale est une géodésique périodique. Cela, à cause de la propriété de minimalité des géodésiques.»

L'analyste ponctue : «Reste à prouver l'existence d'une tapisserie économique.»

«C'est là qu'intervient la deuxième idée de Birkhoff : passer d'un espace de dimension infinie à un espace de dimension finie. Il a cherché à fabriquer, à partir d'une tapisserie quelconque, une deuxième tapisserie dont toutes les courbes sont strictement plus petites, à l'exception des géodésiques périodiques. Une tapisserie minimale contient toujours une géodésique périodique ; sinon, la transformation de nouveau appliquée fournirait une tapisserie dont toutes les courbes seraient plus courtes, et la tapisserie précédente ne serait donc pas minimale. Le procédé de Birkhoff

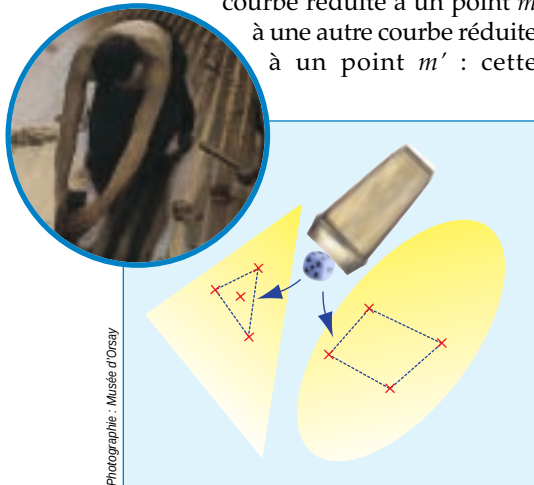
consiste à découper les courbes des tapisseries de la surface en N morceaux de longueurs inférieures à une distance donnée. Puis on remplace chaque morceau par le chemin le plus court, toujours sur la surface, joignant les deux extrémités du morceau. La deuxième opération consiste à prendre ensuite les milieux des côtés des polygones géodésiques obtenus et de les joindre en un nouveau polygone strictement plus petit, car, sur n'importe quelle surface, la longueur d'un côté d'un triangle est toujours inférieure à la somme des deux autres. Le nouveau polygone aura une longueur strictement inférieure à celle de la courbe initiale, à moins que tous les angles aux sommets soient égaux à π ; or cela arrive exactement pour une géodésique périodique. Voilà comment le passage d'une dimension infinie à une dimension finie résout le problème.

Ce procédé ne fournit qu'une seule géodésique périodique, et ce n'est qu'en 1992 que J. Franks, de l'Université de Fribourg, en Allemagne, et Victor Bangert ont montré qu'il en existe une infinité.»

L'analyste ne put qu'approuver : «Admirable !»

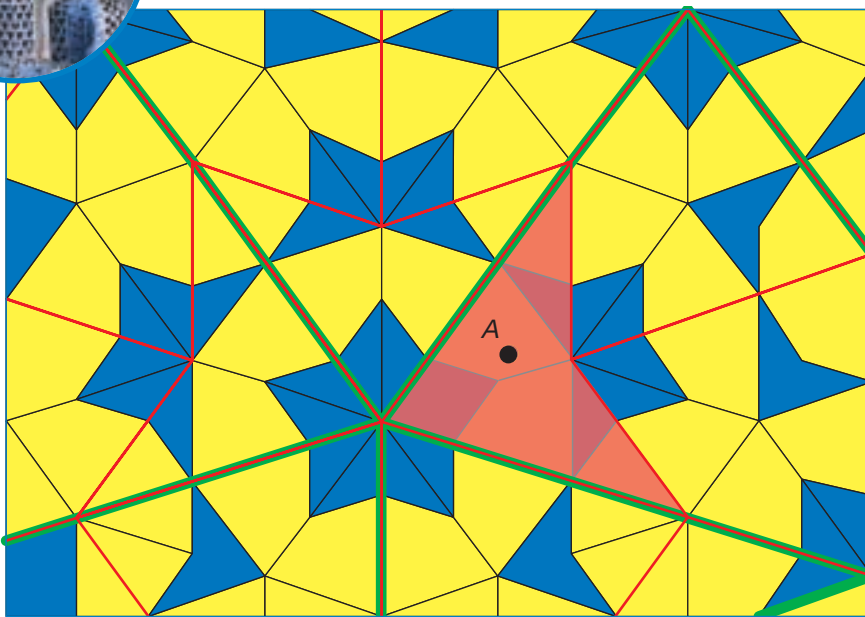
L'algébriste, après avoir fini son dessert, reprend la parole : «Face à un espace géométrique de dimension infinie, Birkhoff se tire d'affaire, certes de façon élégante, en approximant cet espace – l'espace de toutes les courbes planes d'une surface – avec un espace de dimension finie. Mais je subodore que ce n'est pas toujours possible.»

«Vrai, examinons par exemple celui des pavages de Penrose, c'est-à-dire ceux qui, construits à partir de deux types de pavés, recouvrent le plan à l'infini, mais de façon nécessairement non périodique, grâce aux règles de construction (voir la figure 12). En 1975, Robinson a montré que l'on peut complètement les classer. Il a remarqué qu'à partir d'un pavage donné P , on peut grouper les pavés afin qu'ils forment des pavés plus grands de deux types seulement. Le nouveau pavage obtenu P^1 est également non périodique. De P^1 , par le même procédé, on obtient P^2 , puis P^3 , jusqu'à P^n , n tendant vers l'infini. Il s'agit ensuite de prendre un point quelconque à l'intérieur du pavage initial et de regarder dans quel

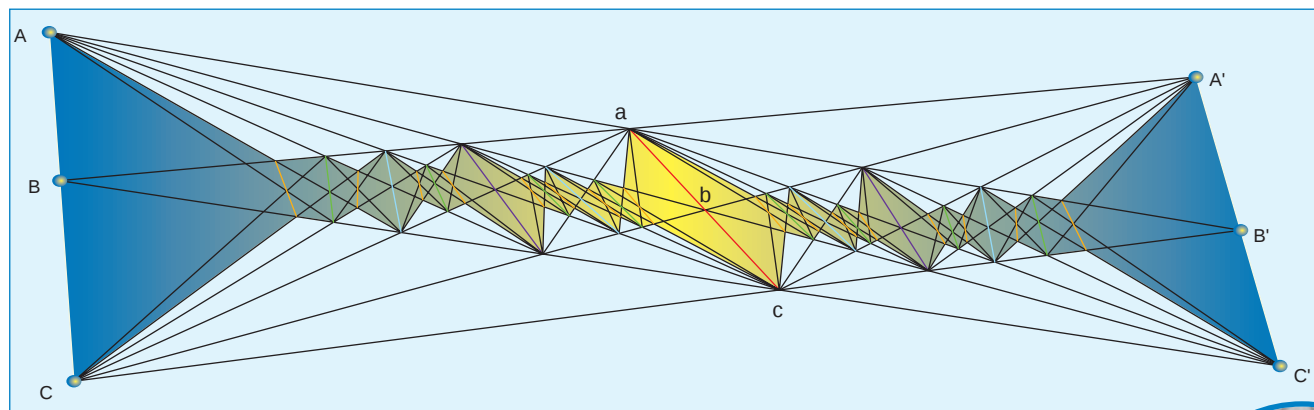


Photographie : Musée d'Orsay

11. QUATRE POINTS jetés dans un plan forment soit un quadrilatère, soit un triangle contenant le quatrième point. La probabilité d'obtenir le premier dépend de la géométrie du domaine : elle est maximale dans un domaine en ellipse, et minimale dans un domaine en triangle. La premier exemple de probabilités géométrique est le fameux problème de Buffon qui étudiait la probabilité d'une aiguille lâchée d'être à cheval sur deux lattes d'un plancher (en haut, à droite).



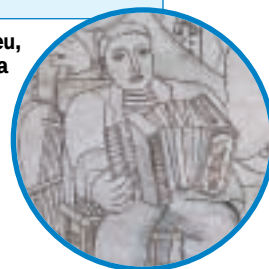
12. UN PAVAGE DE PENROSE est un pavage du plan non périodique avec deux types de pavés (en bleu et jaune). Groupés, les pavés forment un second pavage non périodique constitué de deux types de pavés plus grands (en rouge). À partir de ce pavage, on en construit un nouveau d'ordre supérieur (en vert). Dans chaque pavage, on associe à chaque type de pavés un chiffre (par exemple, 1 pour les pavés convexes et 0 pour les pavés concaves). On montre que l'espace des pavages de Penrose est similaire à celui des suites de 0 ou 1 qui coïncident à partir d'un certain rang et dans lesquelles un 1 est toujours suivi d'un 0. Ces suites sont construites en identifiant pour chaque ordre de pavage, le type de pavés dans lequel un point initial quelconque se trouve. Par exemple, le point A donne la suite 101...



Photographie : Musée Fernand Léger

13. LE THÉORÈME DE PAPPUS associe à deux triplets de point alignés (A, B, C) et (A', B', C') un troisième (a, b, c) : à partir des deux paires de triplets ainsi formées, le même procédé crée deux autres triplets. L'ensemble obtenu après l'itération à l'infini (seules cinq

sont représentées ici : dans l'ordre, rouge, violet, bleu, vert et orange) de cette opération, est fractal et a la structure de groupe modulaire que l'on retrouve dans de nombreux domaines des mathématiques.



type de pavé il est dans P^1 , dans P^2 , jusqu'à P^n : quand, à chaque type de pavé, à un niveau d'itération donné, correspond un chiffre 0 ou 1, on obtient une série de 0 et de 1. Ainsi, les pavages de Penrose forment un espace infini qui est similaire – en mathématiques, on dit isomorphe – à celui de l'ensemble de toutes les suites infinies de 0 ou de 1 dans lesquelles tout 1 est nécessairement suivi d'un 0. Des pavages sont identiques quand leurs suites coïncident à partir d'un certain rang, et ce jusqu'à l'infini : cette caractéristique signifie que, dans un même pavage, deux points initiaux différents sont, à partir d'un certain ordre, dans des pavés de même type.»

Le dîner touchait à sa fin. Au moment du café, l'algébriste et l'analyste sont presque convaincus. Dans leur élan, un cure-dents négligé pendant l'apéritif, roule et tombe sur plancher... à cheval sur deux lattes.

Des probabilités géométriques

Le géomètre reprend derechef : «Vous connaissez sans doute le problème des aiguilles de Buffon ? Il s'agit de calculer la probabilité qu'une aiguille a de tomber à cheval sur deux lattes d'un plancher. Ce problème est le premier exemple connu de probabilités géométriques. Nous ne sommes pas éloignés des itérations, puisqu'une probabilité dans la pratique signifie répéter une opération un nombre infini de fois. Je vous soumetts un autre exemple. Il est intéressant, car enfantin d'énoncé, et cependant en partie incompris aujourd'hui.

Lorsqu'on jette quatre points au hasard dans un domaine du plan, deux

cas sont possibles : ou les quatre points forment un vrai quadrilatère, ou bien l'un est à l'intérieur du triangle formé par les trois autres (voir la figure 12). Quelle est la probabilité d'avoir un vrai quadrilatère?»

L'algébriste et l'analyste réfléchissent un instant, puis le premier répond : «De nouveau intuitivement, je pense que la réponse dépend de la forme du domaine ; si c'est un triangle, il sera plus difficile d'avoir un quadrilatère près des sommets que dans une autre région. À l'inverse, un quadrilatère est plus probable près du bord d'un domaine dont la frontière est une courbe lisse.»

Le géomètre poursuit : «Exact, en 1930, Blaschke a montré que la probabilité cherchée est maximale quand le domaine est une ellipse, et minimale pour un triangle. En revanche, dès que l'on s'intéresse à des dimensions plus grandes, par exemple au jet de cinq points dans l'espace tridimensionnel, et plus généralement de $n+2$ points dans un espace à n dimensions, les géomètres sont perdus. Pour la dimension trois, on montre, avec difficulté, que les ellipsoïdes sont les domaines les plus favorables à la formation d'un pentaèdre. En revanche, pour la plus mauvaise probabilité, on ne sait pas si c'est celle d'un tétraèdre. Plus grave encore : on ne connaît pas la valeur exacte de la probabilité pour le tétraèdre, c'est-à-dire que l'on ne sait pas calculer le volume moyen de l'infinité de tétraèdres contenu dans un tétraèdre donné. Les géomètres ont encore du pain sur la planche.»

Les mathématiciens s'apprentent à partir, quand dehors, un accordéoniste passe avec son instrument.

De nouveau Pappus

«Messieurs, avant de nous quitter, et puisqu'il a beaucoup été question de cercles dans mon exposé, je terminerais par où j'ai commencé, à savoir avec le théorème de Pappus. En 1993, R. Schwarz a remarqué que Pappus associe, à toute paire de triplets de points alignés sur une droite, un troisième triplet du même type. Il est donc naturel d'itérer l'opération et, à partir de chacune des deux paires créées, de construire de nouveaux triplets (voir la figure 13). Que se passe-t-il quand cette opération est répétée à l'infini ? L'ensemble des droites obtenues forme une fractale qui, de plus, a une structure de groupe modulaire. Cette structure que l'on retrouve dans tant de branches des mathématiques : la théorie des nombres, la théorie des groupes, la géométrie des surfaces, la théorie des nombres complexes, etc.

Vous voyez, messieurs, conclut le géomètre alors que trois taxis se garent devant l'hôtel des Trois Faisans, à l'inverse de ce que vous pouviez penser, l'infini est omniprésent en géométrie, et il permet même de relier nos trois disciplines.»

Et les trois de se quitter... maintenant tous convaincus que les mathématiques, toutes les mathématiques, sont la science de l'infini.

Marcel BERGER est l'ancien directeur de l'Institut des hautes études scientifiques.

Marcel BERGER, *Géométrie vivante ou l'échelle de Jacob de la géométrie*, Cassini, à paraître.

L'analyse non standard

JEAN-MICHEL SALANSKIS



En introduisant la notion d'objet insaisissable, l'analyse non standard a conféré aux infiniment petits un statut honorable et justifié des opérations qui, bien que pratiquement satisfaisantes, posaient des problèmes de principe.

Qui ne s'est amusé à calculer sur sa calculatrice e^{-x} pour une valeur de x positive et assez grande ? Le résultat affiché est alors zéro. De même, lorsque nous appliquons la fonction racine carrée à un nombre positif très grand, nous constatons qu'après un certain nombre d'applications successives, nous obtenons 1. En gros, la calculatrice a calculé $e^{-\infty} = 0$ et, en fin de compte, $\sqrt{1+\epsilon} = 1$. Étranges égalités qui nous paraissent être fausses ou exprimer de manière inadéquate une propriété de limite. Comme la calculatrice calcule, néanmoins, et comme nous avons spontanément traduit son calcul par des écritures reflétant ce qu'elle fait, il semble bien qu'elle traite certains nombres comme infiniment grands ou infiniment petits, et en tire des conséquences pour le calcul.

Ces conséquences évoquent une algèbre des calculs sur les limites. Par exemple, nous avons appris que l'inverse d'une fonction tendant vers l'infini tend vers zéro. Bien sûr, nous l'avons appris de façon correcte et nous avons été mis en garde contre l'équation illicite : $1/\infty = 0$. Toutefois, en pratique, quand nous résolvons des exercices,

nous passons parfois mentalement par l'une de ces équations illicites.

L'analyse non standard est une méthode mathématique rigoureuse qui justifie les équations illicites : elle autorise le discours mathématique à faire usage des nombres infiniment petits ou grands, réputés inexistants. Georg Cantor nous a appris à poser, comme des objets respectables de la théorie des ensembles, des cardinaux infinis associés à des ensembles infinis ; de même, l'analyse non standard a défini des nombres réels infiniment grands ou infiniment petits.

En examinant l'histoire de l'infini et de l'infiniment petit ou grand, nous ferons ressortir cinq étapes : le calcul infinitésimal des pionniers où l'on néglige les grandeurs «très petites» ; la conception de l'infini issue de la théorie des ensembles ; l'analyse non standard du mathématicien américain Abraham Robinson, où l'infiniment grand et l'infiniment petit apparaissent par l'intermédiaire de «l'axiome du choix» ; la théorie des ensembles internes, une extension de la théorie des ensembles qui permet d'apercevoir des infiniment petits ou grands au

1. LES MESUREURS, tableau du XVII^e siècle attribué à Hendrick Van Balen (*ci-dessus*), illustre l'aphorisme du poète romain Horace : «Il y a une mesure pour toute chose.» Quelle que soit la précision des mesures, les quantités infinitésimales définies par l'analyse non standard resteront à jamais hors de portée.

milieu des nombres usuels ; enfin, l'analyse non standard d'inspiration «constructiviste» suggérée par le mathématicien strasbourgeois récemment disparu, Georges Reeb, qui conduit à envisager des entiers infiniment grands, au-delà des entiers saisissables.

La contradiction infinitésimale

Revenons à l'origine du calcul infinitésimal. Étant donné la relation entre variables $y = x^2$, nous voulons calculer l'accroissement dy de la variable y lorsque x est modifié de la quantité infiniment petite dx . Nous calculons alors $y + dy = (x+dx)^2$ soit, $dy = 2x dx + dx^2$; de là, nous passons à $dy = 2x dx$, en négligeant dx^2 comme carré d'un infiniment petit. Ainsi, dans le calcul infinitésimal, nous considérons comme licite d'affirmer que la différentielle de x^2 est $2x dx$.

Cela paraît une erreur, une «équation imparfaite», mais au fond, le concept de quantité infiniment petite n'étant pas clairement défini, nous ne savons pas très bien jusqu'à quel point il y a erreur. «Infiniment petit», pour les mathématiciens qui élaborent ce calcul, cela signifie «plus petit que toute quantité finie donnée». Chacun s'accorde, au fond, à considérer que la notion même de quantité réelle inclut la propriété que l'on dénomme «archimédianité» : tant qu'une quantité positive n'est pas nulle, on peut toujours trouver une autre quantité, du type $1/n$ (où n est un entier naturel non nul), plus petite qu'elle. Une quantité ne peut alors pas être «plus petite que toute quantité finie donnée», les nombres du type $1/n$ étant de (petites) quantités données. Autrement dit, un réel strictement positif ne peut être infiniment petit au sens strict : on en produit facilement un plus petit.

Pourtant, les pionniers ont développé un calcul où l'on utilise de telles quantités et où l'on en néglige certaines. Toutefois, en dépit de problème de principe, ce calcul s'applique à la fois au sein des mathématiques, puisqu'il permet de caractériser localement les courbes, notamment la pente de la tangente en un point, de calculer les longueurs, des aires, etc. – et hors des mathématiques – puisqu'il autorise la modélisation des phénomènes, par intégration des «équations différentielles» établies sur des données élémentaires de la situation. Il doit donc y avoir une rationalité derrière l'usage des infiniment petits et des «équations imparfaites» et la «pensée infinitésimale» des pionniers était une démarche mal fondée, mais qui allait dans la bonne voie.

La reformulation ensembliste

La pensée ensembliste contemporaine, c'est-à-dire la conception de l'infini issue de la théorie des ensembles, supprime les contradictions en traduisant le discours des pionniers dans son nouveau langage. Ainsi, l'équation précédemment citée, $dy = 2xdx + dx^2$, devient $f(x+dx) - f(x) = 2xdx + dx^2$, « f » étant le nom de la fonction qui à x associe son carré x^2 . On la considère comme une équation exacte où dx est une quantité ordinaire, finie. On divise alors les deux membres par dx et l'on aboutit à : $dy/dx = (f(x+dx) - f(x))/dx = 2x + dx$. On constate maintenant que, «lorsque

dx tend vers 0, le taux d'accroissement $(f(x+dx) - f(x))/dx$ tend vers $2x$ ». Ce n'est pas $2xdx$ qui est assimilable à dy , c'est dy/dx qui «tend» vers $2x$. Cela suppose que l'on définisse la notion de «tendre vers», ce qu'ont fait Karl Weierstrass, Georg Cantor et Richard Dedekind.

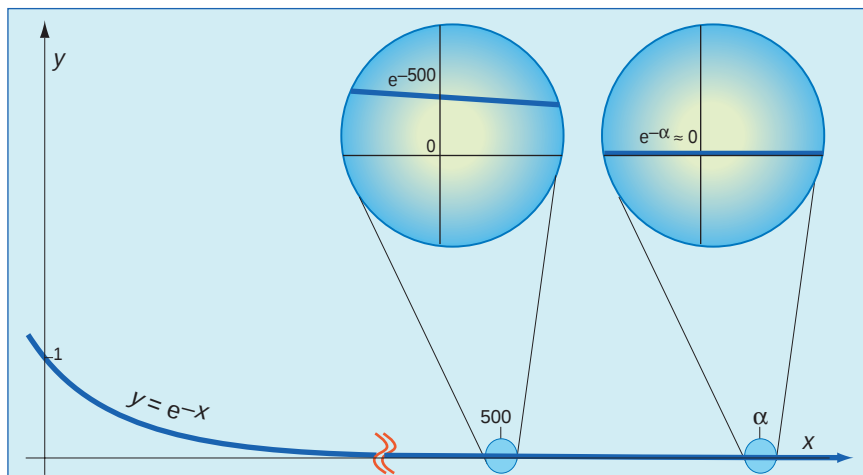
En substance, dire que le taux d'accroissement tend vers le nombre l , c'est dire que l'on peut rendre sa différence avec l inférieure en valeur absolue à une valeur strictement positive arbitraire, quelle qu'elle soit, à condition de choisir un dx inférieur en valeur absolue à une valeur frontière strictement positive adaptée à la première. Cette définition en termes de contrôle satisfait les mathématiciens. Toutefois, le théoricien des fondements remarque, qu'avec cette définition, une propriété de convergence vers une limite apparaît comme établie à l'issue de la résolution d'une infinité de petits problèmes posés par chaque degré de petitesse que l'on veut réaliser. Plus sobrement, elle est la propriété d'un objet supposé bien défini qu'est une fonction, c'est-à-dire une application de l'ensemble des réels dans lui-même, ou encore, selon la conception ensembliste, un ensemble fabriqué à partir d'ensembles infinis.

Ainsi, l'interprétation «moderne» du discours des pionniers de l'analyse infinitésimale le sauve de l'équation imparfaite, qui était l'un des problèmes, mais pas de la référence à une quantité douteuse. Seulement celle-ci, au lieu d'être infiniment petite, est infinie «tout court». À vrai dire, une difficulté de l'ordre de celle de l'équation imparfaite surgit aussi, mais elle

ne le fait que pour un spécialiste soucieux des problèmes de fondement. En acceptant les ensembles infinis actuels de points, le mathématicien est obligé d'accepter une «logique de l'infini», de prolonger l'usage de ses règles logiques de discours de manière systématique à l'infini.

Notamment, il décide que la phrase exprimant que telle fonction tend vers l en zéro, phrase qui évoque une infinité de problèmes paramétrés par un nombre réel strictement positif (pour tout réel strictement positif ε , il existe un réel strictement positif δ tel que, quel que soit x , le fait que la valeur absolue de x soit inférieure à δ implique que la valeur absolue de $f(x) - l$ est inférieure à ε) est ou bien vraie ou bien fausse, selon le principe du tiers exclu. (Le principe du tiers exclu, voir *L'infini, pierre de touche du constructivisme*, pages 74 à 81, énonce que si l'on considère une affirmation P portant sur des objets mathématiques, alors soit P soit non- P est vraie, et l'on exclut toute autre possibilité.)

Or cela ne va pas sans difficulté théorique, ne serait-ce que parce que l'on est amené à formaliser les mathématiques, et, par suite, à adopter un concept de déductibilité plus essentiel que le concept de vérité, concept pour lequel, on le sait depuis la démonstration par Gödel de son fameux théorème d'incomplétude, le tiers exclu ne vaut précisément pas (en 1930, Gödel a montré que dans tout système formel assez riche, on peut construire une proposition P telle que ni P ni non- P ne sont déductibles dans le système formel : une proposition «indécidable»).



2. LA FONCTION e^{-x} décroît très rapidement lorsque x augmente. Toutefois, si l'on prend un x fini, par exemple 500, en grossissant la courbe on notera une différence avec 0. En revanche, si l'on choisit un α infiniment grand, même en grossissant la courbe, la différence restera insaisissable. Selon cette intuition infinitésimale, on écrira $e^{-\alpha} \approx 0$.



La théorie des ensembles

C'est principalement Georg Cantor (ci-contre) qui a eu l'idée d'introduire en mathématiques la notion d'ensemble quelconque d'objets absolument quelconques et de démontrer des résultats généraux sur ces d'ensembles. La conception ensembliste stipule que tous les objets et toutes les notions des mathématiques peuvent être exposés en termes d'ensembles, de sorte que la «théorie des ensembles» devient la théorie sur laquelle se fondent toutes les mathématiques.

Les raisonnements sur les ensembles que l'on fait naturellement en partant de l'intuition de ce qu'est une collection et de ce qu'est, pour un objet, que d'appartenir à une telle collection, conduisent à des paradoxes (antinomie de Russell, paradoxe de Cantor, etc.). Les mathématiciens du début du xx^e siècle ont alors été conduits à formaliser la théorie des ensembles, à la présenter comme une théorie logique abstraite gouvernant l'emploi de la relation d'appartenance notée $\in$, symbole qui, désormais, n'a plus un sens intuitif, mais qui est soumis à des règles d'emploi consignées dans des axiomes. La théorie formelle ainsi obtenue, et qui constitue la référence implicite des mathématiques depuis près d'un siècle, est la théorie ZFC (Zermelo-Fraenkel avec axiome du choix).

Dans cette théorie, on définit la notion d'*ensemble infini* et d'*ensemble fini* comme Cantor en avait déjà eu l'idée : est infini un ensemble qui peut être mis en bijection avec une partie de lui différente de lui (ainsi, l'ensemble N des entiers naturels est mis en correspondance bijective avec l'ensemble P des nombres pairs par l'application $n \rightarrow 2n$) ; est fini un ensemble qui n'est pas infini. À tout ensemble, même infini, est associé un «cardinal», qui «compte» le nombre de ses éléments : ainsi N a «aleph 0» éléments, comme $\{0,1\}$ en a deux. Une application $f : E \rightarrow F$, de source l'ensemble E et de but l'ensemble F , est définie comme un triplet (E, G, F) où $G \subset E \times F$ est constitué de couples du type (x, y) avec $x \in E$ et $y \in F$, le graphe G ayant été choisi tel qu'à tout $x \in E$ peut être associé un unique $y \in F$ tel que $(x, y) \in G$ (le y en question s'appelle alors *image* de x). Si E est infini, le graphe G est alors, lui aussi, infini.

La reconquête de l'indéfini

Cette situation laissait soupçonner que l'on aurait pu employer le formalisme issu de la théorie des ensembles et les formalismes associés pour justifier le mieux possible le discours des pionniers, plutôt que pour le disqualifier, comme le fait dans une certaine mesure la conception ensembliste que nous venons d'évoquer. Or, dans les années 1960, une telle réhabilitation a été proposée par le mathématicien américain Abraham Robinson, spécialiste de la branche alors triomphante de la logique mathématique nommée théorie des modèles.

La réinterprétation de Robinson du calcul et du discours des pionniers est en substance la suivante : elle sauve l'idée d'infiniment petit, en demandant aux mathématiciens de penser qu'il y a plus de nombres réels qu'ils ne le croyaient. Ou plutôt, les réels connus sont plongés dans un système plus vaste, qu'on peut nommer système des «hyperréels», et contenant des éléments strictement positifs plus petits que tout élément strictement positif du système de base, qu'on dénommera désormais système des

réels standard. Toute la subtilité de l'opération réside alors dans ceci que le nouveau système de nombres est une «extension élémentaire» du premier, c'est-à-dire qu'il en est jusqu'à un certain point un double, un élargissement non essentiel.

En d'autres termes, les nouveaux objets ne peuvent être caractérisés en fonction des réels standard et des moyens discursifs qui les accompagnent. Ils sont des entités supplémentaires, indéterminées, simplement choisies de façon à procurer, à toute une classe de relations définies sur le système de base, des objets-limite dominant tous les objets standard suivant ces relations. Pour donner un exemple, disons que la relation $\leq$ a la propriété suivante : pour tout ensemble fini $\{x_1, \dots, x_n\}$ de réels standard, on peut trouver un réel standard y majorant chacun des x_i . Le processus d'élargissement de Robinson nous garantit alors dans le système $^*\mathbb{R}$ des hyperréels l'existence d'un μ tel que $x < \mu$ pour tout x standard, c'est-à-dire d'un μ «infiniment grand».

Pour construire les hyperréels, Robinson se fonde sur l'axiome du choix. Cet axiome – introduit explicitement par Zermelo en 1908 – pose,

pour le résumer de manière commode, qu'«il existe» une fonction choisissant exactement un élément dans tout ensemble d'une famille infinie d'ensembles non vides donnés. Autrement dit, si l'on considère une famille d'ensembles non vides $(X_i)_{i \in I}$ variant dans l'ensemble I des indices, il y a une fonction qui «prend» exactement un élément dans chaque ensemble X_i , qui à chaque indice i de l'ensemble d'indices I associe un élément choisi dans X_i . Toutefois, le «choix» ainsi exprimé par la fonction censée exister n'est ni déterminé ni déterminable au-delà de l'axiome qui pose son existence. On ne peut pas décrire cette fonction autrement que par la sélection qu'elle opère, on ne peut pas l'exprimer par rapport aux objets déjà connus de la situation où elle intervient.

Néanmoins, la simple existence de la fonction est intéressante pour le mathématicien : elle lui garantit que certaines constructions ensemblistes sont possibles, que certains objets commodes peuvent être admis. Par exemple, elle permet l'introduction d'une base à laquelle rapporter tous les vecteurs de n'importe quel espace vectoriel.

Malgré les services rendus par cet axiome, plusieurs mathématiciens l'ont contesté au début du XX^e siècle. Ils estimaient que l'axiome de choix affirmait une existence trop idéale et trop peu contrôlée. Robinson justifie *a posteriori* cet axiome en montrant que le caractère peu contrôlé et indéterminé des objets introduits par l'axiome du choix peut être intéressant en lui-même.

L'axiome du choix amène la mathématique à jouer avec des objets dont elle n'a pas suivi le protocole de création, des objets qui s'ajoutent à son paysage, qui lui sont uniquement raccordés par une propriété miraculeuse. Les infiniment petits de Robinson sont justement de tels objets : leur statut se définit en termes de l'infinité des objets standard, mais ils sont pour le reste indéterminés. Robinson justifie le discours des pionniers en donnant un crédit ensembliste irréfutable aux infiniment petits (et, du coup, aux infiniment grands), mais en ajoutant à la figure de l'infini de totalité celle d'un infini-indéfini, d'un infiniment indéfini en quelque sorte (ou plutôt, cette figure, il ne l'ajoute pas, il la met en relief). C'est d'ailleurs en termes de cette figure que l'équation imparfaite est aussi interprétée : on écrira avec Robinson $dy/dx \approx 2x$, le signe $\approx$ exprimant que

la différence est infiniment petite, c'est-à-dire qu'elle est inférieure en valeur absolue à tout nombre standard. De manière philosophique, on peut dire que dy/dx ne diffère de $2x$ que d'un «indéfinissablement petit».

La syntaxe de l'indéfini

Dans les années 1970, le mathématicien américain Edward Nelson a proposé une nouvelle théorie des ensembles qu'il nomme théorie interne des ensembles, ou IST (pour *Internal Set Theory*). La théorie IST étend la théorie formelle des ensembles usuels, la théorie Zermelo-Fraenkel avec choix, ou ZFC. Cette nouvelle théorie des ensembles fait intervenir, en plus des symboles fondamentaux de l'égalité ($=$) et de l'appartenance ($\in$), un nouveau prédicat nommé st (comme standard). À l'inverse de $=$ ou de $\in$, qui sont des symboles à deux places ($x = y$, $x \in y$), le prédicat st est un prédicat à une place, c'est-à-dire qu'il affirme une propriété sur un objet (on écrit $st(x)$ pour x est standard). Cela signifie qu'on distinguera désormais, parmi tous les ensembles de l'univers des ensembles, ceux qui sont standard et ceux qui ne le sont pas (c'est un peu comme si la télévision ensembliste était passée du noir et blanc à la couleur, nous révélant des distinctions nouvelles).

Le prédicat st ne peut être réduit aux notions ensemblistes classiques. Sa signification est régie par trois nouveaux schémas d'axiome qui ont pour effet de lui faire exprimer plus ou moins la notion d'objet assignable, saisissable, logiquement connu dans le paysage ensembliste. Ce prédicat permet de distinguer les entités qui tiennent le rôle des infiniment petits chez Robinson. Ainsi, dans la mathématique nelsonienne, les réels infiniment petits non nuls sont des réels ne satisfaisant pas le prédicat st , des réels non standard.

Cette approche de Nelson met en évidence un rapport entre le fini et le

standard. Elle laisse entendre que les objets non standard sont là «à cause» de l'infini de totalité, et que, en un certain sens, tout ce qui est assignable tient dans le fini. Bien que la théorie IST qu'il propose soit une extension de la théorie classique, et puisse donc être considérée comme porteuse du même message logique et ontologique, elle nous engage à voir autrement dans l'univers des ensembles la répartition du fini et de l'infini, et leur rapport.

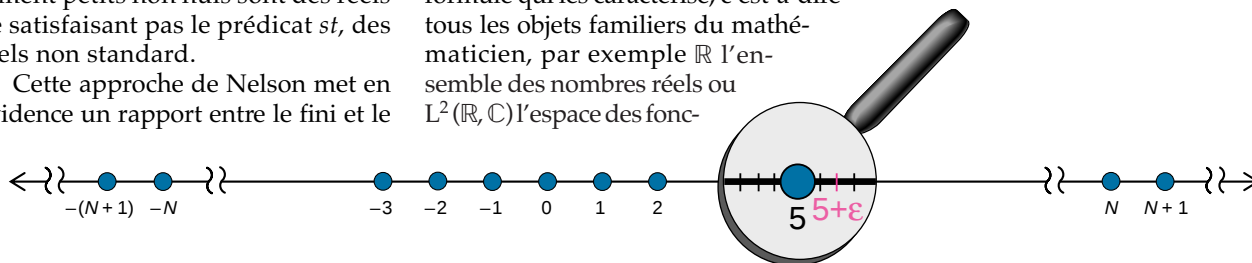
Les deux enseignements de la théorie IST sur le fini et l'infini sont les suivants : premièrement, tout ensemble infini contient un élément non standard, deuxièmement il y a un ensemble fini qui contient tous les objets standard de l'univers. Ces deux assertions sont des théorèmes qui découlent à peu près directement des axiomes spécifiques imposés au prédicat st par lesquels E. Nelson a complété la théorie des ensembles.

La première assertion se comprend au fond très bien : dans un ensemble que nous tenons pour infini, comment pourrait-il ne pas apparaître des objets échappant à notre contrôle logique, c'est-à-dire des objets non assignables, non standard ? La conclusion que tire le théorème semble être une simple conclusion raisonnable, dictée par la connaissance de notre finitude (bien que, naturellement, une telle considération métaphysique ne figure nullement parmi les justifications du dispositif de Nelson ; des justifications, à vrai dire, il n'en donne pas, c'est en cela notamment que consiste la démarche axiomatique).

La seconde assertion est plus surprenante : parmi les objets standard figurent, Nelson prend soin de nous le faire remarquer, tous les objets définissables dans la théorie de base des ensembles (classique) au moyen d'une formule qui les caractérise, c'est-à-dire tous les objets familiers du mathématicien, par exemple $\mathbb{R}$ l'ensemble des nombres réels ou $L^2(\mathbb{R}, \mathbb{C})$ l'espace des fonc-

tions complexes de carré sommable sur $\mathbb{R}$ (les fonctions de $\mathbb{R}$ sur $\mathbb{C}$ dont l'intégrale sur $\mathbb{R}$ du module au carré est finie) ; mais aussi, bien sûr, les entiers «connus», 0, 1, 2, et ainsi de suite. La collection des objets standard contient donc toute une faune, apparemment infinie, d'objets qui sont eux-mêmes susceptibles d'être idéaux (c'est-à-dire que l'on peut penser, mais non construire, parce qu'ils se situent par principe au-delà de toute saisie immédiate, comme $\mathbb{N}$ ou $\mathbb{R}$) et du coup infinis. Néanmoins, ce que dit le théorème, c'est que la collection totale des objets de cette sorte peut être intégrée à un ensemble fini, et donc, en un sens, demeure justiciable d'un traitement analogue à celui des objets finis. Par exemple, on pourra former la somme finie d'une classe de nombres réels contenant tous les réels standard.

Cette affirmation contredit l'intuition qui résulte de la construction de Robinson : pour lui, les nombres réels standard sont les nombres usuels, et l'on n'a obtenu un système avec des nombres infiniment petits et grands qu'en ajoutant de nouveaux nombres à ces nombres usuels ; il y a donc une infinité continue de réels standard. Le point de vue de Nelson est autre : nous avons des nombres infiniment petits uniquement parce que nous reconnaissons que, parmi les nombres réels, figurent forcément des éléments insaisissables, non assignables, parce que nous entrons dans la logique de cette distinction entre assignables et non assignables. C'est cette distinction qui différencie fonctionnellement les nombres réels entre eux et qui dégage la possibilité «grammaticale» de l'infiniment petit ou grand. Elle fait aussi apparaître, à l'intérieur des réels, une collection de réels standard qui se laisse inclure dans une partie de $\mathbb{R}$ que l'on peut traiter comme une partie finie.



3. LES NOMBRES RÉELS INCLUENT les nombres entiers, positifs et négatifs, les nombres rationnels et les nombres irrationnels. Ces nombres réels peuvent être représentés par des points sur une droite nommée la droite réelle. L'analyse non standard nous conduit à envisager, autour de 0, tout un «halo» de nombres non standard infiniment petits. Plus généralement, chaque nombre réel usuel, comme $\sqrt{2}$ ou 5 sur le schéma, est environné d'un tel halo, consti-

tué de nombres de type $5 + \varepsilon$ où ε est infiniment petit. Mais l'analyse non standard nous fait voir aussi, sur la droite réelle des nombres non limités, comme N , $-N$, $N + 1$, $-(N + 1)$ sur le schéma, qui sont supérieurs à tout nombre standard (infiniment grands positifs) ou inférieurs à tout nombre standard (infiniment grands négatifs) ; néanmoins aucun de ces nombres ne peut être dit «compter» l'infinité des points de la droite.

Dans ce contexte de la théorie dénommée IST, gardons à l'esprit que fini ne signifie pas intuitivement fini, absolument fini, fini au sens des ordinateurs ou même de l'intuition ordinaire des nombres entiers. Fini a le sens que lui donne la théorie des ensembles ZFC, c'est-à-dire le sens formel de ce qui peut être mis en bijection avec une partie propre de soi. Un ensemble fini en ce sens peut ne pas paraître fini à notre intuition. Ainsi, lorsque μ est un entier infiniment grand, le segment $[0, \mu]$ est infini à nos yeux, mais fini du point de vue des manipulations mathématiques. C'est toute la difficulté qu'avaient déjà explorée les pionniers de la théorie des modèles et que Robinson avait exploitée pour faire émerger des réels infiniment petits. On peut, si l'on désire se raccrocher à une signification intuitive absolue du fini, nommer «hyperfini» ce fini qui n'est fini que formellement, au sens de la théorie des ensembles. Les objets standard restent dans l'orbe de l'hyperfini.

Assignable et constructif

La manière dont la théorie IST relie l'intention justificatrice de l'analyse non standard à l'égard du discours des pionniers du calcul infinitésimal et la question générale du fini et de l'infini prépare et éclaire l'étape suivante. Celle-ci consiste à mettre directement en rapport l'analyse non standard et les idées «constructivistes». Le constructivisme accorde une importance décisive à la distinction entre ce qui peut être effectivement construit, ou devrait pourvoir l'être en principe, et ce qui se situe au-delà, ce qui est «idéel». En effet, nous commençons à le comprendre, ce qui permet d'avoir des nombres infiniment petits dans la mathématique du continu sans engager les calculs sur les nombres réels dans des paradoxes, c'est de disposer d'un statut d'entité non assignable pour certains objets, certains nombres.

Ce statut, Robinson l'obtient en faisant jouer ensemble deux modèles de la même théorie dont l'un est une extension élémentaire de l'autre, et E. Nelson l'obtient en installant dans le langage une notion d'entité assignable. Le résultat de cette démarche est de lier l'assignable au fini, le non-assignable à l'infini.

Pourquoi, dans ces conditions, ne pas tenter de se fonder sur ce qui semble être la figure contemporaine de la notion d'entité assignable, et

RÉCURRENCE INTERNE ET EXTERNE

Le principe de récurrence nous dit que si une propriété $P(n)$ est vraie pour $n = 0$ et si elle est héréditaire (la vérité de $P(n)$ entraîne celle de $P(n + 1)$, alors elle est vraie pour tout n . Ce principe reste valable dans l'analyse non standard «nelsonienne», mais seulement pour les propriétés dites *internes*, qui ne font pas intervenir la notion d'objet standard (le prédicat *st*). Pour une propriété externe jouant sur la nouvelle notion, sachant que $P(0)$ est vraie et que, pour le standard, $P(n)$ entraîne $P(n + 1)$, on pourra seulement conclure que $P(n)$ est vraie pour tout entier standard, la valeur de vérité de la propriété P pour les entiers infiniment grands restant parfaitement indéterminée.

poser qu'est assignable ce qui est représentable par l'informatique? Parce que cette notion est trop technique et contingente : elle semble dépendre d'un état actuel de la puissance des machines et de la taille des mémoires, par exemple. Les notions plus théoriques d'objets constructibles au sens de Brouwer (c'est-à-dire d'objets dont l'élaboration apparaît comme réalisable à l'intuition qui accompagne la succession indéfinie des nombres entiers) ou de calculabilité (dans les années 1930, des mathématiciens ont introduit cette notion en cherchant à préciser la notion intuitive de calcul) fournissent exactement ce qu'il est naturel de rechercher : des concepts d'entité assignable assez généraux pour dépasser les limites contingentes de la technique, mais qui n'entrent pourtant pas dans un type d'idéalité comparable à celui des ensembles infinis de nombres $\mathbb{N}$ ou $\mathbb{R}$. D'où l'idée de fonder un discours d'analyse non standard dont l'équation de base serait *standard = constructif*.

Il revient à Georges Reeb d'avoir eu l'idée générale et systématique d'une telle stratégie, et également d'avoir compris que c'était toujours de cela qu'il s'était agi dans l'analyse non standard historique.

Il s'agit de comprendre que tout ce qui est nécessaire au développement de l'analyse non standard est un entier infiniment grand. On doit s'être également muni, bien entendu, d'une règle du jeu instruisant sur la façon de raisonner et de calculer avec lui, sur le type d'entités que l'on peut construire à partir des nombres entiers divisés en finis et infinis dont on dispose ; enfin, on doit se munir d'un bon motif «phi-

losophico-fondationnel» pour ajouter crédit à tout ce montage intellectuel faisant intervenir un nombre entier infiniment grand. Dans ces conditions, on peut introduire un système de nombres comportant des infiniment petits et grands qui permet de faire de l'analyse. Le texte qui formule le programme d'une mathématique non standard fondée sur l'entier infiniment grand tout en interprétant ce dernier comme entier non constructif est l'article *La mathématique non standard, vieille de soixante ans?* de Georges Reeb, publié en 1979.

Pour Reeb, nous devons reconnaître que la boîte rassemblant l'ensemble des entiers du formaliste, le $\mathbb{N}$ obéissant au principe du tiers exclu, que l'on définit comme le plus petit ensemble infini en théorie des ensembles, n'a aucune raison de contenir seulement les entiers intuitifs, les entiers qui viennent à leur heure au fil de l'énumération 0, 1, 2, etc. C'est ce que les inlassables objections de Brouwer au formalisme naissant, au début de XX^e siècle, auraient montré sans ambiguïté. Brouwer a critiqué la mathématique ensembliste et formelle qui s'affirmait sous ses yeux, et a proposé une autre orientation, qu'il nomma «intuitionniste». Selon Reeb, Brouwer avait raison de considérer qu'une preuve formelle de la fausseté de la proposition «quel que soit l'entier n , la propriété $P(n)$ est vraie» – le plus probablement obtenue au moyen d'un raisonnement par l'absurde – n'apportait pas la garantie que l'on sache indiquer la construction d'un entier N tel que non $P(N)$ soit vraie. Toutefois, dans ces conditions, l'entier n tel que non $P(n)$ est vraie, dont la théorie affirme l'existence via la preuve formelle, semble avoir le statut d'un entier «au-delà» des entiers effectifs : un tel entier est réputé «exister» parce qu'une phrase de la théorie affirmant son existence est démontrée, mais on ne peut l'identifier avec aucun entier que nous sachions construire. Il est donc «plus grand» que tous ceux-ci.

Dans le raisonnement de Reeb, l'entier infiniment grand est introduit par un étrange raisonnement, qui accepte l'existence de la théorie formelle des entiers et envisage même la mythique collection de nombres satisfaisant cette théorie, mais qui, en même temps, débat sur ce que cette collection peut et doit contenir comme éléments, en comparant la notion formelle d'entier à sa notion intuitive. Le résultat

auquel on arrive (que le putatif $\mathbb{N}$ doit contenir un entier infiniment grand) a beau être nécessaire au gré du raisonnement tenu, il ne peut pas être assumé dans le formalisme : si le prédicat «infiniment grand» était disponible dans le formalisme, on obtiendrait un paradoxe en appliquant le schéma d'induction à la classe des entiers finis ou à la classe des entiers infiniment grands. Par exemple, on dirait en substance que l'ensemble des entiers infiniment grands a un plus petit élément α , donc $\alpha - 1$ est fini, mais alors $\alpha - 1$ est aussi, d'où contradiction.

Tous les montages de l'analyse non standard nouvelle, jouant sur un entier infiniment grand sans passer sous les fourches Caudines de $\mathbb{R}$ ou de la théorie des ensembles, consistent à transformer le formalisme de la théorie des entiers pour que disparaissent ces paradoxes. L'usage du principe de récurrence devra, par exemple, être limité à des prédicats non bâtis au moyen de la notion d'entier «fini», «standard» ou «limité». Il résulte alors de tels dispositifs que les entiers infiniment grands apparaissent comme une classe insituable, intervenant sur la droite entière «après» ceux qui sont finis. Ils se distinguent des entiers standard, que l'on interprète désormais comme les «entiers constructifs», disons les entiers que l'on peut construire avec des ordinateurs, en imaginant ceux-ci suffisamment puissants.

L'infiniment grand

L'intention de ces théories non standard est ainsi de saisir l'infiniment grand comme l'excédent insituable au-delà du constructif. Ces théories idéalisent l'intuition des rapports entre entiers constructifs et entiers énormes, inaccessibles, rapports qui sont «concrètement» assez imprévisibles. Comme l'explique Jacques Harthong, les entiers constructibles ou accessibles sont nécessairement définis relativement à un langage de programmation choisi comme logiciel de production ou d'écriture de référence. Ainsi le nombre $10^\wedge (10^\wedge (10^\wedge (10^\wedge (10^\wedge 10)))$, où $a^\wedge b$ désigne a^b , peut être calculé par un programme très simple, long d'une vingtaine de pas sur une calculatrice programmable : ce nombre est néanmoins gigantesque.



Des infinis aux visages différents

La notion d'infini a évolué avec les conceptions de l'analyse non standard. Abraham Robinson (à gauche) se fonde sur les ensembles infinis de la théorie des ensembles et sur l'axiome du choix pour produire un ensemble des hyper-réels dans lequel on trouve, en sus des nombres réels connus, des nombres infiniment grands et infiniment petits.

En reformalisant la théorie des ensembles, Edward Nelson (à droite) introduit la notion d'entité standard : les infiniment petits ou grands existent non pas comme nouveaux objets «ajoutés», mais parce que nous reconnaissons que, parmi les réels, existent nécessairement des éléments insaisissables.

Si Nelson relie l'infini à l'insaisissable et le fini à l'assignable, il revient au mathématicien Georges Reeb (à gauche) de fonder l'analyse non standard sur une sensibilité constructiviste. Autrement dit, à la suite de Reeb, on essaie de rapprocher le concept d'infini du concept d'incalculable. Les entiers standard sont alors les entiers «constructifs», ceux qu'on peut calculer à l'aide d'ordinateurs ou qui apparaissent au fil de l'énumération naïve.



Un nombre de même taille, mais n'ayant pas la même structure simple n'aurait aucune chance d'être programmable dans le même langage (à défaut de la structure, il ne resterait que la ressource d'écrire ses chiffres, mais il faudrait un nombre de symboles supérieur à celui des particules de l'Univers, et un temps dépassant la capacité de survie de tout vivant). La distribution des nombres «accessibles à la calculatrice programmable» est non régulière, tributaire des particularités du matériel symbolique et des opérations de base de la calculatrice considérée. La distribution des nombres qui apparaissent comme constructibles au sens de telle calculatrice programmable est donc clairsemée.

Dans le nouveau point de vue, l'analyse non standard est une systématisation théorique, indépendante de toute machine et de tout langage particuliers, des phénomènes que nous citons en exemple au début de l'article, lorsque nous remarquons que, pour les calculatrices, certains entiers finis jouent le rôle de l'infini. Si l'on comprend bien le concept de nombre fini accessible, constructif, on remarque que la pensée arithmétique contemporaine rend plausible une théorisation qualitative des entiers énormes qui leur donne un statut d'infini. Toutefois, dès que l'on dispose d'un contexte théorique d'arithmétique avec des entiers infiniment grands, on peut introduire un calcul sur des rationnels à dénominateur infiniment grand qui permet de retrouver, dans une certaine mesure, la mathématique du continu à laquelle nous sommes habitués. Alors, nous pourrions

faire de l'analyse non standard en justifiant les infiniment petits et les équations imparfaites des pionniers.

Pour y parvenir, il faut, encore une fois, faire valoir certains objets comme non assignables ou relativement mal assignés, des entiers énormes cette fois. Selon les stratégies intellectuelles, la justification de ce caractère non assignable peut être plus ou moins subtile et utiliser soit la philosophie intuitionniste de Brouwer, soit le théorème d'incomplétude de Gödel, soit un examen précis des possibilités des langages de programmation.

Les conceptions successives de l'analyse non standard ont en commun d'avoir tenté de faire jouer en plus, à côté ou à la place de la notion cantorienne d'infini de totalité, une notion d'indéfini, d'entité non assignable. Cette notion rend jusqu'à un certain point des services techniques comparables à ceux que rendaient les notions infinitésimales des pionniers. En tout cas, elle justifie *a posteriori* un usage intuitif de l'infinitésimal auquel les mathématiciens et les physiciens sont restés attachés.

Jean-Michel SALANSKIS est professeur de philosophie des sciences à l'Université Paris X (Nanterre).

J.-M. SALANSKIS, *Le constructivisme non standard*, Presses du Septentrion, 1999.

M. PANZA et J.-M. SALANSKIS, *L'objectivité mathématique*, Dunod-Masson, 1995.

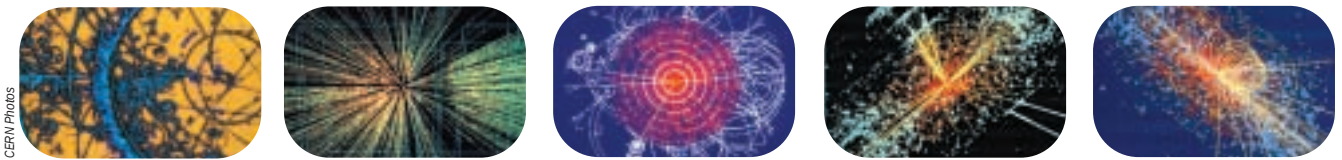
R. LALEMENT, *Logique, réduction, résolution*, Hermès, 1990.

F. DIENER et G. REEB, *Analyse non standard*, Hermann, 1989.

Les géométries non commutatives

DANIELA BIGATTI

Les espaces d'opérateurs de dimensions infinies étudiés à l'aide des géométries non commutatives servent à élaborer de nouvelles théories en physique.



CERN PHOTOS

Les espaces qui nous sont le plus familiers, ceux qui correspondent à notre expérience quotidienne, sont les espaces euclidiens. Au III^e siècle avant notre ère, Euclide définit les notions de point, de droite, de distance et d'angle. Il énonce les «règles», ou postulats, auxquels les objets présents dans ces espaces obéissent. Cette géométrie correspond si bien à notre monde que, pendant plus de 2 000 ans, personne ne pense à la remettre en cause.

Cependant, à partir du XVIII^e siècle, des mathématiciens tels que Hugo Riemann ou Nikolai Lobatchevski s'aperçoivent qu'ils peuvent faire de la géométrie dans des espaces où l'un au moins des postulats d'Euclide n'est plus vrai. Par exemple, Lobatchevski propose de considérer des espaces où, par un point donné, passe une infinité de droites parallèles à une droite donnée (dans les espace euclidiens, il n'y en a qu'une) : il fonde ainsi une «extension» de la géométrie euclidienne, nommée géométrie hyperbolique.

Certaines extensions de la géométrie euclidienne, reposant sur des postulats qui semblent au départ arbitraires, trouvent des applications utiles en physique. Par exemple, Hermann Minkowski (qui fut le professeur d'Albert Einstein) ajoute une dimension supplémentaire aux espaces euclidiens, le temps. Ce nouveau concept, appelé espace-temps, est à la base de la relativité restreinte, dont les applications pratiques sont

nombreuses, notamment en électromagnétisme et en astronomie.

Pour élaborer la théorie de la relativité générale, Einstein utilise une autre de ces extensions : il considère la gravitation, force universelle qui s'applique sur toute chose. Le seul moyen de s'y soustraire est d'y opposer une force d'inertie. Cette dernière se manifeste soit lors des accélérations et des décélérations sur les trajectoires rectilignes, mais qui apparaissent courbées dans l'espace-temps (sur une fusée qui décolle, la force de gravitation est compensée par l'accélération), soit lors d'un mouvement sur une trajectoire courbe (la rotation des stations spatiales sur elles-mêmes crée une force d'inertie, qui joue le rôle de force de gravitation).

La courbure de l'espace-temps

Einstein a ainsi l'idée de modéliser la gravitation comme une courbure de l'espace. Grâce à cette théorie, il prévoit des effets qui n'existent pas dans les espaces euclidiens, mais que l'on observe dans notre monde. Par exemple, un rayon lumineux passant à proximité du Soleil est dévié par la force de gravitation, tandis que, dans un espace euclidien, il se propage en ligne droite. La force de gravitation induit ainsi une courbure de l'espace (voir la figure 1). On le vérifie lors des éclipses de Soleil : les étoiles dont les coordonnées sont connues et dont les rayons passent près du Soleil

(celles qui sont à quelques minutes d'angle du Soleil) semblent légèrement décalées, par rapport à la position qu'elles devraient occuper. Cette observation montre que la lumière est effectivement déviée par le champ de gravitation du Soleil.

Puis les physiciens s'intéressent à la matière aux échelles microscopiques et élaborent la théorie de la mécanique quantique. Dans les espaces euclidiens, il y a une infinité de points sur la droite qui relie deux points donnés, si proches soient ils l'un de l'autre. Ce concept a-t-il un sens en physique ? Pour le savoir, les physiciens essaient de mesurer les positions des particules, indivisibles, qui sont l'équivalent des points dans les espaces euclidiens. Pour mesurer les dimensions d'une pièce, on la compare aux objets qu'elle contient, par exemple à un mètre. Plus la dimension que l'on veut mesurer est grande par rapport à l'objet qui sert de référence, plus la précision est grande (en admettant que l'objet de référence soit indivisible). Ainsi, pour mesurer la position d'une particule, on utilise l'objet de référence le plus petit possible, soit une autre particule. La précision est médiocre. En outre, la mesure, c'est-à-dire l'impact de la particule ou du faisceau de particules censé mesurer la position de la particule à laquelle on s'intéresse, modifie les coordonnées de cette dernière. Ainsi, connaître exactement la position et la vitesse d'une particule est impossible. Cette propriété a été énoncée par Werner Heisenberg.

Des coordonnées non commutatives

La mécanique quantique, notamment la théorie quantique des champs, qui inclut les interactions quantiques et relativistes, ne tient pas compte de la gravitation. Aujourd'hui, on essaie d'élaborer des théories où la gravitation devrait obéir aux lois de la mécanique quantique. Dans ces théories, le principe de Heisenberg s'applique à la gravitation : la courbure de l'espace ne peut pas être complètement déterminée. Dès lors que l'on mesure l'une des coordonnées, on modifie cette courbure, donc les autres coordonnées. Ainsi, si l'on mesure les coordonnées d'un objet dans un certain ordre, par exemple x , y puis z , et qu'ensuite on les mesure dans un ordre différent, par exemple z , x , puis y , on n'obtient pas le même résultat : la mesure des coordonnées ne commute pas. Pour décrire cet état de fait, on utilise les géométries non commutatives.

Les espaces non commutatifs sont construits en découpant un espace de départ en régions, ou classes d'équiva-

lence (dont les points ont une propriété en commun) et en identifiant chacune d'elles à un point du nouvel espace (voir la figure 2). L'opération qui transforme une classe d'équivalence en un point du nouvel espace est dite opération de passage au quotient. Considérons notre monde comme le quotient d'un espace de départ (que l'on peut choisir parmi une multitude d'espaces distincts) : le flou existant sur les coordonnées en physique quantique (correspondant aux différents états possibles d'un système à un instant donné) laisse supposer que l'espace de départ a plus de dimensions que celles que nous connaissons. Par exemple, il pourrait avoir sept dimensions spatiales et trois dimensions temporelles.

En mathématique, on utilise couramment des fonctions, qui assignent à chaque point un nombre, par exemple l'une de ses coordonnées. De même, en physique, la fonction température associe à chaque point de l'atmosphère un nombre qui exprime, en degrés, la température en ce point. La fonction que l'on peut nommer « fonction inter-

action » associe à chaque point un vecteur force à trois composantes, auquel serait soumise une particule supposée ponctuelle placée en ce point. Par définition, les fonctions à valeurs numériques sont commutatives : l'attribution d'un nombre ou d'un vecteur à chaque point est indépendante des autres attributions effectuées. On ne peut donc pas décrire notre monde (en englobant la mécanique quantique et la gravitation, c'est-à-dire en considérant qu'il est non commutatif, comme nous l'avons vu précédemment sur les coordonnées d'une particule) avec des fonctions commutatives.

Puisque l'on ne peut pas utiliser de fonction, on se sert d'opérateurs, bien adaptés, car ils ne commutent pas. Un opérateur est une procédure qui associe à chaque point ou entité de l'espace de départ une autre entité appartenant à l'espace d'arrivée (éventuellement, le même que l'espace de départ). Une rotation de 30 degrés du plan cartésien est un exemple d'opérateur d'un espace dans lui-même. Si un espace est invariant par rotation autour d'un point, on



1. CINQ IMAGES d'une même galaxie (indiquées par les flèches rouges) sont visibles sur cette photographie prise par le Télescope spatial *Hubble*. Cette observation illustre la courbure de l'espace-temps prévue par Einstein dans sa théorie de la relativité générale : les rayons lumineux sont déviés par les amas de galaxies (tel celui qui apparaît en jaune, au centre). Si l'on intègre cette propriété à la mécanique quantique, alors l'ordre dans lequel on mesure les coordonnées d'une particule n'est plus indifférent : la

mesure de chaque coordonnée modifie la courbure de l'espace-temps, donc la valeur des autres coordonnées. Pour modéliser ce phénomène, on a élaboré des géométries non commutatives. Pour choisir celle qui correspond le mieux à notre monde, on confronte les prévisions de chacune d'elles à nos observations. De même qu'Einstein a prévu la courbure de l'espace-temps avant de l'avoir observée, les géométries non commutatives révéleront peut-être des phénomènes insoupçonnés.

NASA, W. Colley et E. Turner, Princeton University ; J. Tyson, Bell Laboratories

exprime ses propriétés de symétrie en affirmant simplement que l'opérateur de rotation transforme l'espace en lui-même. L'intérêt de l'approche opérationnelle réside dans son caractère synthétique.

Même dans un espace euclidien, une telle approche se justifie, car elle évite de considérer les points un à un et permet pourtant de retrouver les mêmes informations que par une approche classique. Prenons l'exemple d'une biopsie : elle consiste à prélever un échantillon en chaque point ou chaque endroit sur lequel on veut obtenir une information. Avec une radiographie, au contraire, on obtient une « synthèse » des données que l'on aurait récoltées en accumulant les biopsies, mais certaines informations sont perdues (la radiographie donne une information ou un niveau de gris pour un segment entier de points alignés dans la direction des rayons X, tandis que les biopsies donnent une information en chaque point). En mathématique, une radiographie peut être considérée comme un opérateur, défini sur le volume étudié (la par-

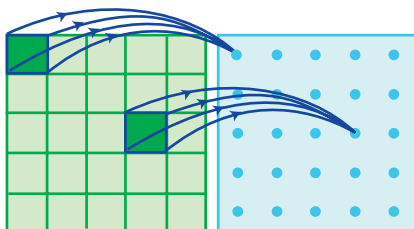
tie du corps radiographié) qui constitue l'espace de départ. L'espace d'arrivée est la plaque photographique. L'opérateur radiographie associe à chaque ensemble de points alignés dans la direction des rayons X, un point de la plaque photographique dont la nuance de gris est définie comme l'intégrale de l'intensité du rayon X qui a traversé le volume étudié (voir la figure 3). Dans ce cas, tous les points alignés dans la direction des rayons X sont « écrasés » en un seul point de la plaque photographique.

Des espaces d'opérateurs de dimension infinie

La mécanique quantique interdit toute description du monde en termes de fonctions et exige d'avoir recours à des opérateurs. Aujourd'hui, pour élaborer de nouvelles théories, les physiciens envisagent notre monde comme le quotient d'un espace de départ, à imaginer. Ils construisent des géométries non commutatives pour étudier des espaces d'opérateurs qui feraient

passer des espaces de départ possibles à notre monde. Ils ont déjà démontré que la partie non commutative de cet espace de départ (qui coexiste avec une partie commutative correspondant aux phénomènes ordinaires que nous connaissons) ne peut pas être décrite avec un espace d'opérateur de dimension finie. Cette restriction est imposée par la théorie de la relativité restreinte. Cependant, ils n'ont que quelques indices sur l'espace de départ et procèdent un peu au hasard pour en choisir un parmi tous les espaces possibles, c'est-à-dire pour choisir la « jauge » : si l'on considère l'opération de passage au quotient comme une projection, pendant laquelle certaines informations présentes dans l'espace de départ sont perdues, la jauge serait l'antiprojection qui ferait passer de notre monde à l'espace de départ.

Comment accéder aux informations perdues, alors que l'on ne connaît même pas leur nature ? Les physiciens sont un peu dans la même situation que quelqu'un qui regarde les photographies d'une maison (une représentation à deux



2. DANS LES NOUVELLES THÉORIES DE PHYSIQUE, notre monde est considéré comme une projection d'un espace de départ sur un espace d'arrivée. De même, une pipe (à trois dimensions) est représentée par une projection sur un plan (la représentation est à deux dimensions ; ci-dessous, *La trahison des images*, de Magritte, 1929). Les géométries non commutatives étudient l'opération dite de passage au quotient, qui transformeraient un espace de départ (*ci-contre en vert*) en notre monde physique (*en bleu*). Chaque classe d'équivalence (où les points ont une propriété commune et schématisée ici par un carré vert) est identifiée à un point de notre monde, où les propriétés de l'espace de départ sont « agglutinées » en un point.





3. L'«OPÉRATEUR RADIOGRAPHIE» associe à chaque ensemble de points alignés dans la direction des rayons X qui traversent l'espace de départ (le volume du corps, à *gauche*), un point de la plaque photographique (à *droite*) ; la nuance de gris est définie comme l'intégrale de l'intensité du rayon X qui a traversé l'espace de départ. Les géométries non commutatives étudient les opérateurs qui font passer d'un espace de départ à notre monde. Lors de la transformation d'un espace de départ vers un espace d'arrivée, des informations sont perdues, de même qu'une partie de l'information sur les organes est perdue lorsque l'on effectue une radiographie : non seulement on ne visualise que les os, mais on passe de structures tridimensionnelles à une photographie bidimensionnelle.

dimensions) sans savoir à quoi ressemble effectivement une maison (c'est-à-dire sans savoir qu'elle existe dans un espace à trois dimensions). Pour décrire la maison, on pourrait croire qu'il faut accumuler le plus de photographies possible, car elles sont toutes différentes. En revanche, si l'on sait que la maison est un objet à trois dimensions, on se rend compte que trois photographies, prises dans les trois directions de l'espace, suffisent pour la décrire.

Les physiciens considèrent que la façon dont ils décrivent le monde avec la théorie quantique des champs est redondante, de même que la description de la maison avec un grand nombre de photographies. En quoi est-elle redondante? Ils ne le sauront que lorsqu'ils auront fait un choix de jauge qui donnera un espace quotient dont les propriétés prédiront des phénomènes existant dans notre Univers, de même que la théorie d'Einstein prédisait la déviation des rayons lumineux par la force de gravitation, hypothèse confirmée ultérieurement.

Les choix de jauges

Pour l'instant, cette approche est une impasse. Les choix de jauge ont abouti à deux types de situations. Dans le premier, les descriptions des espaces de départ envisagés reproduisent les descriptions classiques de notre monde (notamment des interactions électromagnétiques), à quelques variantes près, mais ne prédisent que des phénomènes déjà connus. Ce résultat est aussi peu

surprenant que l'habileté d'un joueur de fléchettes qui vise toujours dans le mille, si ce joueur dessine la cible autour de la fléchette après l'avoir lancée. Dans le deuxième type de situation, les physiciens ont décrit l'espace de départ, dont notre monde serait le quotient, en s'affranchissant de certaines contraintes qui y existent. Leurs théories prédisent alors une foule de phénomènes qui n'existent pas dans notre monde.

Ils devront affiner le choix des contraintes à abandonner, mais ne pourront pas le faire à l'aide des seules géométries non commutatives. Les accélérateurs leur donneront peut-être des indices : à mesure que les énergies employées augmentent, elles révéleront probablement des phénomènes insoupçonnés. Ces phénomènes désigneront alors la géométrie non commutative qui les prévoit comme celle qui fera progresser la physique.

Daniela BIGATTI est mathématicienne à l'Institut Weizmann, en Israël. Elle a récemment travaillé à l'Université catholique de Louvain sur les géométries non commutatives et leurs applications à la théorie des cordes.

Daniela BIGATTI, *Noncommutative spaces in physics and mathematics*, in *Class. Quant. Grav.*, n° 17, pp. 3403-3428, 2000.

A. CONNES, *Noncommutative geometry*, Academic Press, 1995.

A. CONNES, *Geometrie non commutative*, InterÉditions, Paris, 1990.
