

que sur notre propre espèce. Nous sommes obligés de constater que la caractéristique qui semblait nous distinguer du monde animal, nos capacités culturelles, n'est plus si... caractéristique.

Pourtant, les coutumes et les traditions humaines, enrichies et véhiculées par le langage, sont beaucoup plus variées que celles que nous avons observées chez le chimpanzé. Nos connaissances nouvelles des comportements culturels des chimpanzés précisent la compréhension que nous avons de notre unicité, plutôt qu'elle ne la menace.

Au fil des générations, les civilisations humaines ont fait des progrès considérables en accumulant et en améliorant le savoir au fil des générations. Par exemple, le concept de marteau, d'abord un simple caillou, a été modifié et amélioré à plusieurs reprises jusqu'à aujourd'hui, où des marteaux robotisés à commande électronique fonctionnent dans les usines modernes. De même, les chimpanzés qui utilisaient des pierres comme enclumes ont progressé, à Bossou, en apprenant à les caler avec une autre pierre quand le sol n'est pas plan, mais ce comportement est peu répandu et semble bien rudimentaire à côté des progrès accomplis par l'homme.

Les capacités culturelles que nous partageons avec les chimpanzés indiquent aussi que la mentalité qui les sous-tend a une origine commune. Nos capacités ne sont pas apparues du jour au lendemain, mais ont évolué progressivement. Un apprentissage social semblable à celui des chimpanzés aurait probablement suffi pour véhiculer les cultures lithiques de nos ancêtres, qui vivaient il y a deux millions d'années.

On ignore encore si d'autres espèces ont des capacités culturelles proches de celles de l'homme. Les recherches n'ont pas encore été assez poussées pour le savoir, mais quelques indices semblent l'indiquer. Carel van Schaik et ses collègues de l'Université de Duke ont trouvé des orangs-outans à Sumatra qui ont l'habitude d'utiliser au moins deux types d'outils. On n'a jamais vu ce genre de comportement chez les orangs-outans observés pendant des années sur d'autres sites.

Hal Whitehead, de l'Université Dalhousie, et ses collègues commencent à décrire les comportements de certaines populations de baleines qui chantent en différents dialectes et

chassent également de diverses façons. Nous espérons que notre étude des cultures des chimpanzés servira de modèle pour l'étude de ces autres espèces.

Quelles sont les conséquences de nos recherches pour les chimpanzés eux-mêmes? Des populations entières de ces animaux sont en train de disparaître, alors que nous commençons tout juste à mieux les connaître. Depuis un siècle, leur nombre diminue de façon alarmante à cause du braconnage, de l'exploitation forestière, de la construction de routes dans la forêt et, plus récemment, du commerce de la «viande de brousse», c'est-à-dire de la viande d'animaux sauvages (dont celle de chimpanzé), consommée sur place ou exportée. Les animaux sont menacés et, avec eux, tous leurs fascinants comportements culturels.

La richesse culturelle des singes les sauvera peut-être. Des programmes de conservations ont déjà modifié l'attitude de certaines populations humaines locales. Quelques organisations ont commencé à montrer des cassettes vidéo présentant les prouesses cognitives des chimpanzés, et l'attitude de l'homme face au chimpanzé commence à évoluer.

Andrew WHITEN est professeur de psychologie de l'évolution à l'Université de Saint Andrews, en Écosse. Christophe BOESCH est codirecteur de l'Institut Max Planck d'anthropologie de l'évolution et professeur à l'Université de Leipzig.

William C. McGREW, *Chimpanzee Material Culture*, Cambridge University Press, 1992.

A. WHITEN, J. GOODALL, W. McGREW, T. NISHIDA, V. REYNOLDS, C. TUTIN, Y. SUGIYAMA, R.W. BRANGHAM et C. BOESCH, *Cultures in Chimpanzees*, in *Nature*, vol. 399, pp. 682-685, 1999.

Christophe BOESCH et Hedwige BOESCH-ASCHERMANN, *Chimpanzees of the Tai Forest : Behavioral Ecology and Evolution*, Oxford University Press, 2000.

Andrew WHITEN, *Primate Culture and Social Learning*, in *Cognitive Science*, Special issue on primate cognition, vol. 24, pp. 477-508, 2000.

Chimpanzee Cultures Web site: <http://chimp.st-and.ac.uk/cultures/WildChimpanzeeFoundationWebSite>: <http://www.wild-chimps.org>

Variations sur une grille

On observe d'énigmatiques effets d'extinction ou de scintillation en dehors du point de fixation du regard : le cerveau alterne entre des interprétations à faible et à forte résolution.

Dans l'illusion de la grille de Hermann, on distingue de petites taches grises aux intersections des allées blanches qui séparent les carrés noirs (1). Elles disparaissent au point de fixation du regard et reparaissent en périphérie du champ visuel. Les spécialistes de la perception ont repéré dans cette grille d'autres anomalies, moins flagrantes. Par exemple, le blanc des allées semble légèrement plus gris que celui qui borde l'image, et l'on voit parfois des lignes grises s'étendre au centre des allées. Le phénomène est plus fort – toujours en vision périphérique – là où les carrés portent des encoches. Quand on fait subir un quart de tour à la grille, on voit un réseau de lignes sombres, horizontales et verticales, traverser les carrés selon leurs diagonales. Tous ces phénomènes sont également observés en inversant le contraste (carrés blancs sur fond noir) ; les taches ou lignes illusoires présentent alors le contraste inverse de celui décrit plus haut. Ces phénomènes sont tenus pour révélateurs des mécanismes par lesquels les réseaux de neurones calculent le niveau de gris local, et le corrigent en fonction du contraste avec le pourtour.

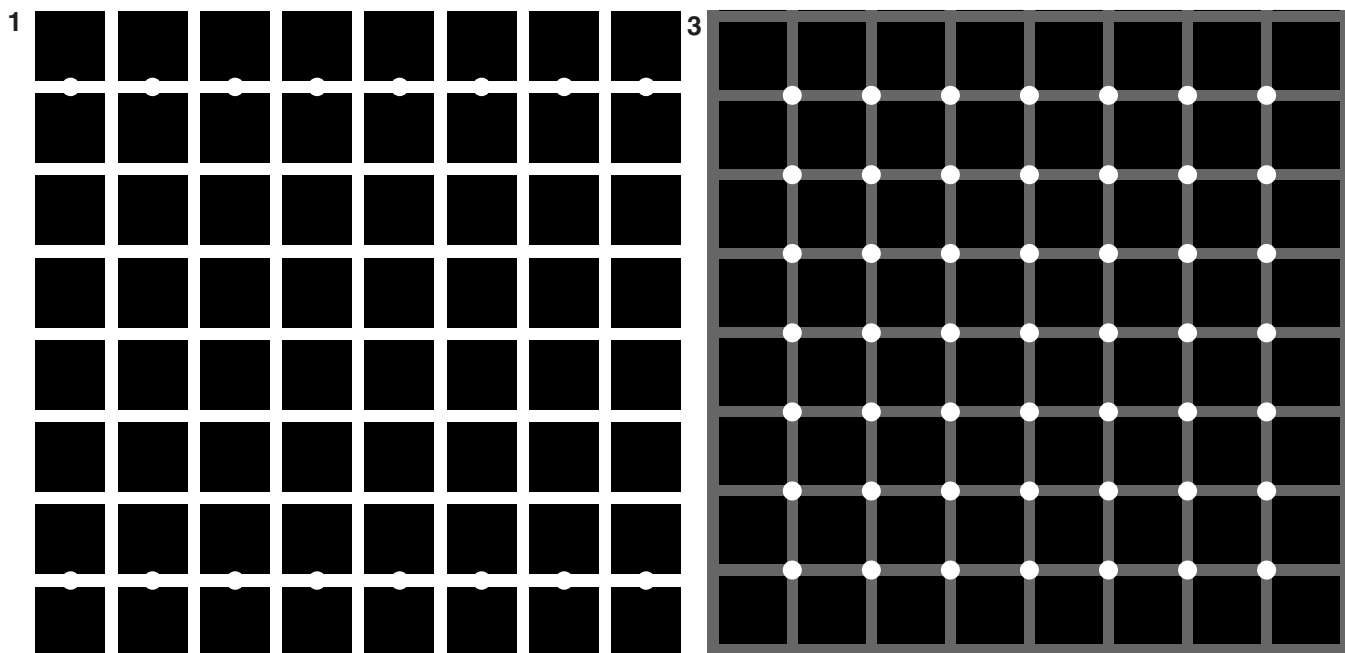
En faisant varier de manière systématique la géométrie de la grille de Hermann, de nouveaux effets se manifestent. En 1985, James Bergen, des Laboratoires RCA à Princeton, a observé une scintillation au croisement des allées, quand la grille est rendue floue (2). L'effet est spectaculaire sur écran, plus difficile à rendre sur papier. La scintillation se manifeste d'abord à l'occasion d'un saut du regard, d'un point à l'autre de l'image. Puis, avec l'habitude, on la voit sans mouvement volontaire. C'est comme si le cerveau oscillait entre deux estimations du niveau de gris aux intersections : l'estimation obtenue à forte résolution, quand on fixe une intersection, et qui la fait voir noire (pour l'image présentée ici), et celle, à faible résolution, de la vision périphérique, où les corrections de contraste local feraient apparaître du blanc, comme sur une grille de Hermann à carrés blancs sur fond noir.

Notons qu'en rendant floue la grille de Hermann, on a rendu grises les allées, et fait apparaître aux intersections des disques plus sombres que les allées. En supprimant le flou, et en réduisant à trois valeurs les niveaux de gris, Michael Schrauf, Bernd Lingelbach et Eugène Wist, de l'Uni-

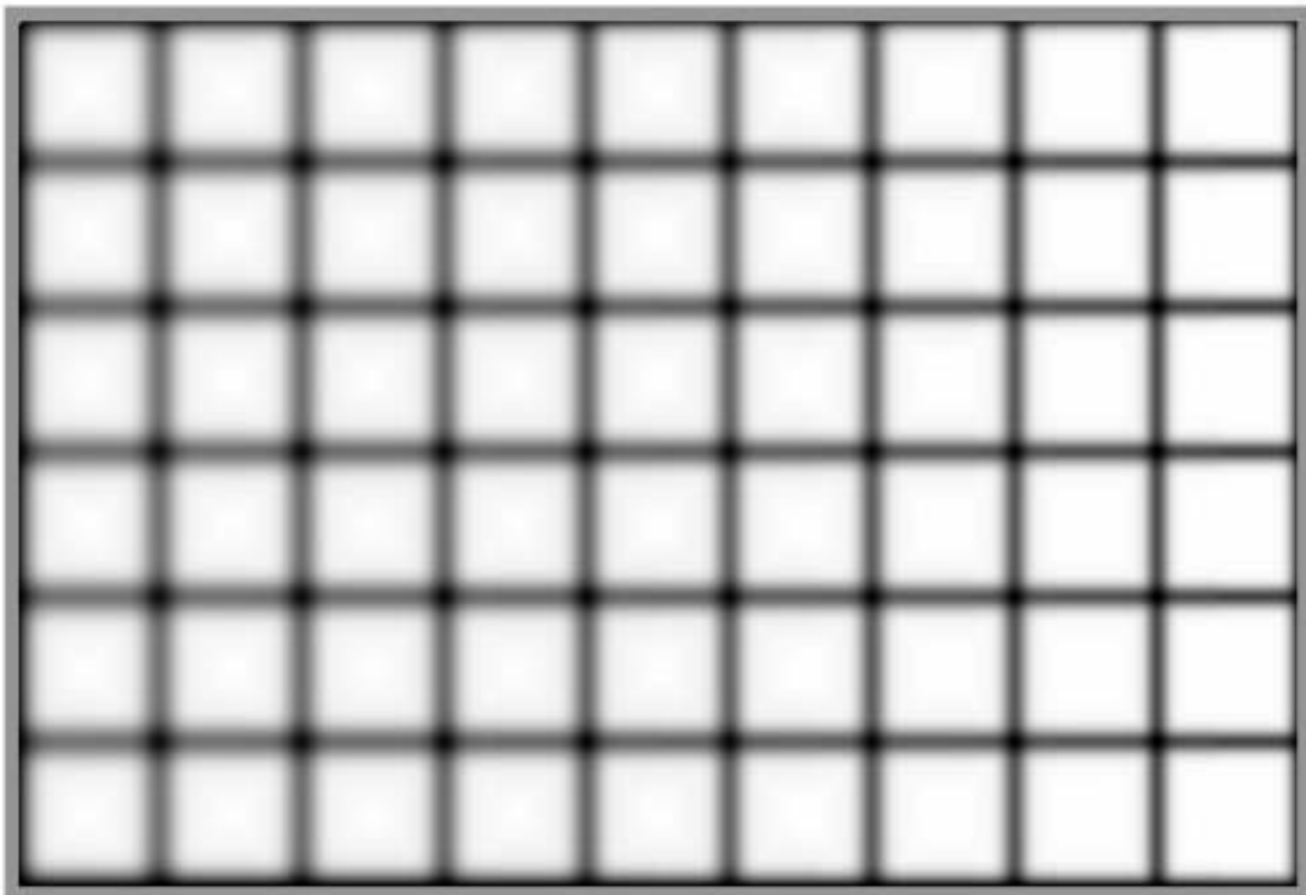
versité Heinrich Heine, à Düsseldorf, ont obtenu une grille de Hermann scintillante, moins spectaculaire, mais plus facile à réaliser que celle de Bergen (3). L'image a été reprise pour illustrer la difficulté d'interprétation des bulletins de vote de Floride ! Prenant cette nouvelle grille pour point de départ et la faisant varier, nous avons mis en évidence avec Kent Stevens, de l'Université d'Oregon à Eugene, un nouvel effet, tout aussi spectaculaire que la scintillation, récemment publié dans la revue *Perception*. La modification essentielle est de cercler avec du blanc les disques noirs, et réduire leur taille. La plupart des disques disparaissent alors de la page. On voit quelques disques là où le regard se pose. Ailleurs, les allées paraissent continues et démunies de disques.

Ce phénomène de disparition – ou d'extinction – est encore plus fort quand les disques sont aux intersections d'un maillage triangulaire (4). L'illusion, dans cette figure, porte sur les grands disques, alors que les petits disques qui sont placés hors des intersections sont vus en totalité. L'illusion d'extinction serait un phénomène attentionnel. Le gris moyen d'un disque noir cerclé de blanc est peu différent de celui de son environnement, aux intersections de trois allées. Il se pourrait qu'à la périphérie du champ visuel, ne soient retenus comme significatifs que les accidents pour lesquels le contraste local dépasserait une certaine valeur. D'où la négligence vis-à-vis des grands disques et, en fin de compte, leur disparition.

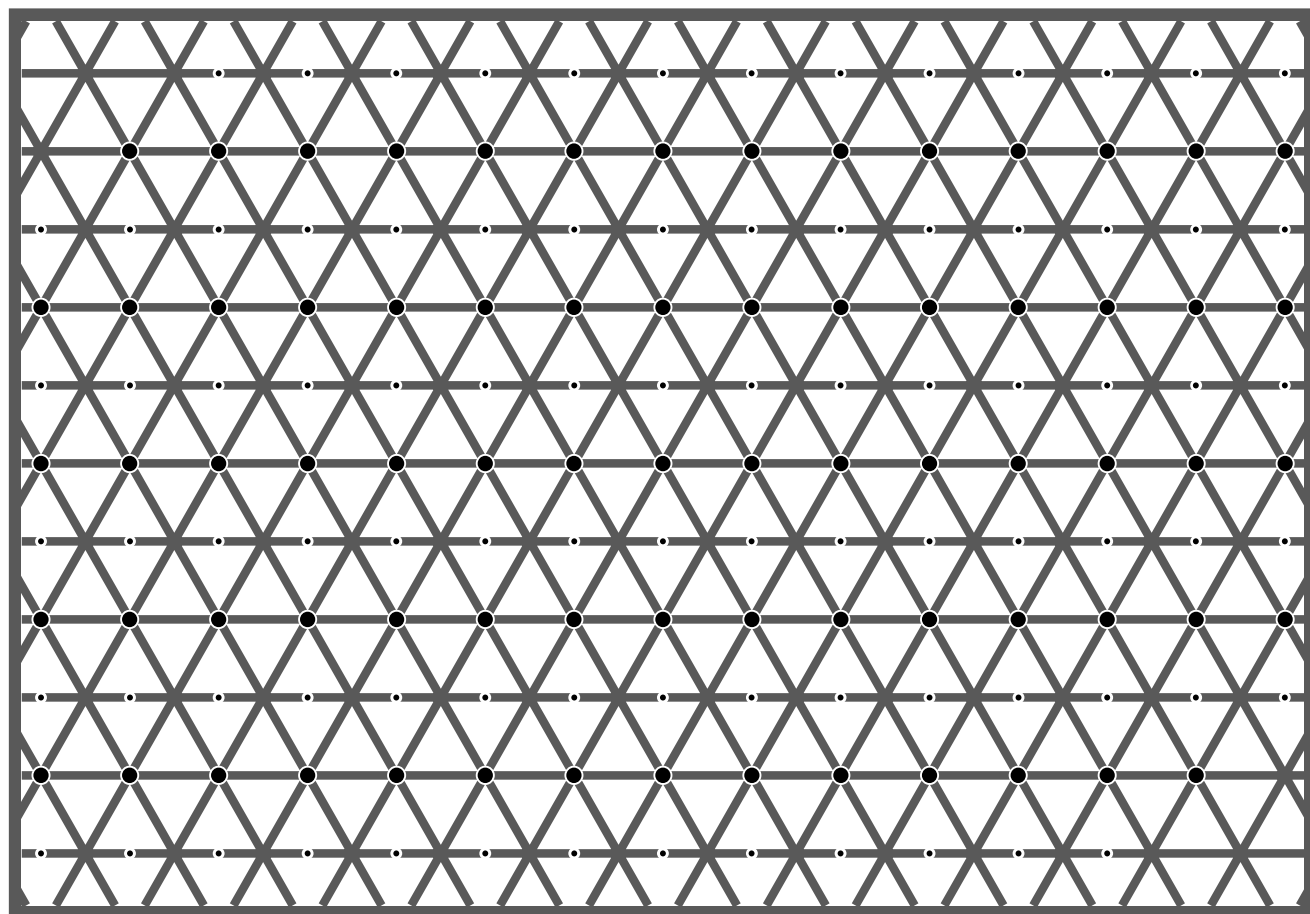
Jacques NINIO,
Laboratoire de physique statistique,
École normale supérieure



2



4



Le jeu des erreurs séduisantes

JEAN-PAUL DELAHAYE

De même qu'il y a de belles escroqueries, il y a aussi de magnifiques erreurs mathématiques : la beauté réside dans l'apparente vérité du faux.



L'art de bien raisonner et l'art de reconnaître les raisonnements fautifs sont deux facettes d'une même médaille : pour raisonner juste, il faut savoir identifier les erreurs et inversement. Certaines démonstrations erronées sont des pièges où on se laisse entraîner, à moins d'être particulièrement attentif ou averti.

Dans la rubrique du mois dernier, nous avons examiné des erreurs mathématiques commises par les meilleurs savants et montré que les mathématiciens, plus peut-être que les autres hommes de sciences, commettent d'énormes bourdes, aux conséquences ennuyeuses ou bénéfiques. L'abstraction, la variété, la complication et la longueur des argumentations mathématiques sont l'explication la plus simple de ce phénomène attesté par les volumineux errata des manuels de mathématiques.

Ce mois-ci, nous allons examiner une série de raisonnements pièges conduisant, par des cheminements d'une rigueur en apparence irréprochable, à d'évidentes absurdités. Les exemples que nous citerons ne proviennent pas d'études historiques, mais sont des pièges conçus dans des buts didactiques ou tout simplement des erreurs commises par des élèves.

SOPHISMES ANTIQUES

L'étude des raisonnements fautifs est antique : Aristote, dans son traité des Réfutations des arguments sophistiques conçut, le premier, l'art d'identifier et de démontrer les arguments erronés que les philosophes sophistes professaient pour ridiculiser leurs adversaires. Voici un exemple des raisonnements qu'analyse Aristote :

Est-ce que le malade qui est guéri se trouve en bonne santé? Oui? Donc le malade est en bonne santé!

Qu'est-ce qui cloche?

Solution : Aristote explique que l'absurdité vient d'une ambiguïté dans le sens

d'un mot. Le mot malade signifie, selon les contextes, personne qui est malade en ce moment (exemple : le malade souffre) ou personne qui a été malade (exemple : quand le traitement a fait son effet, le malade rentre chez lui). La conclusion du sophisme semble paradoxale, car on y utilise le mot malade pour désigner une personne malade (premier sens de malade), alors que, dans la question initiale, qui appelle un oui, c'est le second sens qui est utilisé.

Voici d'autres exemples de sophismes proposés par Aristote que je vous laisse décortiquer vous-même.

Est-ce qu'il faut décider en faveur de celui qui dit des choses justes ou des choses injustes? En faveur de celui qui dit des choses justes. Pourtant il est juste que celui qui a subi une injustice dise d'une manière complète ce dont il a souffert ; or il a souffert de choses injustes. Il nous faut donc décider en faveur de celui qui dit des choses injustes.

Et encore : *Si le non-être est objet d'opinion, le non-être est.*

La classification des erreurs proposée par Aristote expose l'art de la déduction et la manière de ridiculiser un adversaire en l'amenant à se contredire, ce qui est utile dans les débats politiques. Il s'agit toutefois essentiellement de jeux linguistiques, le plus souvent élémentaires. Dès qu'on entre dans le domaine mathématique, des pièges de natures bien différentes surgissent.

ÉVIDENT... MAIS FAUX

Parfois les pièges mentaux sont si simples qu'un petit effort permet à chacun de se corriger. Les exemples suivants, destinés aux enfants, sont dus à E. Northrop.

Une horloge sonne 6 heures en 5 secondes. Combien de temps lui faut-il pour sonner midi?

Non, la réponse n'est pas 10 secondes. Une bouteille et son bouchon coûtent

ensemble 110 F. La bouteille coûte 100 F de plus que le bouchon. Combien coûte la bouteille?

Non, la réponse n'est pas 10 F.

Une grenouille est au fond d'un puits de 30 mètres. En une heure, elle gravit 3 mètres, puis glisse de 2 mètres. Combien d'heures lui faut-il pour sortir?

Non, la réponse n'est pas 30 heures.

Un TGV part de Paris pour Orléans en même temps qu'un omnibus part d'Orléans pour Paris. Le TGV roule à 250 km à l'heure, l'omnibus roule à 120 km à l'heure. Lequel sera le plus loin de Paris quand ils se croiseront?

Non, la réponse n'est pas le TGV.

L'exemple suivant est très classique, mais n'a pas piégé que des enfants : *Le père de Toto a trois fils : Pim, Pam et ...? Non, la réponse n'est pas Poum.*

RETOUR VERS LE PASSÉ

Dans le numéro de l'été 1997 d'Adobe Magazine, Apple faisait la publicité pour son ordinateur Power Mac en indiquant : *En fait, Adobe Photoshop fonctionne 50 pour cent plus rapidement avec un Power Mac. Cela se traduit par un gain de 50 pour cent sur le temps que vous passez devant votre écran à attendre que votre machine retouche les photos, transforme les images ou applique des filtres.*

Fonctionner 50 pour cent plus rapidement signifie que, pendant un temps donné, par exemple 10 secondes, le Power Mac traitera 50 pour cent de segments ou de surface d'image en plus que le modèle d'ordinateur précédent. Si ce que dit la publicité est vrai, alors se pose le problème suivant : avec un

1. Dans son autoportrait, Chagall, pour attirer l'attention, dessine plusieurs doigts à ses mains. L'erreur, voulue, illustre l'habileté du peintre polydactyle. En mathématiques, les erreurs ne sont pas destinées à attirer l'attention sur les qualités du mathématicien.



Autoportrait aux couleurs de la France, Marc Chagall, 1944. © Adagp, Paris 2001. Stedelijk Museum Amsterdam on loan from Netherlands Institute for Cultural Heritage

2. CLASSIFICATION DES ERREURS

Tout le monde commet des erreurs, et tout auteur de livre de mathématiques est étonné quand, malgré le soin qu'il prend à relire les épreuves, il découvre, quelques mois après la parution, une multitude d'imperfections dans les exemplaires mis en vente. Les longs *errata* étonnent les non-mathématiciens et ont parfois donné l'impression que les auteurs des livres de mathématiques ne sont pas sérieux ! La difficulté d'écrire un texte mathématique sans erreur provient de ce que celles-ci peuvent être de natures très différentes et qu'en effectuant une relecture il est impossible de penser simultanément à tout. Outre les fautes de « typo », d'orthographe et de grammaire qui sévissent en mathématiques comme partout, un document de mathématiques peut contenir une grande variété d'erreurs de gravité variée. Par ordre, de gravité voici quelques erreurs typiques.

– L'erreur d'inattention faite au dernier moment, qui conduit à écrire une formule fautive ou un énoncé faux, car on recopie mal son brouillon (qui, lui, était correct) : avec un peu de chance le lecteur attentif « rectifiera de lui-même ». Dès qu'on s'en aperçoit, l'erreur peut être corrigée.

– Plus ennuyeuse, l'erreur vraie de calcul ou de raisonnement conduit à une affirmation fautive dont on n'utilise qu'une partie, elle, vérifiée. Par exemple, on sait que $x \leq 0$ et l'on démontre que $A = 2x$, ce qu'on utilise ensuite pour conclure que $A \leq 0$. Le calcul exact devait conduire à $A = x$ qui, lui aussi, a pour conséquence que $A \leq 0$; l'erreur s'est effacée d'elle-même et n'a pas laissé de trace. Un célèbre cas de deux erreurs s'annulant l'une l'autre est dû à l'astronome anglais Edmond Halley lorsqu'il calcula l'orbite de la comète de 1682 qui porte son nom. Sa double erreur fut notée lorsque Alexis Clairaut, Joseph Lalande et Nicole Lepaute reprirent ses calculs en 1757 pour obtenir une plus grande précision (ce qui leur permit finalement de prédire le retour de la comète avec une exactitude prouvant que la mécanique newtonienne « voit l'avenir » et que la loi d'attraction universelle en $1/r^2$ ne nécessite pas de terme correcteur).

– Plus grave est l'erreur provenant d'un oubli qui conduit à une démonstration trouée : un cas n'a pas été traité. S'il l'avait été, on aurait pu conclure comme dans les autres cas, et donc l'oubli est à nouveau sans gravité. Le raisonnement est faux, mais se complète sans trop de mal. Bien sûr, le mathématicien qui commet souvent cette faute risque un jour de laisser un trou qui ne se rebouchera pas !

– Plus grave encore est l'erreur réparable, mais qui demande un effort important pour être effacée. C'est ce type d'erreur que commit Andrew Wiles quand il proposa sa première démonstration du théorème de Fermat en 1993 (annoncée en première page du *New York Times* du 24 juin). L'effort pour réparer exigera plus d'un an de travail et la participation de certains des collègues de Wiles. Dans le film réalisé par Simon Singh sur cette démonstration, on voit Wiles sur le point de pleurer à l'évocation de la pénible période d'incertitude qu'il vécut à la recherche du contournement de son erreur.

– Une autre catégorie d'erreurs plus graves est celle qui conduit à un énoncé irrémédiablement faux, mais qu'on peut modifier un peu pour en tirer un énoncé vrai. C'est ce type d'erreur qui fut commis par Cauchy quand il crut que toute limite de fonctions continues est continue. Ce type d'erreur est généralement très instructif, car il permet d'isoler une condition oubliée qui n'allait pas de soi et qui parfois fait découvrir un concept important.

– Le dernier cas d'erreur est bien sûr celui qui met tout définitivement par terre, sans qu'aucune leçon puisse en être tirée.

ordinateur procurant un gain de 150 pour cent (ce qui est largement atteint par les machines actuelles), est-ce que vous gagnez 150 pour cent du temps ? Auquel cas, lorsque vous utilisez votre ordinateur, vous retournez dans le passé !

Où est l'erreur du texte publicitaire ?
Solution : Il s'agit là d'un piège classique dans lequel tout le monde tombe un jour : une hausse de X pour cent n'est pas l'inverse d'une baisse de X pour cent et réciproquement. Cette erreur est à rapprocher de celle, tout aussi courante, qui

consiste à croire qu'une hausse de Y pour cent suivie d'une hausse de Z pour cent est équivalente à une hausse de $(Y + Z)$ pour cent.

En réalité, une hausse de X pour cent fait passer de A à $A(1 + X/100)$ et donc, pour la compenser, il faut passer de $A(1 + X/100)$ à A , ce qui correspond à une multiplication par $1/(1 + X/100)$ et non pas à une multiplication par $(1 - X/100)$. Dans notre exemple, une hausse de 50 pour cent (qui fait passer de 100 à 150) est compensée par une

baisse de 33,33 pour cent (qui fait passer de 150 à 100). La publicité aurait dû dire que le temps économisé était de 33,33 pour cent. Avec un ordinateur procurant un gain de 150 pour cent, le gain est 60 pour cent, et non pas un voyage dans le temps passé.

Une erreur analogue à celle d'*Apple* a été commise dans les années 1990 par *France Télécom*, qui utilisait un slogan stupide : « Après 18 heures, le téléphone est 30 pour cent moins cher, soit 30 pour cent de temps en plus. » Ce slogan était pourtant moins favorable que la réalité, car si, après 18h, le téléphone est 30 pour cent moins cher, alors le temps en plus est de 42,8 pour cent (ce qui vous coûtait 100 F vous coûte maintenant 70 F, et donc, si vous utilisez les 30 F gagnés pour prolonger votre coup de téléphone, vous disposez d'une durée supplémentaire de 30/70 de votre coup de téléphone initial, soit 42,85 pour cent). Reconnaissons toutefois que le slogan juste aurait été incongru et incompréhensible pour la majorité des gens.

Venons-en maintenant à des erreurs de raisonnement faites dans le cours d'une démonstration, si possible brève.

EXISTENCE DANGEREUSE

Théorème : 1 est le plus grand entier.

Démonstration. Soit A le plus grand entier. Pour tout entier A supérieur à 1, le carré (entier) A^2 de cet entier A , est plus grand que A , donc A n'est pas le plus grand entier. Sauf si A est égal à 1, le carré de 1 étant égal à 1. Donc 1 est le plus grand entier.

Où est l'erreur ?

Solution : Supposer l'existence d'objets qui n'existent pas conduit à la catastrophe, et c'est ce qu'illustre l'exemple. Dans le raisonnement, nous faisons l'hypothèse qu'il existe un plus grand entier, ce qui est faux. Puisque, à part cela, le raisonnement est entièrement correct et aboutit à une contradiction, il montre rigoureusement qu'il n'existe pas de plus grand entier.

MAIS QU'AI-JE FAIT D'INTERDIT ?

Considérons l'équation :

$$(x + 5)/(x - 7) - 5 = (4x - 40)/(13 - x).$$

Elle peut s'écrire :

$$[x + 5 - 5(x - 7)]/(x - 7) = (4x - 40)/(13 - x),$$

ou encore :

$$(4x - 40)/(7 - x) = (4x - 40)/(13 - x).$$

Les numérateurs étant égaux, les dénominateurs le sont aussi : donc $7 - x$ est égal à $13 - x$, et donc $7 = 13$.

Où est l'erreur ?

Solution : Le piège est subtil. L'affirmation : « les numérateurs étant égaux,

les dénominateurs le sont aussi» n'est vraie que si les numérateurs ne sont pas nuls : par exemple, $0/12 = 0/7$, et 12 n'est pas égal à 7!

Dans notre exemple, $4x - 40$ est nul, car x n'est pas n'importe quel nombre, mais la solution de l'équation de départ, qui justement est 10, ce que l'on vérifie par les calculs suivants.

L'équation de départ donne :

$(13 - x)(4x - 40) = (7 - x)(4x - 40)$,
ce qui, après regroupement dans le premier membre, conduit à $24(x - 10) = 0$,
c'est-à-dire $x = 10$.

PLUS QUE LUI-MÊME?

Prouvons que tout nombre est plus grand que lui-même! Soit a et b deux nombres positifs tels que : $a > b$.

On multiplie les deux nombres par b : $ab > b^2$, on soustrait a^2 de chaque côté et l'on factorise : $a(b-a) > (b-a)(b+a)$.

On divise maintenant par $(b-a)$, qui n'est pas nul, et l'on obtient : $a > b+a$. Puisque b est positif, a est donc plus grand que lui-même et même plus grand que tout nombre plus grand que lui-même. Où est l'erreur?

Solution : Simplement dans le fait (à ne jamais oublier) que, lorsqu'on multiplie ou divise les deux membres d'une inégalité, celle-ci n'est préservée que si ce par quoi on multiplie ou divise est positif. Si ce par quoi on multiplie ou divise est négatif, l'inégalité est inversée. Ici, en divisant par $(b-a)$, qui est négatif, on obtient : $a < b+a$, ce qui n'est nullement paradoxal.

TROIS FOIS RIEN

L'exemple suivant est encore élémentaire, mais, à la fin d'un cours, il ne manquera pas de faire son petit effet.

On démontre que $1+2+3+\dots+n = n(n+1)/2$ par la méthode que l'on veut (par exemple en écrivant $1+2+3+\dots+n$, et juste en dessous $n+ \dots+3+2+1$).

L'égalité étant vraie pour tout n , écrivons-la pour $n-1$:

$$1 + 2 + 3 + \dots + (n-1) = (n-1) n/2.$$

Ajoutons 1 de chaque côté de cette égalité :

$$1 + 2 + 3 + \dots + (n-1) + 1 = (n-1) n/2 + 1.$$

Arrangeons tout cela en sachant que $(n-1) + 1 = 0n$ obtient :

$$1 + 2 + 3 + \dots + n = (n-1) n/2 + 1.$$

On a donc : $n(n+1)/2 = (n-1) n/2 + 1$ ce qui donne : $n^2 + n = n^2 - n + 2$, et finalement $n=1$. Tout entier est égal à 1!

Que doivent vous dire les élèves endormis que la conclusion étrange a peut-être sortis de leur torpeur?

Solution : Si $(n-1) + 1 = n$ est vrai, il n'est pas exact que $1 + 2 + 3 + \dots + (n-1) + 1 = 1 + 2 + 3 + \dots + n$, car, l'avant-dernier

3. LES PAGES D'ERRATA DES LIVRES DE MATHÉMATIQUES

Les livres de mathématiques ont une particularité remarquable : ils sont bourrés de fautes. Les auteurs consciencieux collectent les erreurs que leur signalent les lecteurs et publient alors des «errata» qui sont parfois étonnants.

Chaque chapitre du traité général de mathématiques de Bourbaki était lu et relu par les membres du collectif qui se cachait sous le nom de Bourbaki (voir *Bourbaki*, par Maurice Mashaal, collection *Les génies de la science*). Après un tel passage au crible, on s'attend à ce qu'aucune erreur ne subsiste. Les fascicules d'*errata* publiés par le vrai-faux Général Bourbaki montrent que c'est loin d'être le cas. On a reproduit ci-dessous les premières pages du fascicule d'*errata* numéro 13 (il y en a donc déjà eu 12 auparavant!), qui comporte 7 pages.

Aujourd'hui l'Internet rend le travail plus simple. Outre qu'il permet à l'auteur d'être rapidement informé de ses erreurs par courrier électronique, il permet aussi la publication de volumineux textes d'*errata* qu'aucun éditeur n'aurait accepté d'imprimer, de peur d'y perdre tout crédit. À titre d'exemple, les *errata* de l'*Encyclopédie concise des mathématiques* d'Éric Weisstein (*Concise Encyclopedia of Mathematics*, CRC Press, 1999) comportent 69 pages bien serrées (on n'ose pas penser ce qu'auraient été les errata si l'*Encyclopédie* n'avait pas été «concise»). Son adresse est : www.treasure-troves.com/buy/math/book/errata/

Donald Knuth, le mathématicien créateur du traitement de texte *Tex* (utilisé par la communauté mathématique) et auteur du *Traité de référence sur les mathématiques de la programmation* (*The Art of Computer Programming*) a lui aussi mis à la disposition des lecteurs une série de fascicules d'*errata* qui comportent 80 pages pour le volume 1, 142 pages pour le volume 2, 109 pour le volume 3. Adresse internet : www-cs-faculty.stanford.edu/~knuth/taocp.html

Bon joueur, D. Knuth offre 2,56 dollars pour chaque nouvelle erreur qu'on lui signalera. (Il n'explique pas pourquoi 2,56, mais, n'en doutons pas cela résulte d'un savant calcul – exact espérons-le – rattaché à ses droits d'auteur).



Théorie des ensembles

CHAPITRE III

Ensembles ordonnés. Cardinaux. Nombres entiers

ÉLÉMENTS DE MATHÉMATIQUE, FASCICULE XX TROISIÈME ÉDITION

- P. 40, ligne 9 du haut, au lieu de : $r \in E$, lire : $r \in E$.
P. 77, ligne 15 du bas, au lieu de : $(E_\alpha)_{\alpha \in I}$, lire : $(E_\alpha)_{\alpha \in I}$. Ligne 3 du bas, au lieu de : *ordonné*, lire : *préordonné*.
P. 82, ligne 9 du haut, au lieu de : de E_α dans E_β , lire : de F_α dans F_β .
P. 83, ligne 6 du haut, au lieu de : prop. 2, lire : prop. 1. Ligne 15 du bas, au lieu de : $\lambda \geq \mu$, lire : $\lambda \leq \mu$.
P. 86, ligne 13 du haut, au lieu de (ii), lire (ii').
P. 90, ligne 6 du haut, au lieu de : $f_{\alpha\beta}(u)$, lire : $f_{\beta\alpha}(u)$.
P. 95, ligne 7 du haut, au lieu de : $f_{\beta\alpha}(u_\alpha(x_\alpha))$, lire : $f_{\beta\alpha}(u_\alpha(x_\alpha))$. Ligne 8 du haut, au lieu de : $u_\alpha(f_{\beta\alpha}(x_\alpha))$, lire : $u_\beta(f_{\beta\alpha}(x_\alpha))$.
P. 96, ligne 8 du bas, au lieu de : $\lim_{\alpha} E_\alpha^1$, lire : $\lim_{\alpha} E_\alpha^1$. Ligne 1 du bas, au lieu de F_α^1 , lire F_α^1 .
P. 97, ligne 14 du haut, au lieu de g^α , lire g_α^1 (deux fois). Ligne 5 du bas, au lieu de : $\lim_{\alpha} (\lim_{\beta} u_\beta^1)$, lire : $\lim_{\alpha} (\lim_{\beta} u_\beta^1)$.
P. 103, ligne 14 du bas, au lieu de A, lire X; au lieu de B, lire Y.

terme (non écrit) à gauche est $n-2$, alors qu'à droite c'est $n-1$. Les trois petits points, selon les contextes, ne désignent pas la même chose.

PEUT-ON ÉCRIRE $\sqrt{-1}$?

Les nombres complexes sont définis à partir de l'idée que -1 possède une racine carrée que l'on considère comme un nombre nouveau et que l'on note souvent par la lettre i . Par définition, $i^2 = -1$. Pour plus de commodité on peut noter ce nombre supplémentaire par $\sqrt{-1}$.

On obtient alors de drôles de choses, comme cette suite d'égalités :

$$\begin{aligned} -1 &= \sqrt{-1} \sqrt{-1} = \sqrt{(-1) \cdot (-1)} = \sqrt{1} = 1, \\ \text{ou ce raisonnement :} \\ \sqrt{-1} &= \sqrt{-1}, \text{ donc } \sqrt{1/(-1)} = \sqrt{(-1)/(1)}, \text{ donc} \\ \sqrt{1x/1} &= \sqrt{-1x/-1}, \text{ donc } 1 = -1. \end{aligned}$$

Comment expliquer cela? Les nombres complexes sont-ils contradictoires, et les mathématiciens qui les utilisent sont-ils nécessairement conduits à écrire des choses absurdes?

Solution : Les nombres complexes ne sont pas plus contradictoires que les nombres réels (cela se démontre en faisant la

$b - 1$, d'où l'on déduit que $a = b$, ce que nous souhaitons.

Maintenant, soit un entier positif ou nul quelconque b . Soit n le maximum de 0 et b (bien sûr $n = b$). Le raisonnement précédent a montré que 0 et b étaient égaux, donc $b = 0$. En conclusion, tous les entiers positifs ou nuls sont nuls.

Si tel était le cas, là aussi les entiers seraient majorés par 0! Où est l'erreur?

Solution : Simplement dans le fait que, même si a et b sont des entiers positifs ou nuls, $a - 1$ et $b - 1$ ne le sont pas nécessairement (car $0 - 1 = -1$). On ne peut donc pas appliquer l'hypothèse de récurrence au couple $a - 1, b - 1$.

DÉSOLANTES DÉRIVÉES

Calculons la dérivée de x^3 de deux façons différentes :

(a) Par la formule habituelle $(x^n)' = nx^{n-1}$. Pour tout x , on a donc $(x^3)' = 3x^2$

(b) Par le raisonnement suivant. Pour x entier, on écrit :

$x^3 = x^2 + x^2 + x^2 + \dots + x^2$ (la somme comporte x fois le terme x^2)

On dérive de chaque côté :

$(x^3)' = (x^2)' + (x^2)' + (x^2)' + \dots + (x^2)'$ (la dérivée d'une somme est la somme des dérivées).

$(x^3)' = 2x + 2x + 2x + \dots + 2x$ (application de la formule générale rappelée plus haut).

$(x^3)' = 2x^2$ (car, sur la ligne précédente il y a x fois le terme $2x$, ce qui permet de remplacer le second membre par $2x^2$).

Pour tout x entier, on a donc :

$(x^3)' = 3x^2 = 2x^2$.

En prenant pour x un entier non nul et en simplifiant par x^2 , on a : $2 = 3$.

Où est l'erreur?

Solution : Le raisonnement (b) suppose que x est entier (sinon, on ne pourrait pas écrire : $x^3 = x^2 + x^2 + x^2 + \dots + x^2$). Appelons n l'entier auquel on s'intéresse.

Les fonctions $x \rightarrow x^3$ et $x \rightarrow x^2 + x^2 + x^2 + \dots + x^2$ (n fois le terme x^2) ont effectivement la même valeur au point $x = n$, et cette valeur est n^3 . Mais, autour de la valeur n , ces deux fonctions n'ont plus les mêmes valeurs, car : $(n+\epsilon)^3 \neq (n+\epsilon)^2 + (n+\epsilon)^2 + \dots + (n+\epsilon)^2 = n(n+\epsilon)^2$. Puisque les deux fonctions sont différentes, il est sans surprise que leur dérivée le soit.

L'endroit où se produit l'erreur est lorsqu'on écrit «on dérive de chaque côté» car, l'égalité $x^3 = x^2 + x^2 + x^2 + \dots + x^2$ désigne une égalité valable entre deux fonctions différentes en un point précis, et non pas une égalité entre deux fonctions sur tout leur ensemble de définition. L'erreur est en fait la même que celle qu'on commettrait en disant :

$\sin 45^\circ = \cos 45^\circ = \sqrt{2}/2$, et donc que $(\sin 45^\circ)' = (\cos 45^\circ)'$, ce qui est absurde car en utilisant que $\sin x' = \cos x$ et que $\cos x' = -\sin x$, on en déduirait que $\cos 45^\circ = -\sin 45^\circ$, soit $\sqrt{2}/2 = -\sqrt{2}/2$.

5. LA PLUS BELLE PROPRIÉTÉ DE π

L'importance mathématique du nombre π n'est plus à démontrer. Pourtant la propriété suivante de π ne manquera pas de vous étonner : toute série converge vers π .

En effet, soit une série quelconque (par prudence, vous pouvez considérer une série absolument convergente) :

$a_1 + a_2 + a_3 + a_4 + a_5 + a_6 + a_7 + \dots$

Les égalités (finies) suivantes sont immédiates :

$a_1 = \pi + (a_1 - \pi)$

$a_2 = -(a_1 - \pi) + (a_1 + a_2 - \pi)$

$a_3 = -(a_1 + a_2 - \pi) + (a_1 + a_2 + a_3 - \pi)$

$a_4 = -(a_1 + a_2 + a_3 - \pi) + (a_1 + a_2 + a_3 + a_4 - \pi)$. Etc.

En ajoutant toutes ces égalités, nous obtenons :

$a_1 + a_2 + a_3 + a_4 + a_5 + a_6 + a_7 + \dots = \pi + (a_1 - \pi) - (a_1 - \pi) + (a_1 + a_2 - \pi) - (a_1 + a_2 - \pi) + (a_1 + a_2 + a_3 - \pi) - (a_1 + a_2 + a_3 - \pi) + (a_1 + a_2 + a_3 + a_4 - \pi) - \dots$

Toutes les parenthèses se simplifient deux à deux et disparaissent. Donc :

$a_1 + a_2 + a_3 + a_4 + a_5 + a_6 + a_7 + \dots = \pi$.

Pensez-vous vraiment que cela soit normal?

Solution

Non, bien sûr, et d'ailleurs le même raisonnement, s'il était juste, s'appliquerait à n'importe quelle constante. Le raisonnement est faux, car l'expression obtenue au second membre après l'addition infinie est une série qui n'est pas même convergente (ni les parenthèses, ni leurs composants ne tendent vers zéro en général). Or avec une série de nombres réels qui n'est pas convergente, toute manipulation, même les simplifications les plus simples, est dangereuse et susceptible de nous entraîner dans l'absurde. Si vous en doutez, considérez l'exemple classique suivant :

$1 - 1 + 1 - 1 + 1 - 1 + 1 - 1 + \dots = (1 - 1) + (1 - 1) + (1 - 1) + (1 - 1) + (1 - 1) + \dots = 0$
 $1 - 1 + 1 - 1 + 1 - 1 + 1 - 1 + \dots = 1 + (-1 + 1) + (-1 + 1) + (-1 + 1) + \dots = 1$

SOMMES INFINIES

En analyse, on démontre que la série : $\ln(2) = 1 - 1/2 + 1/3 - 1/4 + 1/5 - 1/6 + \dots$ est une série convergente, c'est-à-dire qu'en prenant 1, puis 2, puis 3, ... termes de la série, on obtient des nombres qui s'approchent de plus en plus d'une constante, qui est ici le logarithme népérien de 2 : $\ln 2 = 0,6931\dots$

Groupons les termes positifs d'une part et les termes négatifs de l'autre :

$\ln 2 = (1 + 1/3 + 1/5 + \dots) - (1/2 + 1/4 + 1/6 + \dots)$.

Par ailleurs nous avons certainement :

$0 = (1/2 + 1/4 + 1/6 + \dots) - (1/2 + 1/4 + 1/6 + \dots)$.

ce qui, en ajoutant membre à membre avec l'égalité au-dessus, donne :

$\ln 2 = (1 + 1/2 + 1/3 + 1/4 + 1/5 + 1/6 + \dots) - 2(1/2 + 1/4 + 1/6 + \dots)$,

c'est-à-dire :

$\ln 2 = (1 + 1/2 + 1/3 + 1/4 + 1/5 + 1/6 + \dots) - (1 + 1/2 + 1/3 + 1/4 + 1/5 + 1/6 + \dots) = 0$,

Donc, $\ln 2 = 0$.

D'où vient cette absurdité?

Solution : L'absurdité provient du fait général qu'on ne peut manipuler les séries sans précaution et qu'en particulier on ne peut changer l'ordre des termes sans risque. Il existe deux sortes de séries convergentes :
 - celles (dénommées «absolument convergentes») dont la série des valeurs absolues est convergente,
 - celles (dénommées «semi-convergentes») dont la série des valeurs absolues est infinie. C'est le cas de :

$\ln(2) = 1 - 1/2 + 1/3 - 1/4 + 1/5 - 1/6 + \dots$

Pour les premières, on peut changer l'ordre des termes, additionner termes à termes, etc, comme nous l'avons fait ici. Pour les autres, Riemann a démontré au contraire que changer l'ordre des termes pouvait conduire à n'importe quelle valeur entre $-\infty$ et $+\infty$, et en fait toutes sortes d'absurdités peuvent être déduites de la manipulation sans précaution de séries semi-convergentes. L'exemple proposé est l'un des plus brefs qui le montre.

ARISTOTE, Organon VI, Réfutation des arguments sophistiques. Traduction J. Tricot. Librairie philosophique J. Vrin, Paris, 1987.

Edward BARREAU, Mathematical Fallacies, Flaws, and Flimflam, édité par The Mathematical Association of America, 2000, (ISBN 0-88385-529-1).

E.P. NORTHROP, Fantaisies et paradoxes mathématiques, Dunod, Paris, 1956, (tra-

duction de : **Riddles in Mathematics**, D. Van Nostrand Company Inc., New York, 1944).

Sylviane GASQUET-MORE, Plus vite que son nombre. Déchiffrer l'information. Éditions du Seuil, Paris 1999.

Massimo PIATELLI-PALMARINI, Inevitable Illusions : How Mistake of Reason Rule our Minds, John Wiley and Sons, New York, 1994 (ISBN 0-471-58126-7).



Un démineur millionnaire

IAN STEWART

Ou comment un jeu informatique peut faire de vous un homme riche.

En manque d'argent? L'Institut mathématique Clay, une fondation éducative à but non lucratif de Cambridge, dans le Massachusetts, offre des prix de un million de dollars pour la résolution de sept problèmes. L'un d'eux est le célèbre problème de l'égalité de la classe P et de la classe NP , dont nous éclaircirons plus loin l'énoncé hermétique. Bien qu'il s'agisse d'une question très difficile, un amateur astucieux peut y répondre grâce au démineur, un jeu populaire accessible à la plupart des ordinateurs personnels.

Richard Kaye, de l'Université de Birmingham, a récemment mis en évidence les relations entre ce jeu et la question « $P = NP?$ » dans un article intitulé «le démineur est-il NP -complet?». Rappelons d'abord les règles du démineur. L'ordinateur commence la partie en présentant une grille de carrés dont certains dissimulent des mines explosives. Il s'agit de les localiser sans en faire exploser une seule. Au premier coup, vous

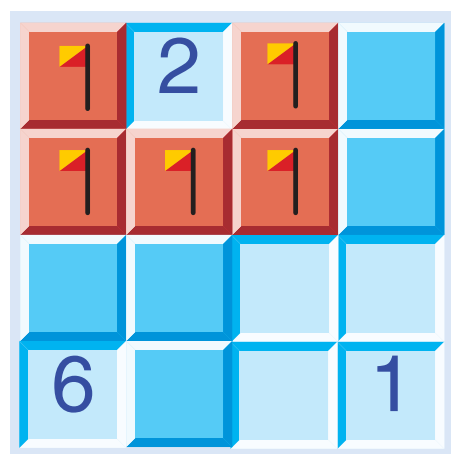
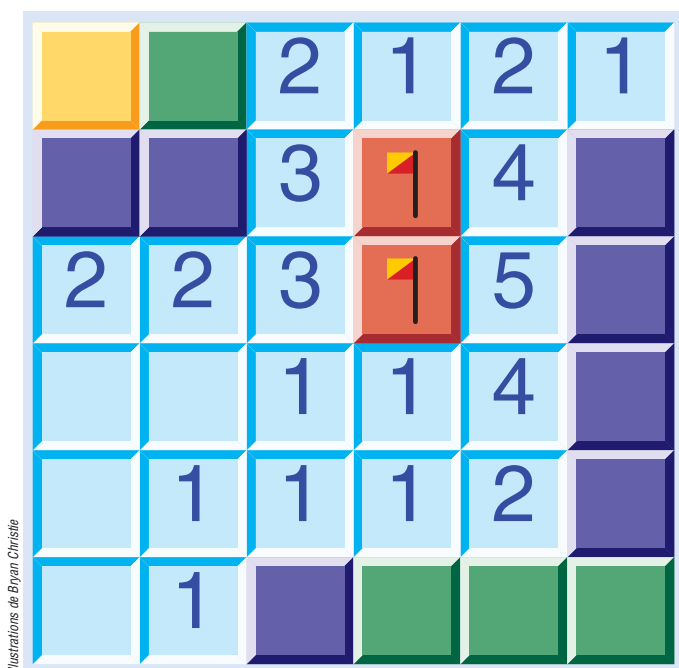
découvrez l'un des carrés : s'il cachait une mine, vous n'avez pas de chance et vous perdez, sinon l'ordinateur affiche dans le carré que vous avez choisi le nombre de mines dans les trois, cinq ou huit carrés adjacents.

Vous utilisez alors cette information pour choisir un autre carré, et de nouveau une mine explose ou un nombre est affiché. Lorsque vous pouvez déduire de ces informations la présence d'une mine dans un carré, vous marquez ce dernier d'un petit drapeau (voir la figure 1). Vous gagnez quand toutes les mines de la grille sont localisées.

Revenons à la question « $P = NP?$ ». Un algorithme est une procédure pas à pas qu'utilise un ordinateur pour résoudre un problème. Pour les informaticiens, il est crucial de savoir avec quelle efficacité un algorithme résout un problème donné, ou, en d'autres termes, comment le temps de traitement, c'est-à-dire le nombre de calculs exigés pour la réponse, varie suivant la taille des don-

nées initiales. On distingue principalement les problèmes de classe P (P comme polynôme) et les autres. Un problème est de classe P quand on peut le résoudre avec un algorithme dont le temps de traitement est majoré par un polynôme de la taille des données, c'est-à-dire par une somme de différentes puissances de la taille des données. Lorsqu'une telle solution n'existe pas, le problème est de classe non- P .

À l'inverse des problèmes de classe P , les problèmes de classe non- P ne sont pas efficacement traités par les ordinateurs : tout algorithme mis au point pour les résoudre requiert un temps trop long. On peut prouver qu'un problème est de classe P en trouvant un algorithme qui le résolve en un temps polynomial, par exemple, ordonner par ordre croissant un ensemble d'entiers est un problème de classe P que des programmes peuvent traiter sur des banques de données commerciales, même très importantes. En revanche, trouver les facteurs premiers



1. Une position typique du démineur (à gauche). Après les premiers coups, le joueur a localisé deux mines (sous les drapeaux). L'examen des chiffres montre que les carrés violets dissimulent d'autres mines tandis que les verts sont sans danger. L'état du carré jaune est encore inconnu, mais peut être déterminé après un coup judicieux. En revanche, la position à droite n'est pas consistante : aucune disposition de mines dans la grille ne peut respecter les chiffres indiqués.

d'un entier et le problème du voyageur de commerce, qui consiste à déterminer le plus court chemin reliant plusieurs villes données, sont des problèmes de classe non- P (bien que cela n'ait pas encore été démontré).

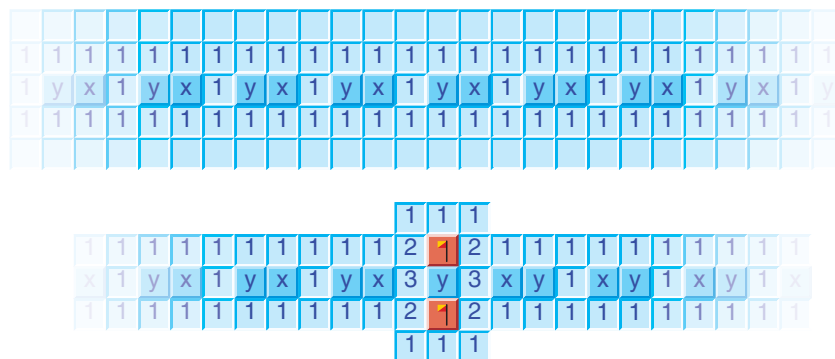
Pourquoi prouver qu'un problème est de classe non- P est-il si difficile ? Parce qu'il faut traiter tous les algorithmes possibles – et non pas un seul – et montrer qu'aucun ne résout le problème en un temps polynomial. Aujourd'hui, on sait qu'une vaste catégorie de problèmes candidats à la classe non- P sont sur un pied d'égalité : si l'un d'eux peut être résolu en un temps polynomial, alors tous peuvent l'être. Ces problèmes ont un comportement polynomial non déterministe, c'est-à-dire qu'avec de la chance – l'algorithme inclut une étape aléatoire –, on peut les résoudre en un temps polynomial ; ils sont de la classe NP .

NON- P , NP ET NP -COMPLET

Les classes non- P et NP sont différentes : un problème est NP dès qu'une solution proposée est bien en temps polynomial. Il est plus facile de vérifier qu'une solution proposée est bien en temps polynomial que de trouver la solution elle-même. Mon exemple favori qui illustre cette différence est le puzzle : en venir à bout peut être difficile, mais vérifier une solution proposée ne requiert souvent que quelques secondes ; le temps d'examiner que chaque pièce s'ajuste aux pièces voisines. Ce temps, environ proportionnel au nombre de pièces, est polynomial. En revanche, on ne peut effectuer un puzzle en un temps polynomial : le nombre de solutions croît plus vite que toute puissance donnée du nombre de pièces.

Beaucoup de problèmes NP ont des temps de traitement équivalents. Un problème NP est dit NP -complet quand l'existence d'une solution en un temps polynomial déterministe, c'est-à-dire sans processus aléatoire, entraîne que tout problème NP a une solution en un temps polynomial. Ainsi, si on peut résoudre un problème NP -complet en un temps polynomial, on peut résoudre tous les problèmes NP en un temps polynomial. Et la question « $P = NP$?» ? Bien que les apparences soient contraires, on peut s'interroger sur l'égalité des classes P et NP . La réponse attendue est négative, cependant, au cas où tous les problèmes NP -complets ont une solution en temps polynomial, alors les deux classes P et NP sont égales.

Nombre de problèmes sont NP -complets : un des plus simples est le problème SAT, qui inclut des circuits booléens.



2. L'électronique du démineur convertit des circuits booléens en position de jeu. Un seul type de carrés (x ou y) cache des mines. Des connexions propagent un signal (y suivi de x) le long de la grille (en haut). Un opérateur *non* inverse l'ordre des x et des y (en bas).

Ces circuits sont construits avec des opérateurs logiques, tel «*et*», «*ou*» et «*non*», qui s'appliquent à des valeurs soit vraies (V), soit fausses (F). Chaque opérateur fournit à partir de valeurs d'entrée une valeur de sortie : par exemple, l'opérateur *non* transforme la valeur V en la valeur F . Le problème SAT s'intéresse à l'existence pour un circuit booléen donné de valeurs d'entrée tel que la valeur de sortie soit V . Pour des circuits simples, le problème est un jeu d'enfants ; il devient ardu lorsque le circuit a un grand nombre d'opérateurs et d'entrées.

LA CONSISTANCE DU DÉMINEUR

Le lien avec le démineur apparaît quand on s'intéresse à la consistance du jeu, c'est-à-dire à la vraisemblance d'une position du jeu. Par exemple, la figure 1 montre une position inconsistante : aucune répartition des mines ne peut mener aux nombres affichés dans les carrés de la grille.

R. Kaye a montré que le problème SAT pour un circuit booléen donné peut-être transformé en un problème de consistance pour une certaine grille du jeu du démineur. En outre, il montre que la conversion des circuits booléens en positions du démineur s'exécute en un temps polynomial. Dans cette conversion, les entrées et les sorties du circuit deviennent une répartition des mines sur la grille : à un carré miné correspond une valeur V , et à un carré sans mine, une valeur F . Les connexions des opérateurs du circuit sont représentées sur la figure 2 : les carrés marqués y ont des valeurs opposées de celles des carrés marqués x . Notons que les chiffres des autres carrés sont corrects que x soit V ou F . Tels des «câbles électroniques», les connexions propagent chaque valeur (V ou F) sur la longueur de la grille.

On représente les opérateurs logiques par le même type de procédé. Par

exemple, la position du démineur de la figure 2 illustre l'opérateur *non* : le bloc du milieu échange les positions de x et de y ; x suit y avant le bloc et précède y après le bloc. Les valeurs des carrés sont inversées, que x soit V ou F , comme après un opérateur *non* dans un circuit booléen.

L'électronique du démineur est plus complexe : elle requiert par exemple de plier ou de scinder des «fils». R. Kaye franchit tous ces obstacles et conclut que quiconque prouve la consistance d'une position donnée du démineur en un temps polynomial, résout aussi le problème SAT du circuit équivalent en un temps polynomial. En d'autres termes, le démineur est un problème NP -complet : une solution en un temps polynomial à la consistance du démineur entraîne que tous les problèmes NP -complets ont des solutions en un temps polynomial : les classes P et NP sont égales. À l'inverse, si l'on montre qu'il n'existe pas de telle solution au problème du démineur, alors les classes P et NP sont différentes. Dans les deux cas, la question serait résolue.

Avant de vous lancer sur votre ordinateur, rappelez-vous que le problème de la consistance du démineur n'est pas une mince affaire : déterminer si une position de jeu est consistante devient vite difficile quand la taille de la grille augmente et la plupart des mathématiciens et des informaticiens pensent qu'il n'y a pas de solution générale traitable en un temps polynomial. De surcroît, l'Institut Clay a imposé des règles strictes au concours : avant d'accepter une solution, elle doit être publiée dans un journal important et être «généralement acceptée» par la communauté mathématique dans les deux années de sa publication.

Richard Kaye, *Minesweeper is NP-Complete*, in *Mathematical Intelligencer*, vol. 22, n° 2, 2000.



Chaleur en perte

ROLAND LEHOUCQ • JEAN-MICHEL COURTY

La chaleur quitte la maison par 1 000 chemins, mais 5 sont plus fréquents.

Avant tout, les murs d'une habitation sont faits pour être solides. Toutefois, quiconque paie des factures de chauffage sait qu'ils doivent être aussi isolants, et déplore que cela ne suffise pas à garder la chaleur dans la maison. Un peu de bon sens physique explique pourquoi, et identifie les cinq déperditions thermiques fréquentes dans une maison.

«Ferme la porte : Il fait froid dehors !» Mille fois entendue, cette exclamation démontre que tout le monde connaît la plus banale des pertes thermiques de la maison : le remplacement de l'air chaud intérieur par de l'air extérieur froid. Un calfeutrage soigneux règle-t-il le problème ? Oui, mais il en crée un autre. Pour que l'habitation soit saine, l'air intérieur chargé en dioxyde de carbone doit souvent être remplacé par de l'air extérieur, riche en oxygène. Selon les normes actuelles pour les habitations, le volume d'air des pièces principales doit être renouvelé une fois par heure. Ainsi 10 kilowattheures par jour sont nécessaires pour réchauffer de 15 °C l'air d'un appartement de 50 mètres carrés si l'on renouvelle l'air une fois par heure.

C'est pour restreindre cette perte, que les bureaux et autres grands ensembles sont souvent équipés d'une ventilation mécanique assortie d'un échangeur de chaleur. Avant d'être intro-

duit dans les pièces, l'air entrant est réchauffé par l'air chaud qui circule dans des tuyaux. Dans certains échangeurs, l'air chaud est évacué par un tuyau qui chemine à l'intérieur du circuit par lequel est introduit l'air froid. De telles «géométries à contre-courant» sont si efficaces que l'air entrant y atteint presque la température de l'air intérieur.

Une fois les portes closes, et les enfants assez disciplinés pour les fermer dès qu'elles sont ouvertes, c'est par les parois que fuit la chaleur. La conduction thermique à travers les murs, le toit, les portes, les fenêtres ou encore le sol, est la deuxième des principales formes de déperdition de chaleur. Expression macroscopique de l'agitation microscopique désordonnée des atomes et des molécules qui constituent la matière, la chaleur se propage de proche en proche par l'intermédiaire des particules qui s'entrechoquent.

Ainsi, quand un milieu plus ou moins conducteur relie deux corps, il transmet la chaleur du corps le plus chaud vers le corps le plus froid. Les températures des deux corps ainsi reliés s'uniformisent en un temps inversement proportionnel à la conductivité thermique du milieu de transmission. Notons toutefois que la géométrie du milieu de transmission influe aussi sur la durée d'uniformisation thermique : quand son épaisseur double ou que sa surface de contact avec l'un des deux corps est réduite de moitié, cette durée double.

La conductivité thermique d'un matériau dépend de sa densité, car plus les particules par unité de volume sont nombreuses, plus les chocs sont fréquents.

De fait, les métaux sont de bien meilleurs conducteurs thermiques que le béton, qui lui-même conduit mieux que le bois. L'acier transmet 52 watts par mètre et par degré, le béton 1,75 et le bois, un bon isolant, seulement 0,06. Les gaz arrêtent aussi très bien la chaleur. L'air, par exemple, a une conductivité de 0,025 watt par mètre et par degré. Toutefois le meilleur de tous les isolants est tout simplement le vide. La raison en est simple : dans une région sans matière, il ne peut pas y avoir de chocs entre particules. C'est pourquoi les bouteilles *Ther-*



Les parois en contact avec l'extérieur, notamment les vitrages, sont une source importante de déperdition de chaleur.

mos, où l'on fait le vide entre les doubles parois, retiennent si bien la chaleur. En fait, le vide serait un isolant parfait, s'il ne laissait pas passer une autre forme d'énergie thermique : le rayonnement infrarouge qu'émet tout corps chaud.

LIMITER LES PERTES PAR CONDUCTION

Que faire pour limiter les pertes par conduction ? Isoler ! Le mieux n'est-il pas de construire sa maison en matériaux isolants ? C'est difficile, car les matériaux isolants sont rarement assez solides pour cela, sauf à prévoir une épaisseur de mur digne d'un château fort. Seul le bois est un isolant à la fois léger et solide, raison pour laquelle il est très utilisé dans les pays nordiques, où abondent les forêts. Très employé, le béton conduit malheureusement bien la chaleur. S'il est massivement utilisé dans une maison, les couches d'isolants sont indispensables. Les isolants courants sont tous à base d'air, le meilleur isolant après le vide.

Ainsi même si le verre est un bon conducteur thermique, la laine de verre est un meilleur isolant, car elle contient 99 pour cent d'air ! Autre exemple, le polystyrène expansé, lui aussi une excellente barrière thermique, renferme une



L'intrusion d'air froid dans les habitations est le principal facteur de déperdition thermique.