

appareils mis en vente reposent sur l'idée que la mécanique quantique produit un hasard qui, une fois bien contrôlé (pour en équilibrer les productions), serait le hasard véritable caractérisé en 1965 par Martin-Löf et que le déterminisme de la mécanique newtonienne n'est pas en mesure d'atteindre.

Cette idée est-elle justifiée et fondée théoriquement ? La réponse n'est pas tranchée, car il n'est pas possible de tirer, des principes de la mécanique quantique, l'affirmation que les suites produites par exemple par les photons du dispositif vendu par *ID Quantique* sont aléatoires au sens de Martin-Löf. De l'indétermination concernant le passage ou non du photon à travers le miroir que la mécanique quantique exprime comme un axiome, on peut sans doute conclure que les suites produites par un tel dispositif sont semblables à celles que produisent des variables aléatoires uniformes indépendantes. Cependant, rien ne permet d'affirmer qu'une suite particulière tirée des photons (ou d'un autre procédé microscopique) est une suite aléatoire au sens de Martin-Löf : les suites provenant de tirages indépendants équilibrés ne sont pas toutes aléatoires au sens mathématique.

Il n'existe ainsi aucune méthode dont on puisse dire avec certitude qu'elle produit des suites aléatoires, et celles de la mécanique quantique ne font pas exception.

Les méthodes algorithmiques n'en produisent certainement pas et cela concerne :

- les chiffres des nombres irrationnels tels que π , $\sqrt{2}$, etc.

- les méthodes proposées dans les langages de programmation qui produisent un hasard assez bien équilibré, mais parfois prédictible et souvent imparfait quand on y regarde de près.

- les méthodes cryptographiques, qui sont conçues pour ne pas être aussi facilement prédictibles, mais qui, du fait de leur nature algorithmique, n'en donnent pas pour autant des suites aléatoires au sens absolu.

Pour les procédés mécaniques, la physique newtonienne affirme que les tirages successifs sont déterministes et, par conséquent, n'ont aucune raison de donner des suites incompressibles ou imprévisibles.

Quant aux méthodes microscopiques, rien dans les principes mêmes de la mécanique quantique ne garantit la production de véritables séquences aléatoires au sens fort. À moins de compléter les axiomes de la théorie, il n'est aujourd'hui pas justifié de dire qu'un procédé quantique produit avec une certitude absolue de l'aléa fort comme Martin-Löf l'a défini. Il faut peut-être reformuler la théorie quantique pour que cela change, mais personne pour l'instant ne le propose. Il faut donc cesser d'affirmer que les méthodes quantiques de production d'aléa sont bien fondées, à l'opposé des méthodes mécaniques ou algorithmiques.

Insaisissable hasard

Malgré tout cela, et c'est ici qu'il y a un paradoxe, pour les tests conçus depuis un siècle et qui sont soigneusement collectés, par exemple par le NIST aux États-Unis, tout semble aller très bien. Il n'y a pas de différences repérables entre les moins assurées (théoriquement) des suites pseudo-aléatoires (comme la suite des chiffres de π), les suites aléatoires mécaniques (produites avec précautions) et les suites aléatoires quantiques.

Régulièrement, certains chercheurs ont prétendu repérer et mesurer des différences entre diverses suites aléatoires (par exemple entre les chiffres de π et ceux d'autres constantes). Cependant, jamais ces affirmations n'ont été confirmées et, au contraire, on a en général découvert des erreurs méthodologiques de la part de ceux qui annonçaient avoir repéré des nuances mesurables entre suites aléatoires. Un article de George Marsaglia, grand spécialiste de l'aléa informatique, publié en 2005 conclut de manière formelle après l'utilisation des meilleures batteries de tests disponibles que les chiffres des développements de nombres irrationnels tels que π , e , $\sqrt{2}$, aussi bien d'ailleurs que ceux des nombres rationnels k/p , où p est un nombre premier assez grand, semblent se comporter comme s'ils provenaient d'une suite de tirages indépendants équitables : la théorie et la pratique divergent au maximum !

L'AUTEUR



Jean-Paul DELAHAYE est professeur à l'Université de Lille et chercheur au Laboratoire d'informatique fondamentale de Lille (LIFL).

✓ BIBLIOGRAPHIE

H. Zenil, *Randomness Through Computation : Some Answers, More Questions*, World Scientific Publishing, 2011.

B. Hayes, *Excursions quasi-aléatoires*, *Pour la Science*, décembre 2011.

A. Dasgupta, *Mathematical foundation of randomness*, dans *Philosophy of Statistics*, North-Holland, pp. 641-710, 2011 (<http://dasgupta.faculty.udmercy.edu/Dasgupta-JSfinal.pdf>).

R. Downey et D. Hirschfeld, *Algorithmic Randomness and Complexity*, Springer-Verlag, 2010.

N. Gauvrit, *Vous avez dit hasard ?*, Belin/Pour la Science, 2009.

Hasard et incertitude, *Pour la Science*, numéro spécial, novembre 2009.

J. Strzalko *et al.*, *Dynamics of Gambling : Origins of Randomness in Mechanical Systems*, Springer, 2009.

P. Diaconis *et al.*, *Dynamical bias in the coin toss*, *SIAM Rev.*, vol. 49, pp. 211-235, 2007.

J.-P. Delahaye, *Information, complexité et hasard*, Hermès, 2^e édition, 1999.

ART & SCIENCE

Recette pour une mosaïque romaine

Comment les motifs complexes d'une mosaïque de la villa romaine de Chedworth, en Grande-Bretagne, ont-ils été obtenus ? Par une méthode simple, également à l'œuvre dans la tradition picturale en usage dans d'autres pays.

Loïc MANGIN

En 1864, une villa romaine fut découverte à Chedworth, dans le Sud-Est de la Grande-Bretagne. Ce complexe de bâtiments, érigés autour d'une cour, a été construit à partir de l'an 120 et fut remanié jusqu'au IV^e siècle. L'ensemble est bien conservé et plusieurs pièces abritent encore de riches mosaïques. Intéressons-nous à l'une d'elles, située dans la salle à manger (*voir la figure g, page ci-contre*).

Un enchevêtrement de lignes géométriques noires sur fond blanc sinuent autour de plusieurs blocs marron et dessinent des svastikas. Ce dernier motif (une croix dont les quatre branches sont « cassées » en angle droit) est fréquent dans l'art romain. On le trouve aussi en Grèce, en Afrique du Nord, chez les Indiens Navajos, aux États-Unis, en Chine, en Inde...

Comment s'y sont pris les artisans pour composer le décor complexe ? Grâce à une analyse géométrique, Godfried Toussaint, de l'Université Harvard, aux États-Unis, et son collègue Yang Liu ont montré que les artisans ont pu obtenir les lignes noires intriquées à partir de seulement quatre courbes déformées selon un procédé simple.

L'étude est fondée sur la reconstitution de la mosaïque par l'écrivain Lewis Day, en 1903 (*voir la figure a*). Les lignes noires serpentent autour de trois types de blocs ornés d'entrelacs : un grand carré au centre, quatre rectangles et 20 petits carrés répartis en quatre groupes de cinq aux quatre coins de la mosaïque.

Au centre du grand carré, un « nœud de Salomon » enferme un svastika. Ce motif est cerné d'un guillochis (un ornement composé de lignes s'entrelaçant avec symétrie) à trois brins. Dans les carrés, trois courbes fermées s'entrecroisent, tandis que dans les rectangles, une seule courbe dessine l'entrelacs.

L'usage de courbes fermées dont l'ordonnement obéit à des contraintes est répandu dans le monde. Ainsi, les motifs de

Les motifs romains ressemblent aux figures de Kolam d'Inde du Sud et au sona d'Angola.

Chedworth ressemble aux figures de Kolam, que les femmes de l'État du Tamil Nadu, en Inde du Sud, dessinent sur le pas de leur porte pour accueillir le visiteur ou protéger le foyer. De même, les entrelacs romains évoquent les *sona* d'Angola, qui consistent en une ou plusieurs courbes tracées par des conteurs dans le sable, autour de points régulièrement disposés. Par exemple, le nœud à trois brins de chacun des 20 carrés correspond à un *lusona* (le singulier de *sona*) à trois courbes (*voir la figure b*). Notons que l'on peut voir dans ces dessins les trajets de rayons lumineux dans une enceinte tapissée de miroirs.

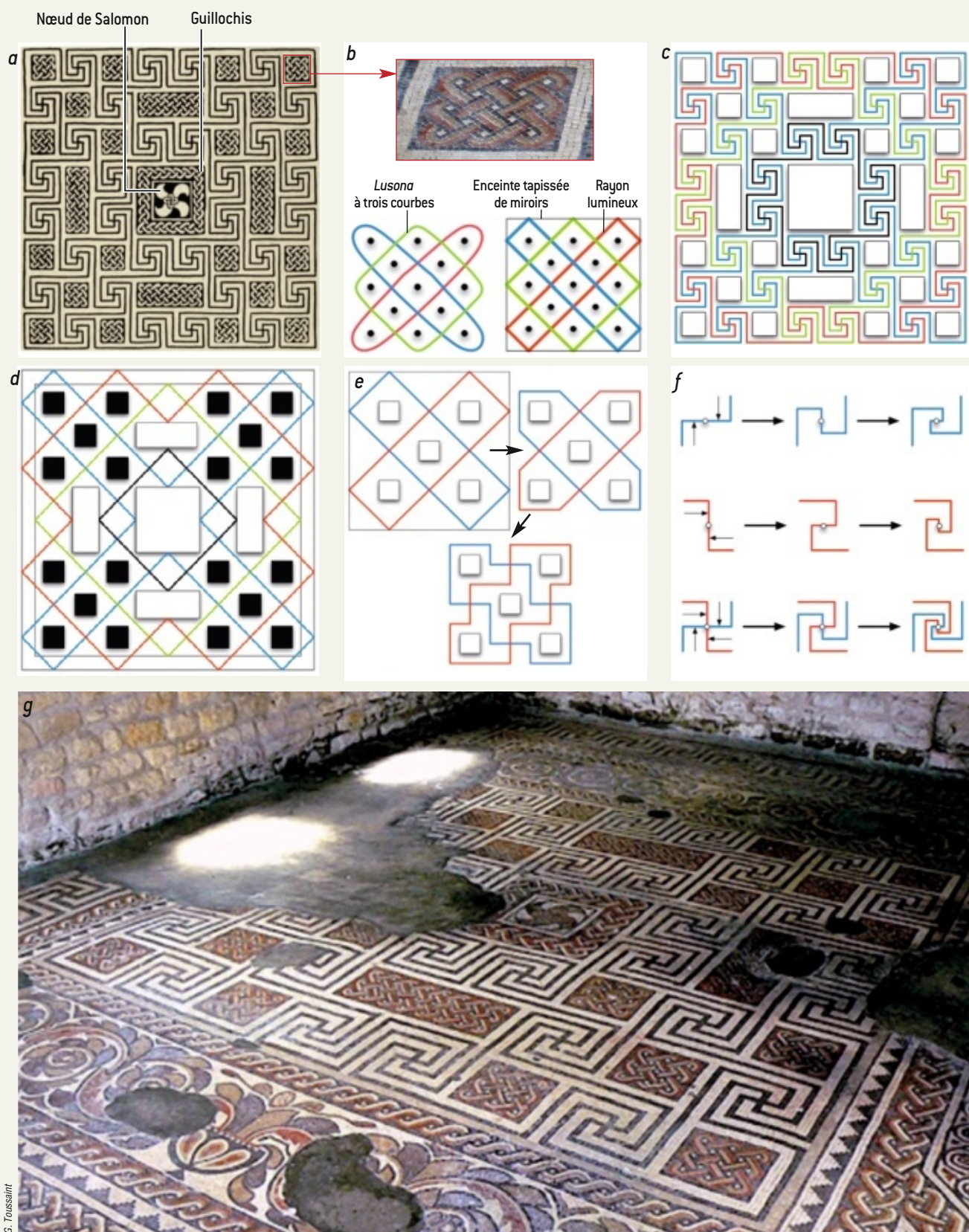
Venons-en aux lignes noires qui parcourent l'ensemble de la mosaïque. Un examen attentif révèle que quatre courbes fermées suffisent à toutes les dessiner (*voir la figure c*).

L'une (*en noir sur la figure*) reste autour du centre, une autre (*en bleu*) associe les quatre coins de la mosaïque, et deux autres encore (*en rouge et en vert*) s'entrelacent pour relier les côtés au carré central. La méthode de construction proposée par G. Toussaint et Y. Liu est fondée sur l'idée des miroirs : ils considèrent que le carré central et les rectangles sont aussi réfléchissants (*voir la figure d*). À chaque croisement de deux courbes (ou rayons lumineux) correspond, sur le motif final, un svastika.

On obtient ainsi de nombreux segments, orientés selon des diagonales, et dont les extrémités sont marquées soit par un changement de direction, soit par un croisement. Chacun de ces segments est ensuite remplacé par un angle droit (*voir la figure e*). Enfin, tous les croisements subissent une rotation (*voir la figure f*, où le motif est d'abord décomposé) de façon à obtenir le svastika souhaité. Le sens de rotation n'a pas d'importance et l'orientation d'un svastika est indépendante de celles des autres.

G. Toussaint et Y. Liu pensent que les motifs intriqués de la mosaïque de Chedworth ont pu être construits selon ce procédé. En outre, leur méthode serait utile pour classer les motifs connus, mais aussi pour restaurer les œuvres endommagées.

Y. Liu et G. Toussaint, *Unravelling Roman mosaic meander patterns : a simple algorithm for their generation*, *Journal of Mathematics and the Arts*, vol. 4, n° 1, pp. 1-11, 2010.
Mathématiques exotiques, Dossier Pour la Science, n° 47, avril-juin 2005.



IDÉES DE PHYSIQUE

Comment éviter le coup de foudre

On ne peut empêcher la foudre de tomber, mais il existe un dispositif simple qui permet de guider en partie les éclairs : le paratonnerre.

Jean-Michel COURTY et Édouard KIERLIK

La foudre est un phénomène naturel qu'il est impossible de prévenir. Ces décharges électriques intenses peuvent tuer et occasionner des dommages considérables. Pour s'en protéger, on utilise de longues tiges métalliques reliées au sol : les paratonnerres. Ils sont censés attirer la foudre et canaliser le courant hors de la zone à protéger. Comment fonctionnent ces dispositifs ? Quelle zone protègent-ils ?

Répondre à ces questions suppose de comprendre comment se forment les éclairs. Les éclairs sont de nature électrique, avait deviné Benjamin Franklin. Le naturaliste français Thomas-François Dalibard fut le premier à confirmer cette hypothèse, en 1752, grâce à l'observation des étincelles créées par temps d'orage par un mât métallique de dix mètres, relié à un condensateur électrique.

En fait, ni le mât de Dalibard ni plus tard le cerf-volant de Franklin n'étaient touchés par la foudre, mais ils captaient les charges électriques de l'air environnant. Franklin eut ainsi l'idée brillante des paratonnerres, dispositifs qui étaient censés décharger les nuages et empêcher la foudre de tomber. Mais les paratonnerres ne font pas exactement cela. Il est impossible de capter sans violence l'électricité des nuages et, aujourd'hui comme hier, rien ne semble empêcher la formation des éclairs. Il est en revanche possible, grâce aux paratonnerres, d'orienter partiellement leurs frappes.

La formation d'un éclair demande d'abord un nuage orageux, c'est-à-dire un nuage de grande taille où, à la suite de divers processus pas toujours bien élucidés, la partie supérieure du nuage se charge positivement tandis que sa base se charge négativement.

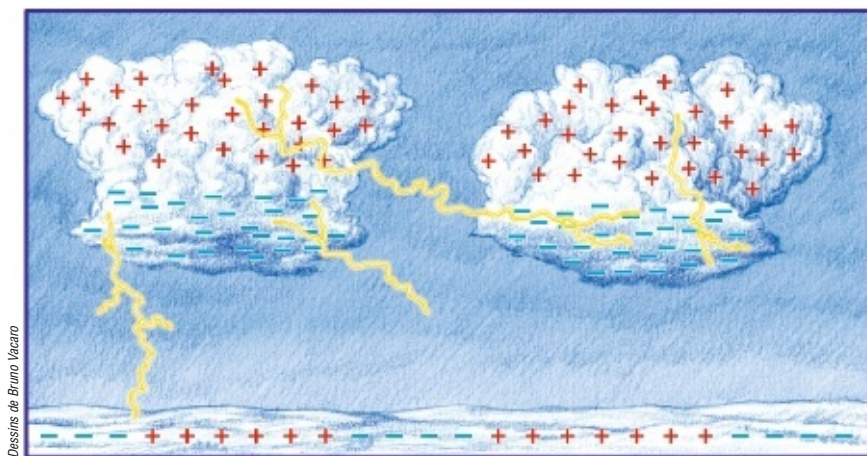
Par influence électrostatique, cette dernière attire au niveau du sol, qui est légèrement conducteur, des charges positives (voir la figure 1). Il se forme ainsi une sorte de condensateur géant où règnent des différences de potentiel de plusieurs dizaines de millions de volts. Avec une différence de potentiel de 100 millions de volts et le bas du nuage situé à 1 000 mètres d'altitude, on obtient des champs électriques de près de 100 000 volts par mètre.

Nuage-sol : un condensateur géant

Ce champ est 1 000 fois plus intense que par temps calme, mais reste très inférieur au champ disruptif de l'air (3,6 millions de volts par mètre, seuil au-delà duquel des arcs électriques naissent). Néanmoins, en raison des instabilités atmosphériques à la base du nuage, les charges peuvent former localement des structures conductrices courbes qui, par effet de pointe, créent des champs électriques bien supérieurs : l'air peut s'ioniser et une décharge s'amorce.

La décharge prend la forme d'un filament hautement conducteur, porteur d'une charge électrique élevée, suffisante pour entretenir l'ionisation. Environ neuf fois sur dix, ce « traceur » se perd dans un nuage voisin ; sinon, il descend vers le sol. Comme le vent chasse l'air ionisé, le traceur s'arrête quand les ions viennent à manquer, puis repart lorsque l'ionisation reprend : il avance par bonds, en zigzags, à une vitesse de quelque 200 kilomètres par seconde.

En se rapprochant du sol, les charges négatives situées à l'extrémité du traceur



Dessins de Bruno Vacaro

1. LES ÉCLAIRS SONT DUS À DE FORTES DIFFÉRENCES DE CHARGE ÉLECTRIQUE qui apparaissent entre le haut et le bas d'un nuage, sous l'effet de processus complexes. La plupart des éclairs sont internes au nuage ou atteignent un nuage voisin, mais quelques-uns atteignent le sol.