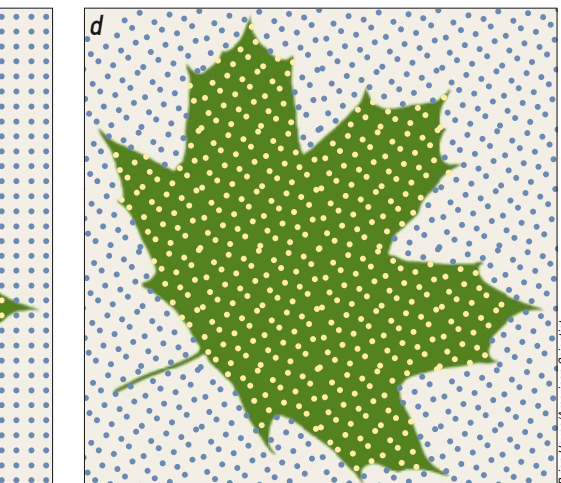




Brian Hayes/American Scientist

à l'intérieur du petit cube. En d'autres termes, quand nous comptons les coups gagnants, il faut que nous en comptons au moins un. Pour un motif de quadrillage à 20 dimensions, cela signifie que nous avons besoin d'un million de points (ou, plus précisément, de $2^{20} = 1\,048\,576$ points). Avec l'échantillonnage aléatoire, l'exigence de taille est probabiliste et donc un peu floue, mais le nombre de points nécessaires pour nous attendre à un seul coup gagnant est là encore 2^{20} . L'analyse pour l'échantillonnage quasi-aléatoire donne le même résultat. En effet, si vous tâtonnez à l'aveuglette pour trouver un objet de volume $1/2^d$, peu importe l'agencement du motif de recherche; il vous faudra chercher à 2^d endroits.

Prévisibles, corrélés, mais équitables

Si les problèmes réels étaient aussi compliqués que celui-ci, les perspectives seraient assez sombres. Il n'y aurait pas le moindre espoir d'estimer une intégrale à 360 dimensions. Mais nous savons que les techniques de Monte Carlo viennent à bout de certains problèmes de cette taille; c'est une conjecture raisonnable de penser que les problèmes solubles ont une structure interne qui accélère la recherche. De plus, le choix du motif d'échantillonnage semble faire une différence, si bien qu'il y a une distinction significative à faire entre toutes les gradations allant de l'aléatoire vrai à l'ordonné, en passant par le pseudo-aléatoire et le quasi-aléatoire.

L'idée de hasard dans un ensemble de nombres a au moins trois aspects. Tout

d'abord, des nombres choisis au hasard sont imprévisibles: il n'existe pas de règle fixe présidant à leur sélection. Deuxièmement, ces nombres sont indépendants, ou non corrélés: le fait de connaître l'un des nombres ne vous aidera pas à en deviner un autre. Enfin, les nombres aléatoires sont non biaisés, ou uniformément distribués: peu importe comment vous découpez l'espace des valeurs possibles, chaque région peut s'attendre à en avoir sa part.

Ces concepts fournissent une clef utile pour distinguer entre les ensembles véritablement aléatoires, pseudo-aléatoires, quasi-aléatoires et ordonnés. Les nombres véritablement aléatoires présentent les trois caractéristiques énoncées: ils sont imprévisibles, non corrélés et non biaisés. Les nombres pseudo-aléatoires abandonnent l'imprévisibilité; ils sont engendrés par une règle arithmétique bien définie, et si vous connaissez la règle, vous pouvez reproduire la séquence entière. Mais les nombres pseudo-aléatoires restent non corrélés et non biaisés (au moins avec une bonne approximation).

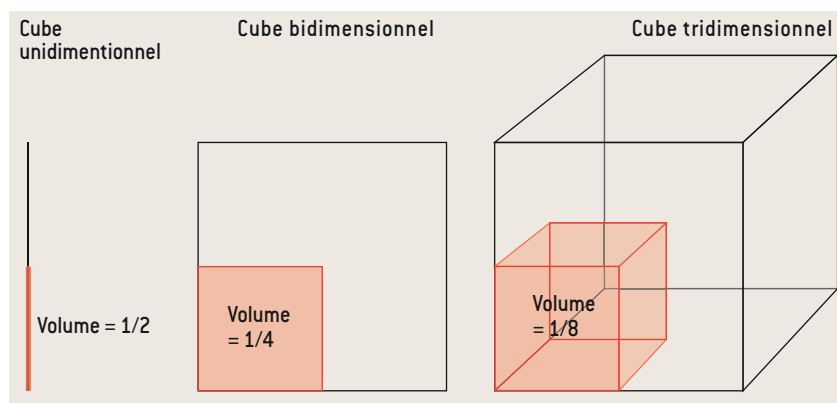
Les nombres quasi-aléatoires sont à la fois prévisibles et fortement corrélés. Il existe une règle bien définie pour les engendrer, et les motifs qu'ils forment, bien que pas aussi stricts que ceux d'un réseau cristallin, présentent une grande régularité. Par rapport au hasard, le seul élément que les nombres quasi-aléatoires préservent est la distribution uniforme ou équitable: ils sont distribués aussi équitablement et également que possible.

Un ensemble très ordonné, comme les points d'un réseau cubique, ne préserve aucune des propriétés du hasard. Il est évident que ces points échouent aux tests d'imprévisibilité et d'indépendance. Il n'est peut-être pas aussi clair que leur distribution n'est pas uniforme. Après tout, il est possible de découper un réseau de N points en N petits cubes, qui contiennent chacun exactement un point. Mais cela ne suffit pas pour qualifier le réseau d'équitabilité et également distribué. L'idéal de la distribution équitable exige un arrangement de N points qui fournit un point par région quand l'espace est divisé en n'importe quel ensemble de N régions identiques. Le réseau rectiligne échoue à ce test quand les régions sont des tranches parallèles aux axes.

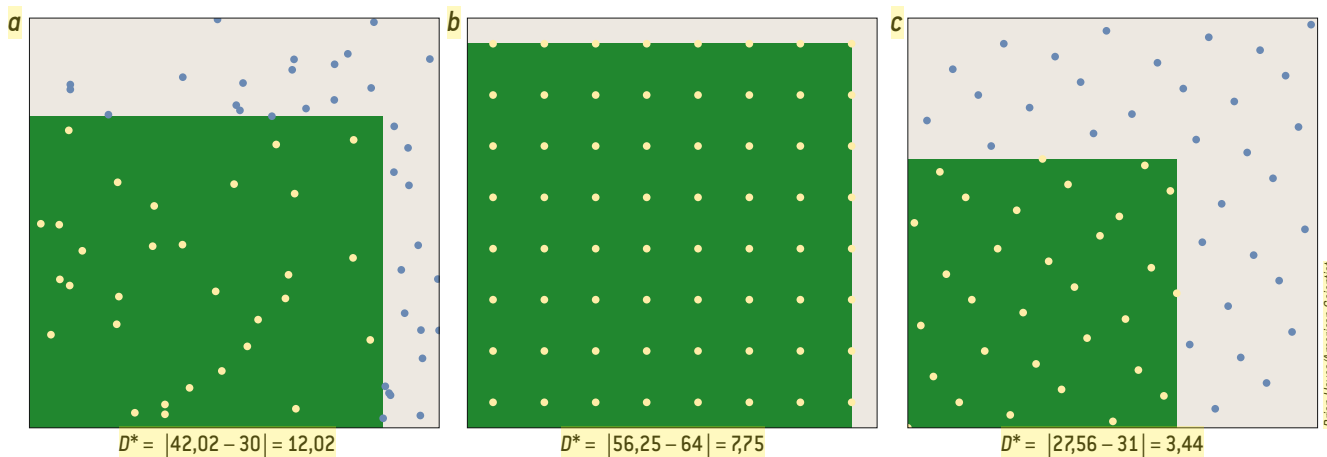
Un léger désaccord: la discrédance

Ces trois aspects du hasard sont importants dans différents contextes. Dans le cas de l'autre Monte Carlo (le casino de Monaco), l'imprévisibilité est tout. Il en va de même pour les applications cryptographiques du hasard. Dans les deux cas, un adversaire essaie activement de détecter et d'exploiter toute ébauche de motif.

Certains types de simulations informatiques sont très sensibles aux corrélations entre des nombres aléatoires successifs, si bien que l'indépendance est importante. Mais pour les estimations de volume que nous discutons ici, c'est l'uniformité de la



3. L'ESTIMATION DES VOLUMES devient de plus en plus difficile à mesure que le nombre de dimensions de l'espace augmente, effet qualifié parfois de « malédiction de la dimension ». Supposons qu'un cube à d dimensions, dont les côtés ont pour longueur 1, renferme un cube de côtés $1/2$. Quand $d = 1$, le « cube » est en fait un segment de droite, et le volume est sa longueur; le petit « cube » [de longueur $1/2$] remplit la moitié du volume total. Pour $d = 2$ (un carré, pour lequel le volume correspond à la superficie), le volume du petit cube est de $1/4$, et, pour $d = 3$, il est égal à $1/8$. Ainsi, le volume fractionnaire diminue comme $1/2^d$, ce qui signifie que pour simplement détecter l'existence du petit cube, il faut déjà environ 2^d points d'échantillonnage.



Brian Hayes/American Scientist

4. LA DISCRÉPANCE mesure de combien un ensemble de points s'écarte d'une distribution spatiale uniforme. La mesure illustrée ici, appelée *discrépance étoile* et notée D^* , est définie en termes de rectangles (*en vert*) ayant un sommet ancré dans le coin inférieur gauche. Si la distribution des points était parfaitement uniforme, le nombre de points que renferme l'un quelconque de ces rectangles serait proportionnel à la superficie du rectangle ; la différence entre ce nombre attendu de

points et le nombre réel est la *discrépance* associée au rectangle. La *discrépance étoile* pour l'ensemble du motif est donnée par le rectangle pour lequel cette différence est maximale. La *discrépance étoile* est donnée ici pour trois motifs de 64 points : un réseau pseudo-aléatoire (*a*), un réseau régulier (*b*) et un réseau quasi-aléatoire (*c*). La configuration quasi-aléatoire est, par définition, une configuration conçue pour minimiser la *discrépance*.

distribution qui importe le plus. Et les nombres quasi-aléatoires, en abandonnant toute prétention à l'imprévisibilité ou à l'absence de corrélation, parviennent à atteindre une plus grande uniformité.

L'uniformité de la distribution est mesurée en termes de son opposé, la « *discrépance* ». Pour les points d'un carré bidimensionnel, la *discrépance* est mesurée comme suit. Considérez tous les rectangles aux côtés parallèles aux axes des coordonnées qu'il est possible de dessiner à l'intérieur du carré. Pour chacun de ces rectangles, comptez le nombre de points qu'il contient, et calculez aussi le nombre de points qui seraient contenus, en fonction de la superficie du rectangle, si la distribution était parfaitement uniforme. L'écart maximal entre ces nombres, calculés pour tous les rectangles possibles, est une mesure de la *discrépance*.

Une autre mesure de la *discrépance*, nommée *discrépance étoile* (parce qu'on la note D^*), ne considère que le sous-ensemble de rectangles parallèles aux axes qui ont un coin ancré à l'origine du carré unité. Et il n'y a aucune raison de ne considérer que les rectangles ; on peut définir d'autres versions de la *discrépance* avec des cercles, des triangles ou d'autres figures. Les mesures ne sont pas toutes équivalentes ; les résultats varient selon la forme. Il est intéressant de noter que le processus pour mesurer la *discrépance* a une grande ressemblance avec le processus Monte Carlo lui-même.

Les grilles et les réseaux ne sont pas très performants lorsqu'il est question de *discrépance*, parce que les points sont disposés en longues rangées et colonnes. Un rectangle infiniment étroit, qui ne devrait contenir aucun point, pourrait en englober toute une rangée. À cause de ces rectangles extrêmes, la *discrépance* d'un réseau carré de N points vaut $\sqrt{N}$. Il est intéressant de noter que la *discrépance* d'un réseau aléatoire ou pseudo-aléatoire vaut encore à peu près $\sqrt{N}$. En d'autres termes, dans un réseau d'un million de points aléatoires, il existe vraisemblablement au moins un rectangle qui a soit 1 000 points en trop, soit 1 000 points en moins.

Quasi-Monte Carlo en dimension élevée ?

Les motifs quasi-aléatoires sont délibérément conçus pour éliminer toutes les possibilités de dessiner de tels rectangles à *discrépance* élevée. Pour les points quasi-aléatoires, la *discrépance* peut ne pas dépasser le logarithme de N , ce qui est beaucoup plus petit que $\sqrt{N}$. Par exemple, pour $N = 10^6$, $\sqrt{N} = 1 000$, mais $\log N = 20$ (avec les logarithmes de base 2).

La *discrépance* est un facteur clef pour évaluer les performances des simulations de Monte Carlo et quasi-Monte Carlo. Elle détermine le degré d'erreur ou d'imprécision statistique à attendre d'une taille d'échantillon donnée. Pour le Monte Carlo classique, avec un échantillonnage aléa-

toire, l'erreur attendue diminue en $1/\sqrt{N}$. Pour le quasi-Monte Carlo, la « *vitesse de convergence* » correspondante vaut $(\log N)^d/N$, où d est la dimension de l'espace (divers détails et facteurs constants étant négligés ici, ces vitesses de convergence sont approximatives).

La vitesse de convergence en $1/\sqrt{N}$ pour l'échantillonnage aléatoire peut être extrêmement lente. Pour gagner une décimale de précision, il faut augmenter N d'un facteur 100, ce qui signifie qu'un calcul d'une heure devient un calcul de quatre jours. La raison de cette lenteur provient des agrégats présents dans une distribution aléatoire de points. En effet, lorsque deux points sont très proches, les ressources de calcul sont gaspillées, puisque la même région est échantillonnée deux fois ; inversement, les grands vides entre les points laissent certaines zones non échantillonnées et contribuent donc à alourdir le poids des erreurs.

En compensation de cet inconvénient, l'échantillonnage aléatoire présente un avantage important : la vitesse de convergence ne dépend pas de la dimension de l'espace. En première approximation, le programme tourne aussi vite en dimension 100 qu'en dimension 2.

La méthode de quasi-Monte Carlo est différente à cet égard. Si seulement nous pouvions ignorer le facteur $(\log N)^d$, les performances de l'échantillonnage quasi-aléatoire seraient radicalement meilleures que celles de l'échantillonnage aléatoire. Nous aurions une convergence

en $1/N$, ce qui multiplie seulement par 10 l'effort pour chaque décimale de précision supplémentaire. Mais nous ne pouvons pas ignorer le terme $(\log N)^d$. Ce facteur croît exponentiellement avec la dimension d . L'effet est petit en dimension 1 ou 2, mais il finit par devenir énorme. En dimension 10, par exemple, $(\log N)^d/N$ reste supérieur à $1/\sqrt{N}$ tant que N est inférieur à 10^{43} .

C'est en raison de la lenteur connue de cette convergence en dimension élevée que S. Paskov et J. Traub ont créé la surprise avec le succès des calculs financiers pour $d = 360$. La seule explication plausible est que la « dimension effective » du problème serait en fait très inférieure à 360. En d'autres termes, le volume mesuré ne s'étendrait pas vraiment à toutes les dimensions de l'espace. De façon similaire, une feuille de papier occupe trois dimensions, mais on ne perd pas grand-chose à considérer qu'elle n'a pas d'épaisseur.

Cette explication pourrait sembler condescendante: le calcul a été un succès non pas parce que l'outil était plus puissant, mais parce que le problème était plus simple qu'il n'y paraissait. Mais il faut noter que cette réduction effective de la dimension fonctionne même quand nous

ignorons lesquelles des 360 dimensions peuvent être négligées sans risque. C'est presque magique.

Dans quelle mesure ce phénomène est-il fréquent? Est-ce juste un coup de chance, ou bien est-il confiné à une classe étroite de problèmes? La réponse n'est pas encore totalement claire, mais une notion portant le nom de « concentration de mesure » donne des raisons d'être optimistes. Elle suggère que les univers de grande dimensionnalité sont essentiellement des mondes assez plats et lisses, analogues à une planète où régnerait une forte gravité et où la formation de paysages escarpés et accidentés serait trop coûteuse en énergie.

Regain d'intérêt pour une idée née à Los Alamos

La méthode de Monte Carlo n'est pas une idée nouvelle, et la variante de quasi-Monte Carlo non plus. Des simulations fondées sur un échantillonnage aléatoire ont été tentées il y a plus d'un siècle par Francis Galton, lord Kelvin et d'autres. À l'époque, ces scientifiques travaillaient

L'AUTEUR



Brian HAYES tient depuis 1993 la chronique *Computing Science* dans la revue *American Scientist*.

Article publié avec l'aimable autorisation de *American Scientist*.

UNE RECETTE DE NOMBRES QUASI-ALÉATOIRES

Les nombres nécessaires pour former des motifs de discrédance faible ont été étudiés pour l'intérêt qu'ils présentaient en eux-mêmes longtemps avant l'invention de la méthode de quasi-Monte Carlo. Dans les années 1930, le mathématicien néerlandais Johannes van der Corput s'est demandé si des nombres distincts en séquence infinie pouvaient être placés un par un sur l'intervalle unité $[0, 1]$ de telle sorte que la discrédance mesurée à chaque étape ne dépasse jamais une borne finie donnée. La réponse, prou-

vée une décennie plus tard par Tatyana van Aardenne-Ehrenfest, une compatriote de van der Corput, est négative: la discrédance croît sans limite. Néanmoins, les travaux de van der Corput ont mis en lumière toute une famille d'ensembles et de suites quasi-aléatoires à discrédance faible.

La procédure pour construire une séquence de van der Corput est particulière en cela qu'elle mélange des nombres (des entités mathématiques) avec des numéros (la représentation des nombres sous

forme de listes de chiffres). L'idée de base consiste à prendre un entier, d'inverser ses chiffres, de mettre une virgule devant, et enfin de traiter le résultat comme une fraction comprise entre 0 et 1. Les opérations peuvent être effectuées en n'importe quelle base. En base 2, par exemple, le nombre 100 (égal à 4 en base 10) devient 001 puis 0,001, qui est la représentation binaire de la fraction $1/8$.

Pour créer ainsi un motif quasi-aléatoire de N points en dimension 2, on part de la séquence de nombres entiers $i = 0, 1, 2, \dots, N-1$. Puis, pour chaque point, on fixe la coordonnée x égale à i/N , tandis que la coordonnée y correspondante est donnée par le processus d'inversion des chiffres de van der Corput pour i . Le résultat pour $N = 8$ points est présenté ci-contre. Le motif quasi-aléatoire de la figure 2d a été créé de la même façon avec $N = 1\,024$. De nombreux autres algorithmes pour

suites à discrédance faible ont été développés, mais la plupart d'entre eux reposent aussi sur le mélange des chiffres d'un nombre. Certaines de ces suites se généralisent facilement à des dimensions arbitrairement grandes.

Les motifs créés par de tels mécanismes maintiennent un équilibre délicat entre ordre et désordre. L'espacement entre les points est raisonnablement uniforme, si bien que n'apparaissent pas de vides ou de regroupements, qui augmentent la discrédance des ensembles aléatoires. Mais la distribution uniforme doit être obtenue sans laisser les points former des rangées ou autres structures régulières qui augmenteraient elles aussi la discrédance. Trouver le juste équilibre entre ces objectifs un peu contradictoires n'est pas difficile en deux dimensions, mais les compromis sont inévitables en dimension supérieure.

