

Visions en mathématiques

W. T. Gowers

GAFA 2000

Introduction

Lorsque l'on m'a demandé pour la première fois de prendre la parole à la conférence "Visions en mathématiques", j'ai eu ce qui me semble être une réaction typique. J'ai voulu essayer d'imiter Hilbert un siècle plus tôt, mais sachant que je ne pourrais pas égaler sa largeur de vue, j'ai été contraint de faire une sorte de compromis. Dans cet article, je discuterai de plusieurs problèmes ouverts, pas toujours dans des domaines que je connais bien, mais ils ne sont pas destinés à être une liste des questions les plus importantes en mathématiques, ni même des questions les plus importantes dans les domaines des mathématiques dans lesquels j'ai travaillé. Il s'agit plutôt d'une sélection personnelle de problèmes qui, pour une raison ou une autre, ont retenu mon attention au fil des ans.

Le titre de cet article fait référence à un thème général qui relie plusieurs problèmes non résolus en combinatoire. Je pense que ce qui les rend difficiles est, dans de nombreux cas, l'absence de théorèmes de "classification approximative" adéquats. J'expliquerai ce que j'entends par là au §3 et je donnerai des exemples de problèmes pour lesquels une classification pourrait s'avérer très utile. Au §4, je donne une liste plus hétéroclite de questions, avec de brefs exposés des raisons pour lesquelles je les trouve intéressantes. La plupart, mais pas toutes, des questions de cette liste sont très connues.

À mesure que la conférence approchait, j'ai décidé que je ne pouvais pas résister à au moins une tentative de "vision des choses" [Bu]. Je ne suis pas favorable à la prédiction des domaines des mathématiques qui seront à la mode dans quinze ans (et encore moins dans cinquante), mais au §2 j'ai fait une prédiction plus générale sur les mathématiques dans leur ensemble. Brièvement, ma thèse est qu'au cours du prochain siècle, les ordinateurs deviendront suffisamment compétents pour démontrer des théorèmes au point que la pratique de la recherche mathématique pure sera complètement révolutionnée, très probablement au point que nous serions réticents à reconnaître cela comme de la recherche mathématique. Les raisons de cette opinion sont exposées dans la section suivante.

Je ne revendique aucune originalité pour mes vues sur ce sujet. En effet, au moment de ma conférence, certaines des idées avaient été abordées par d'autres orateurs et lors des sessions de discussion. De plus, l'intérêt pour la démonstration automatique de théorèmes existe depuis de nombreuses années, et depuis la conférence, on m'a signalé des vues similaires exprimées par des personnes travaillant dans ce domaine. (Une référence est donnée à la fin de la section suivante.)

1. Les mathématiques existeront-elles en 2099 ?

Depuis au moins un siècle, les mathématiciens réfléchissent à la possibilité d'automatiser les mathématiques. Hilbert a célèbrement posé dans son dixième problème s'il existait un algorithme pour résoudre les équations diophantiennes, et a ensuite étendu cette question à celle de savoir s'il existait un algorithme capable de trouver une démonstration pour tout énoncé mathématique en ayant une. En 1936, Turing, tout aussi célèbrement, a formalisé la notion d'algorithme et a peu après démontré l'insolubilité du problème de l'arrêt, montrant ainsi qu'un tel algorithme n'existait pas. Bien que ce résultat, et les démonstrations ultérieures que plusieurs problèmes naturels et bien connus en mathématiques étaient également impossibles à résoudre systématiquement, aient pu sembler initialement quelque peu négatifs, ils avaient aussi un côté positif. L'idée que toute notre créativité et notre intuition pourraient être réduites à quelque chose de mécanique n'était, après tout, pas très attrayante. Le résultat de Turing fut donc un soulagement, puisqu'il laissait aux mathématiciens quelque chose à faire.

Plus récemment, avec l'essor de la théorie de la complexité computationnelle et de la NP-complétude, on a largement compris que ce soulagement était quelque peu mal placé, car, comme Gödel l'a d'ailleurs compris, un algorithme pour trouver des démonstrations n'est pas très utile s'il prend un temps prohibitif. Bien plus menaçant pour les mathématiciens serait un algorithme capable de déterminer, pour tout énoncé mathématique donné, s'il possède une démonstration de longueur au plus n , et ce en temps polynomial (idéalement avec un polynôme de petit degré). L'existence d'un tel algorithme est l'un des principaux problèmes ouverts des mathématiques, car elle est équivalente à $P = NP$, mais de nombreux mathématiciens se réconfortent du fait que personne n'a trouvé d'algorithme court pour résoudre un problème NP-complet et que la plupart des experts pensent qu'un tel algorithme n'existe pas.

Cependant, ce réconfort, même si vous êtes convaincu que $P \neq NP$, est toujours mal placé, car la recherche de démonstrations assez courtes d'énoncés mathématiques arbitraires est un très mauvais modèle de ce que font réellement les mathématiciens. En effet, il n'y a absolument aucune preuve que les mathématiciens surpassent les ordinateurs dans leur capacité à résoudre rapidement des questions NP-complètes. Si quelqu'un déterrait un savant idiot capable de déterminer dans sa tête si des collections de formules booléennes étaient satisfaisables, ou même de factoriser de grands entiers, on serait alors forcé de se poser des questions profondes, mais cela n'arrivera pas (une prédiction que je fais avec une grande confiance). Par conséquent, la seule explication possible du grand succès des mathématiciens au fil des siècles est que notre recherche de démonstrations de théorèmes est très restreinte. En effet, cela est évidemment vrai – nous avons tendance à préférer les questions intéressantes, compréhensibles, et qui ne semblent pas totalement hors de portée, et nous aimons que nos démonstrations fournissent des explications plutôt que de simples garanties formelles de vérité.

Cette ligne de pensée conduit à la question suivante : la notion de “bonne démonstration” peut-elle être formalisée de manière utile ? Si c'est le cas, alors l'impact sur les mathématiques pourrait bien être comparable à l'impact de la formalisation des notions de démonstration et d'algorithme. Il est peut-être peu probable qu'une définition émerge qui soit aussi incontestablement correcte que celle d'une fonction récursive ou d'une machine de Turing, mais il pourrait être réaliste d'essayer d'isoler certaines propriétés de ce que nous considérons généralement comme de bonnes démonstrations

et de les formaliser. Voici une liste incomplète des caractéristiques souvent associées aux bonnes démonstrations.

- (i) Les démonstrations sont généralement plus claires si elles ont une structure hiérarchique. Par exemple, si le théorème principal découle de trois lemmes qui ont tous des énoncés clairs, alors on peut isoler des parties de la démonstration, y réfléchir séparément et comprendre l'argument sans avoir à garder trop de choses en tête à la fois.
- (ii) Beaucoup de bonnes démonstrations sont des variantes d'arguments existants mieux connus. Cela les rend plus faciles à comprendre, car tout ce que l'on a à faire est de se concentrer sur les parties qui sont nouvelles.
- (iii) Il est plus facile de comprendre une démonstration qui peut être réduite à quelques idées de base. (Inutile de dire qu'il n'est pas facile de formaliser ce que l'on entend par "idée de base".)
- (iv) Parfois, la forme d'une démonstration est dictée dans une certaine mesure par l'existence de certains exemples et contre-exemples. Ceux-ci clarifient la démonstration en suggérant des approches qui pourraient fonctionner et en démontrant que d'autres approches ne peuvent pas fonctionner.
- (v) Une bonne définition peut grandement améliorer une démonstration. Dans des cas extrêmes, le choix de la bonne définition peut réduire un problème difficile à un simple exercice de vérification.
- (vi) Beaucoup des meilleures démonstrations, ou du moins les mieux rédigées, sont accompagnées d'une certaine indication sur la façon dont elles ont été découvertes. Pour dire les choses autrement, les démonstrations qui reposent sur des tours de magie, bien qu'elles puissent être impressionnantes et élégantes, ne sont pas complètement satisfaisantes en tant qu'explications.

Les caractéristiques ci-dessus sont interdépendantes, et elles ont toutes en commun qu'une démonstration possédant l'une de ces caractéristiques est, toutes choses étant égales par ailleurs, plus facile à comprendre qu'une démonstration qui n'en possède pas. Je voudrais me concentrer sur (vi) et l'adopter comme définition de travail de "l'explication" par opposition à la "garantie de vérité". C'est-à-dire que je considérerai une explication comme une démonstration accompagnée d'un récit de la façon dont la démonstration a été découverte. Ce récit ne devrait pas être quelque chose comme "J'étais dans mon bain et j'ai soudainement eu l'idée brillante suivante", mais plutôt une démystification de la démonstration : c'est-à-dire une démonstration convaincante de la façon dont un simple être humain (formé en mathématiques) pourrait chercher une démonstration et trouver celle-ci, sans compter sur une quantité excessive de chance.

Je suis conscient que cette définition (qui sera rendue plus précise dans un instant) ne capture pas toutes nos associations avec le mot "explication". En particulier, il pourrait être plus facile de montrer comment on est arrivé à un argument de force brute que de montrer comment on a pensé à une définition astucieuse qui, après l'avoir vue, rendait le résultat évident. Néanmoins, si nous souhaitons apprendre aux ordinateurs à trouver des démonstrations, il est probablement bon de réfléchir à la façon dont nous le faisons nous-mêmes.

Une façon d'expliquer une démonstration, au sens ci-dessus, est de montrer qu'elle peut être générée en utilisant ce que l'on pourrait considérer comme des réflexes mathématiques standard. Certains d'entre eux pourraient être des méthodes d'approche très générales, et d'autres pourraient consister à essayer d'appliquer des techniques bien connues du domaine des mathématiques en question. Voici quelques exemples de ce que j'entends par méthodes d'approche générales. Pour ma conférence, j'ai dessiné des organigrammes, mais mes compétences en TeX sont très basiques, je vais donc les représenter ici comme des pseudo-algorithmes. Une discussion beaucoup plus détaillée et sophistiquée des réflexes mathématiques peut être trouvée dans les livres classiques de Polya, tels que [P, vols. I et II].

A. Généralisation.

1. p est-il un cas particulier de q ? Si cela ne semble pas être le cas, passer à 2, sinon passer à 3.
2. Peut-on modifier q de manière appropriée ? Si oui, passer à 1, sinon passer à 5.
3. q est-il vrai ? Si cela pourrait être le cas, passer à 4, sinon passer à 2.
4. Essayer de prouver q à la place de p .
5. Essayer une approche différente.

B. Modification d'un argument connu.

1. Trouver un énoncé connu q qui ressemble à p .
2. La démonstration de q peut-elle être facilement modifiée pour donner p ? Si oui, FIN, sinon passer à 3.
3. Identifier un énoncé r faisant partie de la démonstration de q , tel que l'énoncé analogue s , que l'on souhaiterait comme partie de la démonstration de p , n'est pas facile à prouver.
4. Si s est manifestement faux, passer à 6, sinon passer à 5.
5. Essayer de prouver s .
6. Essayer une approche différente.

C. Repérer des motifs dans les cas particuliers.

1. Trouver un cas particulier r de p .

2. Si r est trivial, passer à 3, sinon passer à 4.
3. Puis-je identifier le prochain cas particulier le plus simple ? Si oui, l'appeler r et passer à 2, sinon passer à 8.
4. Puis-je voir comment prouver r ? Si oui, passer à 5, sinon passer à 7.
5. La démonstration de r se généralise-t-elle en une démonstration de p ? Si oui, FIN, sinon passer à 6.
6. Puis-je trouver une démonstration différente, plus instructive, de r ? Si oui, passer à 5, sinon passer à 3.
7. Essayer de prouver r à la place de p .
8. Essayer une approche différente.

D. Habituation, ou s'ennuyer.

1. Ai-je réussi à penser à des idées intéressantes liées à p au cours des deux dernières heures ? Si oui, passer à 2, sinon passer à 3.
2. Penser à l'une des idées intéressantes.
3. Se faire une bonne tasse de café et penser à autre chose pendant un moment.

Quant aux réflexes de nature plus spécifique, je veux dire des choses comme essayer d'appliquer des théorèmes ou des techniques bien connus. Pour cela, il faut avoir une certaine capacité à classer les problèmes, afin d'avoir une idée des théorèmes ou des techniques qui peuvent être utiles. Plutôt que de donner plusieurs exemples d'utilisation de méthodes standard pour résoudre des problèmes, je reviens à la question de l'automatisation des mathématiques et présente un dialogue imaginaire entre un mathématicien et un ordinateur dans deux ou trois décennies. L'idée du dialogue est que l'ordinateur est très utile au mathématicien, sans rien faire de particulièrement intelligent. Cela représente une étape intermédiaire peu menaçante entre ce que nous avons maintenant, des ordinateurs qui agissent comme des esclaves effectuant des calculs incroyablement ennuyeux pour nous, et l'automatisation complète des mathématiques. J'ai écrit le dialogue en anglais, mais il est censé être la traduction d'un langage plus formel qui n'a pas encore été inventé. (Je discuterai un peu plus de ce point plus tard.)

MATHÉMATICIEN. L'affirmation suivante est-elle vraie ? Soit $\delta > 0$. Alors pour N suffisamment grand, tout ensemble $A \subset \{1, 2, \dots, N\}$ de taille au moins δN contient un sous-ensemble de la forme $\{a, a + d, a + 2d\}$?

ORDINATEUR. Oui. Si A est non vide, choisir $a \in A$ et poser $d = 0$.

MATHÉMATICIEN. D'accord, d'accord, mais que se passe-t-il si d n'est pas autorisé à être nul ?

ORDINATEUR. Avez-vous essayé une récurrence sur N , avec un certain $\delta = \delta(N)$ tendant vers zéro ?

MATHÉMATICIEN.. Cette idée n'est d'aucune aide. Donnez-moi quelques exemples, s'il vous plaît.

ORDINATEUR. L'algorithme glouton évident donne l'ensemble

$$\{1, 2, 4, 5, 10, 11, 13, 14, 28, 29, 31, 32, 37, 38, 40, 41, \dots\}.$$

Je remarque que de grandes parties de l'ensemble sont des translations d'autres parties. En fait, cet ensemble ressemble beaucoup à l'ensemble de Cantor, ce qui donne une borne de $\delta \geq N(\log 2 / \log 3) - 1$.

Quant aux méthodes aléatoires, choisissons chaque élément de $\{1, 2, \dots, N\}$ pour être dans A avec probabilité δ , les choix étant indépendants. Le nombre attendu de sous-ensembles de A de la forme $\{a, a+d, a+2d\}$ est $\delta^3 N^2$, à une constante absolue près. Par conséquent, si $\delta < N^{-2/3}$, il est possible qu'il n'y en ait aucun. La taille attendue de A est δN , et des estimées standard nous disent qu'il est très probable que A soit à peu près de cette taille.

En appliquant une autre idée standard, nous remarquons que si le nombre d'ensembles dans A de la forme $\{a, a+d, a+2d\}$ est au plus la moitié du nombre de points dans A , alors nous pouvons supprimer un élément de chacun d'eux et nous retrouver avec une proportion substantielle de A . Cela nous indique que nous pouvons choisir n'importe quel δ satisfaisant $C\delta^3 N^2 \leq \delta N$, donc en particulier nous pouvons trouver un certain δ de la forme $cN^{-1/2}$.

MATHÉMATICIEN. Eh bien, les méthodes aléatoires donnent souvent la meilleure réponse pour ce genre de problèmes. Essayons de prouver que $cN^{-1/2}$ est le meilleur possible.

ORDINATEUR. [Pause de 0,001 seconde] En fait, ce n'est pas le cas. Behrend a trouvé une bien meilleure borne en 1946. [Télécharge l'article]

MATHÉMATICIEN. Oh là là, je n'ai plus d'idées alors. Pourriez-vous me faire une suggestion, par hasard ?

ORDINATEUR. Nous avons un ensemble A . Nous voulons prouver qu'un sous-ensemble d'une certaine forme existe. La meilleure façon de prouver l'existence est souvent de compter.

MATHÉMATICIEN. [Intrigué] Oui, mais qu'est-ce que cela signifierait pour un problème comme celui-ci ?

ORDINATEUR. Ici, nous souhaitons compter le nombre de solutions (x, y, z) de l'équation linéaire unique $2y = x + z$. Un outil standard dans de telles situations est l'analyse de Fourier, les sommes

exponentielles, la méthode du cercle, appelez-la comme vous voulez.

MATHÉMATICIEN. Où est la structure de groupe ?

ORDINATEUR. Identifier $\{1, 2, \dots, N\}$ avec $\mathbb{Z}/N\mathbb{Z}$. Ce n'est pas joli mais ça marche parfois. J'écris $\mathbb{Z}_N$ pour $\mathbb{Z}/N\mathbb{Z}$ et $\hat{A}(r)$ pour $\sum_{s \in A} \exp(2\pi i r s/N)$. Alors le nombre de triplets $(a, a+d, a+2d) \in A^3$ est $N^{-1} \sum_{r \in \mathbb{Z}_N} \hat{A}(r)^2 \hat{A}(-2r)$.

MATHÉMATICIEN. Nous essayons de montrer que cette expression ne peut pas être nulle à moins que δ ne soit très petit.

ORDINATEUR. [Entraîné à amuser les mathématiciens] C'est en effet presque équivalent à notre problème initial. Merci. Remarquez que $N^{-1} \hat{A}(0)^3 = \delta^3 N^2$, exactement le nombre que nous avons auparavant lorsque A était choisi aléatoirement. Ce nombre est grand et positif – un bon signe.

MATHÉMATICIEN. Nous voudrions donc que le reste de la somme soit suffisamment petit pour ne pas annuler le terme nul.

ORDINATEUR. Oui. J'essaie quelques idées évidentes comme l'inégalité de Cauchy-Schwarz. Malheureusement, il ne semble pas y avoir de raison pour que le reste de la somme soit petit.

MATHÉMATICIEN. Montrez-moi vos calculs, je vous prie.

ORDINATEUR. Les voici. [Les affiche.] Ils me permettent, ou plutôt nous permettent, de prouver un résultat partiel. Il stipule que si $\max_{r \neq 0} |\hat{A}(r)| \leq c\delta^2 N$, alors l'image de A dans $\mathbb{Z}_N$ contient de nombreux ensembles de la forme souhaitée.

Ce dialogue pourrait être poursuivi pour montrer ce qu'il faut faire lorsque $\hat{A}(r)$ est grand pour un certain r . Cependant, l'ordinateur a déjà guidé le mathématicien vers une idée importante qui est la base de la démonstration par Roth de son théorème sur les progressions arithmétiques [Ro1]. Terminer la démonstration n'est pas très difficile, bien que cela puisse l'être pour le mathématicien ci-dessus, qui semble un peu trop dépendant de son ordinateur.

Remarquez que l'ordinateur essaie à chaque étape des idées standard : récurrence, algorithme glouton, méthodes aléatoires, dénombrement, prise de la transformée de Fourier et application de Cauchy-Schwarz à une somme de produits. Qu'est-ce qui lui fait penser à ces idées standard, plutôt qu'à d'autres complètement inappropriées ? Une partie de la réponse réside dans la façon dont le problème est initialement présenté à l'ordinateur. J'envisage non pas l'énoncé formel donné au début du dialogue, mais quelque chose de plus interactif. Par exemple, cette étape du processus pourrait être guidée par des menus. On choisirait d'abord Combinatoire, ou peut-être Théorie combinatoire des nombres. Cela ferait apparaître un autre menu parmi lequel on choisirait Problèmes extrémaux. Pour le menu suivant, on choisirait le mot Maximiser. Le choix suivant serait Sur des ensembles A d'entiers. Puis on choisirait Cardinalité de A . Ensuite, un menu apparaîtrait avec un choix de contraintes, ou peut-être faudrait-il les saisir : $A \subset \{1, 2, \dots, N\}$ et A ne contient pas de progression arithmétique de longueur trois.

À la fin d'un processus comme celui-ci, l'ordinateur aurait de nombreuses idées sur la façon dont le problème était conventionnellement classé. Par exemple, il serait naturel pour l'ordinateur de penser à la méthode du cercle parce que de nombreux problèmes en théorie additive des nombres impliqueraient des concepts tels que le nombre de progressions arithmétiques de longueur trois dans un ensemble A . L'ordinateur pourrait alors chercher une liste d'astuces standard et trouver le slogan dans le dialogue ci-dessus concernant le dénombrement des solutions d'une équation linéaire.

Remarquez qu'il est important pour l'ordinateur qu'il ait accès à une base de données mathématiques beaucoup plus sophistiquée que tout ce que nous avons actuellement. On pourrait se demander comment l'ordinateur a pu découvrir l'article de Behrend si rapidement, car la compréhension du texte ordinaire est notoirement difficile pour les ordinateurs. J'imaginais qu'un travail considérable avait été consacré à la classification des articles de sorte que l'article de Behrend apparaisse dans la base de données non pas comme Proc. Nat. Acad. Sci. 23 (1946), 331-332, mais plutôt comme quelque chose comme Maximiser - Taille de A - Sous réserve de - (1) $A \subset \{1, 2, \dots, N\}$ et (2) A ne contient pas de 3-AP - Borne inférieure. En plus de l'article, il y aurait un résumé du résultat principal - Borne inférieure au moins $n \exp(-c\sqrt{\log n})$.

Je ne vois aucun obstacle à la création d'une base de données utile du type ci-dessus, autre que le temps et l'effort qu'elle impliquerait. La structure de la base de données serait très bien adaptée à Internet, ce qui faciliterait également la répartition du travail entre de nombreuses personnes, donc peut-être que même cet obstacle n'est pas si grand. Il y a bien sûr des parties des mathématiques qui ne se prêtent pas particulièrement facilement à une classification du type ci-dessus, mais ce n'est pas une raison pour rejeter toute l'idée. Même une base de données ne couvrant que certains domaines des mathématiques serait utile.

Encoder et classer les méthodes générales, les astuces, les règles empiriques et ainsi de suite serait évidemment plus difficile. L'une des principales raisons en est qu'elles sont souvent formulées de manière imprécise, et le traitement des énoncés vagues est un problème bien connu en informatique. Cependant, difficile ne signifie pas impossible et des progrès importants ont déjà été réalisés dans ce domaine. De plus, au moins certains des principes vagues des mathématiques ne semblent pas présenter de difficultés sérieuses. Considérons par exemple la conséquence simple mais extrêmement utile suivante de l'inégalité de Cauchy-Schwarz :

$$\left(\sum_{i=1}^n |a_i| \right)^2 \leq n \sum_{i=1}^n |a_i|^2.$$

Cette inégalité est utile car la norme ℓ_2 est souvent plus facile à manier que la norme ℓ_1 . Cependant, dans de nombreux contextes, l'inégalité est bien trop faible, on a donc tendance à ne l'utiliser que lorsqu'il y a une raison de croire que les a_i sont à peu près de même taille. Cela ressemble à une idée vague, et de nombreux mathématiciens, lorsqu'ils réfléchissent à un problème, n'auront pas besoin de la clarifier davantage, mais une clarification supplémentaire est parfaitement possible. Par exemple, une autre inégalité très simple mais étonnamment utile est

$$\sum_{i=1}^n |a_i|^2 \leq \max_{1 \leq i \leq n} |a_i| \sum_{i=1}^n |a_i|,$$

qui nous dit que la première inégalité sera précise à une constante C près si $\max_{1 \leq i \leq n} |a_i| \leq C n^{-1} \sum_{i=1}^n |a_i|$.

Ce n'est pas une condition nécessaire, mais elle pourrait former l'un d'un certain nombre de tests qu'un ordinateur pourrait appliquer (ou plutôt, tenter d'appliquer) chaque fois qu'il utilise la première inégalité, pour voir si elle est susceptible d'être coûteuse.

S'il y avait une tentative sérieuse de codifier les nombreux principes et méthodes de travail des mathématiciens, ce serait un exercice extrêmement utile de revenir aux mathématiques que nous connaissons déjà et d'essayer de réécrire les démonstrations comme des explications (au sens restreint que j'ai défini plus tôt). De nombreux manuels commencent par des définitions qui se justifient plus tard en menant à des résultats intéressants : cela serait interdit dans le genre de manuel que j'ai en tête. Au lieu de cela, chaque définition devrait être présentée comme découlant naturellement de la recherche de la solution d'un problème. Bien sûr, le problème lui-même devrait également être justifié ; ou bien on pourrait prendre comme point de départ que le problème mérite d'être étudié, laissant la justification à quelqu'un d'autre. (La question de la justification des problèmes ouverts a été discutée par Gromov lors de la conférence. Voir sa contribution à ce volume, en particulier les fins des §1 et §2. Par souci de brièveté, j'ai ignoré de nombreux aspects de la pratique mathématique, comme la formulation de conjectures et la proposition de définitions. Celles-ci doivent également être prises au sérieux si l'on veut automatiser les mathématiques.)

Même si l'automatisation des mathématiques s'avère beaucoup plus difficile que je ne l'ai supposé, cet exercice de réécriture pourrait avoir des avantages énormes, car, sans surprise, la question de savoir comment nous pourrions enseigner les mathématiques aux ordinateurs est étroitement liée à la question de savoir comment nous devrions les enseigner aux êtres humains. Je n'ai que faire de ceux qui croient que le talent pour la recherche mathématique est une qualité mystérieuse et inenseignable que l'on a ou que l'on n'a pas. Bien sûr, il existe une chose telle que les capacités naturelles, mais la plupart des mathématiciens développent leurs capacités en mûrissant. À l'heure actuelle, ils ont tendance à le faire avec très peu de conseils extérieurs et très peu d'aide de la littérature. Une série de livres visant à rendre transparents les processus de pensée qui ont conduit à certains résultats pourrait montrer aux jeunes mathématiciens pourquoi la recherche n'est pas aussi impossible qu'elle en a l'air et transmettre beaucoup plus efficacement les astuces utilisées par ceux qui ont plus d'expérience.

Deux exemples de ce que je veux dire par la justification des définitions sont fournis par l'homologie simpliciale et la fonction zêta de Riemann. Lorsque l'on m'a enseigné l'homologie simpliciale, on m'a présenté une définition compliquée (comme elle semblait à l'époque) qui, miraculeusement, conduisait ensuite à des résultats intéressants et bien plus concrets comme les théorèmes de point fixe. De même, une introduction typique à la fonction zêta commence par la formule du produit d'Euler, qui conduit rapidement à des théorèmes intéressants sur les nombres premiers, mais n'est pas présentée comme le résultat inévitable d'une réflexion sur les nombres premiers. Pour ces deux concepts, je peux donner des récits plausibles (mais totalement anhistoriques et nullement uniques) de la façon dont ils ont pu être inventés. J'ai l'intention de les mettre sur ma page web personnelle (<http://www.dpmms.cam.ac.uk/~wtg10/>), mais je ne prévois pas de m'y mettre avant environ janvier 2001. J'espère aussi donner quelques exemples de la façon dont des idées apparemment brillantes peuvent être expliquées comme le résultat naturel de méthodes familières de recherche de démonstrations.

Je suis pleinement conscient que ce que j'ai décrit jusqu'ici est loin de l'automatisation complète des mathématiques. En particulier, j'ai ignoré certains problèmes notoires en intelligence artificielle. J'ai brièvement discuté du vague, mais l'exemple que j'ai donné était simple et ne se généralise pas de manière évidente. Je n'ai pas mentionné la difficulté de la reconnaissance de formes, bien que la reconnaissance de la similarité soit une partie très importante de la recherche mathématique. Ce que j'ai essayé de montrer, c'est que ces difficultés ne sont pas une bonne raison pour ne rien faire. Réfléchir à ses propres processus de pensée ne nécessite aucune compétence en informatique et est potentiellement très bénéfique pour les autres ; et un ordinateur disposant d'une grande quantité d'informations mais d'une intelligence très modeste pourrait encore être extrêmement utile.

Je n'ai pas non plus discuté des nombreux efforts qui ont déjà été faits dans le domaine général de la démonstration automatique de théorèmes, un domaine sur lequel (comme c'est probablement évident) je sais peu de choses. Mon impression est que la plupart des progrès jusqu'à présent ont été réalisés par une approche ascendante – en essayant d'amener les ordinateurs à découvrir des théorèmes simples mais intéressants en partant presque de zéro. Ce dont j'ai parlé est davantage un projet descendant auquel tout mathématicien pourrait contribuer. Il se pourrait qu'éventuellement une telle approche descendante puisse réduire le problème général de l'automatisation des mathématiques à un grand nombre de problèmes plus petits et plus faciles à gérer, de sorte que peut-être la communauté mathématique pourrait rencontrer les démonstrateurs automatiques de théorèmes à mi-chemin. (Une façon de se renseigner sur l'état de l'art en démonstration automatique de théorèmes est de visiter le site web <http://www.math.temple.edu/~zeilberg/OPINIONS.html> et de suivre les liens, en particulier à partir de l'opinion assez extrême n° 36.)

Enfin, voici donc ma supposition sur l'impact des ordinateurs sur les mathématiques au cours du prochain siècle. Bien que certaines des difficultés impliquées dans la démonstration automatique de théorèmes soient formidables, un siècle est une très longue période en informatique, donc je m'attends à ce que les ordinateurs soient meilleurs que les humains pour démontrer des théorèmes en 2099. Je ne me sens pas particulièrement heureux à ce sujet, mais je m'attends à ce que l'impact sur ma propre vie mathématique soit entièrement positif, car une base de données semi-intelligente du type que j'ai décrit serait une ressource merveilleuse et éliminerait une grande partie du travail fastidieux de la recherche. Il pourrait même y avoir un âge d'or où les ordinateurs seraient bons pour les exercices mais pas encore bons pour avoir des intuitions profondes. Alors, au lieu de perdre une semaine à ne pas remarquer qu'un lemme espéré avait un simple contre-exemple, on pourrait faire vérifier cela par l'ordinateur. En d'autres termes, les ordinateurs feraient encore les parties ennuyeuses pour nous, mais elles ne seraient pas tout à fait aussi ennuyeuses qu'elles le sont maintenant. Cependant, un tel âge d'or, s'il se produit, est peu susceptible de durer longtemps. L'étape suivante pourrait être celle où seuls quelques mathématiciens exceptionnels pourraient découvrir des démonstrations inaccessibles aux ordinateurs. Peut-être que d'autres se concentreraient sur l'enseignement, mais les ordinateurs deviendraient probablement meilleurs que nous dans ce domaine également, du moins s'ils avaient eux-mêmes été formés d'une manière à peu près similaire à celle que j'ai esquissée. En fin de compte, le travail du mathématicien consisterait simplement à apprendre à utiliser efficacement les machines à démontrer des théorèmes et à trouver des applications intéressantes pour elles. Ce serait une compétence précieuse, mais ce ne serait guère des mathématiques pures telles que nous les connaissons aujourd'hui.

2. Structure approximative et classification, et problèmes connexes

Dans son exposé, Gromov a attiré l'attention sur le fait que peu de façons étaient connues de définir des variétés. (Voir la sous-section “Randomisation” au §5 de sa contribution à ce volume.) Pour une raison différente, quelque chose de similaire est vrai pour les graphes. Cela peut sembler à première vue une affirmation ridicule – tout ce que vous avez à faire pour définir un graphe est de prendre l'ensemble de toutes les paires d'éléments de l'ensemble $\{1, 2, \dots, n\}$ et de passer à un sous-ensemble arbitraire. Cependant, bien qu'il soit vrai qu'un objet mathématique n'a pas besoin de satisfaire des contraintes délicates pour être un graphe, on est toujours confronté au problème purement logique que le nombre de descriptions compréhensibles est limité. De plus, bien que les théoriciens des graphes parlent souvent comme s'ils considéraient des graphes individuels, la plupart des problèmes de théorie des graphes concernent, strictement parlant, des familles $(G_n)_{n=1}^\infty$ de graphes, où le nombre de sommets de G_n est généralement n , ou peut-être une fonction simple de n qui tend vers l'infini.

Voici une liste de quelques façons courantes de définir des familles de graphes.

(i) Motifs répétitifs. Quelques exemples en sont

- (a) K_n , le graphe complet d'ordre n ;
- (b) P_n , le chemin à n sommets ;
- (c) C_n , le cycle à n sommets ;
- (d) $K_{1,n}$, l’“étoile” avec un sommet relié à n autres, et aucune autre arête ;
- (e) $\{0, 1\}^n$, le cube discret de dimension n ;
- (f) un arbre dyadique avec une racine et n couches supplémentaires.

(ii) Constructions simples à partir de graphes connus. Il s'agit notamment des produits, des quotients, des unions disjointes, de la suppression de quelques arêtes, etc.

(iii) Utilisation de l'algèbre et de la théorie des nombres. Deux exemples notables sont le graphe de Paley et les graphes de Ramanujan de Lubotzky, Phillips et Sarnak [LPS]. En général, les graphes de Cayley ont souvent des propriétés intéressantes.

(iv) Graphes aléatoires. On ne sait pas très bien en quel sens cette expression définit une famille de graphes : voir la discussion ci-dessous.

Cette liste pourrait sans aucun doute être poursuivie, mais je serais surpris si quelqu'un pouvait me donner dix autres classes de définitions vraiment distinctes et intéressantes.

Considérons un moment ce que signifie dire qu'un graphe est “aléatoire”. Nous aimerions faire des déclarations telles que “Un graphe aléatoire sur n sommets avec une probabilité d'arête de $\frac{1}{3}$ a environ $\frac{1}{3} \binom{n}{2}$ arêtes.” Parfois, on peut donner un sens à une telle déclaration en insérant les mots “avec une probabilité élevée” à l'endroit approprié. Cependant, il existe des contextes où cela n'est pas adéquat. Par exemple, il existe une classe de graphes, connus sous le nom d'expandeurs, qui

ont des applications pratiques directes dans la conception de certains algorithmes. On sait que les graphes choisis aléatoirement sont, avec une probabilité élevée, de bons expandeurs. D'un autre côté, face à un graphe particulier choisi aléatoirement, comment être sûr qu'il ne fait pas partie des exceptions ? Ce n'est peut-être pas un problème pratique réel, car un graphe qui fonctionne avec une probabilité $1 - 10^{-8}$ est tout aussi utile qu'un graphe qui fonctionne avec une certitude totale (bien que j'aie entendu dire que les personnes de l'industrie qui achètent des systèmes informatiques se méfient beaucoup des éléments aléatoires dans la conception). Néanmoins, un développement très important dans la théorie des graphes a été la prise de conscience que plusieurs propriétés partagées par presque tous les graphes sont équivalentes et servent de définition utile du "caractère aléatoire".

Ces propriétés s'appliquent à un graphe G avec n sommets tel que presque chaque sommet a un degré d'environ $n/2$. Elles montrent que G ressemble à un graphe où chaque arête est choisie indépendamment avec une probabilité $1/2$. (Des résultats similaires sont valables pour d'autres probabilités.) En voici trois.

- (i) Pour deux grands ensembles A et B de sommets, le nombre $e_G(A, B)$ de paires $(x, y) \in A \times B$ telles que xy est une arête de G est $|A||B|/2 + o(n^2)$ (et donc approximativement le nombre que l'on attendrait si G était un graphe aléatoire).
- (ii) Pour tout graphe (petit) fixé H , le nombre de sous-graphes (induits) de G isomorphes à H est approximativement ce que le nombre attendu aurait été si G avait été choisi aléatoirement.
- (iii) Définir la matrice d'adjacence A de G en posant $A(x, y)$ à 1 si xy est une arête et 0 sinon. Alors la deuxième plus grande valeur propre (en module) de A est $o(n)$.

En ce qui concerne la propriété (i) ci-dessus, définissons la densité $\delta_G(A, B)$ comme $e_G(A, B)/|A||B|$. Une autre façon d'énoncer la propriété est de dire que $\delta_G(A, B)$ est approximativement $1/2$ pour tous les grands ensembles A et B . Fait intéressant, on peut remplacer (ii) par l'hypothèse apparemment faible que G contient approximativement le nombre attendu de cycles de longueur quatre.

L'un des premiers articles dans lesquels il a été montré que la petitesse de la deuxième plus grande valeur propre impliquait des propriétés de caractère aléatoire est dû à Alon et Milman [AM]. Le premier article traitant explicitement de l'idée de graphes pseudo-aléatoires est dû à Thomason [T1]. L'article de Chung, Graham et Wilson [CGW] a rendu tout le domaine beaucoup plus systématique, donnant une longue liste de propriétés de caractère aléatoire équivalentes. Ils ont appelé un graphe satisfaisant l'une (et donc toutes) de ces propriétés quasi-aléatoire, peut-être pour clarifier la distinction entre cette idée et le concept plutôt différent de pseudo-aléatoire en informatique.

Je voudrais tirer deux leçons des résultats ci-dessus. La première est que dans de nombreux contextes, ils permettent de donner une signification claire au concept de graphe aléatoire, même lorsque le graphe a été explicitement choisi. (Par exemple, le graphe de Paley est connu pour être quasi-aléatoire en ce sens.) La seconde est qu'ils soutiennent notre intuition selon laquelle deux graphes aléatoires sont fondamentalement les mêmes. Bien sûr, ils ne seront pas isomorphes, mais ils partagent toutes les propriétés qui nous intéressent, donc la différence entre eux n'est pas importante.

Je peux maintenant expliquer ce que j'entends par "classification approximative". Il est clair qu'il

il y a trop de graphes pour qu'il soit raisonnable d'essayer de les classer tous à isomorphisme près. Cependant, comme je viens de le commenter, nous aimons souvent considérer des graphes non isomorphes comme fondamentalement les mêmes. Une métrique était implicitement présente dans cette discussion, qui peut être définie comme suit. Étant donné deux graphes G et H , tous deux avec n sommets, et une bijection ϕ entre $V(G)$ et $V(H)$, définir

$$d_\phi(G, H) = n^{-2} \max_{A, B \subset V(G)} |e_G(A, B) - e_H(\phi A, \phi B)|$$

et soit $d(G, H)$ le minimum de tous les $d_\phi(G, H)$. Si G et H sont proches dans cette métrique, cela signifie qu'il existe une correspondance biunivoque entre leurs sommets telle que les densités $\delta_G(A, B)$ dans G sont approximativement les mêmes que les densités correspondantes dans H , au moins lorsque A et B sont grands. Un graphe est quasi-aléatoire si et seulement s'il est proche d'un graphe aléatoire typique dans cette métrique.

Par une classification approximative des graphes, je veux dire un choix d'un petit nombre de graphes représentatifs, tels que tout graphe est proche dans la métrique ci-dessus de l'un des représentants. Une telle classification sera utile si les graphes représentatifs peuvent être facilement décrits et analysés. Il n'est pas du tout évident qu'une classification utile doive exister, mais un résultat important, le lemme de régularité de Szemerédi [S1], montre qu'elle existe.

Pour énoncer le résultat, introduisons une définition supplémentaire. Si G est un graphe et A et B sont des ensembles de sommets dans G , disons que la paire (A, B) est ϵ -régulière si, chaque fois que A' et B' sont des sous-ensembles de A et B avec $|A'| \geq \epsilon|A|$ et $|B'| \geq \epsilon|B|$, on a $|d_G(A', B') - d_G(A, B)| \leq \epsilon$. Si A et B sont disjoints, cette condition, qui est très similaire à la première propriété des graphes quasi-aléatoires listée ci-dessus, stipule que les arêtes entre A et B se comportent beaucoup comme celles d'un graphe biparti aléatoire sur A et B avec probabilité d'arête $d_G(A, B)$. Le lemme de régularité de Szemerédi est le suivant.

Lemme 1. *Soit $\epsilon > 0$. Alors il existe un entier positif $k = k(\epsilon)$ tel que, étant donné un graphe fini quelconque G , les sommets de G peuvent être partitionnés en k classes $V_1, \dots, V_k$, toutes de taille $\lfloor n/k \rfloor$ ou $\lceil n/k \rceil$, telles que toutes les paires (V_i, V_j) (avec $i < j$) sauf au plus $\epsilon \binom{k}{2}$ sont ϵ -régulières.*

Une façon courante de définir des graphes est de choisir quelques ensembles $V_1, \dots, V_k$ et une matrice symétrique $(p_{ij})_{i,j=1}^k$ à valeurs dans $[0, 1]$, et de relier $x \in V_i$ à $y \in V_j$ avec probabilité p_{ij} , tous les choix étant indépendants. Le lemme de régularité de Szemerédi stipule que tout graphe est comme cela, du moins à proximité dans la métrique définie précédemment. Comme les graphes aléatoires sont assez bien compris, cela nous donne une classification utile.

Une conséquence de ce résultat est que si G est un graphe, et H est un sous-graphe induit de G formé en choisissant la moitié des sommets de G au hasard, alors H ressemble beaucoup à G . La raison en est que, avec une probabilité élevée, l'ensemble des sommets de H sera composé d'environ la moitié de chacun des ensembles V_i de la partition donnée par le lemme de Szemerédi, et la régularité d'une paire (V_i, V_j) est héritée par les grands sous-ensembles (V'_i, V'_j) , au prix d'une légère détérioration du paramètre ϵ . Ce principe général concernant les sous-graphes aléatoires est difficile à démontrer sans utiliser un résultat du type du lemme de régularité (bien qu'il existe des

versions plus faibles suffisantes).

Un inconvénient majeur du lemme de régularité est que la dépendance de k par rapport à ϵ est très mauvaise – une tour de deux de hauteur proportionnelle à ϵ^{-5} . De plus, il a été montré que cette situation fâcheuse est nécessaire. Dans [G1] un graphe est construit qui donne une borne inférieure pour k d’une tour de deux de hauteur proportionnelle à $\epsilon^{-1/16}$. Cela limite l’utilisation du lemme aux situations où la dépendance par rapport à ϵ n’est pas importante.

Avant de passer aux problèmes ouverts, permettez-moi de discuter brièvement de l’autre moitié de mon titre : la structure approximative. Les mathématiciens varient considérablement dans la quantité de structure qu’ils aiment, et cela permet de les classer approximativement de manière utile. À une extrémité du spectre se trouvent ceux qui partent d’autant d’hypothèses que même trouver un nouvel objet satisfaisant les hypothèses est un grand accomplissement. Ceux qui y parviennent sont souvent récompensés par la découverte de motifs beaux et détaillés. À l’autre extrémité se trouvent les mathématiciens qui aiment plus de liberté. Ils opèrent sous moins de contraintes, donc les objets qu’ils étudient existent en grande profusion et les problèmes qui se posent sont d’un caractère très différent.

Comme ma discussion sur la classification l’indique déjà, cette distinction, bien qu’elle existe indubitablement, n’est pas aussi grande qu’elle pourrait le paraître. Je voudrais l’estomper un peu plus en montrant que des motifs peuvent émerger de manière inattendue même dans des problèmes très combinatoires.

Il existe un théorème de Freiman [Fr1,2] qui illustre extrêmement bien cela. Soit A un ensemble d’entiers, et écrivons $A + A$ pour l’ensemble $\{x + y : x, y \in A\}$. Il est facile de montrer que si $|A| = n$ alors $2n - 1 \leq |A + A| \leq \frac{1}{2}n(n + 1)$ et que les deux bornes sont optimales. Le théorème de Freiman donne une description complète de A lorsque $A + A$ est petit, ou pour être précis, lorsque $|A + A| \leq Cn$ pour une certaine constante fixée C . Une façon évidente de construire un tel ensemble est de laisser P être une progression arithmétique de taille au plus $Cn/2$ et de laisser A être un sous-ensemble arbitraire de P de taille n . Cependant, ce n’est pas la seule source d’exemples : on peut prendre une progression arithmétique “bidimensionnelle”, qui est un ensemble de la forme

$$\{a + r_1d_1 + r_2d_2 : 0 \leq r_1 < m_1, 0 \leq r_2 < m_2\}.$$

On peut vérifier que si A est une progression bidimensionnelle, alors $|A + A| \leq 4|A|$. Plus généralement, $A + A$ sera petit si A est un grand sous-ensemble d’une progression arithmétique de dimension d pour un certain petit d . (La définition devrait être évidente.)

Un autre exercice facile est de montrer que si $|A + A| = 2n - 1$, alors A doit être une progression arithmétique de longueur n . Plus difficile que cela, mais encore pas trop difficile, est de montrer que si $|A + A|$ est juste un peu plus grand que $2n$, au plus $2.1n$ par exemple, alors il existe une progression arithmétique de longueur légèrement supérieure à n qui contient A . Le résultat le plus fort de ce type est dû à Freiman.

Théorème. Soit A un ensemble d’entiers de taille n . Si $0 \leq b \leq n - 3$ et $|A + A| \leq 2n - 1 + b$, alors A est contenu dans une progression arithmétique de longueur $n + b$.

Ce résultat montre que A est forcé d'être unidimensionnel si $|A + A| \leq 3n - 4$. Cette borne est optimale, car si m est un entier supérieur à $2n - 3$ et A est l'ensemble $\{1, 2, \dots, n - 1, m\}$, alors $|A + A| = 3n - 3$.

Le théorème principal de Freiman est beaucoup plus profond et nous dit ce qui se passe lorsque l'on relâche progressivement la condition sur la taille de l'ensemble somme. Il stipule que les seuls ensembles avec de petits ensembles sommes sont de grands sous-ensembles de progressions arithmétiques de basse dimension.

Théorème. Pour tout C , il existe des constantes $K = K(C)$ et $d = d(C)$ telles que si A est un ensemble quelconque de n entiers et $|A + A| \leq Cn$, alors A est un sous-ensemble d'une progression arithmétique généralisée de dimension au plus d et de cardinalité au plus Kn .

Remarquez que l'hypothèse du théorème de Freiman est combinatoire, mais la conclusion est très définitivement structurelle, dans un sens approximatif approprié. Ce n'est pas particulièrement surprenant lorsque C est proche de la valeur extrême de 2, lorsque A est, comme on pourrait s'y attendre, une petite perturbation d'un ensemble extrémal facilement décrit, mais pour des valeurs plus grandes de C , le résultat est beaucoup moins attendu et n'est pas du tout facile à démontrer. Une démonstration belle et complètement différente du théorème de Freiman a été obtenue par Ruzsa [Ru], et une version considérablement nettoyée de l'argument original de Freiman a récemment été donnée par Bilu [Bi].

J'ai présenté le lemme de régularité de Szemerédi et le théorème de Freiman de manière très positive, et ils ont tous deux des applications importantes. Cependant, ils ne résolvent pas tout. Pour le reste de cette section, je discuterai de problèmes qui semblent difficiles à résoudre sans une sorte de théorème de classification approximative. Dans chaque cas, le théorème de classification devrait être meilleur que tout ce que nous avons actuellement. Je commence par trois problèmes en théorie élémentaire de Ramsey.

1. Bornes pour les nombres de Ramsey diagonaux. Le théorème de Ramsey affirme que pour tout entier positif k , il existe un entier positif N tel que si les arêtes du graphe complet sur N sommets sont colorées en rouge et bleu, alors il doit y avoir k sommets dont toutes les arêtes les reliant sont de la même couleur. (Le graphe couvert par ces sommets est appelé monochromatique.) Le plus petit N ayant cette propriété est noté $R(k, k)$. On sait que $R(3, 3) = 6$ et $R(4, 4) = 18$, mais $R(k, k)$ n'est connu pour aucun $k \geq 5$. Il est peu probable qu'une formule raisonnable pour $R(k, k)$ existe, donc une question beaucoup plus intéressante est de savoir comment il se comporte asymptotiquement. Des arguments simples montrent que $2^{k/2} \leq R(k, k) \leq 2^{2k}$, mais, malgré les efforts acharnés de nombreux mathématiciens, aucune amélioration significative de l'une ou l'autre borne n'a été obtenue. (Par significative, j'entends une borne inférieure de $2^{\alpha k}$ avec $\alpha > 1/2$ ou une borne supérieure de $2^{\beta k}$ avec $\beta < 2$.)

Il existe une définition équivalente de $R(k, k)$, que j'utiliserai parfois sans autre commentaire. C'est le plus petit N tel que tout graphe à N sommets contienne une clique ou un ensemble indépendant à k sommets (c'est-à-dire k sommets soit complètement reliés, soit pas reliés du

tout).

2. Recouvrement par des graphes sans triangles. Le nombre de Ramsey $R(3, 3, \dots, 3)$ (où 3 apparaît r fois) est le plus petit entier N tel que si le graphe complet sur N sommets est coloré avec r couleurs, alors il doit y avoir un triangle dont toutes les arêtes sont de la même couleur. Des arguments relativement simples montrent que ce nombre se situe entre e^r et r^r . Il est ouvert de savoir s'il existe une borne supérieure de la forme C^r pour une certaine constante absolue C .

Une formulation équivalente de ce problème est la suivante : si le graphe complet sur n sommets doit être écrit comme une réunion de k graphes sans triangles, quelle doit être la taille de k ? On peut le faire si $k = \lceil \log_2 n \rceil$. Est-il nécessaire d'utiliser au moins $c \log n$ graphes ?

Il existe une autre formulation équivalente du problème en termes de capacité de Shannon – voir la contribution d'Alon dans ce volume pour une discussion de ce concept.

3. Nombres de Ramsey sous contraintes supplémentaires. Voici un beau problème d'Erdős et Hajnal. Soit H un graphe fini fixé. Si l'on suppose que G ne contient pas de copie induite de H , alors de combien la borne pour $R(k, k)$ s'améliore-t-elle ? En particulier, existe-t-il une constante C dépendant uniquement de H telle que tout graphe avec au moins k^C sommets ne contenant pas de H induit ait un sous-graphe complet ou vide à k sommets ?

Qu'est-ce qui rend ces trois problèmes difficiles ? La raison la plus importante semble être qu'il n'est pas évident de savoir quels doivent être les graphes extrémaux. Pour le premier problème, la borne inférieure $R(k, k) \geq 2^{k/2}$ a été prouvée par Erdős [E] en 1947. Sa démonstration peut être résumée comme suit : colorer chaque arête en rouge ou bleu aléatoirement et indépendamment avec une probabilité d'un demi, et vérifier que le nombre attendu de sous-graphes complets monochromatiques à k sommets est inférieur à un. Pour le dire encore plus librement : une coloration aléatoire fonctionne.

Une fois que l'on a vu cette démonstration, elle semble si naturelle qu'il est tentant de conjecturer que la meilleure borne possible, ou quelque chose qui s'en approche, est atteinte par un graphe aléatoire, et en effet beaucoup de ceux qui ont réfléchi au problème croient qu'il en est ainsi. Cependant, un résultat troublant de Thomason [T2] suggère que la démonstration d'un tel fait ne sera pas facile. Une approche naturelle pour estimer $R(k, k)$ consiste à considérer le problème plus général suivant. Soient k et n des entiers positifs, et supposons que les arêtes d'un graphe complet sur n sommets sont colorées en rouge et bleu. Quel est le plus petit nombre possible de sous-graphes complets monochromatiques d'ordre k ? Erdős a conjecturé qu'une coloration aléatoire devrait être la meilleure possible, ce qui suggérerait une borne d'environ $2.2^{-\binom{k}{2}} \binom{n}{k}$. Ce résultat est vrai lorsque $k = 3$, mais Thomason a montré qu'il est faux lorsque $k = 4$.

Il est important de souligner en quel sens les graphes aléatoires ne sont pas les meilleurs. Certains graphes, comme les remarquables graphes de Ramanujan, améliorent les constructions aléatoires en étant "plus aléatoires que le hasard". Par exemple, les graphes de Ramanujan sont d -réguliers et ont une deuxième plus grande valeur propre de $2\sqrt{d-1}$, qui est asymptotiquement aussi petite que possible. Bien que les graphes aléatoires aient également tendance à avoir une deuxième

plus grande valeur propre de taille $C\sqrt{d-1}$, la constante C est plus grande que 2. Ce n'est pas la nature de la non-aléatoire de l'exemple de Thomason, et en effet cela ne peut pas l'être, car tout graphe meilleur que le hasard sera quasi-aléatoire au sens défini précédemment, et donc lui et son complément contiendront à peu près le même nombre de copies de K_4 qu'un graphe aléatoire.

L'exemple de Thomason est construit comme suit. Il choisit d'abord un graphe fixé H avec k sommets. (Un exemple qui fonctionne est un certain produit de sept triangles.) Il partitionne ensuite les sommets d'un graphe complet en ensembles de taille égale $V_1, \dots, V_k$ et colore une arête entre $x \in V_i$ et $y \in V_j$ en bleu si ij est une arête de H et en rouge sinon. Les détails se trouvent dans [T2,3]. Remarquez que le lemme de régularité de Szemerédi et les faits sur les graphes quasi-aléatoires impliquent que tout exemple est nécessairement similaire à celui de Thomason en ce qu'il impliquera de partitionner les sommets en quelques ensembles et de changer les probabilités d'arête entre les ensembles de $1/2$ à autre chose, bien qu'il soit assez surprenant que Thomason y parvienne avec des probabilités de zéro et un.

Quant au deuxième problème, une façon évidente de couvrir n arêtes avec $\lceil \log_2 n \rceil$ graphes sans triangles est d'utiliser des graphes bipartis. On peut laisser les sommets de G être toutes les suites $01-(\epsilon_1, \dots, \epsilon_k)$ de longueur k , et pour chaque i définir un graphe biparti H_i en reliant $(\epsilon_1, \dots, \epsilon_k)$ à $(\eta_1, \dots, \eta_k)$ si et seulement si $\epsilon_i \neq \eta_i$. Une fois de plus, nous avons un exemple très naturel qui semble pourrait être le meilleur possible. Cependant, ce n'est pas le cas.

Pour démontrer cela, faisons une observation simple. Supposons que G est un graphe complet sur m sommets qui peut être couvert par r graphes sans triangles $H_1, \dots, H_r$. Maintenant, soit $k > 1$ et considérons le graphe complet dont l'ensemble des sommets est G^k . Pour $1 \leq i \leq r$ et $1 \leq j \leq k$, définissons un graphe H_{ij} en reliant $(x_1, \dots, x_k)$ à $(y_1, \dots, y_k)$ si et seulement si $x_j y_j$ est une arête de H_i . Il est facile de vérifier que les graphes H_{ij} sont sans triangles et couvrent toutes les arêtes de G^k . Le nombre de graphes H_{ij} est rk et G^k a m^k sommets. Si $n = m^k$ alors $rk = \log n / ((\log m)/r)$, donc pour battre la première borne, tout ce que nous avons à faire est de trouver un exemple où $(\log m)/r$ est supérieur à $\log 2$. Un tel exemple existe lorsque $m = 5$. Le graphe complet sur cinq sommets peut être couvert avec deux cycles de cinq, et $(\log 5)/2 > \log 2$. Un meilleur exemple encore est obtenu comme suit. Soit G le graphe complet K_{41} , et identifions les sommets de G avec les entiers modulo 41. Les résidus quatrièmes forment un sous-groupe Q du groupe multiplicatif de $\mathbb{F}_{41}$. Ce sous-groupe est d'indice quatre, et on peut vérifier qu'aucun de ses deux éléments ne diffère de un. Nous colorons maintenant les arêtes xy de G avec quatre couleurs selon la classe de $x - y$. (Puisque -1 est un résidu quatrième, la classe de $x - y$ est égale à la classe de $y - x$.) Il est facile de voir qu'un triangle monochromatique impliquerait l'existence de quatre éléments non nuls a, u, v, w de $\mathbb{F}_{41}$ tels que $au^4 + av^4 = aw^4$. Cela impliquerait à son tour l'existence de p et q non nuls tels que $1 + p^4 = q^4$, dont on a vérifié qu'ils n'existaient pas. Par conséquent, chaque classe de couleur représente un graphe sans triangle, et dans ce cas $(\log m)/r$ est $(\log 41)/4$, qui est plus grand que $(\log 5)/2$.

Même cela n'est pas le meilleur exemple connu, qui (comme je l'ai appris de Noga Alon lors de la conférence) est lié à un certain graphe de Cayley avec plus de trois cents sommets et a été découvert par un ordinateur.

Pour le troisième problème, si H est un graphe modérément compliqué avec sept sommets, par exemple, il n’y a aucun espoir de deviner quel sera le graphe extrémal G . Le résultat est plausible, car la condition que G ne contienne pas de copie induite de H exclut presque toute utilisation du hasard, ce qui semble très important si un graphe ne doit pas contenir de grandes cliques ou ensembles indépendants.

Il est intéressant de réfléchir à la raison pour laquelle pour certains problèmes extrémaux, il est facile d’identifier les meilleurs exemples possibles, alors que pour d’autres, c’est beaucoup plus difficile. Les trois problèmes ci-dessus sont tous difficiles car il existe deux techniques complètement différentes pour produire des exemples, et on peut produire tout un spectre d’exemples supplémentaires en les mélangeant. Par exemple, le graphe G suivant montre que $R(k, k) \geq k^2$. Soient les sommets de G $V_1 \cup \dots \cup V_k$, où les V_i sont des ensembles disjoints de taille k . Relier $x \in V_i$ à $y \in V_j$ si et seulement si $i \neq j$. À une époque, Turán croyait que la borne de k^2 était optimale, et bien que nous sachions maintenant qu’il avait complètement tort, l’exemple de Thomason montre qu’il est difficile d’écarter des exemples qui mélangent ce type d’approche explicite avec une approche probabiliste. Cela explique également la difficulté du troisième problème. Un phénomène similaire se produit avec le deuxième, en ce qu’il existe un spectre de méthodes de construction de graphes sans triangles. À une extrémité se trouve un graphe biparti complet et à l’autre un graphe aléatoire avec des probabilités d’arête appropriées (et quelques suppressions). Entre les deux, on peut construire un graphe H sans triangles sur r sommets en utilisant des méthodes aléatoires, puis le “dilater” en un graphe G sur n sommets en partitionnant les sommets de G en ensembles $V_1, \dots, V_r$ et en reliant $x \in V_i$ à $y \in V_j$ si et seulement si ij est une arête de H .

C’est parce que les graphes extrémaux pour ces problèmes ne sont pas évidents qu’il semble y avoir un besoin d’une sorte de classification. Revenant au problème de l’estimation de $R(k, k)$, bien que le graphe aléatoire puisse bien donner la bonne réponse, il semble nécessaire de prendre en compte des graphes tels que celui construit par Thomason, ne serait-ce que pour les écarter. On pourrait imaginer que le lemme de régularité de Szemerédi était idéalement adapté à cet usage, mais les mauvaises bornes dans le lemme le rendent beaucoup trop grossier. En fait, même la définition des graphes quasi-aléatoires est trop grossière, car il est parfaitement possible qu’un graphe G sur n sommets soit $n/2$ -régulier et ait une deuxième plus grande valeur propre proportionnelle à $\sqrt{n}$, mais contienne un sous-graphe complet avec $\sqrt{n}$ sommets. En d’autres termes, la distinction entre deux graphes quasi-aléatoires telle que définie précédemment peut être très importante dans le contexte du théorème de Ramsey. Cela suggère la nécessité d’une définition plus forte : à ma connaissance, rien de convenable n’a été proposé.

Je me suis concentré jusqu’à présent sur la théorie des graphes, mais des difficultés similaires se posent dans de nombreux autres contextes. Voici deux problèmes en théorie combinatoire des nombres.

4. Théorème de Roth sur les progressions arithmétiques. En 1974, Szemerédi [S2] a prouvé le théorème suivant, qui avait été conjecturé par Erdős et Turán en 1936 [ET].

Théorème. Soit $k \geq 3$ un entier et soit $\delta > 0$. Alors si N est suffisamment grand, tout sous-ensemble A de $\{1, 2, \dots, N\}$ de cardinalité au moins δN contient une progression arithmétique

de longueur k .

La dépendance asymptotique correcte de N en fonction de k et δ est encore un problème ouvert majeur, mais le problème sur lequel je souhaite me concentrer ici est plus spécifique. Il s'agit de déterminer la dépendance de N en fonction de δ lorsque $k = 3$.

Les premiers progrès vers la conjecture d'Erdős-Turán sont dus à Roth [Ro1], qui a montré que la conjecture était vraie lorsque $k = 3$, et que N pouvait être pris comme étant $\geq \exp \exp(C/\delta)$. Cela a ensuite été amélioré à $\exp(1/\delta^C)$ par Szemerédi [S3] et Heath-Brown [He]. Très récemment, il a été montré par Bourgain [Bo3] que C peut être $2 + \epsilon$. Plus précisément, si $\delta \geq c \log \log N (\log N)^{-1/2}$, alors A doit contenir une progression arithmétique de longueur trois.

La meilleure borne connue dans l'autre direction est donnée par un exemple ingénieux découvert par Behrend en 1946 [B]. Il montre que l'inégalité $N \geq \exp(C \log(\delta^{-1})^2)$, ou de manière équivalente $\delta \geq \exp(-c\sqrt{\log N})$, n'est pas suffisante pour garantir une progression arithmétique de longueur trois.

Cette situation est intéressante pour plusieurs raisons. Premièrement, les bornes laissent ouverte la question de savoir si une densité $\delta = (\log N)^{-1}$ est suffisante. C'est une question importante car $(\log N)^{-1}$ est la densité des nombres premiers dans l'intervalle $\{1, 2, \dots, N\}$. Un résultat de van der Corput, basé sur des techniques de Vinogradov, montre que les nombres premiers contiennent une infinité de progressions arithmétiques de longueur trois. Il est cependant possible qu'il existe une démonstration de ce fait qui ne dépende que de la densité des nombres premiers et non de faits plus détaillés sur leur distribution.

Une deuxième raison est que bien que l'exemple de Behrend donne une borne très faible par rapport à ce qui est connu dans l'autre direction, il montre au moins que les méthodes purement aléatoires ne donnent pas un ensemble extrémal. (Ce point a été discuté par le mathématicien et l'ordinateur dans le dialogue de la section précédente.)

5. Bornes pour le théorème de Freiman. Dans l'énoncé du théorème de Freiman donné plus tôt, aucune mention n'a été faite de la dépendance de K et d par rapport à C . Cela est encore loin d'être complètement résolu. Voir l'article de Bilu [Bi] pour plus de détails sur ce qui est connu. L'amélioration des bornes connues aurait des applications importantes. Pour donner une idée de ce qui n'est pas connu, voici deux questions spécifiques. Si A est un ensemble de n entiers et $|A + A| \leq C|A|$, peut-on dire quelque chose sur la structure de A si $C = n^\alpha$ pour un certain $\alpha > 0$, ou même si $C = \log n$? Les démonstrations existantes ne donnent aucune information dans cette situation. Pour la deuxième question, soit A satisfaisant les mêmes conditions. A a-t-il un sous-ensemble B de taille cn tel que B soit contenu dans une progression arithmétique généralisée P de cardinalité K et de dimension d , où c et K ont une dépendance raisonnable en fonction de C et $d = A \log C$ pour une certaine constante absolue A ? Pour les applications, il serait encore très intéressant si d pouvait être borné supérieurement par $(\log C)^A$.

Ces deux problèmes soulèvent des difficultés similaires dans l'esprit à celles des problèmes de théorie des graphes. L'exemple de Behrend montre qu'il n'y a pas d'ensemble extrémal évident pour le problème 4, et une fois de plus il y a un conflit entre les techniques aléatoires et explicites qui explique en partie cette situation. De même, les ensembles qui se présentent dans toute considération du théorème de Freiman sont très variés. Donc pour les deux problèmes, une classification serait presque certainement utile.

Plus spécifiquement, pour les deux problèmes, on se retrouve à vouloir comprendre des sous-ensembles $A \subset \mathbb{Z}_N$ de taille αN , où α est une constante, ou peut-être un nombre qui tend vers zéro très lentement lorsque N tend vers l'infini. (Ici $\mathbb{Z}_N$ est le groupe des entiers modulo N .) Une bonne façon de faire est de regarder la transformée de Fourier de la fonction caractéristique de A , définie par $\hat{A}(r) = \sum_{s \in A} e(rs/N)$, où $e(x)$ est une notation pour $\exp(2\pi i x)$. Le coefficient $\hat{A}(0)$ est simplement la cardinalité de A . Si $\hat{A}(r)$ est significativement plus petit que $\alpha^2 N$ pour tout r non nul, alors A se comporte à bien des égards comme un sous-ensemble aléatoire de $\mathbb{Z}_N$, où chaque élément est choisi avec probabilité α . Nous avons donc une bonne définition du caractère quasi-aléatoire pour les sous-ensembles de $\mathbb{Z}_N$. Voir [CG] pour plus de détails à ce sujet.

Remarquez que l'information véhiculée par les coefficients de Fourier sur les sous-ensembles de $\mathbb{Z}_N$ est très similaire à l'information véhiculée par les valeurs propres sur les graphes. Ce n'est pas une coïncidence : si A est un sous-ensemble de $\mathbb{Z}_N$, on peut définir un graphe orienté G_A avec pour ensemble de sommets $\mathbb{Z}_N$ en ayant une arête de x à y si et seulement si $y - x \in A$, et il est facile de vérifier que les valeurs propres de G_A sont les coefficients de Fourier de A . Les analogies vont plus loin. Tout comme on peut essayer de classer les graphes, on peut aussi essayer de classer les sous-ensembles de $\mathbb{Z}_N$. Pour cela, on a besoin d'une notion de "différant pas de manière intéressante" : il s'avère que la différence entre deux ensembles A et B est pour de nombreux objectifs sans importance si $\hat{A}(r) - \hat{B}(r)$ est petit pour tout r .

L'analyse de Fourier a été une partie essentielle des meilleurs arguments connus pour analyser les sous-ensembles A de $\mathbb{Z}_N$ en rapport avec les problèmes 4 et 5. La raison en est que, parce que perturber les coefficients de Fourier d'un ensemble A ne fait pas une différence importante, nous n'avons à nous soucier que des grands coefficients. Il y en a très peu, donc nous avons la façon efficace suivante d'encoder l'information essentielle sur A : choisir une constante β , soit $K = \{r : |\hat{A}(r)| \geq \beta N\}$ et pour $r \in K$ soit $f_A(r) = \hat{A}(r)$. Il pourrait sembler que le fait de faire correspondre des ensembles A à leurs fonctions f_A est un schéma de classification tout fait, mais je considérerais "semi-classification" comme une meilleure description, car il n'est pas facile d'utiliser toutes les informations contenues dans f_A . Les arguments connus ont tendance à jeter toute information sur l'ensemble K sauf sa cardinalité, mais la structure de K est également très importante. En particulier, le comportement de A dépend beaucoup de l'existence ou non de relations arithmétiques simples entre les éléments de K , comme deux d'entre eux s'ajoutant à un troisième. Malheureusement, il n'est pas facile de voir comment incorporer cette information dans les arguments, et cela semble être un obstacle à de nouveaux progrès sur les deux problèmes. Une fois de plus, ce qui semble manquer, c'est un meilleur schéma de classification. (Des remarques similaires s'appliquent aux valeurs propres des graphes. Bien qu'utiles, pour de nombreux problèmes, il est important de savoir comment les vecteurs propres correspondants s'agencent, et cette information est difficile à utiliser.)

Les deux problèmes suivants sont conçus pour montrer que des difficultés similaires à celles discutées jusqu'à présent peuvent se poser dans de nombreux contextes.

6. Un problème sur les sous-ensembles de la sphère n -dimensionnelle. Notons par S_{n-1} la sphère unité euclidienne de $\mathbb{R}_n$. Soit $m < n$, soit $\epsilon > 0$ et soit A un sous-ensemble de S_{n-1} qui intersecte tout sous-espace de dimension m de $\mathbb{R}_n$. Écrivons A_ϵ pour l'ensemble $\{x : d(x, A) \leq \epsilon\}$. Il y a des questions intéressantes que l'on peut se poser sur l'ensemble A_ϵ . Par exemple, quelle peut être la petitesse de sa mesure ? Pour des valeurs données de m et ϵ , quel est le plus grand r tel que A_ϵ est forcé de contenir la sphère unité d'un sous-espace de dimension r ?

Ces questions sont difficiles, car il y a deux façons naturelles de construire un ensemble A ayant la propriété donnée. La première est de prendre la sphère unité d'un sous-espace de dimension $(n - m)$, et la seconde est de prendre une réunion de calottes sphériques placées aléatoirement d'un rayon donné. Nous sommes à nouveau dans une situation où de nombreuses autres constructions peuvent être imaginées qui mélangent ces deux idées (par exemple, au lieu de calottes, qui sont des dilatations d'ensembles de dimension zéro, on pourrait prendre des dilatations d'ensembles de dimension k pour un petit k). Il semble probable que si l'on exige que A soit centralement symétrique, alors la mesure de A_ϵ est minimisée lorsque A est la sphère d'un sous-espace de dimension $(n - m)$. D'un autre côté, il ne semble pas nécessaire que A_ϵ contienne la sphère d'un sous-espace aussi grand. Peut-être peut-on faire mieux en permettant à A_ϵ d'être plus grand, mais moins régulier.

Pour certaines valeurs de m et ϵ , une classification approximative des sous-ensembles de S_{n-1} pourrait aider à déterminer une bonne borne pour r . On voudrait exploiter le fait que les ensembles A pour lesquels A_ϵ est petit ont une certaine structure, et en déduire que r est plus grand que ce qu'impliquent les seules considérations de théorie de la mesure.

Une méthode possible pour fournir des bornes supérieures pour r est de laisser B être la sphère unité d'un sous-espace de dimension $(n - m)$, de laisser ϕ être une application antipodale continue de S_{n-1} dans elle-même et de poser $A = \phi(B)$. Cela fait de A une sorte de sphère topologique de dimension $n - m$. Malheureusement, il est difficile de penser à un choix pour ϕ qui fournira une bonne borne et qui soit facile à analyser.

7. Un problème sur les couches du cube discret. Soient $r < s < t < n$ des entiers positifs, écrivons $[n]^{(s)}$ pour l'ensemble de tous les sous-ensembles de $[n] = \{1, 2, \dots, n\}$ de taille s et soit $\mathcal{A}$ un sous-ensemble de $[n]^{(s)}$. C'est-à-dire que $\mathcal{A}$ est une collection d'ensembles de taille s . Écrivons $\partial_r \mathcal{A}$ pour la collection de tous les ensembles $B \in [n]^{(r)}$ tels que $B \subset A$ pour un certain $A \in \mathcal{A}$, et $\partial^t \mathcal{A}$ pour la collection de tous les ensembles $C \in [n]^{(t)}$ tels que $A \subset C$ pour un certain $A \in \mathcal{A}$. Supposons que tout ensemble de $[n]^{(t)}$ contienne un ensemble de $\mathcal{A}$, donc $\partial^t \mathcal{A} = [n]^{(t)}$. Quelle peut être la petitesse de $\partial_r \mathcal{A}$?

La nature de ce problème dépend beaucoup des valeurs de r, s, t et n . Remarquez qu'une fois de plus, il y a deux façons complètement différentes de construire un ensemble $\mathcal{A}$ qui satisfait les conditions du problème. Si l'on souhaite minimiser la cardinalité de $\mathcal{A}$, alors il est préférable de le choisir aléatoirement. Cependant, parmi tous les ensembles $\mathcal{A}$ d'une

cardinalité donnée, les ensembles aléatoires sont ceux pour lesquels $\partial_r \mathcal{A}$ est le plus grand. Une construction alternative exploite le principe des tiroirs de manière simple. Pour plus de commodité, supposons que s divise t et que $k = t/s$ divise n . Partitionner l'ensemble $[n]$ en k sous-ensembles de taille égale $V_1 \cup \dots \cup V_k$ et soit $\mathcal{A}$ la collection de tous les ensembles A de taille r tels que $A \subset V_i$ pour un certain i . Tout ensemble de taille t intersecte au moins l'un des V_i en au moins $n/k = s$ points, donc $\mathcal{A}$ satisfait la condition requise.

Supposons un instant que s est beaucoup plus petit que t . Si r n'est que très légèrement plus petit que s , alors il semble qu'un choix aléatoire de $\mathcal{A}$ soit le meilleur, puisque chaque ensemble de $\mathcal{A}$ ne contiendra que quelques ensembles dans $[n]^{(r)}$. Si en revanche r est beaucoup plus petit que s , alors la construction explicite s'avère meilleure. Entre ces deux extrêmes, il y a une plage de valeurs de r où il n'est pas évident de savoir quelle construction est la meilleure, ou si quelque chose "d'intermédiaire" est plus approprié. En fait, je ne sais même pas comment prouver ce qui semble être le cas dans les situations extrêmes – que les exemples que j'ai esquissés sont les meilleurs possibles.

Ce problème est peut-être légèrement artificiel, mais c'est un exemple simple d'un type de difficulté que j'ai rencontré plusieurs fois en réfléchissant à d'autres questions, et je doute d'être le seul. Quelle que soit l'opinion que l'on a sur l'intérêt ou non du problème, il est difficile d'imaginer qu'une solution soit ennuyeuse. Une fois de plus, il semble qu'une classification approximative serait très utile, cette fois des sous-ensembles de $[n]^{(s)}$. (Il existe des versions du lemme de régularité de Szemerédi lorsque s est fixé, mais elles ne seraient pas utiles si l'on prenait $s = \sqrt{n}$, par exemple.)

8. Trouver des classifications approximatives. Découlant des problèmes précédents, il y a le problème général de trouver de nouvelles classifications approximatives utiles. Parmi les objets que l'on pourrait vouloir classer, citons les suivants.

(i) Graphes. Le problème ici est de trouver un remplacement utile du lemme de régularité de Szemerédi. Parce que la mauvaise borne dans le lemme est (dans un sens très faible!) optimale, remplacer le lemme signifierait remplacer la description qu'il donne d'un graphe général par quelque chose de plus compliqué, en échange de bien meilleures bornes. Un autre projet intéressant est de trouver des démonstrations alternatives pour des résultats qui utilisent le lemme (tout comme il est intéressant de remplacer une démonstration qui repose sur la classification des groupes simples finis par une qui n'en repose pas). Voir [KoS], [GrRR] et [G2, Prop. 12] pour quelques exemples de la façon de l'éviter. Pour plus d'informations générales sur le lemme de régularité et ses variantes, voir [KoS].

(ii) Graphes sans triangles. Deux façons de les construire sont d'utiliser des méthodes aléatoires, ou de prendre un graphe sans triangles donné H avec des sommets $x_1, \dots, x_k$ et de définir un graphe G avec un ensemble de sommets partitionné en ensembles $V_1, \dots, V_k$ en reliant $x \in V_i$ à $y \in V_j$ si et seulement si $x_i x_j$ est une arête de H . La combinaison de ces méthodes, et d'autres opérations simples comme les sous-graphes et les unions disjointes, donne une grande classe de graphes sans triangles. Des techniques comme celles-ci s'approchent-elles de la génération de tous les graphes sans triangles ?

(iii) Graphes clairsemés. Il est depuis longtemps reconnu qu'un inconvénient majeur du lemme de régularité de Szemerédi est que les seuls graphes pour lesquels il donne des informations sont ceux qui ont un nombre d'arêtes proportionnel à $\binom{n}{2}$, alors que de nombreux problèmes d'un grand intérêt concernent des graphes avec, disons, n^α arêtes pour un certain $\alpha \in (1, 2)$. Il existe des variantes du lemme de régularité pour de tels graphes clairsemés, mais elles sont loin de ce que l'on pourrait souhaiter appeler une classification.

(iv) Sous-ensembles de groupes abéliens, en particulier $\mathbb{Z}_N$ et $\mathbb{Z}$, avec deux ensembles considérés comme similaires si leurs transformées de Fourier sont proches.

(v) Sous-ensembles de $\mathbb{Z}_N$, avec deux ensembles considérés comme similaires si leurs fonctions caractéristiques sont proches dans la norme suivante :

$$\|f\| = N^{-4} \sum_{x,a,b,c \in \mathbb{Z}_N} \prod_{\epsilon \in \{0,1\}^3} f_\epsilon(x - \epsilon_1 a - \epsilon_2 b - \epsilon_3 c).$$

Ici f_ϵ désigne f si $\epsilon_1 + \epsilon_2 + \epsilon_3$ est pair et le conjugué complexe de f sinon. Il s'avère que, bien que les grands coefficients de Fourier d'un ensemble A contiennent toutes les informations sur les progressions arithmétiques de longueur trois, ils ne le font pas pour les progressions plus longues. Si A et B sont proches dans la norme ci-dessus, alors ils contiendront approximativement le même nombre de progressions arithmétiques de longueur quatre – d'où son importance. Pour plus de détails, voir [G2]. Une classification par rapport à cette norme serait nécessairement beaucoup plus sophistiquée qu'une classification par rapport à la proximité des coefficients de Fourier.

(vi) Sous-ensembles d'espaces métriques tels que $[0, 1]^2$ et S_{n-1} avec deux ensembles considérés comme similaires si leurs propriétés métriques sont à peu près les mêmes. Il existe de nombreuses façons de rendre cette question précise, dont certaines ont été largement étudiées. Une possibilité est de dire que les sous-ensembles A et B d'un espace métrique X sont similaires s'il existe un petit $\epsilon > 0$ et des applications 1-Lipschitz $\phi : A \rightarrow X$ et $\psi : B \rightarrow X$ telles que tout point de B soit à ϵ d'un point de ϕA , tout point de A soit à ϵ d'un point de ψB , et que $d(x, \psi \phi x)$ et $d(y, \psi \psi y)$ soient au plus ϵ pour tout $x \in A$ et $y \in B$.

(vii) Sous-ensembles de structures combinatoires telles que le cube discret $\{0, 1\}^n$ et $[n]^{(r)}$.

3. Une sélection supplémentaire de problèmes

Chaque problème de cette section devait satisfaire à deux critères. Premièrement, j'aurais moi-même dû y consacrer au moins une heure de réflexion, et deuxièmement, j'aurais dû avoir au moins l'intention d'y consacrer plus de temps à l'avenir. Cela a suffi à éliminer la plupart des problèmes intéressants en mathématiques, me laissant une liste gérable.

9. L'hypothèse de Riemann. Elle est si connue que je ne prendrai même pas la peine de l'énoncer ici. Je l'inclus parce que je ne souhaite pas faire une sorte de déclaration en ne le faisant pas. Cependant, je n'ai pas grand-chose à dire à son sujet. Je renvoie plutôt le lecteur à la vaste littérature sur le sujet, y compris certaines parties des articles de Connes, Iwaniec et Sarnak dans ce volume.
10. Le problème de Kakeya. Il est étroitement lié à l'hypothèse de Riemann. Il a de nombreuses formulations, dont voici trois.

(i) Soit A un sous-ensemble de $\mathbb{R}^d$ tel que pour toute droite L dans $\mathbb{R}^d$, A contienne une droite L' parallèle à L . Un tel ensemble est souvent appelé un ensemble de Kakeya, ou un ensemble de Besicovitch. Besicovitch a montré que la mesure de A peut être nulle, mais tous les exemples connus ont une dimension de Hausdorff d . Bourgain [Bo4] a montré que la dimension de Hausdorff de A doit être au moins $\frac{13d}{25} + \frac{12}{25}$, ce qui est une amélioration significative d'une borne relativement facile de $\frac{d}{2} + 1$. (Très récemment, cette borne a été encore améliorée par Katz et Tao à $\frac{6d}{11} + \frac{5}{11}$ [KaT].) On sait que la dimension doit être d lorsque $d = 2$, mais pour $d > 2$, le problème est ouvert. La première borne inférieure non triviale lorsque $d = 3$ a été découverte par Bourgain [Bo2]. Son argument a montré que la dimension d'un ensemble de Kakeya tridimensionnel était d'au moins $7/3$. Cela a été amélioré par Wolff à $5/2$ [W]. Cette borne a récemment été légèrement améliorée par Katz, Laba et Tao [KaLT], un résultat intéressant car $5/2$ était une limite naturelle dans la démonstration de Wolff et la dépasser était très difficile.

(ii) Soit S le carré $\{(x, y, z) \in \mathbb{R}^3 : z = 0, 0 \leq x, y \leq 1\}$ et soit P le point $(1/2, 1/2, 1)$. Divisons maintenant S en n^2 sous-carrés S_{ij} de côté n^{-1} de la manière évidente. Pour chacun de ces sous-carrés, formons un cône C_{ij} en prenant l'enveloppe convexe de S_{ij} et P . Traduisons chaque cône C_{ij} dans une direction xy (ou en d'autres termes, glissons sa base ailleurs dans le plan xy) et appelons le cône traduit C'_{ij} . Alors, quelle peut être la petitesse du volume de $\bigcup_{i,j=1}^n C'_{ij}$? Si le plus petit volume possible se comporte comme $n^{-\alpha}$ pour un certain $\alpha > 0$, alors il existe un ensemble de Kakeya dans $\mathbb{R}^3$ de dimension de Hausdorff inférieure à trois. La réciproque est également vraie. L'analogie d -dimensionnelle évidente de cette équivalence est également valable.

(iii) Soient m et r des entiers, et pour tout d entre 1 et r , soit $P_d \subset \mathbb{Z}$ une progression arithmétique de longueur m et de différence commune d . Quelle peut être la petitesse de l'union $\bigcup_{d=1}^r P_d$? En particulier, si $r = m^C$ pour un certain entier fixé $C \geq 2$, peut-on améliorer la borne triviale de mr par un facteur de m^α pour un certain $\alpha > 0$? Une fois de plus, une construction qui y parviendrait pourrait être traduite en un exemple d'ensemble de Kakeya de dimension inférieure à la dimension totale.

La troisième version du problème donne une indication du fait qu'il a un contenu en théorie des nombres. Un lien étroit avec une conjecture de Montgomery (voir [Mo, Chapitre 7]), et donc avec les zéros de la fonction zêta, a été découvert par Bourgain dans son article au titre modeste [Bo1]. Depuis un célèbre article de Fefferman [F], on a réalisé que les ensembles de Kakeya sont également d'une grande importance en analyse harmonique. Une solution au

problème de Kakeya serait donc une avancée majeure pour plusieurs raisons différentes.

11. Le problème $P = NP$? J'ai déjà mentionné ce problème au §2. Comme tous les mathématiciens ne sont pas à l'aise pour l'énoncer, permettez-moi de donner la définition que je trouve la plus claire de la classe des problèmes NP. Tout d'abord, il est commode de penser non pas à des problèmes mais à des fonctions booléennes, c'est-à-dire des fonctions $\phi : \{0, 1\}^n \rightarrow \{0, 1\}$. Un problème tel que le problème du voyageur de commerce peut facilement être considéré comme une telle fonction : coder votre graphe G avec n sommets comme une suite $01\text{-}\eta_G$ de longueur $\binom{n}{2}$ de manière évidente et définir $\phi(\eta_G)$ comme étant 1 si G contient un cycle hamiltonien et 0 sinon. Une fonction booléenne ϕ appartient à P s'il existe un algorithme en temps polynomial pour calculer ϕ . (Strictement parlant, on devrait parler d'une suite ϕ_n de fonctions, où ϕ_n a $f(n)$ variables, et le temps pris pour calculer ϕ_n est borné supérieurement par une fonction polynomiale de $f(n)$.) Elle appartient à NP si elle est une projection d'une fonction dans P, au sens suivant : il existe un entier m borné supérieurement par un polynôme en n et une fonction $\psi : \{0, 1\}^m \times \{0, 1\}^n$ qui appartient à P telle que $\phi(\eta) = 1$ si et seulement s'il existe $\xi \in \{0, 1\}^m$ tel que $\psi(\xi, \eta) = 1$.

Pour le problème du voyageur de commerce, ξ représentera un ordre des sommets, η un graphe, et $\psi(\xi, \eta) = 1$ si et seulement si l'ordre des sommets donné par ξ définit un cycle hamiltonien dans le graphe donné par η . Ceci est très facile à vérifier, donc ψ appartient à P, et la fonction ϕ définie précédemment, qui vaut 1 si et seulement si $\psi(\xi, \eta) = 1$ pour un certain ξ , est dans NP.

Une définition supplémentaire, que je ne donnerai pas précisément, est celle d'une fonction NP-complète. Une découverte importante de Cook [Co] et Karp [K], qui explique l'intérêt pour le problème $P=NP$, est que de nombreux problèmes NP (ou fonctions booléennes correspondantes) sont universels, en ce sens qu'un algorithme en temps polynomial pour résoudre l'un d'eux peut être converti en un algorithme en temps polynomial pour résoudre un problème arbitraire dans NP. Un tel problème est dit NP-complet. En fait, des centaines de tels problèmes sont maintenant connus, et beaucoup d'entre eux sont d'une importance pratique directe.

Le problème $P=NP$ peut sembler rebutant car l'ensemble de toutes les fonctions calculables par machine de Turing en temps polynomial n'est pas très mathématique. J'ai parfois vu écrit que le problème est difficile car il est difficile d'exclure tous les algorithmes en temps polynomial possibles. Cependant, il existe une reformulation très propre et complètement mathématique du problème en termes de complexité de circuit dite qui montre que ce n'est pas vraiment la difficulté. Je donnerai la définition sous une forme équivalente.

Soit $\phi : \{0, 1\}^n \rightarrow \{0, 1\}$ une fonction booléenne et soit $A = \{\eta : \phi(\eta) = 1\}$. Pour chaque $1 \leq i \leq n$, soit $E_i = \{\eta : \eta_i = 1\}$, et appelons les ensembles E_i et leurs compléments ensembles élémentaires. La complexité de circuit de ϕ , ou de A , est la longueur N de la plus courte suite $A_1, A_2, \dots, A_N$ telle que $A_N = A$ et chaque A_i est soit un ensemble élémentaire, le complément d'un A_j antérieur, soit l'union ou l'intersection de deux A_j antérieurs.

Un problème presque, mais pas tout à fait, équivalent à celui de savoir si $P=NP$ est le suivant.

Soit ϕ une fonction booléenne dans NP. La complexité de circuit de ϕ doit-elle être bornée supérieurement par un polynôme en le nombre de variables ? Un algorithme en temps polynomial pour calculer ϕ peut être converti en une suite $A_1, \dots, A_N$ de longueur polynomiale, donc une borne superpolynomiale pour la complexité de circuit impliquerait que P et NP ne sont pas égaux. Il est théoriquement possible que la réciproque de ceci soit fausse, mais il est difficile d'envisager une construction d'une suite $A_1, \dots, A_N$ avec si peu de régularité qu'elle, ou une adaptation de celle-ci, ne puisse pas être générée efficacement par un algorithme. Je suis donc convaincu que la complexité de circuit est la bonne approche du problème. (C'est le point de vue standard des informaticiens théoriciens, mais pas de tous les mathématiciens.)

Malheureusement, cela ne rend pas le problème facile, et la meilleure borne inférieure connue pour la complexité de circuit d'une fonction NP est encore seulement linéaire. Un article récent, beau et important de Razborov et Rudich [RR] va loin pour expliquer pourquoi le problème est difficile, et est une lecture obligatoire pour quiconque voudrait prouver que $P \neq NP$. Une approche naturelle pour prouver des bornes inférieures pour la complexité de circuit est de définir une mesure de simplicité, de prouver que les ensembles A avec une petite complexité de circuit doivent être simples, et de prouver que l'ensemble A correspondant à une fonction NP n'est pas simple. Razborov et Rudich définissent une notion de "preuve naturelle", qui signifie essentiellement une preuve résultant d'une mesure "naturelle" de simplicité. Ils montrent que toutes les bornes connues en complexité de circuit ont été obtenues en utilisant des preuves naturelles dans leur sens, et ils montrent ensuite que toute preuve naturelle que $P \neq NP$ pourrait être convertie en un algorithme en temps polynomial pour factoriser de grands entiers. C'est une preuve convaincante qu'une démonstration de $P \neq NP$ devrait être très étrange en effet. Pour plus de détails, voir les contributions de Razborov et Wigderson dans ce volume, ou l'article de Razborov et Rudich lui-même, qui est bien écrit et pas difficile à lire.

Le problème suivant (mentionné par Wigderson lors de la conférence) a une saveur similaire au problème $P = NP$, mais n'a rien à voir avec les ordinateurs et semble à première vue beaucoup plus simple. Soit M une matrice inversible $n \times n$ sur $\mathbb{F}_2$. L'élimination de Gauss nous dit que M peut être convertie en la matrice identité par une séquence d'au plus n^2 opérations élémentaires sur les lignes. De plus, un simple argument de dénombrement (il y a moins de $2n^2$ opérations élémentaires sur les lignes et presque 2^{n^2} matrices non singulières) montre que presque toutes les matrices ont besoin d'au moins $cn^2/\log n$ opérations. Le problème est de donner un exemple explicite d'une matrice qui a besoin d'un nombre superlinéaire. ("Explicite" ne signifie pas "défini sans ambiguïté", ou on pourrait juste énumérer toutes les matrices et prendre la première qui fonctionne. Une interprétation possible est qu'il devrait être possible de générer une matrice explicite, ou plutôt une suite de matrices, avec un algorithme en temps polynomial.) Ce problème reste obstinément ouvert, et une solution pourrait bien éclairer le problème $P = NP$ lui-même.

12. Une version effective du théorème de Roth. Par théorème de Roth, j'entends maintenant non pas son théorème sur les progressions arithmétiques mais son théorème beaucoup plus célèbre sur l'approximation rationnelle des nombres algébriques [Ro2]. Celui-ci stipule que pour tout nombre algébrique α et tout $\epsilon > 0$, il existe $c > 0$ tel que pour toute paire p, q d'entiers

positifs, on ait l'inégalité $|\alpha - p/q| \geq c/q^{2+\epsilon}$. C'est un renforcement très considérable du théorème de Liouville. Le problème est que la démonstration ne donne aucune information sur la dépendance de c par rapport à α et ϵ . Il en va de même pour les résultats antérieurs, plus faibles, de Thue et Siegel.

Les résultats connus dans cette direction sont très modestes. Par exemple, Baker [B] a prouvé une version effective du théorème de Roth dans le cas $\alpha = 2^{1/3}$, obtenant l'inégalité $|2^{1/3} - p/q| \geq 10^{-6}/q^{2.995}$, qui doit être comparée à la borne inférieure du théorème de Liouville de c/q^3 , où c peut être pris égal à $1/16$.

Un problème étroitement lié, mentionné par Furstenberg dans son exposé à cette conférence, concerne le développement en fraction continue des nombres algébriques. Si α est une irrationnelle quadratique, alors tout est connu, et les quotients partiels sont éventuellement périodiques. Le théorème de Roth implique que les quotients partiels d'un nombre algébrique ne peuvent pas croître trop rapidement, mais, sauf dans le cas quadratique, il n'y a pas d'exemple pour lequel on sache s'ils sont bornés.

13. La conjecture de Littlewood. Ce problème a été discuté par Margulis lors de la conférence. Littlewood a conjecturé que si α et β sont deux nombres réels quelconques et si $\|x\|$ désigne la distance de x à l'entier le plus proche, alors $\liminf_{n \in \mathbb{N}} n\|n\alpha\|\|n\beta\| = 0$.

Pour comprendre le problème, considérons une irrationnelle quadratique α . Le théorème de Liouville est équivalent à l'énoncé selon lequel $\|n\alpha\|$ est toujours borné inférieurement par cn^{-1} . D'un autre côté, généralement $\|n\alpha\|$ est beaucoup plus grand que cela. On peut donc se demander s'il est possible de trouver deux nombres α et β , dont aucun ne peut être bien approché par des rationnels, avec la propriété supplémentaire que $\|n\beta\|$ est grand lorsque $\|n\alpha\|$ est petit et vice-versa. La conjecture de Littlewood est une formulation précise de cette question.

L'existence d'une telle paire de nombres α, β aurait la conséquence plus combinatoire que pour tout n on pourrait trouver un ensemble X de n points dans $[0, 1]^3$ et une constante absolue $c > 0$ tels que l'inégalité $\prod_{i=1}^3 |x_i - y_i| \geq cn^{-1}$ soit vérifiée pour deux points distincts quelconques $x, y \in X$. En effet, en écrivant $\langle t \rangle$ pour la partie fractionnaire de t , un ensemble approprié serait $\{(mn^{-1}, \langle \alpha m \rangle, \langle \beta m \rangle) : 1 \leq m \leq n\}$. À ma connaissance, l'existence d'un tel ensemble X est une question ouverte, il est donc possible que la conjecture de Littlewood, si elle est vraie, ait une démonstration combinatoire. Un argument simple montre que l'on peut trouver $n/\log n$ points avec la propriété donnée. Toute amélioration serait intéressante.

Si la conjecture de Littlewood est fausse, alors une bonne candidate pour un contre-exemple serait la paire $(\sqrt{2}, \sqrt{3})$. Il semble très difficile d'analyser des exemples concrets comme celui-ci, et c'est peut-être la principale difficulté du problème.

14. Marches aléatoires auto-évitantes. Comme beaucoup de mathématiciens, je suis fasciné par la longue liste de résultats en physique statistique qui ont été obtenus par les physiciens en utilisant des arguments non rigoureux. L'une des sources de problèmes les plus attrayantes

est le domaine des marches aléatoires auto-évitantes. Une marche auto-évitante de longueur n est un chemin $x_0, x_1, \dots, x_n$ de sommets dans le réseau $\mathbb{Z}^d$, où $x_0 = 0$, x_{i-1} et x_i sont adjacents pour $1 \leq i \leq n$ et il n'y a pas de répétition. Si l'on place la distribution uniforme sur l'ensemble de toutes les marches auto-évitantes de longueur n dans $\mathbb{Z}^2$ ou $\mathbb{Z}^3$, alors on obtient un modèle probabiliste important sur lequel les physiciens savent beaucoup de choses et les mathématiciens presque rien. Je ne mentionnerai que quelques problèmes.

Presque toutes les questions les plus intéressantes sont des cas particuliers de la question vague suivante : à quoi ressemble une marche auto-évitante typique de longueur n ? La question la plus évidente est peut-être la distance attendue de l'origine à x_n . En deux dimensions, les arguments des physiciens montrent que la réponse devrait être $n^{3/4}$, mais aucune borne supérieure de la forme $n^{1-\epsilon}$ n'a été prouvée rigoureusement. La situation pour les bornes inférieures est encore pire. On pourrait s'attendre à ce qu'un argument trivial donne une borne inférieure de $cn^{1/2}$, puisqu'une marche auto-évitante parcourt sûrement une plus grande distance qu'une marche aléatoire pure. Cependant, on ne sait pas comment prouver rigoureusement qu'une marche auto-évitante typique n'a pas tendance à se replier sur elle-même. Bien sûr, la plupart des points d'une marche auto-évitante doivent être à une distance $cn^{1/2}$ de l'origine par un simple argument de dénombrement, donc cette dernière question concerne spécifiquement le point d'arrivée.

Une autre question concerne le nombre $f(n)$ de marches auto-évitantes de longueur n . Il est facile de voir que f est sous-multiplicative, donc un argument standard montre que $f(n) = C^n g(n)$ pour une certaine constante C , connue sous le nom de constante connective, et une certaine fonction g qui croît de manière sous-exponentielle. La valeur exacte de C n'est pas connue, mais ce n'est pas un problème particulièrement intéressant pour deux raisons. Premièrement, C peut en principe être déterminée avec une précision arbitraire par un dénombrement exact de toutes les marches d'une certaine longueur appropriée n_0 (bien que cela devienne rapidement impossible en pratique). Plus important encore, C ne possède pas une propriété connue sous le nom d'universalité.

On peut modifier le modèle des marches auto-évitantes de nombreuses façons, comme changer le réseau dans lequel il vit. Les physiciens savent que certains paramètres ne sont pas affectés par de tels changements. Par exemple, la distance d'arrivée attendue est toujours proportionnelle à $n^{3/4}$. Ces paramètres inchangés sont appelés universels et sont considérés comme ayant une signification physique. La constante connective n'est pas universelle en ce sens, car elle dépend certainement du réseau sous-jacent.

D'un autre côté, la fonction g semble être universelle. Une fois de plus, les physiciens en savent beaucoup plus que les mathématiciens : ils nous disent que $g(n)$ croît comme $n^{11/32}$. On pourrait d'abord se demander comment on pourrait espérer étudier la fonction g alors que la constante connective C n'est pas connue. Cependant, il est facile de montrer que g a l'interprétation suivante : $g(n)$ croît comme n^α si et seulement si la probabilité que deux marches aléatoires auto-évitantes de longueur n ne s'intersectent pas (sauf à l'origine) se comporte comme $n^{-\alpha}$. Cela montre que g a une certaine signification indépendante. D'un point de vue purement mathématique, c'est un problème ouvert très intéressant de prouver que g est

borné supérieurement par un polynôme. La meilleure borne connue en deux dimensions est due à Hammersley et Welsh [HaW], qui ont montré par un argument élémentaire ingénieux que $g(n) \leq \exp(cn^{1/2})$.

La raison pour laquelle l'hypothèse d'universalité est plausible est que l'on s'attend à ce que le comportement macroscopique d'une très longue marche auto-évitante ne dépende pas trop de manière critique du comportement précis à de très petites échelles. Liée à cette croyance est l'opinion qu'il devrait y avoir une sorte de limite mise à l'échelle de manière appropriée des marches aléatoires auto-évitanes, jouant le rôle que le mouvement brownien et la mesure de Wiener jouent pour les marches aléatoires ordinaires. C'est un autre problème ouvert important.

Ces problèmes sont tous à leur plus intéressant en deux et trois dimensions. Puisque la dimension de Hausdorff d'un chemin brownien est deux, on ne s'attend pas à ce que la contrainte d'auto-évitement soit très significative en quatre dimensions ou plus. Cette intuition a en fait été justifiée rigoureusement lorsque la dimension est au moins cinq, par Hara et Slade [HarS1,2]. Bien que leurs résultats ne soient pas très surprenants, ils constituent un accomplissement majeur et les démonstrations ne sont pas du tout faciles. Pour beaucoup plus de détails sur tous ces problèmes et résultats, voir l'excellent livre de Madras et Slade [MS].

15. Dénouer le nœud trivial. Il y a eu beaucoup d'enthousiasme autour de la théorie des nœuds au cours des vingt dernières années, y compris la découverte de nouveaux invariants importants. Cependant, le problème fondamental suivant est ouvert : existe-t-il un algorithme en temps polynomial pour reconnaître un diagramme du nœud trivial et pour produire une séquence de mouvements de Reidemeister qui le dénoue ? (Le temps est fonction du nombre de croisements dans le diagramme.) Si un tel algorithme pouvait être trouvé, il ouvrirait une toute nouvelle approche de la reconnaissance et de la classification des nœuds. Je connais au moins un expert en théorie des nœuds qui croit fermement en l'existence d'un tel algorithme. Il n'est pas du tout trivial de trouver un algorithme pour reconnaître le nœud trivial, mais cela a été fait par Haken [H]. J'ai moi-même réalisé l'expérience scientifique suivante. Je prends un long morceau de ficelle, j'attache les deux extrémités ensemble avec un petit nœud, puis j'emmêle la boucle résultante, en essayant de la rendre difficile à démêler (mais sans serrer quoi que ce soit). Je n'ai jamais réussi à me causer la moindre difficulté.
16. L'hypothèse de Knaster. C'est la question suivante. Soit $f : S_n \rightarrow \mathbb{R}$ continue et soient $x_1, \dots, x_{n+1}$ $n + 1$ points quelconques dans S_n . Doit-il exister une application orthogonale α telle que les valeurs $f(\alpha x_i)$ soient toutes égales ? Lorsque $n = 1$, une réponse positive découle trivialement du théorème des valeurs intermédiaires. Je pense que la conjecture est connue pour être vraie lorsque $n = 2$ mais ouverte pour $n \geq 3$. Elle est également connue pour certains ensembles de points très particuliers $x_1, \dots, x_{n+1}$.

L'une des motivations de l'hypothèse de Knaster, outre l'attrait intrinsèque de la question, est que, comme l'a observé Milman [Mi], une réponse positive fournirait une démonstration complètement nouvelle de l'un des résultats centraux de la théorie des espaces de Banach – le théorème de Dvoretzky [D]. Celui-ci stipule que pour tout k et $\epsilon > 0$, il existe n tel que tout

espace normé de dimension n ait un sous-espace de dimension k dont la distance de Banach-Mazur à un espace de Hilbert est au plus $1 + \epsilon$. (Géométriquement, il dit qu'un corps convexe centralement symétrique de dimension n a une section k -dimensionnelle presque ellipsoïdale.) Le théorème de Dvoretzky découle de l'hypothèse de Knaster comme suit. Considérons S_{n-1} comme la sphère unité de ℓ_2^n , soient $x_1, \dots, x_n$ un δ -réseau de la sphère unité d'un sous-espace V de dimension k , et soit f la restriction d'une norme arbitraire $\|\cdot\|$ sur $\mathbb{R}^n$ à S_{n-1} . Par l'hypothèse de Knaster, il existe une rotation α telle que toutes les valeurs $\|\alpha x_i\|$ soient égales. Des arguments standard montrent alors que tous les points x dans le sous-espace αV ont à peu près le même rapport $\|x\|/\|x\|_2$. Cet argument, rendu précis, donne la meilleure estimée possible pour la dépendance asymptotique de n en fonction de ϵ pour k fixé, qui est encore un problème ouvert.

Certaines approches naturelles de l'hypothèse de Knaster s'avèrent échouer pour de bonnes raisons, et beaucoup ont le sentiment qu'elle pourrait bien être fausse. Cependant, même un affaiblissement considérable, comme le prouver pour $\sqrt{n}$ points dans S_n , serait utile pour le théorème de Dvoretzky.

17. Le théorème de densité de Hales-Jewett. Le théorème de Hales-Jewett concerne les colorations de l'ensemble $\{1, 2, \dots, k\}^N$. Si x appartient à cet ensemble, A est un sous-ensemble de $\{1, 2, \dots, N\}$ et $1 \leq j \leq k$, alors écrivons $x \oplus jA$ pour le point y obtenu à partir de x en changeant x_i en j si $i \in A$ et en le laissant inchangé sinon. Une ligne combinatoire est un ensemble de la forme $\{x \oplus jA : 1 \leq j \leq k\}$. Par exemple, lorsque $k = 3$ et $N = 5$, le point $x = (3, 3, 1, 1, 2)$ et l'ensemble $\{1, 3, 5\}$ donnent la ligne combinatoire

$$\{(1, 3, 1, 1, 1), (2, 3, 2, 1, 2), (3, 3, 3, 1, 3)\}.$$

Remarquez que si l'ensemble $\{1, 2, \dots, k\}^N$ est imaginé comme une grille N -dimensionnelle, alors les lignes combinatoires sont des cas particuliers de lignes géométriques, c'est-à-dire des ensembles de k points alignés.

Le théorème de Hales-Jewett stipule que pour tout k et r , il existe N tel que si l'ensemble $\{1, 2, \dots, k\}^N$ est coloré avec r couleurs, alors il doit y avoir une ligne combinatoire composée de points de la même couleur. Ce théorème est facilement vu comme un renforcement du théorème de van der Waerden, et comme le théorème de van der Waerden, il a une version en densité : pour tout k et tout $\delta > 0$, il existe N tel que tout sous-ensemble $A \subset \{1, 2, \dots, k\}^N$ avec au moins δk^N éléments contienne une ligne combinatoire.

Ce théorème de densité est dû à Furstenberg et Katznelson [FuK] et fait un usage intensif de la théorie ergodique. De plus, aucune démonstration n'est connue qui n'utilise pas ces techniques. Il serait très intéressant de trouver une démonstration non ergodique, même dans le cas $k = 3$. À partir du théorème dans ce cas, il est facile de déduire le résultat suivant. Pour tout $\delta > 0$, il existe N tel que tout sous-ensemble de la grille bidimensionnelle $\{1, 2, \dots, N\}^2$ de taille au moins δN^2 contienne un triplet de la forme $\{(x, y), (x, y + d), (x + d, y)\}$, c'est-à-dire un triangle isocèle rectangle aligné. Même cet énoncé manque de toute démonstration qui donne une borne supérieure raisonnable pour N en fonction de δ .

18. Le problème de Danzer. Est-il possible de trouver un ensemble X de points dans $\mathbb{R}^2$ tel que tout corps convexe d'aire 1 contienne au moins un point de X et pas plus de C points, où C est une constante absolue? Si K est un corps convexe d'aire 1, alors il existe des rectangles homothétiques R_1 et R_2 tels que $R_1 \subset K \subset R_2$ et tels que l'aire de R_2 soit au plus quatre fois celle de R_1 . Il s'ensuit qu'il suffit de considérer les rectangles dans l'énoncé du problème.

Il y a un petit lien entre ce problème et certaines des questions de la section précédente, car les méthodes aléatoires et explicites évidentes pour construire un tel ensemble X échouent toutes deux, mais pour des raisons très différentes.

La question réelle de Danzer est plus faible que celle ci-dessus. Il demande un ensemble X qui intersecte tout corps convexe d'aire 1 et satisfasse également l'inégalité $|X \cap RD| \leq CR^2$ pour tout R , où D est le disque unité centré à l'origine. Un bel article de Bambah et Woods [BW] montre que X ne peut pas être une union finie de translatés de réseaux, écartant ainsi ce qui semble être une source prometteuse d'exemples potentiels.

19. Billards dans les polygones. Il existe un exemple célèbre, connu sous le nom de l'étude de Kafka, qui montre que l'on peut concevoir une pièce connexe (mais non convexe) telle que, même si les murs sont entièrement recouverts de miroirs, une seule source lumineuse ne peut pas éclairer toute la pièce. L'exemple exploite le fait qu'un rayon lumineux qui passe entre les deux foyers d'une ellipse revient entre les deux foyers après avoir été réfléchi par la limite de l'ellipse. On ne sait pas s'il existe un exemple polygonal.

Une question étroitement liée est la suivante. Supposons qu'une bille de billard infinitésimale soit mise en mouvement sur une table polygonale sans frottement (et sans trous). Reviendra-t-elle arbitrairement près de sa position et direction initiales? On sait que la réponse est positive si les angles entre tous les côtés du polygone sont des multiples rationnels de π , car alors il n'y a qu'un nombre fini de directions à considérer, et le problème peut facilement être réduit à un résultat connu et facile concernant les applications dites d'échange d'intervalles.

Soit $[0, 1)$ partitionnée en m intervalles semi-ouverts de deux manières différentes, $I_1 \cup \dots \cup I_m$ et $J_1 \cup \dots \cup J_m$ de telle sorte que I_r et J_r aient la même longueur pour tout $r \leq m$. Soit ϕ_r la translation de I_r vers J_r . Soit ϕ l'application qui envoie $x \in I_r$ vers $\phi_r(x)$. Une application d'échange d'intervalles de $[0, 1)$ dans $[0, 1)$ est toute application de ce type. Il est assez simple de démontrer que de telles applications sont récurrentes, au sens où pour tout x et tout $\epsilon > 0$, il existe $N > 0$ tel que la distance de $\phi^N x$ à x est au plus ϵ .

J'ai interrogé des gens et je n'ai trouvé personne qui puisse me dire si la généralisation bidimensionnelle évidente du résultat ci-dessus est vraie ou fausse. C'est-à-dire que l'on suppose que l'on découpe un carré en rectangles et que l'on les réarrange (par translation) pour qu'ils forment à nouveau le carré original. L'application d'échange de rectangles résultante est-elle récurrente? Je crois qu'elle l'est, mais certains experts en systèmes dynamiques à qui j'ai posé la question ont l'instinct opposé.

Références

- [AM] N. Alon, V. D. Milman, Eigenvalues, expanders and superconcentrators, Proc. 25th Annual FOCS, Singer Island, FL, IEEE, New York (1984), 320-322.
- [BW] R. P.ambah, A.C. Woods, On a problem of Danzer, Pacific J. Math. 37 (1971), 295-301.
- [B] F. A. Behrend, On sets of integers which contain no three in arithmetic progression, Proc. Nat. Acad. Sci. 23 (1946), 331-332.
- [Bi] Y. Bilu, Structure of sets with small sumset, Asterisque 258 (1999), 77-108.
- [Bo1] J. Bourgain, Remarks on Montgomery's conjectures on Dirichlet series, Geometric Aspects of Functional Analysis (1989-1990), Springer Lecture Notes in Mathematics 1469 (1991), 153-165.
- [Bo2] J. Bourgain, Besicovitch type maximal operators and applications to Fourier analysis, GAFA 1 (1991), 147-187.
- [Bo3] J. Bourgain, On triples in arithmetic progression, GAFA 9 :5 (1999), 968-984.
- [Bo4] J. Bourgain, On the dimension of Kakeya sets and related maximal inequalities, GAFA 9 :2 (1999), 256-282.
- [Bu] George Bush, Interview, Time, 26th Jan. 1987.
- [CG] F. R. K. Chung, R.L. Graham, Quasi-random subsets of $\mathbb{Z}_n$, J. Comb. Th. A 61 (1992), 64-86.
- [CGW] F. R. K. Chung, R. L. Graham, R. M. Wilson, Quasi-random graphs, Combinatorica 9 (1989), 345-362.
- [Co] S. A. Cook, The complexity of theorem proving procedures, Proc. 3rd Annual ACM Symposium on the Theory of Computing (1971), 151-158.
- [D] A. Dvoretzky, Some results on convex bodies and Banach spaces, Proc. Symp. on Linear Spaces, Jerusalem (1961), 123-160.
- [E] P. Erds, Some remarks on the theory of graphs, Bull. Amer. Math. Soc. 53 (1947), 292-294.
- [ET] P. Erds, P. Turán, On some sequences of integers, J. London Math. Soc. 11 (1936), 261-264.
- [F] C. Fefferman, The multiplier problem for the ball, Annals of Math. 94 (1971), 330-336.
- [FuK] H. Furstenberg, Y. Katznelson, A density version of the Hales-Jewett theorem, J. D'Analyse Math. 57 (1991), 64-119.
- [Fr1] G. A. Freiman, Foundations of a Structural Theory of Set Addition (in Russian), Kazan Gos. Ped. Inst., Kazan 1966.
- [Fr2] G. A. Freiman, Foundations of a Structural Theory of Set Addition, Translations of Mathematical Monographs 37, Amer. Math. Soc., Providence, R.I., USA, 1973.
- [G1] W. T. Gowers, Lower bounds of tower type for Szemerédi's uniformity lemma, GAFA 7 (1997), 322-337.
- [G2] W. T. Gowers, A new proof of Szemerédi's theorem for arithmetic progressions of length four, GAFA 8 (1998), 529-551.
- [GrRR] R.L. Graham, V. Rödl, A. Ruciński, On graphs with linear Ramsey numbers, preprint.
- [H] W. Haken, Theorie der Normalflächen, ein Isotopiekriterium für den Kreisnoten, Acta Math. 105, 245-375.
- [HaW] J. M. Hammersley, D. J. A. Welsh, Further results on the rate of convergence to the connective constant of the hypercubical lattice, Quart. J. Math. Oxford Ser. (2) 13 (1962), 108-110.
- [HarS1] T. Hara, G. Slade, Self-avoiding walk in five or more dimensions, Comm. Math. Phys. 147 (1992), 101-136.
- [HarS2] T. Hara, G. Slade, The lace expansion for self-avoiding walk in five or more dimensions, Rev. Math. Phys. 4 (1992), 235-327.
- [He] D. R. Heath-Brown, Integer sets containing no arithmetic progressions, J. London Math. Soc. (2) 35 (1987), 385-394.
- [K] R. M. Karp, Reducibility among combinatorial problems, Complexity of Computer Computations, Proc. Sympos., IBM Thomas J. Watson Res. Centr, Yorktown Heights, N.Y., 1972, (R.E. Miller, J. W. Thatcher, eds.) Plenum Press, New York 1972, 85-103.
- [KaLT] N. H. Katz, I. Laba, T. Tao, An improved bound on the Minkowski dimension of Besicovitch sets in $\mathbb{R}^3$, Annals of Math., to appear.
- [KaT] N. H. Katz, T. Tao, A new bound on partial sum-sets and difference-sets, and applications to the Kakeya conjecture, submitted.

- [KoS] J. Komlós, M. Simonovits, Szemerédi's Regularity Lemma and its applications in Graph Theory, in "Combinatorics, Paul Erdős is 80 (Vol 2)", Bolyai Society Math. Studies 2, 295-352, Kesthely (Hungary) 1993, Budapest 1996.
- [LPS] A. Lubotzky, R. Phillips, P. Sarnak, Explicit expanders and the Ramanujan conjectures, Proceedings of the 18th ACM Symposium on the Theory of Computing 1986, 240-246; also *Combinatorica* 8 (1988), 261-277.
- [MS] N. Madras, G. Slade, The Self-Avoiding Random Walk, Birkhäuser, Boston, 1992.
- [Mi] V. D. Milman, A few observations on the connections between local theory and some other fields, in Geometric Aspects of Functional Analysis, Israel seminar (GAFA) 1986-1987 (J. Lindenstrauss, V.D. Milman, eds.), Springer LNM 1317, (1988), 283-289.
- [Mo] H. L. Montgomery, Ten Lectures on the Interface Between Analytic Number Theory and Harmonic Analysis, CBMS Regional Conference Series in Math. 84, AMS 1994.
- [P] G. Polya, Mathematics and Plausible Reasoning, Vols. I and II, Princeton University Press, 1954.
- [RR] A. A. Razborov, S. Rudich, Natural proofs, in 26th Annual ACM Symposium on the Theory of Computing (STOC '94, Montreal, PQ, 1994); also *J. Comput. System Sci.* 55 (1997), 24-35.
- [Ro1] K. Roth, On certain sets of integers, *J. London Math. Soc.* 28 (1953), 245-252.
- [Ro2] K. Roth, Rational approximations to algebraic numbers, *Mathematika* 2 (1955), 1-20 (with corrigendum p. 168).
- [Ru] I. Z. Ruzsa, Generalized arithmetic progressions and sumsets, *Acta Math. Hungar.* 65 (1995), 379-388.
- [S1] E. Szemerédi, Regular partitions of graphs, *Colloques Internationaux C.N.R.S. 260 - Problemes Combinatoires et Théorie des Graphes*, Orsay 1976, 399-401.S2>
- E. Szemerédi, On sets of integers containing no k elements in arithmetic progression, *Acta Arith.* 27 (1975), 299-345.
- [S3] E. Szemerédi, Integer sets containing no arithmetic progressions, *Acta Math. Hungar.* 56 (1990), 155-158.
- [T1] A. G. Thomason, Pseudo-random graphs, *Proceedings of Random Graphs*, Poznan 1985 (M. Karonski, ed.), *Annals of Discrete Mathematics* 33, 307-331.
- [T2] A. G. Thomason, A disproof of a conjecture of Erdős in Ramsey theory, *J. London Math. Soc.* 39 (1989), 246-255.
- [T3] A. G. Thomason, Graph products and monochromatic multiplicities, *Combinatorica* 17 (1997), 125-134.
- [W] T. H. Wolff, An improved bound for Kakeya type maximal functions, *Revista Mat. Iberoamericana* 11 (1995), 651-674.