

Discret et continu : deux faces d'une même pièce ?

László Lovász

GAFA 2000

Résumé

Dans quelle mesure la frontière entre mathématiques discrètes et continues est-elle profonde ? Les structures et méthodes fondamentales des deux versants de notre science sont assez différentes. Mais à un niveau plus profond, il existe un lien plus solide qu'il n'y paraît.

1. Introduction

Il existe une frontière assez nette entre les mathématiques discrètes et continues. Les mathématiques continues sont classiques, bien établies, avec une grande variété d'applications. Les mathématiques discrètes sont issues de casse-têtes et sont souvent identifiées à des domaines d'application spécifiques comme l'informatique ou la recherche opérationnelle. Les structures et méthodes fondamentales des deux versants de notre science sont assez différentes (les mathématiciens du continu utilisent des limites ; les mathématiciens du discret utilisent l'induction).

Cet état des mathématiques à cette époque particulière, le tournant du millénaire, est le produit d'une logique intrinsèque forte, mais aussi d'une coïncidence historique.

La principale source extérieure de problèmes mathématiques est la science, en particulier la physique. La vision traditionnelle est que l'espace et le temps sont continus, et que les lois de la nature sont décrites par des équations différentielles. Il y a, bien sûr, des exceptions : la chimie, la mécanique statistique et la physique quantique sont basées sur un certain degré de discrétion. Mais ces objets discrets (molécules, atomes, particules élémentaires, graphes de Feynman) vivent dans le continuum de l'espace-temps. La théorie des cordes tente d'expliquer ces objets discrets comme des singularités d'un continuum de dimension supérieure.

En conséquence, les mathématiques utilisées dans la plupart des applications sont l'analyse, le noyau dur de notre science. Mais surtout dans les sciences en plein essor, les modèles discrets deviennent de plus en plus importants.

On pourrait observer qu'il y a un nombre fini d'événements dans tout domaine fini de l'espace-temps. Le reste d'une variété à quatre (ou dix) dimensions a-t-il une signification physique ? L'expression "le point dans le temps à mi-chemin entre deux interactions consécutives d'une particule élémentaire" a-t-elle un sens ? Que la réponse à cette question soit oui ou non, les caractéristiques

discrètes du monde subatomique sont indéniables. Jusqu'où une description combinatoire du monde nous mènerait-elle ? Ou les descriptions du monde comme un continuum, ou comme une structure discrète (mais terriblement grande), pourraient-elles être équivalentes, mettant en évidence deux facettes d'une même réalité ?

Mais laissez-moi m'éloigner de ces questions sans réponse et regarder ailleurs dans la science. La biologie tente de comprendre le code génétique : une tâche gigantesque, qui est la clé de la compréhension de la vie et, en fin de compte, de nous-mêmes. Le code génétique est discret : des questions de base simples comme trouver des motifs correspondants, ou tracer les conséquences du retournement de sous-chaînes, semblent plus familières au théoricien des graphes qu'au chercheur en équations différentielles. Des questions sur le contenu informationnel, la redondance ou la stabilité du code peuvent sembler trop vagues pour un mathématicien classique, mais un informaticien théoricien verra immédiatement au moins quelques outils pour les formaliser (même si trouver la réponse peut être trop difficile pour l'instant).

L'économie est un grand consommateur de mathématiques – et une grande partie de ses besoins ne font pas partie de la boîte à outils des mathématiques appliquées traditionnelles. L'outil peut-être le plus réussi en économie et en recherche opérationnelle est la programmation linéaire, qui vit à la frontière du discret et du continu. L'applicabilité de la programmation linéaire dans ces domaines dépend cependant de conditions de convexité et de divisibilité illimitée ; prendre en compte les indivisibilités (par exemple, les décisions logiques, ou les agents individuels) conduit à la programmation en nombres entiers et à d'autres modèles de nature combinatoire.

Enfin, le monde des ordinateurs est essentiellement discret, et il n'est pas surprenant que tant de mathématiciens du discret travaillent dans cette science. La communication et le calcul électroniques fournissent un vaste éventail de problèmes mathématiques bien formulés, difficiles et importants sur les algorithmes, les bases de données, les langages formels, la cryptographie et la sécurité informatique, le placement de circuits VLSI, et bien plus encore.

Dans tous ces domaines, la compréhension réelle implique, je crois, une synthèse du discret et du continu, et c'est un objectif intellectuellement très stimulant de développer ces outils mathématiques.

Il existe différents niveaux d'interaction entre les mathématiques discrètes et continues, et je les traite (je crois) par ordre d'importance croissante.

1. Nous utilisons souvent le fini pour approximer l'infini. Discrétiser une structure continue compliquée a toujours été une méthode fondamentale – depuis la définition de l'intégrale de Riemann en passant par la triangulation d'une variété en (disons) théorie de l'homologie jusqu'à la résolution numérique d'une équation aux dérivées partielles sur une grille.

C'est une idée un peu plus subtile que l'infini est souvent (ou peut-être toujours ?) une approximation du grand fini. Les structures continues sont souvent plus propres, plus symétriques et plus riches que leurs homologues discrets (par exemple, une grille planaire a un degré de symétrie beaucoup plus petit que le plan euclidien tout entier). C'est une méthode naturelle et puissante d'étudier les structures discrètes en les "plongeant" dans le monde continu.

2. Parfois, l'étape clé dans la démonstration d'un résultat purement "discret" est l'application d'un théorème purement "continu", ou vice versa. Nous avons tous nos exemples préférés ; j'en décris deux dans la section 2.
3. Dans certains domaines des mathématiques discrètes, des progrès clés ont été réalisés grâce à l'utilisation de méthodes de plus en plus sophistiquées de l'analyse. Ceci est illustré dans la section 3 par la description de deux méthodes puissantes en optimisation discrète. (Je n'ai trouvé aucun domaine des mathématiques "continues" où des progrès auraient été réalisés à une échelle similaire grâce à l'introduction de méthodes discrètes. Peut-être que la topologie algébrique s'en rapproche le plus.)
4. Les liens entre le discret et le continu peuvent faire l'objet d'une étude mathématique en soi. L'analyse numérique peut être considérée de cette manière, mais la théorie des discordances en est le meilleur exemple. Dans cet article, nous devons nous limiter à discuter de deux résultats classiques dans ce domaine florissant (section 4) ; nous renvoyons au livre de Beck et Chen [BeC], à l'article de synthèse de Beck et Sós [BeS] et au livre récent de Matoušek [M].
5. Le niveau d'interaction le plus significatif est celui où un même phénomène apparaît à la fois dans le cadre continu et discret. Dans de tels cas, l'intuition et la perspicacité acquises en considérant l'un de ces cadres peuvent être extrêmement utiles dans l'autre. Un exemple bien connu est le lien entre les suites et les fonctions analytiques, fourni par le développement en série entière. Dans ce cas, il existe un "dictionnaire" entre les aspects combinatoires (réurrences, asymptotiques) de la suite et les propriétés analytiques (équations différentielles, singularités) de sa fonction génératrice.

Dans la section 5, je discute en détail la notion discrète et continue de "Laplacien", reliée à une variété de processus de type dispersion. Cette notion relie des sujets allant de l'équation de la chaleur au mouvement brownien, en passant par les marches aléatoires sur les graphes et les plongements de graphes sans entrelacs.

Un parallélisme passionnant mais encore mal compris est lié au fait que l'itération d'étapes très simples donne lieu à des structures très complexes. En mathématiques continues, cette idée apparaît dans la théorie des systèmes dynamiques et l'analyse numérique. En mathématiques discrètes, elle est en un sens la base de nombreux algorithmes, de la génération de nombres aléatoires, etc. Une synthèse des idées des deux côtés pourrait apporter des développements spectaculaires ici.

On pourrait mentionner de nombreux autres domaines avec des composantes discrètes et continues substantielles : les groupes, les probabilités, la géométrie... En effet, mon observation de départ sur cette division devient discutable si l'on pense à certains des développements récents dans ces domaines. Je crois que les développements futurs la rendront totalement dénuée de sens.

Remerciements. Je suis redevable à plusieurs de mes collègues, en particulier à Alain Connes et

Mike Freedman, pour leurs commentaires précieux sur cet article.

2. Discret dans le continu et continu dans le discret

2.1. Mariages et mesures

Une application frappante de la théorie des graphes à la théorie de la mesure est la construction de la mesure de Haar sur les groupes topologiques compacts. Cette application a été mentionnée par Rota et Harper [RoH], qui ont développé l'idée sur l'exemple de la construction d'une intégrale invariante par translation pour les fonctions presque périodiques sur un groupe arbitraire. La démonstration est également liée à la démonstration de Weil de l'existence de la mesure de Haar. (Il existe d'autres applications, peut-être plus significatives, des couplages aux mesures qui pourraient être mentionnées ici, par exemple le théorème d'Ornstein [O] sur l'isomorphisme des décalages de Bernoulli ; cf. aussi [Du], [LoM], [LoP], [St].)

Nous décrivons ici la construction d'une intégrale invariante par translation pour les fonctions continues sur un groupe topologique compact, équivalente à l'existence de la mesure de Haar [LoP]. Une intégration invariante pour les fonctions continues sur un groupe topologique compact est une fonctionnelle linéaire L définie sur les fonctions continues à valeurs réelles sur G , avec les propriétés suivantes :

- (a) $L(\alpha f + \beta g) = \alpha L(f) + \beta L(g)$ (linéarité)
- (b) Si $f \geq 0$, alors $L(f) \geq 0$ (monotonie)
- (c) Si l désigne la fonction identité, alors $L(l) = 1$ (normalisation)
- (d) Si s et t sont dans G et f et g sont deux fonctions continues telles que $g(x) = f(sxt)$ pour tout x , alors $L(g) = L(f)$ (invariance par double translation).

Théorème 1. *Pour tout groupe topologique compact, il existe une intégration invariante pour les fonctions continues.*

Démonstration. Soit f la fonction que nous voulons intégrer. L'idée est d'approcher l'intégrale par la moyenne

$$f(A) = \frac{1}{|A|} \sum_{a \in A} f(a),$$

où A est un sous-ensemble fini "uniformément dense" du groupe G . La question est de savoir comment trouver un ensemble A approprié.

La définition clé est la suivante. Soit U un sous-ensemble ouvert non vide de G (pensez-y comme étant "petit"). Un sous-ensemble $A \subset G$ est appelé un U -réseau si $A \cap sUt \neq \emptyset$ pour tout $s, t \in G$. Il résulte de la compacité de G qu'il existe un U -réseau fini.

Bien sûr, un U -réseau peut être très inégalement réparti sur le groupe, et en conséquence, la moyenne sur celui-ci peut ne pas approcher l'intégrale. Ce que nous montrons, c'est que le simple fait de nous restreindre aux U -réseaux de cardinalité minimale (en abrégé, U -réseaux minimaux),

nous obtenons un ensemble fini suffisamment uniformément dense.

Pour mesurer cette uniformité, nous définissons

$$\delta(U) = \sup \{|f(x) - f(y)| \mid x, y \in sUt \text{ pour un certain } s, t \in G\}.$$

La compacité de G implique que si f est continue, alors elle est aussi uniformément continue au sens où pour tout $\epsilon > 0$, il existe un ensemble ouvert non vide U tel que $\delta(U) < \epsilon$.

Soient A et B des U -réseaux minimaux. L'inégalité suivante est le cœur de la preuve :

$$|f(A) - f(B)| \leq \delta(U). \quad (1)$$

Le noyau combinatoire de la construction réside dans la preuve de cette inégalité. Définissons un graphe biparti H avec bipartition (A, B) en reliant $x \in A$ à $y \in B$ si et seulement s'il existe $s, t \in G$ tels que $x, y \in sUt$. Nous utilisons le théorème de Mariage pour montrer que ce graphe biparti a un couplage parfait. Nous devons vérifier deux choses :

- (I) $|A| = |B|$. Ceci est trivial.
- (II) Tout ensemble $X \subseteq A$ a au moins $|X|$ voisins dans B . En effet, soit Y l'ensemble des voisins de X . Nous montrons que $T = Y \cup (A \setminus X)$ est un U -réseau. Soient $s, t \in G$, nous montrons que T intersecte sUt . Puisque A est un U -réseau, il existe un élément $x \in A \cap sUt$. Si $x \notin X$, alors $x \in T$ et nous avons fini. Sinon, nous utilisons qu'il existe un élément $y \in B \cap sUt$. Alors xy est une arête de H et donc $y \in Y$, et nous avons à nouveau fini.

Ainsi T est un U -réseau. Puisque A est un U -réseau de cardinalité minimale, nous devons avoir $|T| \geq |A|$, ce qui équivaut à $|Y| \geq |X|$.

Ainsi H a un couplage parfait $\{a_1b_1, \dots, a_nb_n\}$. Alors

$$\begin{aligned} |f(A) - f(B)| &= \frac{1}{n} \left| \sum_{i=1}^n (f(a_i) - f(b_i)) \right| \\ &\leq \frac{1}{n} \sum_{i=1}^n |f(a_i) - f(b_i)| \leq \frac{1}{n} (n\delta(U)) = \delta(U). \end{aligned}$$

Le reste de la preuve est assez routinier. D'abord, nous devons montrer que la moyenne sur des U -réseaux minimaux pour différents U donne approximativement le même résultat. Plus exactement, soit A un U -réseau minimal et B un V -réseau minimal, alors

$$|f(A) - f(B)| \leq \delta(U) + \delta(V). \quad (2)$$

En effet, pour tout $b \in B$, Ab est aussi un U -réseau de cardinalité minimale, donc par (1),

$$|f(A) - f(Ab)| \leq \delta(U).$$

Par conséquent,

$$|f(A) - f(AB)| = \left| f(A) - \frac{1}{|B|} \sum_{b \in B} f(Ab) \right| \leq \frac{1}{|B|} \sum_{b \in B} |f(A) - f(Ab)| \leq \delta(U).$$

De même,

$$|f(B) - f(AB)| \leq \delta(V),$$

d'où (2) découle.

Choisissons maintenant une suite U_n d'ensembles ouverts tels que $\delta(U_n) \rightarrow 0$, et soit A_n un U_n -réseau minimal. Alors (2) implique que la suite $f(A_n)$ tend vers une limite $L(f)$, qui est indépendante du choix de U_n et A_n , et donc elle est bien définie. Les conditions (a)–(d) sont triviales à vérifier. $\square$

2.2. Sous-ensembles disjoints et topologie

Nous nous tournons maintenant vers des exemples démontrant la direction inverse : les applications de méthodes continues à des problèmes purement combinatoires. Notre premier exemple est un résultat où la topologie algébrique et la géométrie sont les outils essentiels dans la preuve d'un théorème combinatoire.

Théorème 2. *Partitionnons les sous-ensembles à k éléments d'un ensemble à n éléments en $n - 2k + 1$ classes. Alors l'une des classes contient deux sous-ensembles à k éléments disjoints.*

Ce résultat a été conjecturé par Kneser [K] et prouvé dans [Lo1]. La preuve a été simplifiée par Barany [B], et nous décrivons sa version ici.

Démonstration. Nous invoquons d'abord une construction géométrique due à Gale. Il existe un ensemble de $2k + d$ vecteurs sur la sphère S^d tel que tout hémisphère ouvert en contienne au moins k . (Pour $d = 1$, prendre les sommets d'un $(2k + 1)$ -gone régulier.)

En choisissant $d = n - 2k$, nous obtenons donc un ensemble S de n points. Supposons que tous les sous-ensembles à k éléments sont partitionnés en $d + 1 = n - 2k + 1$ classes $\mathcal{P}_0, \mathcal{P}_1, \dots, \mathcal{P}_d$. Soit A_i l'ensemble des vecteurs unitaires $h \in S^d$ pour lesquels l'hémisphère ouvert centré en h contient un sous-ensemble à k éléments de S qui appartient à $\mathcal{P}_i$. Clairement, les ensembles $A_0, A_1, \dots, A_d$ sont ouverts, et par définition de S , ils recouvrent toute la sphère.

Nous passons maintenant de la géométrie à la topologie : par le théorème de Borsuk-Ulam, l'un des ensembles A_i contient deux points antipodaux h et $-h$. Ainsi, les hémisphères centrés en h et $-h$ contiennent tous deux un sous-ensemble à k éléments de $\mathcal{P}_i$. Puisque ces hémisphères sont disjoints, ces deux sous-ensembles à k éléments sont disjoints. $\square$

Il existe de nombreux autres exemples où des méthodes de topologie algébrique ont été utilisées pour prouver des énoncés purement combinatoires (voir [Bj] pour une synthèse).

3. Optimisation : discrète, linéaire, semi-définie

Notre exemple suivant illustre comment tout un domaine de problèmes mathématiques appliqués importants, à savoir l'optimisation discrète, repose sur des idées qui font appel à des outils de plus

en plus sophistiqués de l'optimisation continue plus traditionnelle.

Dans un problème d'optimisation typique, on nous donne un ensemble S de solutions réalisables, et une fonction $f : S \rightarrow \mathbf{R}$, appelée fonction objectif. Le but est de trouver l'élément de S qui maximise (ou minimise) la fonction objectif.

Dans un problème d'optimisation traditionnel, S est un continuum décent dans un espace euclidien, et f est une fonction lisse. Dans ce cas, l'optimum peut être trouvé en considérant l'équation $\nabla f = 0$, ou par une méthode itérative utilisant une version de la descente de plus grande pente.

Ces deux méthodes dépendent de la structure continue dans un voisinage du point d'optimisation, et échouent donc pour un problème d'optimisation discret (ou combinatoire), lorsque S est fini. Ce cas est trivial d'un point de vue classique : “en principe”, on pourrait évaluer f pour tous les éléments de S , et choisir le meilleur. Mais dans la plupart des cas qui nous intéressent, S est très grand par rapport au nombre de données nécessaires pour spécifier l'instance du problème. Par exemple, S peut être l'ensemble des arbres couvrants, ou l'ensemble des couplages parfaits, d'un graphe ; ou l'ensemble de tous les horaires possibles de tous les trains d'un pays ; ou l'ensemble des états d'un verre de spin. Dans de tels cas, nous devons utiliser la structure implicite de S , plutôt que la force brute, pour trouver l'élément optimisant.

Dans certains cas, des méthodes totalement combinatoires nous permettent de résoudre de tels problèmes. Par exemple, soit S l'ensemble de tous les arbres couvrants d'un graphe G et supposons que chaque arête de G ait une “longueur” non négative qui lui est associée. Dans ce cas, un arbre couvrant de longueur minimale peut être trouvé par l'algorithme glouton (dû à Borouvka et Kruskal) : nous choisissons à plusieurs reprises l'arête la plus courte qui ne forme pas de cycle avec les arêtes choisies précédemment, jusqu'à obtenir un arbre.

Dans de nombreux cas (dans un sens, la plupart), un tel algorithme combinatoire direct n'est pas disponible. Une approche générale consiste à plonger l'ensemble S dans un continuum S' de solutions, et aussi à étendre la fonction objectif f en une fonction f' définie sur S' . Le problème de la minimisation de f' sur S' est appelé une relaxation du problème original. Si nous faisons cela correctement (disons, S' est un ensemble convexe et f' est une fonction convexe), alors le minimum de f' sur S' peut être trouvé par des outils classiques (différentiation, descente de plus grande pente, etc.). Si nous sommes vraiment chanceux (ou habiles), l'élément minimisant de S' appartiendra à S , et alors nous aurons résolu le problème discret original. (Sinon, nous pouvons encore utiliser la solution de la relaxation pour obtenir une borne sur la solution du problème original, ou pour obtenir une solution approchée. Voir plus loin.)

3.1. Combinatoire polyédrique

La première réalisation réussie de ce schéma a été élaborée dans les années 60 et 70, où les techniques de la programmation linéaire ont été appliquées à la combinatoire. Supposons que notre problème d'optimisation combinatoire puisse être formulé de sorte que S soit un ensemble de vecteurs 0–1 et que la fonction objectif soit linéaire. L'ensemble S peut être spécifié par une variété de contraintes

logiques et autres ; dans la plupart des cas, il est assez facile de les traduire en inégalités linéaires :

$$S = \{x \in \{0, 1\}^n : a_1^\top x \leq b_1, \dots, a_m^\top x \leq b_m\}. \quad (3)$$

La fonction objectif est

$$f(x) = c^\top x = \sum_{i=1}^n c_i x_i, \quad (4)$$

où les c_i sont des nombres réels donnés. Une telle formulation est généralement facile à trouver.

Ainsi, nous avons traduit notre problème d'optimisation combinatoire en un programme linéaire avec des conditions d'intégralité. Il est assez facile de résoudre ce problème si nous négligeons les conditions d'intégralité ; le vrai jeu est de trouver des moyens d'écrire ces programmes linéaires de telle manière que la négligence des conditions d'intégralité soit justifiée.

Un bel exemple est le couplage dans les graphes bipartis. Supposons que G soit un graphe biparti, avec un ensemble de nœuds $\{u_1, \dots, u_n, v_1, \dots, v_n\}$ où chaque arête relie un nœud u_i à un nœud v_j . Nous voulons trouver un couplage parfait, c'est-à-dire un ensemble d'arêtes couvrant chaque nœud exactement une fois.

Supposons que G ait un couplage parfait M et soit

$$x_{ij} = \begin{cases} 1, & \text{si } u_i v_j \in M, \\ 0, & \text{sinon.} \end{cases}$$

Alors la propriété définissante des couplages parfaits peut être exprimée comme suit :

$$\sum_{i=1}^n x_{ij} = 1 \quad \text{pour tout } j, \quad \sum_{j=1}^n x_{ij} = 1 \quad \text{pour tout } i. \quad (5)$$

Réciproquement, si nous trouvons une solution de ce système d'équations linéaires avec chaque $x_{ij} = 0$ ou 1 , alors nous avons un couplage parfait.

Malheureusement, la résolubilité d'un système d'équations linéaires (même d'une seule équation) est NP-difficile. Ce que nous pouvons faire, c'est remplacer la condition $x_{ij} = 0$ ou 1 par la condition plus faible

$$x_{ij} \geq 0. \quad (6)$$

Résoudre un système linéaire comme (5) en nombres réels non négatifs n'est toujours pas trivial, mais peut être fait efficacement (en temps polynomial) en utilisant la programmation linéaire.

Nous n'avons pas fini, bien sûr, car si nous trouvons que (5)-(6) a une solution, cette solution peut ne pas être entière et donc peut ne pas "signifier" un couplage parfait. Il existe diverses façons de conclure, en extrayant un couplage parfait d'une solution fractionnaire de (5)-(6). La plus élégante est la suivante. L'ensemble de toutes les solutions forme un polytope convexe. Or, tout sommet de ce polytope convexe est entier. (La preuve de ce fait est amusante, utilisant la règle de Cramer et des calculs de déterminants de base. Voir, par exemple, [LoP].)

Donc, nous exécutons un algorithme de programmation linéaire pour voir si (5)-(6) a une solution en nombres réels. Si ce n'est pas le cas, alors le graphe n'a pas de couplage parfait. Si oui, la plupart des algorithmes de programmation linéaire donnent automatiquement une solution de base, c'est-à-dire un sommet. C'est une solution entière de (5)-(6), et correspond donc à un couplage parfait.

De nombreuses variantes de cette idée ont été développées tant en théorie qu'en pratique. Il existe des moyens de générer automatiquement de nouvelles contraintes et de les ajouter à (3) si elles échouent ; il existe des méthodes plus subtiles, plus efficaces pour générer de nouvelles contraintes pour des problèmes spéciaux ; il existe des moyens de gérer le système (3) même s'il devient trop grand pour être écrit explicitement.

3.2. Optimisation semi-définie

Cette nouvelle technique de production de relaxations de problèmes d'optimisation discrète rend les problèmes continus d'une manière plus sophistiquée. Nous l'illustrons par le problème de la coupe maximum. Soit $G = (V, E)$ un graphe. Nous voulons trouver une partition $(S, \bar{S})$ de V pour laquelle le nombre M d'arêtes reliant S à $\bar{S}$ est maximum.

Ce problème est NP-difficile, donc nous ne pouvons pas espérer trouver un algorithme efficace (en temps polynomial) qui le résolve. Au lieu de cela, nous devons nous contenter de moins : nous essayons de trouver une coupe qui est proche de l'optimale.

Il est facile de trouver une coupe qui ramasse la moitié des arêtes. Il suffit de traiter les nœuds un par un, en les plaçant dans l'ensemble S ou $\bar{S}$, selon celui qui donne le plus grand nombre de nouvelles arêtes entre les deux ensembles. Puisqu'aucune coupe ne peut avoir plus que toutes les arêtes, cet algorithme simple obtient une approximation de la coupe maximum avec une erreur relative d'au plus 50%.

Il s'est avéré assez difficile d'améliorer ce fait simple, jusqu'à ce que Goemans et Williamson [GW] combinent l'optimisation semi-définie avec une technique d'arrondi aléatoire pour obtenir un algorithme d'approximation avec une erreur relative d'environ 12%. D'autre part, il a été prouvé par Hastad [H] (renforçant des résultats antérieurs d'Arora et al. [ArLMSS]) qu'aucun algorithme en temps polynomial ne peut produire une approximation avec une erreur relative inférieure à (environ) 6%, à moins que $P = NP$.

Nous esquissons l'algorithme de Goemans-Williamson. Décrivons toute partition $V_1 \cup V_2$ de V par un vecteur $x \in \{-1, 1\}^V$ en posant $x_i = -1$ ssi $i \in V_1$. Alors la taille de la coupe correspondant à x est $(1/4) \sum_{i,j} (x_i - x_j)^2$. Par conséquent, le problème MAX CUT peut être formulé comme la maximisation de la fonction quadratique

$$\frac{1}{4} \sum_{i,j} (x_i - x_j)^2 \quad (7)$$

sur tout $x \in \{-1, 1\}^V$. La condition sur x peut aussi être exprimée par des contraintes quadratiques :

$$x_i^2 = 1 \quad (i \in V). \quad (8)$$

Réduire un problème d'optimisation discret au problème de maximisation d'une fonction quadratique soumise à un système d'équations quadratiques peut sembler formidable, et on pourrait essayer d'utiliser les multiplicateurs de Lagrange et d'autres techniques classiques. Mais elles ne seront pas utiles ; en fait, le nouveau problème est très difficile (NP-difficile), et il n'est pas clair que nous gagnions quoi que ce soit.

L'astuce suivante consiste à linéariser, en introduisant de nouvelles variables $y_{ij} = x_i x_j$ ($1 \leq i, j \leq n$). La fonction objectif et les contraintes deviennent linéaires dans ces nouvelles variables y_{ij} .

$$\begin{aligned} & \text{maximiser} \quad \frac{1}{4} \sum_{i,j} (y_{ii} + y_{jj} - 2y_{ij}) \\ & \text{sujet à} \quad y_{11} = \dots = y_{nn} = 1. \end{aligned} \quad (9) - (10)$$

Bien sûr, il y a un piège : les variables y_{ij} ne sont pas indépendantes ! En introduisant la matrice symétrique $Y = (y_{ij})_{i,j=0}^n$, nous pouvons noter que

$$Y \text{ est semi-définie positive,} \quad (11)$$

et

$$Y \text{ a un rang 1.} \quad (12)$$

Le problème de la résolution de (9)-(10) avec les contraintes supplémentaires (11) et (12) est équivalent au problème original, et donc NP-difficile en général. Mais si nous supprimons (12), nous obtenons une relaxation traitable ! En effet, ce que nous obtenons est un programme semi-défini : maximiser une fonction linéaire des entrées d'une matrice semi-définie positive, soumise à des contraintes linéaires sur les entrées de la matrice.

Puisque la semi-définition positive d'une matrice Y peut être traduite en (infiniment) d'inégalités linéaires impliquant les entrées de Y :

$$v^T Y v \geq 0 \quad \text{pour tout } v \in \mathbf{R}^{n+1},$$

les programmes semi-définis peuvent être vus comme des programmes linéaires avec une infinité de contraintes. Cependant, ils se comportent beaucoup mieux que ce à quoi on pourrait s'attendre. Entre autres, il existe une théorie de la dualité pour eux (voir, par exemple, Wolkowitz [W]). Il est également important que les programmes semi-définis soient solubles en temps polynomial (à une erreur arbitrairement petite près ; notons que la solution optimale peut ne pas être rationnelle). En fait, la méthode de l'ellipsoïde [GrLS] et, plus important encore d'un point de vue pratique, les méthodes de points intérieurs [Al] s'étendent aux programmes semi-définis.

Revenant au problème de la coupe maximum, soit Y une solution optimale de (9)-(10)-(11), et soit M^* la valeur optimale de la fonction objectif. Puisque toute coupe définit une solution, nous avons $M^* \geq M$.

La deuxième astuce consiste à observer que puisque Y est semi-définie positive, nous pouvons l'écrire comme une matrice de Gram, c'est-à-dire qu'il existe des vecteurs u_i dans $\mathbf{R}^n$ tels que $u_i^T u_j = Y_{ij}$ pour tout i et j . En particulier, nous obtenons que $|u_i|^2 = 1$, et

$$\sum_{i,j} (u_i - u_j)^2 = 4M^*. \quad (13)$$

Choisissons maintenant un hyperplan aléatoire uniformément distribué H passant par l'origine. Cela divise les u_i en deux classes, et définit ainsi une coupe dans G . Le poids attendu de cette coupe est

$$\sum_{i,j} C_{ij} \mathbf{P}(H \text{ sépare } u_i \text{ et } u_j).$$

Mais il est trivial que la probabilité que H sépare u_i et u_j est $(1/\pi)$ fois l'angle entre u_i et u_j . À son tour, l'angle entre u_i et u_j est d'au moins $0.21964(u_i - u_j)^2$, ce qui se voit facilement par un calcul élémentaire. Par conséquent, la taille attendue de la coupe obtenue est d'au moins

$$\sum_{i,j} 0.21964(u_i - u_j)^2 = 0.87856M^* \geq 0.87856M.$$

Ainsi, nous obtenons une coupe qui est au moins environ 88% de l'optimale.

Sauf pour la dernière partie, cette technique est entièrement générale et a été utilisée dans un certain nombre de preuves et d'algorithmes en optimisation combinatoire [LoS1].

4. Théorie de la discordance

Dans la section 3, nous avons commencé avec un problème discret et avons essayé de trouver une bonne approximation continue (relaxation) de celui-ci. Discutons un peu d'un problème inverse (seulement dans la mesure de quelques exemples classiques). Supposons que nous ayons un ensemble muni d'une mesure ; comment peut-il être bien approximé par une mesure discrète ?

Dans notre premier exemple, nous considérons la mesure de Lebesgue λ sur le carré unité $[0, 1]^2$. Supposons que nous voulions approximer celle-ci par la mesure uniforme sur un sous-ensemble fini T . Bien sûr, nous devons spécifier comment l'erreur d'approximation est mesurée : ici, nous la définissons comme l'erreur maximale sur les rectangles parallèles aux axes :

$$\Delta(T) = \sup_R ||T \cap R| - \lambda(R)|T||,$$

où R parcourt tous les ensembles de la forme $R = [a, b] \times [c, d]$ (nous avons mis à l'échelle par $|T|$ pour plus de commodité). Nous cherchons à trouver le meilleur ensemble T , et à déterminer la "discordance"

$$\Delta_n = \inf_{|T|=n} \Delta(T).$$

La question peut être posée pour toute famille d'"ensembles tests" au lieu de rectangles : cercles, ellipses, triangles, etc. La réponse dépend de manière complexe de la structure géométrique de la famille de tests ; voir [BeC] pour un exposé de ces beaux résultats.

La question peut également être posée dans d'autres dimensions. Le cas unidimensionnel est facile, car le choix évident de $T = \{1/(n+1), \dots, n/(n+1)\}$ est optimal. Mais pour des dimensions supérieures à 2, l'ordre de grandeur de Δ_n n'est pas connu.

Revenant à la dimension 2, le premier choix évident est d'essayer une grille $\sqrt{n} \times \sqrt{n}$ pour T . Cela laisse de côté un rectangle de taille d'environ $1/\sqrt{n}$, et a donc $\Delta(T) \approx \sqrt{n}$. Il existe de nombreuses

constructions qui font mieux ; la plus élémentaire est la suivante. Soit $n = 2^k$. Prendre tous les points (x, y) où x et y sont des multiples de 2^{-k} , et en les développant en k bits en binaire, on obtient ces bits dans l'ordre inverse : $x = 0.b_1b_2 \dots b_k$ et $y = 0.b_nb_{n-1} \dots b_1$. C'est un exercice agréable de vérifier que cet ensemble a une discordance $\Delta(T) \approx k = \log_2 n$.

Il est difficile de prouver que l'on ne peut pas faire mieux ; même prouver que $\Delta_n \rightarrow \infty$ était difficile [A]. La borne inférieure correspondant à la construction ci-dessus a finalement été prouvée par Schmidt [S] :

$$\Delta_n = \Theta(\log n).$$

Ce résultat fondamental a de nombreuses applications.

Notre second exemple peut être introduit par une motivation statistique. Supposons que l'on nous donne une séquence 0 – 1 $x = x_1x_2 \dots x_n$ de longueur n . Nous voulons tester si elle provient de lancers de pièces indépendants.

Une approche (liée à la proposition de von Mises pour la définition d'une séquence aléatoire) pourrait être de compter les 0 et les 1 ; leurs nombres devraient être à peu près les mêmes. Cela devrait également être vrai si nous comptons les bits dans les positions paires, ou impaires, ou plus généralement, dans toute progression arithmétique fixée d'indices.

Nous considérons donc la quantité

$$\Delta(x) = \max_A \left| \sum_{i \in A} x_i - \frac{1}{2}|A| \right|,$$

où A parcourt toutes les progressions arithmétiques dans $\{1, \dots, n\}$. La théorie élémentaire des probabilités nous dit que si x est généré par des lancers de pièces indépendants, alors $\Delta(x) \approx \sqrt{n}/2$ avec une probabilité élevée. Donc, si $\Delta(x)$ est plus grand que cela (ce qui est le cas de la plupart des séquences non aléatoires auxquelles on pense), alors nous pouvons conclure qu'elle n'est pas générée par des lancers de pièces indépendants.

Mais x peut-elle échouer à ce test dans l'autre direction, en étant “trop lisse” ? En d'autres termes, qu'est-ce que

$$\Delta_n = \min_x \Delta(x),$$

où x parcourt toutes les séquences 0 – 1 de longueur n ? Nous pouvons considérer la question comme un problème d'approximation de mesure : on nous donne la mesure sur $\{1, \dots, n\}$ dans laquelle chaque atome a une mesure $1/2$. C'est une mesure discrète mais pas à valeurs entières ; nous voulons l'approximer par une mesure à valeurs entières. La famille de tests définissant l'erreur d'approximation est constituée de progressions arithmétiques.

Il résulte de la considération des séquences aléatoires que $\Delta_n = O(\sqrt{n})$. Le premier résultat dans la direction opposée a été prouvé par Roth [Rot], qui a montré que $\Delta_n = \Omega(n^{1/4})$. On s'attendait à ce que la borne soit améliorée à $\sqrt{n}$ éventuellement, montrant que les séquences aléatoires sont les extrêmes, du moins en ordre de grandeur (dans de nombreux cas similaires de problèmes extrémaux combinatoires, par exemple en théorie de Ramsey, le choix aléatoire est le meilleur). Mais de manière

surprenante, Sarkozy a montré que $\Delta_n = O(n^{1/3})$, ce qui a été amélioré par Beck [Be] à la valeur presque optimale $\Delta_n = O(n^{1/4} \log n)$, et même le facteur logarithmique a été récemment supprimé par Matousek et Spencer [MS]. Ainsi, il existe des séquences qui simulent trop bien les séquences aléatoires !

5. Le laplacien

Le Laplacien, tel que nous l'apprenons à l'école, est l'opérateur différentiel

$$\sum_i \frac{\partial^2}{\partial x_i^2}.$$

Quoi de plus lié à la continuité des espaces euclidiens, ou du moins des variétés différentielles, qu'un opérateur différentiel ? Dans cette section, nous montrons que le Laplacien a un sens en théorie des graphes, et qu'il est en fait un outil de base. De plus, l'étude des versions discrètes et continues interagit de diverses manières, de sorte que l'utilisation de l'une ou l'autre est presque une question de commodité dans certains cas.

5.1. Marches aléatoires

L'un des problèmes algorithmiques fondamentaux généraux est l'échantillonnage : générer un élément aléatoire à partir d'une distribution donnée sur un ensemble V . Dans les cas non triviaux, l'ensemble V est soit infini, soit fini mais très grand, et souvent seulement décrit implicitement. Dans la plupart des cas, nous voulons générer un élément à partir de la distribution uniforme, et ce cas particulier sera suffisant pour nos discussions.

J'utiliserai comme exemples deux tâches d'échantillonnage :

- (a) Étant donné un corps convexe K dans $\mathbf{R}^n$, générer un point aléatoire uniformément distribué de K .
- (b) Étant donné un graphe G , générer un couplage parfait uniformément distribué dans G .

Le problème (a) se pose dans de nombreuses applications (calcul de volume, intégration de Monte-Carlo, optimisation, etc.). Le problème (b) est lié au modèle d'Ising en mécanique statistique.

Dans des cas simples, on peut trouver des méthodes spéciales élégantes pour l'échantillonnage. Par exemple, si nous avons besoin d'un point aléatoire uniformément distribué dans la boule unité de $\mathbf{R}^n$, alors nous pouvons générer n coordonnées indépendantes à partir de la distribution normale standard, et "normaliser" le vecteur obtenu de manière appropriée.

Mais en général, l'échantillonnage à partir d'un grand ensemble fini, ou d'un ensemble infini, est un problème difficile, et la seule approche générale connue est l'utilisation de chaînes de Markov. Nous définissons (en utilisant la structure de S) une chaîne de Markov ergodique dont l'espace d'états est S et dont la distribution stationnaire est uniforme. Dans le cas d'un ensemble fini S , il peut être plus facile de penser à cela comme un graphe régulier non biparti connexe G avec un ensemble de nœuds V . En partant d'un nœud arbitraire (état), nous effectuons une marche aléatoire sur le

graphe : à partir d'un nœud i , nous passons à un nœud suivant sélectionné uniformément parmi l'ensemble des voisins de i . Après un nombre T d'étapes suffisamment grand, nous nous arrêtons : la distribution du dernier nœud est approximativement uniforme.

Dans le cas d'un corps convexe, une chaîne de Markov assez naturelle à considérer est la suivante : en partant d'un point pratique dans K , nous nous déplaçons à chaque étape vers un point sélectionné uniformément dans une boule de rayon δ autour du point actuel (si le nouveau point est en dehors de K , nous restons où nous étions, et considérons l'étape comme "gaspillée"). La taille du pas δ sera choisie de manière appropriée, mais typiquement elle est d'environ $1/\sqrt{n}$.

Si nous voulons échantillonner à partir de couplages parfaits dans un graphe G , la marche aléatoire à considérer n'est pas si évidente. Jerrum et Sinclair considèrent une marche aléatoire sur les couplages parfaits et presque parfaits (couplages qui laissent seulement deux nœuds non couplés). Si nous générons un tel couplage, il sera parfait avec une probabilité petite mais non négligeable ($1/n^{\text{const}}$ si le graphe est dense). Nous itérons donc jusqu'à obtenir un couplage parfait.

Une étape de la marche aléatoire est générée en choisissant une arête aléatoire e de G . Si le couplage actuel est parfait et contient e , alors nous supprimons e ; si le couplage actuel est presque parfait et que e relie les deux nœuds non couplés, nous l'ajoutons ; si e relie un nœud non couplé à un nœud couplé, nous ajoutons e au couplage et supprimons l'arête qu'il intersecte.

Une autre façon de voir cela est de faire une marche aléatoire sur un "grand" graphe H dont les nœuds sont les couplages parfaits et presque parfaits du "petit" graphe G . Deux nœuds de H sont adjacents si les couplages correspondants proviennent l'un de l'autre en changeant (supprimant, ajoutant ou remplaçant) une seule arête. Des boucles sont ajoutées pour rendre H régulier. Notons que nous ne voulons pas construire H explicitement ; sa taille peut être exponentiellement grande par rapport à G . L'essentiel est que la marche aléatoire sur H peut être implémentée en utilisant seulement cette définition implicite.

Il n'est généralement pas difficile de faire en sorte que la distribution stationnaire de la chaîne soit uniforme, et que la chaîne soit ergodique. La question cruciale est : quel est le temps de mélange, c'est-à-dire combien de temps devons-nous marcher ? Cette question conduit à l'estimation du temps de mélange des chaînes de Markov (le nombre d'étapes avant que la chaîne ne devienne essentiellement stationnaire). Du point de vue des applications pratiques, il est naturel de considérer des chaînes de Markov finies – un calcul dans un ordinateur est nécessairement fini. Mais dans l'analyse, cela dépend de l'application particulière si l'on préfère utiliser un espace d'états fini ou mesurable général. Toutes les connexions essentielles (et très intéressantes) qui ont été découvertes sont valables dans les deux modèles. En fait, la question mathématique générale est la dispersion : nous pourrions nous intéresser à la dispersion de la chaleur dans un matériau, ou à la dispersion de la probabilité lors d'une marche aléatoire, ou à de nombreuses autres questions connexes.

Une chaîne de Markov peut être décrite par sa matrice de transition $M = (p_{ij})$ où p_{ij} est la probabilité de passer à j , étant donné que nous sommes en i . Dans le cas particulier des marches aléatoires sur un graphe non orienté régulier, nous avons $M = (1/d)A$, où A est la matrice d'adjacence habituelle du graphe G . Notons que le vecteur $\mathbf{1}$ est un vecteur propre de M appartenant à la valeur

propre 1.

Si σ est la distribution de départ (qui peut être vue comme un vecteur dans $\mathbf{R}^V$), alors la distribution après t étapes est $M^t\sigma$. De là, il est facile de voir que la vitesse de convergence vers la distribution stationnaire (qui est le vecteur propre $(1/n)\mathbf{1}$) dépend de la différence entre la plus grande valeur propre 1 et la deuxième plus grande valeur propre λ (le trou spectral). Nous pourrions aussi définir le trou spectral comme la plus petite valeur propre positive de la matrice $L = M - I$.

Nous appelons cette matrice L le Laplacien du graphe. Pour tout vecteur $f \in \mathbf{R}^V$, la valeur de Lf au nœud i est la moyenne de f sur les voisins de i , moins f_i . Cette propriété montre que le Laplacien est en effet un analogue discret de l'opérateur de Laplace classique.

Certaines propriétés de l'opérateur de Laplace dépendent de la structure fine des variétés différentiables, mais de nombreuses propriétés de base peuvent être généralisées facilement à tout graphe. On peut définir des fonctions harmoniques, le noyau de la chaleur, prouver des identités comme la formule analogue de Green :

$$\sum_i (f_i(Lg)_i - g_i(Lf)_i) = 0$$

et ainsi de suite. De nombreuses propriétés de l'opérateur de Laplace "continu" peuvent être dérivées en utilisant de telles formules combinatoires faciles sur une grille, puis en passant à la limite. Voir Chung [Chu] pour un exposé de certaines de ces connexions.

Mais plus significativement, le Laplacien discret est un outil important dans l'étude de diverses propriétés de la théorie des graphes. Pour illustrer ce fait, revenons aux marches aléatoires.

Souvent, des informations sur le trou spectral sont difficiles à obtenir (c'est le cas dans nos deux exemples d'introduction). Un remède possible est de relier la vitesse de dispersion aux inégalités isopérimétriques dans l'espace d'états. Pour être plus précis, définissons la conductance de la chaîne par

$$\Phi = \max_{\emptyset \subset S \subset V} \frac{\sum_{i \in S, j \in V \setminus S} \pi_i p_{ij}}{\pi(S)\pi(V \setminus S)}.$$

(Le numérateur peut être vu comme la probabilité que, en choisissant un nœud aléatoire à partir de la distribution stationnaire, puis en faisant un pas, le premier nœud soit dans S et le second dans $V \setminus S$. Le dénominateur est la même probabilité pour deux nœuds aléatoires indépendants de la distribution stationnaire.) L'inégalité suivante a été prouvée par Jerrum et Sinclair [JS] :

Théorème 3.

$$\frac{\Phi^2}{8} \leq 1 - \lambda \leq \Phi.$$

Cela signifie que le temps de mélange (que nous devons utiliser ici de manière informelle, car la définition exacte dépend de la façon dont nous commençons, de la façon dont nous mesurons la convergence, etc.) est compris entre $1/\Phi$ et $1/\Phi^2$.

Dans le cas des marches aléatoires dans un corps convexe, la conductance est très étroitement liée au théorème isopérimétrique suivant [LoSi], [DyF] :

Théorème 4. *Soit K un corps convexe dans $\mathbf{R}^n$, de diamètre D . Soit la surface F divisant K en deux parties K_1 et K_2 . Alors*

$$\text{vol}_{n-1}(F) \geq \frac{2}{D} \frac{\text{vol}(K_1)\text{vol}(K_2)}{\text{vol}(K)}.$$

(Un long cylindre mince montre que la borne est optimale.)

Du théorème 4, il résulte qu’après un pré-traitement approprié et d’autres difficultés techniques que nous passons sous silence, on peut générer un point d’échantillonnage dans un corps convexe de dimension n en $O^*(n^3)$ étapes.

Dans le cas des couplages, Jerrum et Sinclair prouvent que pour les graphes denses, la conductance de la chaîne décrite est bornée inférieurement par $1/n^{\text{const}}$, et donc le temps de mélange est borné par n^{const} .

Comment établir les inégalités isopérimétriques ? La méthode générique consiste à construire, explicitement ou implicitement, des flots ou des flots multi-produits. Supposons que pour chaque sous-ensemble $\emptyset \subset S \subset V$, nous puissions construire un flot à travers le graphe de sorte que chaque nœud $i \in S$ soit une source qui produit une quantité $\pi_i \pi(V \setminus S)$ de flot, et chaque $j \in V \setminus S$ soit un puits qui consomme une quantité $\pi_j \pi(S)$ de flot. Supposons en outre que chaque arête ij transporte au plus $K \pi_i p_{ij}$ de flot. Alors un calcul simple montre que la conductance satisfait

$$\Phi \geq \frac{1}{K}.$$

Au lieu de construire un flot pour chaque sous-ensemble S de nœuds, on pourrait préférer construire un flot multi-produits, c’est-à-dire un flot de valeur $\pi_i \pi_j$ pour chaque i et j . Si le flot total à travers chaque arête ij est d’au plus $K \pi_i p_{ij}$, alors la conductance est d’au moins $1/K$.

Quelle est la qualité de la borne sur la conductance obtenue par la méthode du flot multi-produits ? Un résultat important de Leighton et Rao [LR] implique que pour le meilleur choix des flots, elle donne un facteur de $O(\log n)$. L’une des raisons pour lesquelles je mentionne ceci est que ce résultat, à son tour, peut être dérivé du théorème de Bourgain [Bo] sur la plongeabilité des espaces métriques finis dans les espaces ℓ_1 avec une petite distorsion. Cf. aussi [LiLR], [DeL].

Construire les “meilleurs” flots ou flots multi-produits est souvent la principale difficulté mathématique. Dans le cas de la preuve du théorème 4, cela peut être fait par une méthode utilisée antérieurement par Payne et Weinberger [PW] dans la théorie des équations aux dérivées partielles. Nous n’en donnons qu’une esquisse grossière.

Supposons que nous voulions construire un “flot” de K_1 à K_2 . Par le théorème du “Sandwich” (Ham-Sandwich), il existe un hyperplan qui coupe à la fois K_1 et K_2 en deux parties de volumes égaux. Nous pouvons construire séparément les flots de part et d’autre de cet hyperplan. En répétant cette procédure, nous pouvons découper K en “aiguilles” de sorte que chaque aiguille soit divisée par la partition $(S, V \setminus S)$ dans la même proportion, et il suffit donc de construire un flot entre K_1 et K_2 à l’intérieur de chaque aiguille. Cela se fait facilement en utilisant le théorème de

Brun-Minkowski.

Pour la marche aléatoire sur les couplages, la construction du flot multi-produits (appelé chemins canoniques par Jerrum et Sinclair) remonte également à l’une des plus anciennes méthodes de la théorie des couplages. Étant donné (disons) un couplage parfait M et un couplage presque parfait M' , nous formons leur union. C’est un sous-graphe du “petit” graphe G qui se décompose en un chemin (avec des arêtes alternant entre M et M'), quelques cycles (encore une fois alternant entre M et M') et les arêtes communes de M et M' . Il est alors facile de transformer M en M' en parcourant chaque cycle et le chemin, et en remplaçant les arêtes une par une. Ceci équivaut à un chemin dans le “grand” graphe H , reliant les nœuds M et M' . Si nous utilisons ce chemin pour transporter le flot, alors on peut montrer qu’aucune arête du graphe H n’est surchargée, à condition que le graphe G soit dense.

5.2. Le théorème de la cage et les applications conformes

Il a été prouvé par Steinitz que tout graphe planaire 3-connexe peut être représenté comme le squelette d’un polytope (tridimensionnel). En fait, il y a beaucoup de liberté dans le choix de la géométrie de ce polytope représentant. Parmi diverses extensions, la plus intéressante pour nous est la construction classique remontant à Koebe [Ko] (prouvée pour la première fois par Andre’ev [An] ; cf. aussi [T]) :

Théorème 5 (Théorème de la Cage). *Soit H un graphe planaire 3-connexe. Alors H peut être représenté comme le squelette d’un polytope tridimensionnel dont toutes les arêtes sont tangentes à la sphère unité.*

Nous pouvons ajouter que le polytope représentant est unique à une transformation projective de l’espace qui préserve la sphère unité près. En considérant l’“horizon” de chaque sommet du polytope, nous obtenons une représentation des nœuds du graphe par des disques circulaires disjoints dans le plan de sorte que les nœuds adjacents correspondent à des cercles tangents.

Le théorème de la Cage peut être considéré comme une forme discrète du théorème de l’application conforme de Riemann, en ce sens qu’il implique le théorème de l’application conforme de Riemann. En effet, supposons que nous voulions construire une application conforme d’un domaine simplement connexe K dans le plan complexe vers le disque unité D . Pour simplifier, supposons que K soit borné. Considérons une grille triangulaire dans le plan (un graphe infini 6-régulier), et soit G le graphe obtenu en identifiant tous les points de la grille à l’extérieur de K en un seul nœud s . Considérons la représentation de Koebe par des cercles tangents sur la sphère ; nous pouvons supposer que le cercle représentant le nœud s est l’extérieur de D . Ainsi, tous les autres cercles sont à l’intérieur de D , et nous obtenons une application de l’ensemble des points de la grille dans K vers D . En laissant la grille devenir arbitrairement fine, il a été montré par Rodin et Sullivan [RS] qu’à la limite nous obtenons une application conforme de K sur D .

Le théorème de la Cage est donc un magnifique pont entre les mathématiques discrètes (théorie des graphes) et les mathématiques continues (analyse complexe). Mais où est le Laplacien ? Nous verrons une connexion dans la section suivante. Une autre connexion se produit dans le travail de Spielmann et Teng [SpT], qui utilisent le théorème de la Cage pour montrer que le trou spectral

du Laplacien d'un graphe planaire est d'au plus $1/n$. Cela implique, entre autres, que le temps de mélange de la marche aléatoire sur un graphe planaire est au moins linéaire en le nombre de nœuds.

5.3. L'invariant de Colin de Verdière

En 1990, Colin de Verdière [Co] a introduit un paramètre $\mu(G)$ pour tout graphe non orienté G . La recherche concernant ce paramètre implique un mélange intéressant d'idées issues de la théorie des graphes, de l'algèbre linéaire et de l'analyse.

La définition exacte dépasse les limites de cet article ; approximativement, $\mu(G)$ est la multiplicité de la plus petite valeur propre positive du Laplacien de G , où les arêtes de G sont pondérées de manière à maximiser cette multiplicité.

Le paramètre a été motivé par l'étude de la multiplicité maximale de la deuxième valeur propre de certains opérateurs différentiels de type Laplacien, définis sur des surfaces de Riemann. Il a approximé la surface par un graphe G suffisamment densément plongé, et a montré que la multiplicité de la deuxième valeur propre de l'opérateur peut être bornée par cette valeur $\mu(G)$ ne dépendant que du graphe.

L'invariant de Colin de Verdière a suscité beaucoup d'intérêt parmi les théoriciens des graphes, en raison de ses propriétés théoriques des graphes étonnamment agréables. Entre autres, il est monotone par mineurs, de sorte que la théorie des mineurs de graphes de Robertson-Seymour s'y applique. De plus, la planarité des graphes peut être caractérisée par cet invariant : $\mu(G) \leq 3$ si et seulement si G est planaire.

La preuve originale par Colin de Verdière de la partie “si” de ce fait était très inhabituelle en théorie des graphes : essentiellement, en inversant la procédure ci-dessus, il a montré comment reconstruire une sphère et un opérateur différentiel partiel elliptique positif P sur celle-ci de sorte que $\mu(G)$ soit borné par la dimension du noyau de P , et a ensuite invoqué un théorème de Cheng [C] affirmant que cette dimension est au plus 3.

Plus tard, van der Holst [Ho] a trouvé une preuve combinatoire de ce fait. Bien que cela puisse sembler un pas en arrière (après tout, cela a éliminé la nécessité de la seule application d'équations aux dérivées partielles en théorie des graphes que je connaisse), cela a ouvert la possibilité de caractériser le cas suivant. Vérifiant une conjecture de Robertson, Seymour et Thomas, il a été montré par Lovász et Schrijver [LoS2] que $\mu(G) \leq 4$ si et seulement si G est plongeable sans entrelacs dans $\mathbf{R}^3$.

Peut-on revenir à la motivation originale de $\mu(G)$ et trouver une version “continue” de ce résultat ? En quel sens un graphe plongeable sans entrelacs approxime-t-il l'espace ? Ces questions semblent très difficiles.

Il s'avère que les graphes avec de grandes valeurs de μ sont également assez intéressants. Par exemple, pour un graphe G sur n nœuds avec $\mu(G) \geq n - 4$, le complément de G est planaire, à l'introduction de points “jumeaux” près ; et la réciproque de cette affirmation est également vraie

sous des conditions raisonnablement générales. La preuve de ce dernier fait utilise la représentation de Koebe-Andre'ev des graphes.

Ainsi, l'invariant de graphe $\mu(G)$ est lié, à une extrémité de l'échelle, aux équations aux dérivées partielles elliptiques (théorème de Cheng) ; à l'autre, au théorème de Riemann sur les applications conformes.

Références

- [A] T. van Aardenne-Ehrenfest, On the impossibility of a just distribution, *Nederl. Akad. Wetensch. Proc.* 52 (1949), 734-739 ; *Indagationes Math.* 11 (1949) 264-269.
- [Al] F. Alizadeh, Interior point methods in semidefinite programming with applications to combinatorial optimization, *SIAM J. on Optimization* 5 (1995), 13-51.
- [AloS] N. Alon, J. Spencer, *The Probabilistic Method*. With an appendix by Paul Erdős, Wiley, New York, 1992.
- [An] E. Andre'ev, On convex polyhedra in Lobachevsky spaces, *Mat. Sbornik, Nov. Ser.* 81 (1970), 445-478.
- [ArLMSS] S. Arora, C. Lund, R. Motwani, M. Sudan, M. Szegedy, Proof verification and hardness of approximation problems, *Proc. 33rd FOCS* (1992), 14-23.
- [B] I. Barany, A short proof of Kneser's conjecture, *J. Combin. Theory A* 25 (1978), 325-326.
- [Be] J. Beck, Roth's estimate of the discrepancy of integer sequences is nearly sharp, *Combinatorica* 1 (1981), 319-325.
- [BeC] J. Beck, W. Chen, *Irregularities of Distribution*, Cambridge Univ. Press (1987).
- [BeS] J. Beck, V.T. Sos, Discrepancy Theory, Chapter 26, in "Handbook of Combinatorics" (R.L. Graham, M. Grottschel, L. Lovász, eds.), North-Holland, Amsterdam (1995).
- [Bj] A. Bjorner, Topological methods, in "Handbook of Combinatorics" (R.L. Graham, L. Lovász, M. Grottschel, eds.), Elsevier, Amsterdam, 1995, 1819-1872.
- [Bo] J. Bourgain, On Lipschitz embedding of finite metric spaces in Hilbert space, *Isr. J. Math.* 52 (1985), 46-52.
- [C] S. Y. Cheng, Eigenfunctions and nodal sets, *Commentarii Mathematici Helvetici* 51 (1976), 43-55.
- [Ch] A.J. Chorin, *Vorticity and Turbulence*, Springer, New York, 1994.
- [Co] Y. Colin de Verdière, Sur un nouvel invariant des graphes et un critere de planarite, *Journal of Combinatorial Theory Series B* 50 (1990), 11-21.
- [Chu] F. R. K. Chung, *Spectral Graph Theory*, CBMS Reg. Conf. Series 92, Amer. Math. Soc., 1997.
- [DP] C. Delorme, S. Poljak, Combinatorial properties and the complexity of max-cut approximations, *Europ. J. Combin.* 14 (1993), 313-333.
- [DeL] M. Deza, M. Laurent, *Geometry of Cuts and Metrics*, Springer Verlag, 1997.
- [Du] R.M. Dudley, Distances of probability measures and random variables, *Ann. Math. Stat.* 39 (1968), 1563-1572.
- [DyF] M. Dyer, A. Frieze, Computing the volume of convex bodies : a case where randomness provably helps, in "Probabilistic Combinatorics and Its Applications" (Bela Bollobas, ed.), *Proceedings of Symposia in Applied Mathematics*, Vol. 44 (1992), 123-170.
- [DyFK] M. Dyer, A. Frieze, R. Kannan, A random polynomial time algorithm for approximating the volume of convex bodies, *Journal of the ACM* 38 (1991), 1-17.
- [GW] M.X. Goemans, D.P. Williamson, .878-Approximation algorithms for MAX CUT and MAX 2SAT, *Proc. 26th ACM Symp. on Theory of Computing* (1994), 422-431.
- [GrLS] M. Grottschel, L. Lovász, A. Schrijver, *Geometric Algorithms and Combinatorial Optimization*, Springer, 1988.
- [H] J. Hastad, Some optimal in-approximability results, *Proc. 29th ACM Symp. on Theory of Comp.* (1997), 1-10.
- [Ho] H. van der Holst, A short proof of the planarity characterization of Colin de Verdiere, *Journal of Combinatorial Theory, Series B* 65 (1995), 269-272.

- [JS] M. Jerrum, A. Sinclair, Approximating the permanent, *SIAM J. Computing* 18 (1989), 1149-1178.
- [K] M. Kneser, Aufgabe No. 360, *Jber. Deutsch. Math. Ver.* 58 (1955).
- [Ko] P. Koebe, Kontaktprobleme der konformen Abbildung, *Berichte über die Verhandlungen d. Sachs. Akad. d. Wiss., Math.-Phys. Klasse* 88 (1936), 141-164.
- [LR] F.T. Leighton, S. Rao, An approximate max-flow min-cut theorem for uniform multicommodity flow problems with applications to approximation algorithms, *Proc. 29th Annual Symp. on Found. of Computer Science, IEEE Computer Soc.* (1988), 422-431.
- [LiLR] N. Linial, E. London, Y. Rabinovich, The geometry of graphs and some of its algebraic applications, *Combinatorica* 15 (1995), 215-245.
- [Lo1] L. Lovász, Kneser's conjecture, chromatic number, and homotopy, *J. Comb. Theory A* 25 (1978), 319-324.
- [Lo2] L. Lovász, On the Shannon capacity of graphs, *IEEE Trans. Inform. Theory* 25 (1979), 1-7.
- [LoM] L. Lovász, P. Major, A note on a paper of Dudley, *Studia Sci. Math. Hung.* 8 (1973), 151-152.
- [LoP] L. Lovász, M.D. Plummer, *Matching Theory*, Akademiai Kiadó - North Holland, Budapest, 1986.
- [LoS1] L. Lovász, A. Schrijver, Cones of matrices and set-functions, and 0-1 optimization, *SIAM J. on Optimization* 1 (1990), 166-190.
- [LoS2] L. Lovász, A. Schrijver, A Borsuk theorem for antipodal links and a spectral characterization of linklessly embeddable graphs, *Proceedings of the AMS* 126 (1998), 1275-1285.
- [LoSi] L. Lovász, M. Simonovits, Random walks in a convex body and an improved volume algorithm, *Random Structures and Alg.* 4 (1993), 359-412.
- [M] J. Matoušek, *Geometric Discrepancy*, Springer Verlag, 1999.
- [MS] J. Matoušek, J. Spencer, Discrepancy in arithmetic progressions, *J. Amer. Math. Soc.* 9 (1996), 195-204.
- [O] D. Ornstein, Bernoulli shifts with the same entropy are isomorphic, *Advances in Math.* 4 (1970), 337-352.
- [PW] L. E. Payne, H. F. Weinberger, An optimal Poincaré inequality for convex domains, *Arch. Rat. mech. Anal.* 5 (1960), 286-292.
- [RS] B. Rodin, D. Sullivan, The convergence of circle packings to the Riemann mapping, *J. Differential Geom.* 26 (1987), 349-360.
- [RoH] G.-C. Rota, L. H. Harper, Matching theory, an introduction, in "Advances in Probability Theory and Related Topics" (P. Ney, ed.) Vol I, Marcel Dekker, New York (1971) 169-215.
- [Rot] K. F. Roth, Remark concerning integer sequences, *Acta Arith.* 9 (1964), 257-260.
- [S] W. Schmidt, Irregularities of distribution, *Quart. J. Math. Oxford* 19 (1968), 181-191.
- [SpT] D. A. Spielman, S.-H. Teng, Spectral partitioning works : planar graphs and finite element meshes, *Proc. 37th Ann. Symp. Found. of Comp. Sci., IEEE* (1996), 96-105.
- [St] V. Strassen, The existence of probability measures with given marginals, *Ann. Math. Statist* 36 (1965), 423-439.
- [T] W. Thurston, *The Geometry and Topology of Three-manifolds*, Princeton Lecture Notes, Chapter 13, Princeton, 1985.
- [W] H. Wolkowitz, Some applications of optimization in matrix theory, *Linear Algebra and its Applications* 40 (1981), 101-118.

László Lovász,
 Microsoft Research, One Microsoft Way, Redmond, WA 98052, USA
 lovasz@microsoft.com