



Étude de la diffusion des processus déterministes et faiblement aléatoires en environnement aléatoire

Yann Chiffaudel

► To cite this version:

Yann Chiffaudel. Étude de la diffusion des processus déterministes et faiblement aléatoires en environnement aléatoire. Physique mathématique [math-ph]. Université de Paris, 2019. Français. NNT : . tel-02397594

HAL Id: tel-02397594

<https://theses.hal.science/tel-02397594>

Submitted on 6 Dec 2019

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Université de Paris
Ecole doctorale : Sciences Mathématiques de
Paris Centre ED386
Laboratoire : Laboratoire de Probabilités,
Statistique et Modélisation

Étude de la diffusion des processus déterministes et faiblement aléatoires en environnement aléatoire

Par Yann Chiffaudel

Thèse de doctorat de mathématiques appliquées

Dirigée par Raphaël Lefevre

Présentée et soutenue publiquement le 22/10/2019

Devant un jury composé de :

Daniel Ueltschi, Professor, University of Warwick, Rapporteur
François Huveneers, Maître de conférence, Université Paris-Dauphine, Rapporteur
Cristina Toninelli, Directrice de Recherche, Université Paris-Dauphine
Giambattista Giacomini, Professeur, Université Paris Diderot
Raphaël Lefevre, Maître de Conférence, Université Paris Diderot, Directeur de thèse

Déposée le 19/09/2019



Except where otherwise noted, this work is licensed under
<https://creativecommons.org/licenses/by-nc-nd/3.0/fr/>

Titre : Étude de la diffusion des processus déterministes et faiblement aléatoires en environnement aléatoire

Résumé : Cette thèse étudie la diffusion dans le modèle des miroirs, modèle inspiré de la physique et introduit en 1988 par Ruijgrok et Cohen. Ce modèle est déterministe et réversible. Pour traiter ce modèle difficile, initialement défini uniquement en dimension 2, nous l'avons d'abord généralisé pour en faire un modèle en dimension quelconque. De premières études numériques permirent de conjecturer que le modèle est diffusif en dimension supérieure ou égale à 3. Nous avons par la suite exploré une approche perturbative du coefficient de diffusion basée sur la technique de lace expansion développée par Gordon Slade pour l'étude de la marche aléatoire auto-évitante. Face à la difficulté des calculs nous avons légèrement simplifié le modèle en abandonnant la contrainte de réversibilité. Nous avons obtenu ainsi un nouveau modèle que nous nommons le modèle des permutations. Nous avons ensuite transformé ces deux modèles pour en faire des marches aléatoires en milieu aléatoire, et ce via une approche systématique et généraliste. Grâce à ces modifications nous avons pu pousser l'approche perturbative jusqu'à obtenir une approximation satisfaisante de la valeur du coefficient de diffusion dans le modèle des permutations. Le résultat principal est l'existence d'une série dont tous les termes sont bien définis et dont les premiers termes fournissent l'approximation voulue. La convergence de cette série reste un problème ouvert. Les résultats analytiques sont appuyés par une approche numérique de ces modèles, ce qui permet de voir que la lace expansion donne des résultats de qualité. De nombreuses questions restent ouvertes, notamment le calcul des termes suivants du développement perturbatif et la généralisation de cette approche au modèle des miroirs, ce qui ne saurait poser problème, puis à une classe plus large de modèles.

Mots clefs : Diffusion, Loi de Fick, Émergence, Gaz de Lorentz, Modèle des miroirs, Lace expansion, Approche perturbative, Coefficient de diffusion, Diagrammes, Marche aléatoire en milieu aléatoire.

Title : Study of diffusion of weekly random process in random environment

Abstract : This thesis studies the diffusion in the mirrors model, a physics-based model introduced in 1988 by Ruijgrok and Cohen. This model is deterministic and reversible. To treat this difficult model, initially defined only in dimension 2, we first generalized it to a model valid in any dimension. Initial numerical studies suggested that the model is diffusive in dimensions greater than or equal to 3. We then explored a perturbative diffusion coefficient approach based on the lace expansion technique developed by Gordon Slade for the study of self-avoiding random walk. Faced with the difficulty of the calculations, we slightly simplified the model by giving up the reversibility constraint. We thus obtained a new model that we call the permutations model. We then transformed these two models into random walks in random environment using a systematic and general approach. Thanks to these modifications, we were able to push the perturbative approach to obtain a satisfactory approximation of the value of the diffusion coefficient in the permutations model. The main result is the existence of a series in which all terms are well defined and the first terms provide the desired approximation. The convergence of this series remains an open problem. The analytical results are supported by a numerical approach to these models, which shows that the lace expansion gives quality results. Many questions remain open, including the calculation of the following terms of perturbative development and the generalization of this approach to the mirrors model -which should not be a problem- and then to a broader class of models.

Keywords : Diffusion, Fick's law, Emergence, Lorentz gas, Mirror's model, Lace expansion, Perturbative Approach, Diffusion coefficient, Diagrams, Random walks in random environments

Remerciements

Mes remerciements vont d'abord naturellement à mon directeur de thèse, Raphaël Lefevre, pour m'avoir accompagné durant presque 5 ans entre le début de mon stage et ma soutenance. Merci Raphaël, pour toutes ces heures passées à parler de maths, de physique et parfois d'autre chose. Merci de m'avoir fait découvrir Shanghai malgré mon aversion à la consommation de kérosène, merci de m'avoir laissé faire des maths à mon rythme et à ma manière, je sais que ça n'a pas toujours été facile, mon entourage sait que je suis compliqué. Alors merci, et bonne chance pour la poursuite des travaux, on sait bien que le sujet est loin d'être clos.

Sur la fin de ma thèse, j'ai eu la chance de recevoir une aide extrêmement précieuse de la part de Bastien Fernandez, au moment où j'étais perdu et découragé. Merci Bastien, tu as su faire preuve de grandes qualités humaines et ton aide a vraiment fait la différence, merci pour tout !

Merci aussi à mes rapporteurs, Daniel Ueltschi et François Huveneers, qui ont lu et évalué mon manuscrit avec patience et application malgré les calculs longs et difficiles qu'ils y ont trouvé.

Merci à Cristina Toninelli et Giambattista Giacomini d'avoir pris le temps de participer à mon jury, on sait tous combien le temps est précieux pour les chercheurs.

Merci également à Vivien Lecomte pour m'avoir dit "tu devrais contacter Raphaël Lefevre", quand je lui ai demandé s'il connaissait des domaines de recherche à l'interface math-physique. Sans lui rien de tout ça n'aurait été possible.

On ne soulignera jamais assez l'importance d'une bonne équipe administrative dans la vie d'un laboratoire, et le LPSM a cette chance. Merci donc à eux, et en particulier à Nathalie Bergame, Valérie Juvé, Florence Deschamps et Francis Comets, pour leur travail, leur compétence mais aussi pour leur gentillesse et leur bonne humeur ! Merci aussi à Amina Hariti pour avoir grandement facilité mes interactions avec l'école doctorale.

Enfin merci à ceux, familles et amis, mathématiciens ou non qui m'ont personnellement et spécifiquement aidé dans ce travail de thèse. Notamment Laura Green, qui m'a appris des techniques d'auto-manipulation (on dit "productivité" en fait) qui ont été indispensables pour terminer ce travail.

Merci à Tristan Roth pour m'avoir invité à travailler ensemble un dimanche et en présence de qui j'ai écrit ce jour-là la dernière ligne de la preuve qui est le centre de ce travail.

Et merci infiniment à mes relectrices et relecteurs, Isabelle Thomas-Chiffaudel, Adrien Chiffaudel, Eugénie Marescaux et Quentin Didier, vous êtes géniaux !

Une thèse n'est pas juste un travail, c'est une tranche de vie, où à tout instant on a une partie du cerveau qui songe à la recherche. Les personnes qui ont partagé ma vie au labo et en dehors du labo ont donc tous participé à leur manière, et je souhaiterais tous les remercier, même si cela risque d'être terriblement long, d'avance pardon à ceux que je vais oublier.

Merci à mes collègues du LPSM, qui ont embelli mon quotidien pendant toutes ces années : Laure Maréché, pour ta bonne humeur et ta pédagogie, les mathématiques ont besoin de plus de gens comme toi, Clément Cosco pour avoir été le ciment de la bonne entente entre les thésards, Côme Huré pour ton amitié sincère, Thomas Galtier pour les discussions d'écologie, Simon Coste pour les échanges enrichissants, Julien-Pierra Vest pour avoir partagé mon sujet de thèse pendant quelques mois, travailler avec toi a été un plaisir, Romain Mismar, Benjamin Havret, Enzo Miller, Guillaume Conchon-Kerjan, Paul Melotti, Marie Théret, Noufel Frikha, Svetlana Gribkova, Arturo Leos, Verónica Miró Pina pour Saint Flour et Maps.me, Assaf Shapira pour l'arbre AVL (Rien que ça !), Fabio Coppini, dankon pro via amikeco, Michel Pain, Olga Lopusanschi, Clément Bonvoisin, Quentin Didier, Thomas Bourany, Sylvain Wolf, Lucas Benigni, Marc Pegon (φ), Barbara Dembin, Ziad Kobeissi, Hiroshi (little brother), Cyril Benezet, Xiaoli Wei, Tingting et Shanqiu pour les repas en Chinois où j'étais heureux de ne rien comprendre, Houzi Li, Mi-song Dupuy pour m'avoir donné les bases du foot, Anna Ben-hamou, Lorick Huang, Vu Lan Nguyen, et tant d'autres !

Merci également à tous mes enseignants qui m'ont formé tout au long de mes études, jamais je ne pourrai tous les remercier mais être dans cette liste signifie avoir des qualités pédagogiques

exceptionnelles, toutes mes excuses à ceux que je vais inévitablement oublier et merci du fond du cœur à Régis Boulter, Paul Lecabellec, Éric Depagnat, Jacques Taillet, Dominique Durin, Sacha Benneto, Olivier Fouquet, David Papoulat, Aurélien Galateau, Patrice Hello, Nicolas Pavloff, Laurent Rozas, Gilles Abramovici, Anuradha Jagannathan, Jean-Luc Raimbault, Claudie Mory, Dominique Hulin.

Merci aussi à mes étudiants, c'était un plaisir d'enseigner les mathématiques et c'est quelque chose qui me manquera certainement en quittant le monde universitaire, merci notamment à Jeanne Alkala, Elisa Chardon-Legrand, Blanche Bouchard, Quentin Dadvisard et Juste Leblanc.

Koran dankon al la "mardaj esperantistoj" kiuj multege pligrandigis mian feliĉon dum jaroj !
Merci du fond du cœur aux "mardaj esperantistoj" qui ont énormément augmenté mon bonheur pendant des années !

Alizé Ville, Thurian Lefort, Florine Authier, Alice Andrès, Valentin Melot, Lucas Teyssier, Louis Petitcolas, Quentin Didier, Louise Verkin, Chloé Dawib, Louis Noizet (Ludo).

Merci aux altruistes efficace qui ont augmenté le bonheur de l'humanité et qui ont aussi été des amis exceptionnels : Antonin Broi, Olivier Bertrand, Laura Green, Caroline Jeanmaire, Attilio Lanza, Guillaume Corlouer, Guillaume Vorreux, Fabio Coppini, Raphaël Pesah, Quentin Didier (encore lui !), Louise Verkin, Paul Louyot, Lê Nguyen Hoang (joyeux anniversaire !), Kelly Floch, Tristan Roth...

Lucile Goubayon et Natalia Zambrana Prado mes adorables colocataires.

Merci à tout les amis qui ont été présents, merci entre autre à Suzanne Gruca, Émeline Lhoumaud, Pierre Bienvenu, Steven De Olivera, Claude Sanz, Nathanaëlle Ullmo, Carole Harry, Flore Pineau, Anaël Pineau, Alexandre Barzyk, Raphaël Barzyk, Vivien Rieu, Isaac Laurenty, Emily Dieu, Jordan Bouaziz, Cécile Boutaud, Baptiste Rongier, Marion Cottin, Florent Cottin, Emmanuelle Claeys, Jeanne Nguyen, Matthieu Kieffer, Nicolas Macé, Matthieu Épiard, Marion Herry, Géraldine Favre, Théophile Choutri, Pérégrine, Sébastien Carrassou, Charlotte Barbier, André Müller, Ariane Marcon, Nicolas Macé, Margaux Hamelin, Alexis Galen, Thomas Guibentif...

Emmanuel Walter, pour avoir illuminé ma vie pendant plus de 20 ans.

Un grand merci à mes parents Isabelle Thomas-Chiffaudel et Arnaud Chiffaudel, pour m'avoir conçu. Il semblerait qu'ils m'aient élevé avec enthousiasme en plus. Merci à Adrien Chiffaudel et Stéphanie Chiffaudel-Cenciai, David Cenciai, Gabriel Cenciai, Aurèle Cenciai, Delphine Meiling Vika Chiffaudel, Pascal Chiffaudel, Sophie Chiffaudel-Lu Jérôme Chiffaudel. Elisabeth Thomas, partie trop tôt, parfois la vie familiale et le stress de la vie professionnelle sont incompatibles, je regretterai sûrement longtemps de ne pas avoir pu partager tes derniers mois dans ce monde. Merci à Jean-Pierre Thomas et à Claudine Chiffaudel d'être encore là.

Merci à ceux qui ont beaucoup donné de leur personne pour préparer un excellent pot de thèse, principalement Jean-Pierre Thomas, Isabelle Thomas-Chiffaudel, Adrien Chiffaudel et Caroline Cabot. Mais aussi Alice Andrès, Steven De Olivera et Quentin Didier.

Merci à Thomas Guibentif pour avoir traversé la moitié de la Chine pour passer une soirée avec moi à l'autre bout du monde ! Merci aussi à Luis Fredes pour avoir visité Shanghai avec moi qui n'avais jamais voyagé aussi loin.

Et merci infiniment à Lucile Goubayon pour sa confiance et son amitié fraternelle.

Dans les remerciements de thèse, la dernière ligne est souvent dédiée à remercier son amoureux ou son amoureuse, ici ça sera un peu différent, je remercie les personnes qui ont été les plus proches de moi émotionnellement durant ces 4 années chacune de ces personnes est ou a été à mes yeux aussi importante qu'une compagne de vie et il n'est ni possible ni souhaitable d'en privilégier une sur les autres. Donc, dans l'ordre chronologique, merci infiniment à : Maëlle Patyn, Florine Authier, Yana Khusanova, Alice Andrès, Alice Contal, Quentin Didier, Caroline Cabot et Eugénie Marescaux.

Table des matières

0.1	Émergence des lois macroscopiques	5
0.2	Pas un problème unique, de nombreux problèmes.	6
0.3	La loi de Fick	6
0.3.1	Premier axe de simplification : approximation physique	6
0.3.2	Second axe de simplification : créer de nouveaux modèles	7
0.4	Plusieurs définitions de la loi de Fick et de la diffusion.	7
0.5	État de l'art sur la loi de Fick - Motivation	8
0.6	Gaz de Lorentz	9
0.7	Le modèle des miroirs	9
0.8	Marche aléatoire auto-évitante	9
0.9	Travaux effectués	10
1	Le modèle des miroirs dans le tore $\{1, \dots, N\} \times (\mathbb{Z}/N\mathbb{Z})^{d-1}$	11
1.1	Abstract of this chapter	11
1.2	Macroscopic laws as laws of large numbers	11
1.3	The mirrors model	14
2	Nouveaux modèles et approche analytique	22
2.1	Élargissement du problème, marche aléatoire cinématique, environnement aléatoire, marche auto-évitante	22
2.1.1	Le modèle des miroirs : un cas particulier de marche avec vitesse	22
2.1.2	Les marches aléatoires cinématiques	22
2.1.3	Marche aléatoire cinématique en milieu aléatoire	23
2.1.4	(Re-)Définition du modèle des miroirs	24
2.1.5	Un autre modèle naturel : le modèle des permutations	24
2.1.6	Lien entre les modèles étudiés et les marches auto-évitantes	25
2.1.7	Versions faiblement corrélées des modèles, marche faiblement auto-évitante	26
2.2	Approche perturbative du coefficient de diffusion pour le modèle des permutations dans $\mathbb{Z}^d$	27
2.2.1	Étude du déplacement carré moyen	27
2.2.2	Notations matricielles	30
2.2.3	Première approche de la convergence de $K_m(n)$	33
2.2.4	Partitions sans singleton	35
2.2.5	Représentation graphique des partitions sans singleton.	39
2.2.6	Majoration de $\sum_{n=0}^{+\infty} \mathcal{U}_{W,n}$	41
2.2.7	Convergence des graphes connectés	43
2.2.8	Convergence dominée de la série $\mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ pour W connecté.	50
2.2.9	Convergence de la série $\mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ pour une partition W non-connectée.	56
2.2.10	Retour au calcul des κ_m	60
2.3	Valeurs numériques aux petits ordres	62
2.3.1	Calcul de κ_2	63
2.3.2	Calcul de κ_3	67
2.3.3	Résumé des résultats analytiques	71

3	Approche numérique	73
3.1	Simulation des marches aléatoires en milieu aléatoire	73
3.1.1	Algorithme général	73
3.1.2	Structures de données	74
3.1.3	Nombres pseudo-aléatoires	75
3.1.4	Parallélisation	76
3.2	Résultats et comparaisons avec l'approche perturbative	76
3.2.1	Déplacement carré moyen dans le modèle des permutations	76
3.2.2	Modèle des Miroirs	80
3.3	Interprétation heuristique de la limite singulière dans le modèle des miroirs	81

Introduction

0.1 Émergence des lois macroscopiques

L'idée que la matière est constituée de minuscules briques élémentaires existait déjà chez certains philosophes grecs, les atomistes. Le concept est en effet assez séduisant, ainsi pour comprendre toute la complexité du monde qui nous entoure il suffirait de comprendre une poignée de lois élémentaires.

De fait, cette idée est restée purement spéculative jusqu'à la seconde moitié du 19ème siècle. À ce moment, l'apparition de la physique statistique, notamment grâce aux travaux de Ludwig E. Boltzmann, donna corps aux idées des atomistes. En toute rigueur, la théorie cinétique des gaz commence en 1738 avec *Hydrodynamica*, de Daniel Bernoulli [17], mais les idées de Bernoulli ne percèrent pas à son époque.

Mais quand on dit qu'il "suffit" de comprendre le fonctionnement des atomes et des particules pour comprendre la matière on ignore une difficulté majeure : le calcul est affreusement difficile. Et quiconque ayant déjà posé l'équation de Schrödinger pour un système de 10^{23} particules comprend bien ce que "affreusement" signifie ici.

On peut donc dire que la démarche atomiste de compréhension de la matière par la compréhension de ses composants est une démarche de type analyse-synthèse avec une première étape de compréhension détaillée des particules élémentaires et des lois de la physique au niveau microscopique suivie d'une seconde étape de déduction des propriétés de la matière à partir des connaissances acquises à la première étape. Le lecteur intéressé pourra consulter [16] pour une excellente vulgarisation de ces idées.

La première étape d'analyse a été menée avec de belles réussites par les physiciens et les chimistes du 19ème siècle et surtout du 20ème siècle qui ont découvert les molécules, les atomes, l'électron, le noyau atomique, les nucléons, les neutrinos et toutes les particules du modèle standard jusqu'au boson de Higgs.

La démarche de synthèse est par certains côtés bien plus ardue. On notera par exemple que le problème à trois corps en mécanique céleste possède déjà des propriétés chaotiques qui complexifient son étude, et ce sans même parler de mécanique quantique. Et pourtant, l'objectif de la démarche de synthèse est la compréhension du problème à N corps en mécanique quantique avec $N \simeq 10^{23}$. La célèbre anecdote de la découverte des propriétés chaotiques du problème à 3 corps est contée dans de nombreuses sources. J'ai personnellement eu le plaisir de la découvrir dans le livre *Théorème Vivant* de Cédric Villani [23].

C'est au problème à N corps que s'est attaqué Ludwig Boltzmann, à partir des lois de la mécanique classique qui étaient les seules connues à son époque. C'est lui qui a commencé à utiliser la théorie des probabilités qui s'est avérée l'outil idéal pour étudier les grands ensembles de particules, et ce même si ces particules suivent des lois déterministes. C'est ainsi que l'on a nommé cette discipline "Physique statistique", à l'époque où les probabilités et les statistiques n'étaient pas encore deux sciences séparées.

Les premiers travaux en physique statistique étaient encourageants, sous certaines hypothèses on était capable d'expliquer certains phénomènes de la physique macroscopique à partir des lois de la physique microscopique. Cependant, on était très loin d'avoir une compréhension totale et cette quête de la transition micro-macro s'annonçait rude pour les physiciens comme pour les

mathématiciens venus en renfort.

Malgré un siècle de progrès théoriques et la contribution inestimable de l'informatique, qui permet entre autre l'approche par simulation numérique et les preuves assistées par ordinateur, ce problème est loin d'être résolu. En supposant que le mot "résolu" ait un sens ici. Mais nous allons tout de même, avec cette thèse, ajouter une toute petite pierre à l'édifice de la physique statistique construit pour comprendre au mieux la complexité de notre monde issue de la simplicité de ses composants.

0.2 Pas un problème unique, de nombreux problèmes.

Cette thèse porte sur l'explication des lois de la physique macroscopique à partir de celles de la physique microscopique. Il ne s'agit pas du tout d'un problème unique mais de nombreux problèmes différents. Pour chaque loi connue à l'échelle macroscopique on va devoir utiliser des approches mathématiques différentes pour démontrer celle-ci à partir des lois microscopiques. Nous avons donc une collection de nombreux problèmes mathématiques difficiles que nous pouvons étudier en parallèle. Parmi ces problèmes on pourra citer l'équation de Navier-Stokes, évoquée ici [12] qui est d'ailleurs l'un des 7 problèmes du millénaire à 1 000 000 \$ posés par l'Institut de mathématiques Clay. Mais il y a aussi la croissance de l'entropie, la loi d'ohm et la conductivité électrique, la magnétisation et le modèle d'Ising, la loi de Fourier et bien sûr la loi de Fick qui nous occupe dans cette thèse.

0.3 La loi de Fick

Nous allons donc nous concentrer sur la loi de Fick qui est la loi qui régit la diffusion de particules dans un solvant. D'un point de vue phénoménologique, cette diffusion est bien modélisée par la loi de Fick qui s'énonce ainsi :

$$\vec{j} = -\kappa \vec{\nabla} \rho,$$

avec ρ la concentration de particules, $\vec{j}$ le courant de particules et κ le coefficient de diffusion. On précisera cette définition en temps utile.

Notre objectif ultime serait de prouver cette loi à partir des lois connues de la physique atomique. Cependant, force est de constater que dans l'état actuel des connaissances en physique statistique et en mathématiques appliquées, un tel objectif est encore très très inaccessible.

La solution pour aborder ce problème est donc de le simplifier le plus possible, quitte à s'éloigner du problème d'origine. On peut distinguer deux approches pour simplifier le problème, qui peuvent bien sûr être combinées.

0.3.1 Premier axe de simplification : approximation physique

Cette première approche est plutôt celle des physiciens. Elle consiste à poser les équations exactes dans toute leur complexité puis à proposer des approximations pour simplifier les équations le plus possible jusqu'à les rendre solubles analytiquement ou au moins pour avoir des résultats intéressants de simulation numérique. Si les approximations sont bien faites cette méthode est celle qui permet d'obtenir des résultats les plus proches de la réalité expérimentale. De plus, au fur et à mesure que la compréhension du problème avance, on peut essayer de réduire petit à petit le nombre d'approximations pour se rapprocher d'une résolution rigoureuse. C'est typiquement ce genre d'approche qui intéresse à juste titre un ingénieur ou un physicien expérimentateur. Mais l'accumulation de nombreuses approximations parfois difficiles à justifier autrement que par la qualité du résultat final peut être frustrante pour un mathématicien et il peut être très enrichissant de renoncer aux approximations et de se forcer à une approche rigoureuse. Il nous faut alors une autre approche simplificatrice afin de ne pas se noyer dans des équations insolubles en pratique. C'est là qu'intervient la seconde approche.

0.3.2 Second axe de simplification : créer de nouveaux modèles

Cette seconde approche est l'approche mathématique qui consiste à ne faire aucune concession sur la rigueur. Et pour cela on renonce à travailler sur des problèmes concrets de la physique, ce qui est une concession forte. On va donc inventer des problèmes différents, sur des modèles différents, inspirés de la physique mais qui sont plus simples et qui n'ont aucune application directe. Le but étant de créer de nouvelles méthodes mathématiques de résolution de tels problèmes. Le but à terme étant de complexifier progressivement ces modèles afin de les rapprocher de plus en plus de la réalité physique.

On peut voir ces deux approches comme deux approches opposées mais complémentaires. D'un côté travailler sur des équations réelles et améliorer petit à petit la rigueur pour converger vers une résolution rigoureuse des équations réelles. D'un autre côté, travailler parfaitement rigoureusement sur des équations imaginaires puis modifier ces équations pour les rendre de plus en plus proches de la réalité et ainsi converger vers une résolution rigoureuse des équations réelles. C'est cette seconde approche que l'on va utiliser pour étudier la diffusion et la loi de Fick dans cette thèse, on reviendra sur les simplifications utilisées.

0.4 Plusieurs définitions de la loi de Fick et de la diffusion.

Si on cherche dans la littérature scientifique des articles qui parlent de la diffusion, on trouve que plusieurs définitions différentes sont utilisées pour la caractériser. En voici trois principales. On travaille toujours dans des modèles avec des particules qui se déplacent dans un espace métrique. Ces modèles peuvent être fondamentalement aléatoires ou bien déterministes mais en général on aura toujours une dimension aléatoire due à l'incertitude sur les conditions initiales.

La première définition se concentre sur le déplacement carré moyen.

Définition 0.1. *On dira qu'un modèle est diffusif si pour chaque particule, la position $\mathbf{q}_n$ au temps n de la particule est une variable aléatoire respectant le critère suivant :*

$$\exists \kappa \in \mathbb{R}_+^*,$$

$$\mathbb{E} \left[\frac{\|\mathbf{q}_n\|^2}{n} \right] \xrightarrow{n \rightarrow +\infty} \kappa.$$

Pour les définitions 2 et 3 on a besoin de la notion de limite d'échelle diffusive. Lorsqu'un modèle est discret en temps et en espace on peut en trouver une limite continue (si celle-ci existe) en faisant tendre vers zéro la *distance caractéristique* Δx (typiquement le pas du réseau) et le *temps caractéristique* Δt (typiquement le temps entre deux pas), mais il est intéressant de ne pas les faire tendre vers zéro à la même vitesse. Par exemple, dans le cas de la marche aléatoire simple, Δx est le pas du réseau et Δt le temps entre deux sauts. Et si on les fait tendre tous les deux vers zéro en gardant le rapport entre les deux constant alors on converge vers un modèle où la particule est immobile, ce qui n'est pas très intéressant. Il est plus intéressant de prendre la limite suivante appelée *limite d'échelle diffusive*, on prend $\Delta x \rightarrow 0, \Delta t \rightarrow 0$ et $\Delta t^2 / \Delta x \rightarrow cste$, dans ce cas la limite de la marche aléatoire simple est le mouvement Brownien. On va donc utiliser cette limite d'échelle diffusive.

La seconde définition se réfère au mouvement Brownien.

Définition 0.2. *On dira qu'un modèle est diffusif si, dans la limite d'échelle diffusive, chaque particule suit un mouvement Brownien.*

La troisième définition est celle la plus proche de la physique.

Définition 0.3. *On dira qu'un modèle est diffusif si, dans la limite d'échelle diffusive, on peut définir une fonction ρ telle que pour toute position $\mathbf{q}$, $\rho(\mathbf{q})$ représente la densité de particules au point $\mathbf{q}$ et un champ de vecteur $\vec{j}$ tel que pour toute position $\mathbf{q}$, $\vec{j}(\mathbf{q})$ représente le flux de particules au point $\mathbf{q}$ et tels que ρ et $\vec{j}$ respectent la loi de Fick qui s'énonce ainsi : $\exists \kappa \in \mathbb{R}_+^*, \forall \mathbf{q}$*

$$\vec{j}(\mathbf{q}) = -\kappa \vec{\nabla} \rho(\mathbf{q}),$$

On voit bien que ces définitions manquent de précision et devront être bien spécifiées pour chaque modèle étudié. Mais ce sont toujours des variantes de ces définitions qui sont utilisées. Nous allons d'abord travailler avec la troisième définition dans le chapitre 1 dans le but d'être plus proche de la physique. Au niveau du chapitre 2 nous nous restreindrons à la première définition qui simplifie grandement les calculs.

0.5 État de l'art sur la loi de Fick - Motivation

Quelle que soit la définition utilisée pour parler de diffusion, on a déjà évoqué la nécessité de simplifier les modèles pour pouvoir les résoudre analytiquement. Si on part du problème de la diffusion de particules réelles (atomes, molécules) dans un solvant (par exemple, de l'eau), alors il y a plusieurs simplifications classiques pour rendre ce problème abordable mathématiquement. La simplification la plus commune est d'utiliser des modèles de physique classique et non de physique quantique. En effet, les équations de la physique quantique sont très délicates à résoudre et pour traiter un problème complexe comme la diffusion c'est souvent une difficulté de trop. On notera que les travaux de T. Bodineau, I. Gallagher et L. Saint-Raymond [15] se placent dans le cadre de cette simplification-là.

Certains auteurs ont le courage d'étudier des modèles quantiques mais on va plutôt regarder les modèles classiques. Au sein des modèles classiques on distingue 4 axes de simplifications que voici et que nous allons détailler :

- Choisir un modèle déterministe ou intrinsèquement aléatoire.
- Passer du problème à N corps au problème à 1 corps.
- Travailler dans un espace continu ou sur réseau.
- Travailler en temps continu ou en temps discret.

Au regard de ces critères on peut essayer de classer les modèles selon leur complexité, et pour respecter l'idée de partir des modèles les plus simples pour aller vers plus de complexité nous allons présenter le plus simple : la marche aléatoire simple sur $\mathbb{Z}^d$ avec $d \in \mathbb{N}^*$.

En effet, une manière très simple de modéliser un ensemble de particules dans un solvant est de considérer que chaque particule subit des chocs aléatoires indépendants de ceux subis par les autres particules. Cela suppose que les particules n'interagissent pas du tout entre elles, par exemple parce qu'elles sont très diluées. Ainsi chaque particule est indépendante des autres et le problème n'est plus un problème à N corps mais simplement une collection de problèmes à 1 corps indépendants. On peut à ce moment-là considérer arbitrairement que les particules évoluent en temps discret par des sauts sur un réseau discret et obtenir le modèle le plus simple possible de la diffusion. On peut facilement prouver, voir par exemple [19, 18] [Herbert Spohn, Lawler], que ce modèle est diffusif selon les trois définitions de la section précédente. Pour la seconde et la troisième définitions il faut prendre la limite quand le pas du réseau Δx et l'intervalle de temps Δt tendent vers zéro avec $\Delta t^2 / \Delta x \rightarrow cste$ et le nombre de particule N tend vers $+\infty$. Ce modèle fournit une sorte de base sur laquelle les preuves sont aisées et qui peut être complexifiée par la suite. Par exemple on peut faire des marches aléatoires en temps discret sur $\mathbb{R}^d$ ou en temps continu sur $\mathbb{Z}^d$ ou même en temps continu sur $\mathbb{R}^d$. Dans tous ces cas-là le système est diffusif quelle que soit la définition choisie. Le problème se corse lorsque l'on commence à introduire des interactions entre les particules ou du déterminisme. On verra que les modèles étudiés en détail dans cette thèse insisteront sur l'aspect déterministe.

Au niveau des modèles déterministes, le modèle des sphères dures est probablement l'un des meilleurs compromis entre réalisme physique et une relative simplicité mathématique (toute relative). Mais l'un des modèles phares de la physique statistique déterministe est le gaz de Lorentz que nous allons voir en détail. En effet, le modèle des miroirs étudié au chapitre 1 ainsi que le modèle des permutations introduit au chapitre 2 sont des gaz de Lorentz sur réseau.

0.6 Gaz de Lorentz

Le premier Gaz de Lorentz est le modèle du vent dans les arbres ou modèle d'Ehrenfest [14], l'image pédagogique qui sous-tend ce modèle est celle de l'étude du mouvement d'un vent léger dans une forêt [figure]. Dans cette situation, la trajectoire des molécules d'air (ou des "particules de vent") est fortement affectée par la présence des arbres. Mais les arbres restent bien immobiles dans le vent, c'est comme si l'interaction était à sens unique. Il reste à supposer que les particules de vent n'interagissent pas du tout entre elles et on obtient un modèle à la fois simple à étudier analytiquement et numériquement et pourtant étonnamment riche. C'est l'idée du gaz de Lorentz : des particules indépendantes les unes des autres (ici le vent) dont la trajectoire est affectée par des diffuseurs (ici les arbres) immobiles et imperturbables. Ce modèle a été décliné dans de nombreuses versions. On notera que selon la répartition des diffuseurs le modèle peut être diffusif ou non-diffusif et qu'il est difficile de prévoir ce comportement a priori. Par exemple, dans [5] le modèle est diffusif au sens de la définition 0.2 malgré une répartition périodique des diffuseurs. Pour prouver ce résultat, Bunimovich et Sinai s'appuient sur les propriétés chaotiques du billard de Sinai et tournent à leur avantage la périodicité. Dans [2, 3], par contre, c'est le contraire. Le modèle est non-diffusif pour une répartition plus aléatoire des diffuseurs. Le modèle de [5] est très particulier, l'idée de base de cette thèse était de trouver un modèle déterministe plus proche de la physique et de prouver son caractère diffusif. C'est pourquoi nous avons choisi le modèle des miroirs, parfaitement déterministe et de plus réversible, à l'image des lois de la physique microscopique.

0.7 Le modèle des miroirs

Le modèle des miroirs fut introduit en 1988 par Ruijgrok et Cohen [13]. Il s'agit d'un gaz de Lorentz sur réseau en 2 dimensions. L'image est la suivante : des particules se déplacent le long des arêtes de $\mathbb{Z}^2$ et sur les noeuds de $\mathbb{Z}^2$ elles rencontrent des miroirs orientés à 45° qui les font dévier à gauche ou à droite. S'il n'y a pas de miroir alors la particule va tout droit. Une limitation de ce modèle est le fait d'être en 2D, et si on s'attache à l'analogie physique du rebond des particules sur un miroir alors il n'est pas évident de trouver une généralisation en dimensions supérieures. Nous avons d'ailleurs passé un certain temps à imaginer des formes géométriques compliquées de miroirs en 3D ou 4D, sans grand succès. La solution est d'abandonner l'analogie physique et de ne regarder que l'effet du miroir. Cet effet est simplement une déviation instantanée du vecteur vitesse de la particule, qui respecte la notion de réversibilité. Avec cette idée-là en tête on peut facilement généraliser le modèle des miroirs en dimension quelconque. Cette généralisation sera détaillée aux chapitres 1 et 2. On pourra ainsi s'affranchir de la dimension 2 qui est toujours passionnante mais souvent ardue en physique statistique.

0.8 Marche aléatoire auto-évitante

Un point important à signaler est le parallèle entre le modèle des miroirs et la marche aléatoire auto-évitante. Si on regarde attentivement le modèle des miroirs on constate que comme la dynamique est injective on peut calculer l'ensemble des positions passées et futures d'une particule en fonction de sa position présente et de sa vitesse. Et il y a alors deux cas possibles : soit la trajectoire est ouverte et part à l'infini dans le passé comme dans le futur, soit la trajectoire est fermée et c'est une boucle dans laquelle la particule tourne à l'infini. Dans le premier cas chaque arête de la trajectoire ne sera visitée qu'une fois, on a donc affaire à une trajectoire d'une marche aléatoire auto-évitante sur les arêtes de $\mathbb{Z}^d$. Dans le second cas, chaque arête de la boucle n'est utilisée qu'une seule fois par tour. On a donc affaire à une boucle auto-évitante sur les arêtes de $\mathbb{Z}^d$. Dans les deux cas on a une forte proximité avec la marche aléatoire auto-évitante, cela explique pourquoi le modèle des miroirs est ardu à étudier et c'est aussi pour cela que la démarche analytique du chapitre 2 est fortement inspirée de la *lace expansion* de Gordon Slade [8, 9]. On

reviendra sur ce parallèle à la section 2.1.6. Le lecteur intéressé pourra consulter [11] pour une introduction générale au problème de la marche aléatoire auto-évitante.

Sans surprise on peut aussi rapprocher ce travail des travaux sur les marches aléatoires en milieux aléatoires (MAMA). En fait, le modèle des miroirs est une marche déterministe en milieu aléatoire, c'est donc un cas particulier de MAMA, on se référera en particulier aux travaux de Jean Bricmont et Antti Kupiainen [10] et à ceux de Bálint Tóth [7] qui sont particulièrement proches des sujets de cette thèse.

0.9 Travaux effectués

Durant ces 4 années de thèse nous nous sommes attaqués au modèle des miroirs, avec l'idée de prouver que ce modèle était diffusif au sens de la définition 0.3 de la section 0.4. Cet objectif était ambitieux car on peut voir facilement que le modèle des miroirs n'est pas diffusif au sens de la seconde définition. C'était un candidat idéal pour être diffusif au sens physique, c'est à dire respecter la loi de Fick sans pour autant approcher un mouvement Brownien. Cette première idée a débouché sur un article [22], dont le contenu constitue le chapitre 1, et dans lequel nous avons donné deux conditions nécessaires pour que le modèle des miroirs dans un tore de dimension d soit diffusif au sens de la loi de Fick. Nous avons aussi confirmé numériquement que le modèle des miroirs sur le tore de dimension d respecte ces deux conditions pour $d = 3$ mais pas pour $d = 2$. Nous avons par conséquent conjecturé que ce modèle est diffusif pour $d \geq 3$.

Par la suite nous avons cherché à prouver analytiquement que ces deux conditions étaient respectées, et ce à partir d'une approche perturbative. Ce problème s'est avéré ardu. Nous avons donc restreint la question à la question de l'existence d'un coefficient de diffusion. Nous sommes alors passés du modèle des miroirs dans le tore au modèle des miroirs dans $\mathbb{Z}^d$, et cela nous a finalement ramené à étudier la diffusivité du modèle des miroirs selon la définition 0.1 de la section 0.4. Obtenir des résultats rigoureux sur le modèle des miroirs était difficile mais nous avons des calculs intéressants qui donnaient un début d'approche perturbative. Cependant les valeurs calculées n'étaient pas en adéquation avec les simulations numériques. Face à ces difficultés nous avons tenté de simplifier le modèle en abandonnant la condition de réversibilité du modèle des miroirs. Cela nous a mené au modèle des permutations étudié dans le chapitre 2. Nous avons même conçu des versions dite "faiblement aléatoires" du modèle des permutations et du modèle des miroirs. Tous ces modèles sont présentés à la section 2.1.

Le modèle des permutations faiblement aléatoires présentait deux avantages, premièrement il était plus simple à étudier analytiquement que le modèle des miroirs mais en plus les résultats numériques étaient cohérents avec les résultats de l'approche perturbative. Nous avons donc poussé l'approche analytique jusqu'à obtenir une expression du coefficient de diffusion κ sous la forme d'une série, et nous avons pu prouver rigoureusement que chaque terme de cette série est bien défini, ce qui est le résultat principal de cette thèse. La preuve de ce théorème est donnée à la section 2.2. Malheureusement nous n'avons pas pu prouver la convergence de cette série, nous laissons ce travail à d'autres chercheurs courageux.

Nous avons alors pu obtenir les 3 premiers termes κ_0 , κ_2 et κ_3 , du développement perturbatif de κ pour le modèle des permutations. Le terme κ_1 étant nul. Et nous avons la garantie que ces termes avaient un sens mathématiquement et que nous pourrions obtenir les suivants à condition de calculer longtemps. Le calcul de ces termes est présenté à la section 2.3.

Le chapitre 3 quant à lui, est consacré à la présentation de nos algorithmes (section 3.1) et des résultats des simulations numériques (section 3.2) ainsi qu'à la comparaison de ces résultats avec ceux de l'approche perturbative. On y présentera aussi, à la section 3.3, une heuristique permettant de comprendre (au moins partiellement) pourquoi les résultats analytiques et les résultats numériques semblent incohérents pour le modèle des miroirs.

Nous allons maintenant rentrer dans le vif du sujet en nous attaquant à la loi de Fick dans le modèle des miroirs.

Chapitre 1

Le modèle des miroirs dans le tore

$$\{1, \dots, N\} \times (\mathbb{Z}/N\mathbb{Z})^{d-1}$$

Le contenu de ce chapitre est celui de l'article "The Mirrors Model : Macroscopic Diffusion Without Noise or Chaos" [22] que nous avons publié en 2016 dans J-Phys A. Cela explique que l'approche et les notations soient un peu différentes de celles des autres chapitre, les choses ayant évoluée depuis. C'est aussi la raison pour laquelle ce chapitre est en anglais.

1.1 Abstract of this chapter

Before stating our main result, we first clarify through classical examples the status of the laws of macroscopic physics as laws of large numbers. We next consider the mirrors model in a finite d -dimensional domain and connected to particles reservoirs at fixed chemical potentials. The dynamics is purely deterministic and non-ergodic but takes place in a random environment. We study the macroscopic current of particles in the stationary regime. We show first that when the size of the system goes to infinity, the behaviour of the stationary current of particles is governed by the proportion of orbits crossing the system. This allows to formulate a necessary and sufficient condition on the distribution of the set of orbits that ensures the validity of Fick's law. Using this approach, we show that Fick's law relating the stationary macroscopic current of particles to the concentration difference holds in three dimensions and above. The negative correlations between crossing orbits play a key role in the argument.

1.2 Macroscopic laws as laws of large numbers

Take a macroscopic box $\Lambda = [0, L]^d$ that contain N freely moving distinguishable particles and fixed obstacles of arbitrary shapes. N should be thought to be of the order of magnitude of the Avogadro number : 6×10^{23} .

A first experiment is performed on this box. The N particles are initially located in a cube $\Lambda' \subset \Lambda$ of side length $L' < L$, see Figure 1.2. The evolution of the density of the cloud of particles is monitored through a beam of light that crosses the system. We call this density $\rho : \Lambda \times [0, \infty[\rightarrow \rho(\mathbf{x}, t) \in \mathbb{R}^+$. The initial state is described by $\rho(\mathbf{x}, 0) = \mathbf{1}_{\Lambda'}(\mathbf{x})/|\Lambda'|$. The empirical fact that is observed at the macroscopic level is that the density evolves according to the laws of diffusion :

$$\begin{cases} \partial_t \rho(\mathbf{x}, t) = \kappa \Delta \rho(\mathbf{x}, t) \\ n_{\mathbf{x}} \cdot \nabla \rho(\mathbf{x}, t) = 0, \quad \mathbf{x} \in \partial \Lambda \\ \rho(\mathbf{x}, 0) = \rho_0(\mathbf{x}) := \mathbf{1}_{\Lambda'}(\mathbf{x})/|\Lambda'|, \end{cases} \quad (1.1)$$

where $n_{\mathbf{x}}$ is the vector normal to the boundary of the box $\partial \Lambda$ at $\mathbf{x}$ and κ is a strictly positive constant. How can we explain this phenomenon from the motion of the individual atoms ? For each *macroscopic* coordinate $\mathbf{x} = (x_1, \dots, x_d) \in \Lambda$, we define a *microscopic* coordinate $\mathbf{q} = \mathbf{r}/\epsilon_N$

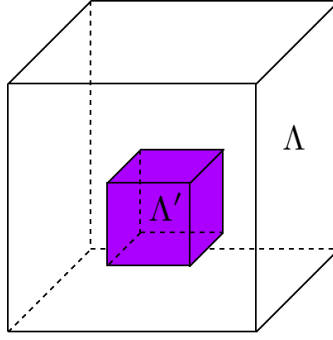


FIGURE 1.1 – In the first experiment, the cloud of particles is initially concentrated in a volume Λ' .

where $\epsilon_N = \frac{1}{N^{1/d}}$. The motion of the particles is entirely determined by the law of Newtonian mechanics. When a particle makes a collision with one of the fixed obstacles or the boundaries of the boxes, its velocity is modified according to the laws of specular reflection. We assume that each particle starts with a speed equal to 1. Since this property is preserved by the dynamics, the microscopic motion of a given particle (with label $i \in \{1, \dots, N\}$) is described by a map $t \rightarrow (\mathbf{q}_i(\epsilon_N^{-2}t), \mathbf{p}_i(\epsilon_N^{-2}t)) \in [0, L/\epsilon_N]^d \times S^{d-1}$ where S^{d-1} is the unit sphere in d dimensions. The coordinates $(\mathbf{q}_i, \mathbf{p}_i)$ are the *microscopic* positions and velocities of the i -th particle. The microscopic time-scale is $\epsilon_N^{-2}t$. The scaling of the time variable is a priori arbitrary but is fixed here by the fact that the solution of the diffusion equation $\rho(\mathbf{x}, t)$ is invariant under the transformation $(\mathbf{x}, t) \rightarrow (\lambda\mathbf{x}, \lambda^2t)$, $\lambda > 0$.

In the absence of any other information, we assume that the initial positions and velocities of the particles are independent and identically uniformly distributed, with density

$$f(\mathbf{q}_i, \mathbf{p}_i) = \frac{\epsilon_N^d}{|\Lambda'| |S^{d-1}|} \mathbf{1}_{\Lambda'}(\epsilon_N \mathbf{q}_i) \mathbf{1}_{S^{d-1}}(\mathbf{p}_i) \quad (1.2)$$

for every $i = 1, \dots, N$. If the position of each particle is chosen independently of the others with that density, the number of particles in a microscopic volume of size of order 1 follows a Poisson distribution with finite mean as $N \rightarrow \infty$. We denote by $\mathbb{P}$ the law of probability of the initial positions and velocities of particles. No information about the spatial location or the shape of the obstacles in Λ is known either. We denote by $\mathbb{Q}$ the probability distribution on those degrees of freedom. It is chosen such that in the limit $N \rightarrow \infty$, the number of obstacles in a microscopic volume of size 1 follows a distribution with a finite mean and such that, almost surely, the dynamics of moving particles is well-defined at all time. The dynamical system defined in this way is an instance of the *random Lorentz gas*.

We define the empirical density of particles :

$$\rho_N(\mathbf{x}, t) = \frac{1}{N} \sum_{j=1}^N \delta(\epsilon_N \mathbf{q}_j(\epsilon_N^{-2}t) - \mathbf{x}) \quad (1.3)$$

The density ρ_N contain all possible information about the density. Indeed, it is easy to see that

$$\rho_N(V, t) := \int_{\mathbb{R}^d} d\mathbf{x} \mathbf{1}_V(\mathbf{x}) \rho_N(\mathbf{x}, t)$$

gives the proportion of particles that belong to any $V \subset \Lambda$ at time t . It is straightforward to see that the following statement holds : if $\{(\mathbf{q}_i(0), \mathbf{p}_i(0)) : i = 1, \dots, N\}$ is a collection of i.i.d variables with marginals given by (1.2) then for any bounded function h and any $\delta > 0$:

$$\lim_{N \rightarrow \infty} \mathbb{P}[|\langle \hat{\rho}_N(0), h \rangle - \langle \rho_0, h \rangle| > \delta] = 0, \quad (1.4)$$

where ρ_0 is given by (1.1), $\langle h, g \rangle = \int_{\mathbb{R}^d} d\mathbf{x} h(\mathbf{x})g(\mathbf{x})$ and we use the notation $\hat{\rho}_N(t) := \rho_N(\cdot, t)$. The goal is to show the

Conjecture 1.1. *There exists a natural¹ distribution $\mathbb{Q}$ such that for any $t > 0$, any bounded function h and any $\delta > 0$:*

$$\lim_{N \rightarrow \infty} \mathbb{P} \times \mathbb{Q} [| \langle \hat{\rho}_N(t), h \rangle - \langle \rho(t), h \rangle | > \delta] = 0. \quad (1.5)$$

where $\rho(t) := \rho(\cdot, t)$ is the solution of (1.1) for some $\kappa > 0$.

The law of ordinary diffusion is therefore understood as a *law of large numbers* : as N becomes very large, the probability that the empirical density $\hat{\rho}_N(t)$ differs significantly from the solution of the diffusion equation goes to zero.

It is natural to consider first a simpler version of the problem in which the randomness of the obstacles is removed, i.e. $\mathbb{Q}$ is taken to be a Dirac distribution δ_C on a special configuration of obstacles giving rise to a chaotic dynamics. This is exactly the result of Bunimovich and Sinai [5]. They consider the 2D case in which obstacles are disks located at the vertices of a regular lattice such that the induced billiard dynamics has a finite horizon². Their result implies that for any $t > 0$, any bounded function h and any $\delta > 0$:

$$\lim_{N \rightarrow \infty} \mathbb{P} \times \delta_C [| \langle \hat{\rho}_N(t), h \rangle - \langle \rho(t), h \rangle | > \delta] = 0. \quad (1.6)$$

Let us sketch how this statement is obtained. Let $h : \mathbb{R}^d \rightarrow \mathbb{R}$ a bounded function. First, one computes :

$$\langle \hat{\rho}_N(t), h \rangle = \frac{1}{N} \int_{\mathbb{R}^d} d\mathbf{x} \sum_{j=1}^N \delta(\epsilon_N \mathbf{q}_j(\epsilon_N^{-2}t) - \mathbf{x}) h(\mathbf{x}) \quad (1.7)$$

$$= \frac{1}{N} \sum_{j=1}^N h(\epsilon_N \mathbf{q}_j(\epsilon_N^{-2}t)). \quad (1.8)$$

Thus, because the initial positions of the particles are identically distributed :

$$\mathbb{E}[\langle \hat{\rho}_N(t), h \rangle] = \mathbb{E}[h(\epsilon_N \mathbf{q}_1(\epsilon_N^{-2}t))].$$

Next, the theorem 2 of [5] implies³ that

$$\lim_{N \rightarrow \infty} \mathbb{E}[h(\epsilon_N \mathbf{q}_1(\epsilon_N^{-2}t))] = \int_{\mathbb{R}^d} \rho(\mathbf{x}, t) h(\mathbf{x}, t) d\mathbf{x} \quad (1.9)$$

where $\rho(\mathbf{x}, t)$ is the solution of (1.1). To derive (1.9), one has to rely on the strong chaotic properties of the billard system under study.

Next, since $\{\mathbf{q}_j(t) : 1 \leq j \leq N\}$ are *independent*, the variance of $\langle \hat{\rho}_N(t), h \rangle$ is

$$\text{Var}[\langle \hat{\rho}_N(t), h \rangle] = \frac{1}{N} \text{Var}[h(\epsilon_N \mathbf{q}_1(\epsilon_N^{-2}t))] = O\left(\frac{1}{N}\right)$$

since h is bounded. Thus, Chebychev inequality allows us to conclude the proof of Conjecture 1.1 in the case where $\mathbb{Q} = \delta_C$. One should note that the proof is made of two steps of very different levels of complexity. The first step is basically given by (1.9) and this is where the whole difficulty is located. The second step is a concentration result of the random variable $\langle \hat{\rho}_N(t), h \rangle$ around the expected value $\mathbb{E}[\langle \hat{\rho}_N(t), h \rangle]$. This part is trivial because when $\mathbb{Q}$ is replaced by δ_C , the positions and velocities of the particles remain independent for all time t . The fact that it is so trivial is probably the reason why it is hard to find a reference where this step is mentioned or even alluded to. It is however essential and when $\mathbb{Q} \neq \delta_C$, the statistical independence of the motions of particles is lost. Controlling the correlations between them to ensure the concentration of $\langle \hat{\rho}_N(t), h \rangle$ around its mean does require some work. We will see below that this issue arises in the mirrors model and how it can be dealt with.

1. By natural we mean as uniform as possible over the locations and shapes of obstacles.
2. For instance, the center of each (sufficiently large) disk is located at a vertex of a triangular lattice.
3. To be more precise Bunimovich and Sinai consider the case $L = \infty$ but their method should apply directly to the finite L case

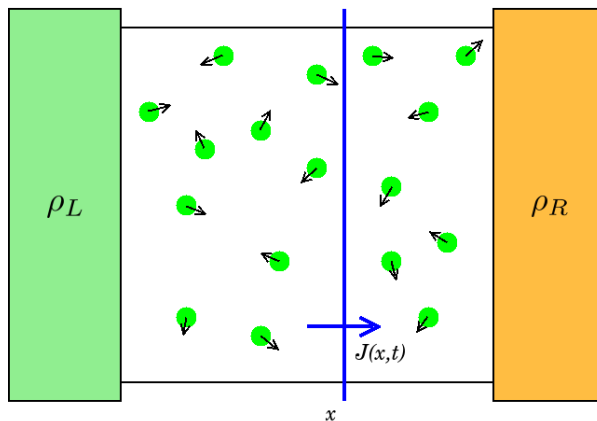


FIGURE 1.2 – At the two sides of the cube Λ , particles reservoirs maintain constant values of the local densities of particles ρ_L and ρ_R .

A second experiment may be performed on the box. At the two sides of the cube Λ perpendicular to $e_1 = (1, 0, \dots, 0)$, particles reservoirs maintain constant values of the local densities of particles ρ_L and ρ_R , respectively on the left and right side, see Figure 2.

A device records the net flux of mass crossing a section of Λ perpendicular to e_1 per unit time. This quantity is denoted by $j(x, t)$ when the section contains the point $(x, 0, \dots, 0)$ for $x \in [0, L]$. After some transient time proportional to L^2 , it is observed that the instantaneous current of particles takes the stationary value :

$$j_s(x) = \frac{\kappa}{L}(\rho_L - \rho_R). \quad (1.10)$$

With respect to the first experiment, the coupling to external reservoirs introduce an additional probabilistic element. The law of the reservoirs and the initial conditions of the particles inside the system is denoted by $\mathbb{P}$. One can introduce an empirical current of particles per unit time $\hat{J}_N(t)$ ⁴. Again, the goal is to show the following conjecture :

Conjecture 1.2. *There exists a natural distribution $\mathbb{Q}$ such that for any bounded continuous function h and any $\delta > 0$:*

$$\lim_{N \rightarrow \infty} \lim_{t \rightarrow \infty} \mathbb{P} \times \mathbb{Q} [| \langle \hat{J}_N(t), h \rangle - \langle j_s, h \rangle | > \delta] = 0. \quad (1.11)$$

While this has never been done explicitly, one should expect that with the choice $\mathbb{Q} = \delta_C$ the result follows from the methods of [5]. In [35], the authors show that in a low density regime limit, the expectation of the stationary current (with respect to $\mathbb{P} \times \mathbb{Q}$) converges to j_s . The statement corresponding to Conjecture 1.2 together with the exponential convergence to the stationary current has been obtained in the case of a discrete space-time dynamics in [31]. We now outline how the problem may be tackled in the context of the mirrors model.

1.3 The mirrors model

The mirrors model was introduced by Ruijgrok and Cohen [32] as a lattice version of the random Lorentz gas or the Ehrenfest wind-tree model. The latter encompasses a large class of models in which obstacles do not induce a chaotic behaviour of the trajectories of the particles. A fundamental question which remains open regarding those models is whether a non-chaotic

4. This quantity will be our main object of study in the next section and will be given a precise definition there.

deterministic dynamics may give rise to a *macroscopic* diffusive behaviour. In the mirrors model, particles travel on the edges of the cubic lattice generated by $\mathbb{Z}^d$. "Mirrors" are located at the vertices of the lattice and deflect the motion of an incoming particle in a new direction, see Figure 1.3 for an illustration in the 2D version of the model. The precise general definition is given below.

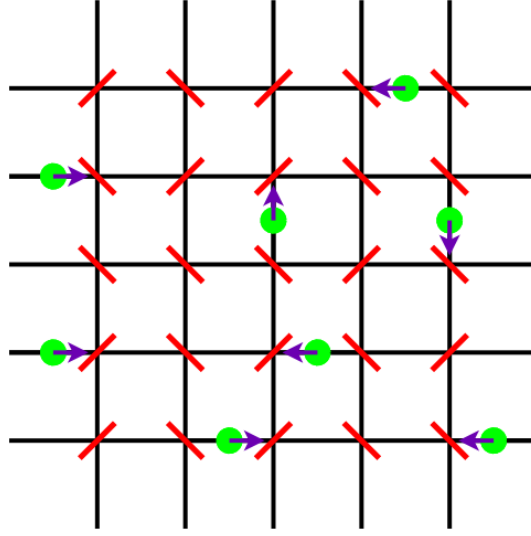


FIGURE 1.3 – In the mirrors model on $\mathbb{Z}^2$, green particles travel on the edge of the lattice and are deflected by mirrors located at the vertices. The green particles do not interact with each other.

A quick look at the structure of the orbits reveals its total lack of ergodicity. Indeed, in Figure 1.5 a sample of orbits in a finite box with periodic boundary conditions in the vertical direction and reflecting boundary conditions in the horizontal direction is pictured with different colors. For almost any configuration of the mirrors, no orbit is able to visit the entire phase space.

A perhaps even more striking fact is that, in any dimension, the motion of a particle in an environment of randomly orientated mirrors is not a gaussian diffusion⁵ [36]. More precisely this means that (1.9) (where the expectation is taken with respect to $\mathbb{P}$ and $\mathbb{Q}$) does not hold.

Our goal is to show that in spite of these unpromising properties, the mirrors model does exhibit normal macroscopic conductive properties when $d \geq 3$ in the sense that the analogue of (1.11) holds. It turns out that quite weak conditions on the statistics of orbits are sufficient to ensure the validity of Fick's law at the macroscopic level. It is therefore not necessary that orbits behave as a Gaussian diffusion to ensure the validity of Fick's law. Thus, the normal *macroscopic* laws of diffusion apply to a much wider class of dynamical systems than generally expected.

The dynamics of the mirrors model is reversible in the usual sense of the word in the context of Hamiltonian dynamics. Namely, under the reversal of the velocities of all particles at a given time $t > 0$, the dynamics brings the system of particles to its initial condition at time 0 (with reversed velocities), see (1.13). This reversibility property of the dynamics will allow us to show that when the system is large, the number of orbits travelling from one side of the system to the other one basically determines the value of the current in the stationary state. This will allow to formulate a condition on the distribution of orbits that is both sufficient and necessary for the validity of Fick's law.

We recall now briefly the set-up of the original mirrors model. Particles travel on the edges of $\mathbb{Z}^2$ with unit speed. Mirrors are located at some vertices of the lattice and take two possible angular orientations : $\{\frac{\pi}{4}, \frac{3\pi}{4}\}$. When a particle hits a mirror, it gets deflected according to the laws of

5. In spite of this, it has been observed numerically [6] that in two dimensions, the mean-square displacement of a given particle is linear in time, allowing the definition of a *microscopic* diffusion coefficient. This is of course a much weaker property than the property (1.11). In particular we will see that one can not define a *macroscopic* diffusion coefficient in $2D$.

specular reflection, see Figure 1.5 for sample trajectories of particles. It is convenient to think that every particle starts at time zero with a given velocity at a vertex of the lattice $\mathcal{Q}$ that is obtained by taking the middle point of every edge of $\mathbb{Z}^2$. As all particles move with unit velocity, one can simply observe the evolution of the system at discrete times $t \in \mathbb{N}$. At those times, the particles will be always located at one of the vertices of the new lattice $\mathcal{Q}$ with a well-defined velocity. In general, the orientation of the mirrors is picked randomly. It is obvious that the motion of a single particle can not be described as a Markov process. When a particle hits a mirror for the second time, no matter how far back in the past the first visit occurred, its reflection is strongly affected by the way its was reflected at the first visit. For instance in Figure 1.5, the two orientations of the mirrors are picked at random, and in that case, at the second visit the reflection is always deterministic.

We come now to a more general definition of the dynamics in d dimensions. We denote by $\mathbf{z} = (z_1, \dots, z_d)$ a generic element of $\mathbb{Z}^d$. As for $\mathbb{Z}^2$, we consider the set of midpoints of edges of an hypercube of $\mathbb{Z}^d$ of side N and with periodic conditions in all but the first direction. We call this set $\mathcal{Q}$. It may be described as follows : $\mathcal{Q} = \bigcup_{i=1}^d L_i$ where

$$L_i = \left\{ \mathbf{z} + \frac{1}{2} \mathbf{e}_i : 0 \leq z_1 \leq N-1, (z_2, \dots, z_d) \in (\mathbb{Z}/N\mathbb{Z})^{d-1} \right\}.$$

Let $(\mathbf{e}_1, \dots, \mathbf{e}_d)$ the canonical basis of $\mathbb{R}^d$, the space of possible velocities is $\mathcal{P} = \{\pm \frac{\mathbf{e}_1}{2}, \dots, \pm \frac{\mathbf{e}_d}{2}\}$ and the phase space of the dynamics is

$$\mathcal{M} = \{(\mathbf{q}, \mathbf{p}) : \mathbf{q} \in \mathcal{Q}, \mathbf{p} \in \mathcal{P} \text{ s. t. if } \mathbf{q} \in L_i \text{ then } \mathbf{p} = \pm \frac{\mathbf{e}_i}{2}\}.$$

We denote a generic point of $\mathcal{M}$ by $(\mathbf{q}, \mathbf{p})$. The set of points in $\mathcal{M}$ whose spatial coordinate belongs to the boundaries of the system is $B = B_- \cup B_+$, with

$$\begin{aligned} B_- &= \{x = (\mathbf{q}, \mathbf{p}) \in \mathcal{M} : \mathbf{q} = (q_1, \dots, q_d) \in L_1, q_1 = \frac{1}{2}\} \\ B_+ &= \{x = (\mathbf{q}, \mathbf{p}) \in \mathcal{M} : \mathbf{q} = (q_1, \dots, q_d) \in L_1, q_1 = N - \frac{1}{2}\}. \end{aligned}$$

See Figure 1.4.

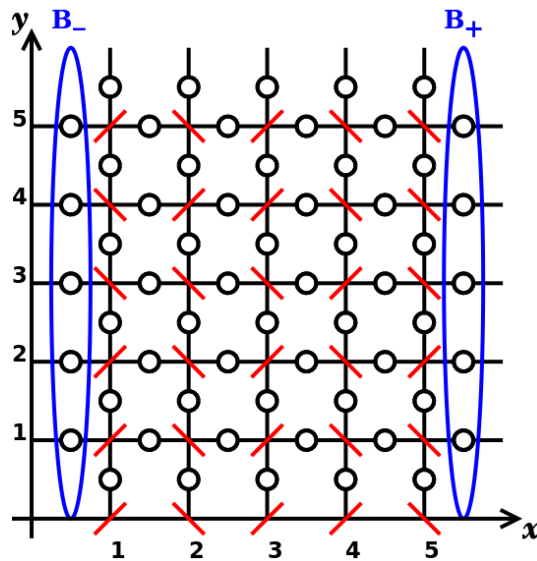


FIGURE 1.4 – The spatial component of the phase space $\mathcal{M}$ and of the sets B_- and B_+ in $2D$.

For each $\mathbf{z} \in \mathbb{Z}^d$, we define randomly the action of a “mirror” on the velocity of an incoming particle by $\pi(\mathbf{z}; \cdot)$ which is a bijection of $\mathcal{P}$ into itself. It satisfies the following conditions :

$$\begin{cases} \pi(\mathbf{z}; -\pi(\mathbf{z}; \mathbf{p})) = -\mathbf{p}, & \forall \mathbf{z} \in \mathbb{Z}^d, \forall \mathbf{p} \in \mathcal{P} \\ \pi(0, z_2, \dots, z_d; -\frac{\mathbf{e}_1}{2}) = \frac{\mathbf{e}_1}{2} \\ \pi(N, z_2, \dots, z_d; \frac{\mathbf{e}_1}{2}) = -\frac{\mathbf{e}_1}{2}, & (z_2, \dots, z_d) \in (\mathbb{Z}/N\mathbb{Z})^{d-1} \end{cases} \quad (1.12)$$

We will take mirrors randomly among the bijections $\pi(\mathbf{z}; \cdot)$ which satisfy (1.12). The law will be given later. The dynamics is defined on $\mathcal{M}$ in the following way. For any $(\mathbf{q}, \mathbf{p}) \in \mathcal{M}$:

$$F(\mathbf{q}, \mathbf{p}) = (\mathbf{q} + \mathbf{p} + \pi(\mathbf{q} + \mathbf{p}; \mathbf{p}), \pi(\mathbf{q} + \mathbf{p}; \mathbf{p})).$$

It is easy to check that the map F is a bijection on $\mathcal{M}$. The two last conditions in (1.12) are just saying that when particles hit the boundaries they are reflected backwards. We define an operator $R : \mathcal{M} \rightarrow \mathcal{M}$ which reverses the velocities by $R(\mathbf{q}, \mathbf{p}) = (\mathbf{q}, -\mathbf{p})$ and we see that the first of the conditions (1.12) ensures that the map F is *reversible*,

$$F^{-1} = RFR. \quad (1.13)$$

We define the orbit of a point $x \in \mathcal{M}$, $\mathcal{O}_x = \{y \in \mathcal{M} : \exists t \geq 0, F^t(x) = y\}$ and its period, $T(x) = \inf\{t \geq 0 : F^t(x) = x\}$. From the fact that F is bijective, one infers that for every $x \in \mathcal{M}$, $\mathcal{O}_x$ is a loop : $T(x) \leq |\mathcal{M}|$ and that orbits are non-intersecting : if $y \notin \mathcal{O}_x$, then $\mathcal{O}_x \cap \mathcal{O}_y = \emptyset$. A given orbit is also non-self-intersecting : if $y \in \mathcal{O}_x$ and $y \neq x$ then $F(y) \neq F(x)$.

As we are interested in the transport of particles, we define occupation variables $\sigma(\mathbf{q}, \mathbf{p}; t) \in \{0, 1\}$ that record the absence or presence of a particle at position $\mathbf{q}$ with velocity $\mathbf{p}$ at time $t \in \mathbb{N}$. When connecting the system to external particles reservoirs, we obtain the following evolution rule : given $\sigma(\cdot; t-1)$, we define $\sigma(\cdot; t)$ for all $t \in \mathbb{N}^*$ by

$$\sigma(x; t) = \begin{cases} \sigma(F^{-1}(x); t-1) & \text{if } x \notin B_- \cup B_+ \\ \sigma_x^-(t-1) & \text{if } x \in B_- \\ \sigma_x^+(t-1) & \text{if } x \in B_+. \end{cases}$$

The families of random variables $\{\sigma_x^-(t) : x \in B_-, t \in \mathbb{N}\}$ and $\{\sigma_x^+(t) : x \in B_+, t \in \mathbb{N}\}$ consist of independent Bernoulli variables with respective parameters ρ_- and ρ_+ . If one chooses $\{\sigma(x; 0) : x \in \mathcal{M}\}$ to be a collection of independent random variables, then it is easy to see by induction that at any $t \geq 0$, $\{\sigma(x; t) : x \in \mathcal{M}\}$ is a collection of i.i.d Bernoulli random variables. To simplify a bit the discussion, we choose an homogeneous initial distribution, i.e. all Bernoulli random variables have a common parameter ρ_I . The distribution of the collection $\{\sigma(x; t) : x \in \mathcal{M}\}$ becomes stationary after a *finite* time. More precisely, for any $t \geq |\mathcal{M}|$, we have the following equality in law :

$$\sigma(x, t) = \begin{cases} \sigma_I & \text{if } \mathcal{O}_x \cap B = \emptyset \\ \sigma_- & \text{if } F^{-t^*}(x) \in B_- \\ \sigma_+ & \text{if } F^{-t^*}(x) \in B_+ \end{cases}$$

where $t^* = \inf\{t : F^{-t}(x) \in B\}$ and $\sigma_{\pm}$ and σ_I are Bernoulli random variables of parameter $\rho_{\pm}$ and ρ_I .

Proceeding as in [31], it is possible to show that when the size of the system goes to infinity, the stationary current converges in probability to the proportion of crossing orbits times the chemical potentials difference. We define the average current of particles that crosses the hyperplane $\mathcal{Q}^l = \{\mathbf{q} \in \mathcal{Q} : q_1 = l + \frac{1}{2}\}$, $l \in \{1, \dots, N-2\}$ during a diffusive time interval N^2 (we will need this time average for kill the natural fluctuations of the Bernoulli variables) :

$$J(l, t) = \frac{1}{N^{d+1}} \sum_{s=t+1}^{t+N^2} \sum_{x \in \mathcal{M}} \sigma(x; s) \Delta(x, l) \quad (1.14)$$

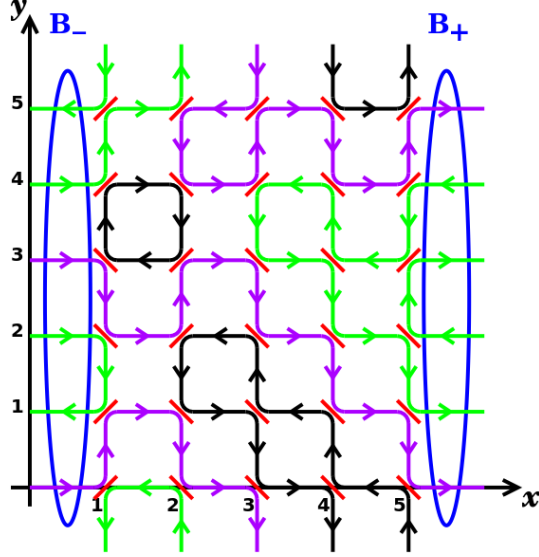


FIGURE 1.5 – $N = 6$. Crossing orbits are coloured in purple, internal loops in black and non-crossing orbits are coloured in green. The travel direction given by the arrows is arbitrary. Each edge of the crossing orbits will be used twice in a given orbit : once in each direction. For this configuration of mirrors $\mathcal{N} = 2$.

where $\Delta(x, l) = 2(\mathbf{p} \cdot \mathbf{e}_1) \mathbf{1}_{\mathbf{q} \in \mathcal{Q}^l}$, with $x = (\mathbf{q}, \mathbf{p})$. Thus $\Delta(x, l)$ takes the value $+1$ (resp. -1) if x crosses the slice $\mathcal{Q}^l$ from left to right (resp. from right to left). We denote by $\mathcal{N}_{\pm}$ the numbers of crossings from $B_{\pm}$ to $B_{\mp}$ induced by F , i.e. $\mathcal{N}_{\pm} = |S_{\pm}|$ where $S_{\pm}$ is given by

$$S_{\pm} = \{x \in B_{\pm} : \exists s > 0, \forall 0 < j < s, F^j(x) \notin B_{\pm}, F^s(x) \in B_{\mp}\}.$$

One notes that $\mathcal{N}_+ = \mathcal{N}_-$. Indeed, since every orbit is closed, it must contain as many left-to-right than right-to-left crossings. Thus, we set $\mathcal{N} = \mathcal{N}_+ = \mathcal{N}_-$. Proceeding as in [31], we get that for any $t \geq |\mathcal{M}|$, $\mathbb{E}[J(l, t)] = \frac{\mathcal{N}}{N^{d-1}}(\rho_- - \rho_+)$.⁶ Moreover, for every $\delta > 0$, any $t \geq |\mathcal{M}|$ and $l \in \{1, \dots, N-2\}$,

$$\mathbb{P} \left[\left| NJ(l, t) - \frac{\mathcal{N}}{N^{d-2}}(\rho_- - \rho_+) \right| \geq \delta \right] \leq 2 \exp(-\delta^2 N^{d-1}). \quad (1.15)$$

Without the time average in the definition of J we would have N^{d-3} in the RHS, and it would be a problem for $d \in \{2, 3\}$. For $d \geq 4$ we don't need this time average.

We take now random configurations of reflectors $\{\pi(\mathbf{z}; \cdot) : \mathbf{z} \in \mathbb{Z}^d\}$. The law of the reflectors is denoted by $\mathbb{Q}$. The map F becomes now a random map.

From the physic's point of view the Fick's law claims that $J(l, t) \sim \kappa/N$. Let's give a mathematical version of this law. The model satisfies *Fick's law* if and only if there exists some $\kappa > 0$ (the conductivity) such that $\forall \delta > 0$,

$$\lim_{N \rightarrow \infty} \lim_{t \rightarrow \infty} \mathbb{P} \times \mathbb{Q} [|NJ(l, t) - \kappa(\rho_- - \rho_+)| > \delta] = 0. \quad (1.16)$$

As in [31], it is easy to infer from (1.15) that the following theorem holds.

Théorème 1.1. Sufficient and Necessary Condition for Fick's law : (1.16) holds if and only if there exists $\kappa > 0$ such that for any $\delta > 0$,

$$\lim_{N \rightarrow \infty} \mathbb{Q} \left[\left| \frac{\mathcal{N}}{N^{d-2}} - \kappa \right| > \delta \right] = 0. \quad (1.17)$$

6. This relation implies that the average current flows in the "right" direction and that when $\rho_- \neq \rho_+$, the average current in the stationary state is different from 0 if and only if $\mathcal{N} \neq 0$.

We see that the central object to study is the distribution of the number of crossing orbits $\mathcal{N}$. The expectation of this quantity is related to the probability that one orbit crosses the system, while the variance is given in terms of the joint probability that orbits with two different starting points cross the system. Indeed by periodicity, we have, using the notations $O = ((\frac{1}{2}, 0, \dots, 0), \frac{e_1}{2})$ and $S = S_-$:

$$\mathbb{E} \left[\frac{\mathcal{N}}{N^{d-2}} \right] = \frac{N}{N^{d-1}} \sum_{x \in B_-} \mathbb{E}[\mathbf{1}_{x \in S}] = N \mathbb{Q}[O \in S] \quad (1.18)$$

and

$$\text{Var} \left[\frac{\mathcal{N}}{N^{d-2}} \right] = \frac{1}{N^{2d-4}} \sum_{x, y \in B_-} \delta(x, y) = \frac{1}{N^{d-3}} \sum_{x \in B_-} \delta(O, x) \quad (1.19)$$

with

$$\delta(x, y) = \mathbb{Q}[x \in S, y \in S] - \mathbb{Q}[x \in S] \mathbb{Q}[y \in S]. \quad (1.20)$$

Thus if the two following conditions, named the **Crossing Conditions**, are satisfied then Fick's law (1.16) holds in the stationary state.

Crossing Conditions. :

1. *There exist $\kappa > 0$ such that the RHS of (1.18) converges to κ as $N \rightarrow \infty$.*
2. *The RHS of (1.19) goes to zero as $N \rightarrow \infty$.*

We note first that when $d = 2$, (1.17) can not hold, whatever the distribution $\mathbb{Q}$ is. To see this, we adapt an argument found in [38]. Indeed, the spatial part of each crossing orbit crosses any "vertical" section $\mathcal{Q}^l$ an odd number of times. On the other hand, the spatial part of any non-crossing orbit must cross any vertical section an even number of times, see Figure 1.5. Thus, N and $\mathcal{N}$ must have the same parity. This implies that there can not exist $\kappa > 0$ such that (1.17) holds when $d = 2$. The origin of this issue lies in the strong correlations between crossing orbits that are present in two dimensions.

We turn now to the higher dimensional case $d \geq 3$ equipped with some natural and spatially homogeneous distribution $\mathbb{Q}$. Now observe that if $\mathbb{Q}[\pi(\mathbf{z}; \frac{e_i}{2}) = -\frac{e_i}{2}] > 0$ for some $i = 1, \dots, d$ then an orbit starting from O will encounter this type of reflecting mirror after an exponential number of steps and therefore $\mathbb{Q}[0 \in S] \leq e^{-cN}$ for some $c > 0$. This, in turn, implies that $\lim_{N \rightarrow \infty} N \mathbb{Q}[0 \in S] = 0$ and that Fick's law can not hold. Thus from now on, we consider maps such that $\pi(\mathbf{z}; \frac{e_i}{2}) \neq -\frac{e_i}{2}$ if $0 < z_1 < N$ and such that the conditions (1.12) are satisfied. We call the set of such maps Π . We take $\mathbb{Q}$ such that the collection of maps

$$\{\pi(\mathbf{z}; \cdot) : 0 < z_1 < N, (z_2, \dots, z_d) \in (\mathbb{Z}/N\mathbb{Z})^{d-1}\}$$

is independent and that each map is uniformly distributed over Π . We note first that if the law of an orbit with respect to $\mathbb{Q}$ was similar to the law of a simple random walk, then there would be a $\kappa > 0$ such that $\lim_{N \rightarrow \infty} N \mathbb{Q}[0 \in S] = \kappa$, this follows from the gambler's ruin argument. Similarly, if the orbits were independent objects, then the RHS of (1.19) would go to zero because the only non-zero term would be the one with $x = O$ and $\mathbb{Q}[O \in S] \sim \kappa/N$. We also note that the average stationary current is identified as the difference between chemical potentials times the probability that a particle crosses the system, an idea that was put forward in [37], in the context of chaotic systems. The law $\mathbb{Q}$ of the mirrors induces a law on the set of orbits which is a priori very far from the distribution of independent simple random walks. The set of orbits is a very interesting lattice object in itself which features some (self-)avoiding properties as we mentioned above.

Fortunately, what is needed to ensure the validity of (1.17) is much less than the full joint distribution of the orbits. Thanks to (1.18) and (1.19), one only has to analyze the marginal of a path starting on the boundary and also the joint probability of two such paths. The distribution

of a path starting at O (i.e. on the boundary) is similar to the one of a “true” self-avoiding random walk [4] but defined on $\mathcal{Q}$ rather than on $\mathbb{Z}^d$ and with further constraints. The diffusive behaviour of those walks for $d \geq 3$ has been conjectured in [4], see also the rigorous results of [39]. It can be expected that as the dimensionality of the system increases, the effect of the revisits of an orbit to the same mirror decreases. In a process where the mirrors are flipped randomly after being used (i.e memory effects are killed), we computed that in $d = 3$ the crossing probability is $\sim 3/2N$. Numerical simulations in $d = 3$ show that this number is indeed a good approximation. The log log plot of the crossing probability $\mathbb{Q}[O \in S]$ is given in Figure 1.6 for N up to 400. The corresponding conductivity is $\kappa = 1.535 \pm 0.005$. As the conductivity measured in simulations is slightly higher than $3/2$, it indicates that recollisions tend to push forward the orbit.

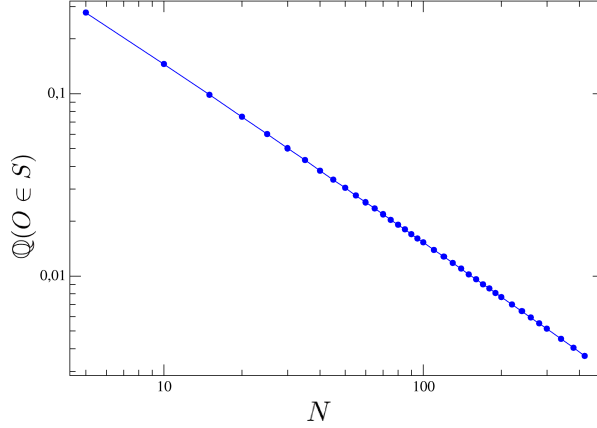


FIGURE 1.6 – $\mathbb{Q}(O \in S)$ for N from 5 to 420. The 95% confidence interval is about half the size of a dot.

We must show now that $\sum_{x \in B_-} \delta(O, x) \rightarrow 0$ as $N \rightarrow \infty$. We know that this sum is positive because it is a variance, and thus it is enough to get an upper bound on the sum. We will use a numerical analysis to show that $\delta(O, x) < 0$ for all but a finite (i.e. independent of N) number $x \in B_-$.

But before doing that and to get a better picture of the origin of the correlations $\delta(0, x)$, we split them in two parts. Given an orbit $\mathcal{O}$, we denote by $\gamma(\mathcal{O})$ the set of edges of $\mathbb{Z}^d$ used by $\mathcal{O}$. For each $\mathbf{z} \in Z_N := \{\mathbf{z} \in \mathbb{Z}^d : 1 \leq z_1 \leq N - 1, (z_2, \dots, z_d) \in (\mathbb{Z}/N\mathbb{Z})^d\}$, let also $b_{\mathbf{z}}(\mathcal{O})$ be the half of the number of times that the orbit $\mathcal{O}$ visits the vertex $\mathbf{z}$. Two *crossing orbits* $\mathcal{O}$ and $\mathcal{O}'$ are incompatible if $\gamma(\mathcal{O}) \cap \gamma(\mathcal{O}') \neq \emptyset$ and compatible otherwise. The law of a given *crossing* orbit is :

$$\mathbb{Q}(\mathcal{O}) = \prod_{\mathbf{z} \in Z_N} \prod_{j=1}^{b_{\mathbf{z}}(\mathcal{O})} \frac{1}{2(d-j)+1}. \quad (1.21)$$

If $\mathbb{Q}(\mathcal{O}, \mathcal{O}')$ is the joint probability of two orbits $\mathcal{O}$ and $\mathcal{O}'$ then $\mathbb{Q}(\mathcal{O}, \mathcal{O}') = 0$ when $\mathcal{O}$ and $\mathcal{O}'$ are incompatible. If they are compatible, the joint law of two crossing orbits is given by

$$\mathbb{Q}(\mathcal{O}, \mathcal{O}') = \prod_{\mathbf{z} \in Z_N} \prod_{j=1}^{b_{\mathbf{z}}(\mathcal{O})+b_{\mathbf{z}}(\mathcal{O}')} \frac{1}{2(d-j)+1}. \quad (1.22)$$

In particular, if they do not share any mirrors, then $\mathbb{Q}(\mathcal{O}, \mathcal{O}') = \mathbb{Q}(\mathcal{O})\mathbb{Q}(\mathcal{O}')$. From those properties, starting from (1.20), we obtain that for $x \in B_-$,

$$\begin{aligned} \delta(O, x) &= \sum_{\mathcal{O}_0, \mathcal{O}_x} (\mathbb{Q}(\mathcal{O}_0, \mathcal{O}_x) - \mathbb{Q}(\mathcal{O}_0)\mathbb{Q}(\mathcal{O}_x)) \\ &\quad - \sum_{\mathcal{O}_0, \mathcal{O}_x}' \mathbb{Q}(\mathcal{O}_0)\mathbb{Q}(\mathcal{O}_x). \end{aligned} \quad (1.23)$$

Both sums run over orbits that cross the box $\mathcal{Q}$. The first sum runs over compatible orbits such that $\gamma(\mathcal{O}_0)$ and $\gamma(\mathcal{O}_x)$ share a vertex of $\mathbb{Z}^d$. The second (prime) sum runs over *incompatible* orbits $\mathcal{O}_0, \mathcal{O}_x$.

Thus, from (1.23), we see that correlations $\delta(O, x)$ are created from two opposite origins, corresponding to each of the two sums in (1.23). If two orbits $\mathcal{O}$ and $\mathcal{O}'$ share some mirrors, then it is easy to see from (1.22) that $\mathbb{Q}(\mathcal{O}, \mathcal{O}') > \mathbb{Q}(\mathcal{O})\mathbb{Q}(\mathcal{O}')$. This in turn implies that the first sum is strictly positive. This is a “cooperative” effect, the orbits help each other crossing the system. The second sum corresponds to a *jamming* effect : an orbit starting from O and crossing the system occupies a certain number of horizontal edges. Because distinct orbits can not share the same edges, the occupied edges are no more available for an orbit starting from $x \in B_-$, this creates negative correlations.

Numerical simulations in $d = 3$ show that the latter effect dominates. For all but a few points, the correlations $\delta(O, x)$ for $x \neq O$ are not only small but negative within confidence intervals, see Figure 1.7. The only exceptions are points $((1/2, 1, 0), \frac{e_1}{2})$, $((1/2, 0, 1), \frac{e_1}{2})$, $((1/2, N-1, 0), \frac{e_1}{2})$ and $((1/2, 0, N-1), \frac{e_1}{2})$ which give clearly positive correlations. However, we checked that for $N = 70$, $\sum_{y=1}^{N-1} \delta(O, ((1/2, y, 0), \frac{e_1}{2})) = -1.360 \times 10^{-04} \pm 1.47 \times 10^{-05}$, i.e. it is negative with a margin of more than 9σ . $\sum_{z=1}^{N-1} \delta(O, ((1/2, 0, z), \frac{e_1}{2}))$ must be equal by symmetry. Increasing values of N do not modify this behaviour. In particular, the number of points with positive correlations do not increase. Since we know already that $\mathbb{Q}[O \in S] \sim \kappa/N$, as $N \rightarrow \infty$, we conclude with the same margin that $\sum_{x \in B_-} \delta(O, x) \leq \kappa/N \rightarrow 0$, as $N \rightarrow \infty$.

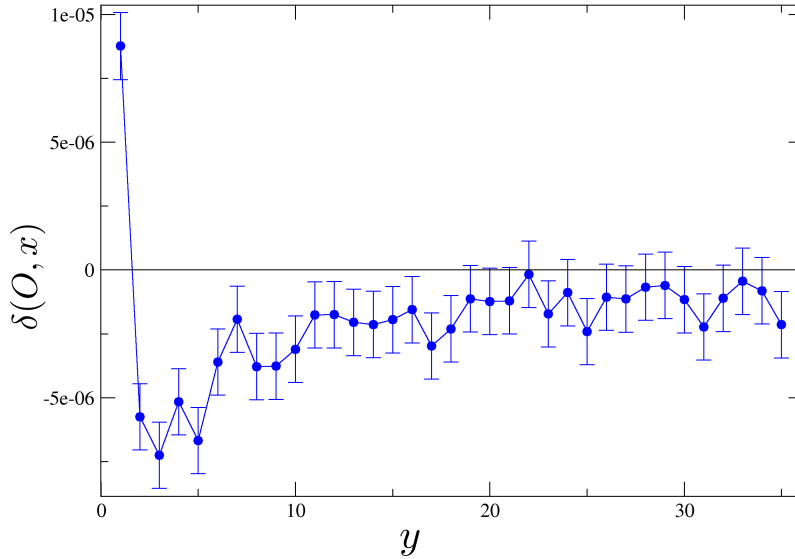


FIGURE 1.7 – $\delta(O, x)$ for $x = ((1/2, y, 0), \frac{e_1}{2})$. $N=70$. We draw the 95% confidence interval.

We expect the same behaviour in $d > 3$. A rigorous proof that the *crossing conditions* introduced above are satisfied seems to be within reach in the present model. Moreover, it is possible to draw a general conclusion from the above discussion. If one is really interested in deriving macroscopic laws from microscopic dynamics, many detailed properties of the latter are irrelevant. Only weaker properties than chaoticity, ergodicity or Gaussian behaviour of the orbits are required. In the present context, the minimal properties necessary to obtain Fick’s law are encapsulated in the crossing conditions. It is of course natural to seek similar weak conditions in different contexts as for instance in the problem of the derivation of Fourier’s law.

Chapitre 2

Nouveaux modèles et approche analytique

2.1 Élargissement du problème, marche aléatoire cinématique, environnement aléatoire, marche auto-évitante

2.1.1 Le modèle des miroirs : un cas particulier de marche avec vitesse

La description que nous avons donnée du modèle des miroirs fait apparaître le vecteur vitesse de la particule qui se déplace de miroir en miroir. Cela n'est pas usuel pour les marches aléatoires et nous pousse à considérer une généralisation possible. L'algorithme de base réglant l'évolution d'une particule dans le modèle des miroirs est le suivant : faire un pas, choisir un nouveau vecteur vitesse, faire un pas, choisir un nouveau vecteur vitesse, etc... Seul le choix d'un nouveau vecteur vitesse est aléatoire. C'est en partant de cette idée que nous allons généraliser.

Remarque : Dans le chapitre 1 nous avons considéré des sites au milieu des arêtes, cela à l'avantage d'être non ambigu, on visualise très bien le prochain pas : avancer d'une demi-arête, changer le vecteur vitesse puis de nouveau avancer d'une demi-arête. Mais dans les calculs cela devient rapidement très très lourd, c'est pourquoi nous adoptons une nouvelle convention. Cette nouvelle convention est de placer les particules sur les points de $\mathbb{Z}^d$ et que chaque pas se décompose ainsi :

- Premièrement changer la vitesse.
- Deuxièmement déplacer la particule.

On aurait pu décréter l'inverse, l'important étant de choisir une convention et de s'y tenir.

2.1.2 Les marches aléatoires cinématiques

On a vu que la dynamique du modèle des miroirs au temps n dépend directement de la position $\mathbf{q}_n$ de la particule et également de sa vitesse $\mathbf{p}_n = \mathbf{q}_n - \mathbf{q}_{n-1}$. Cette description est très naturelle en physique où le couple position-vitesse est fondamental pour décrire l'état d'une particule en mécanique classique. Nous nous sommes donc intéressés de manière générale aux marches aléatoires avec vitesse. Nous les nommons les *marches aléatoires cinématiques*.

On considère une particule se déplaçant sur $\mathbb{Z}^d$ et ayant une vitesse de norme 1 dans $\mathcal{P} := \{\pm \mathbf{e}_1, \dots, \pm \mathbf{e}_d\}$. L'espace des états possibles de cette particule est alors :

$$\mathcal{M} = \mathbb{Z}^d \times \mathcal{P}.$$

On notera en général $x = (\mathbf{q}, \mathbf{p})$ un point de $\mathcal{M}$. Une marche aléatoire *cinématique* est une chaîne de Markov $(x_n)_{n \in \mathbb{N}}$, $x_n = (\mathbf{q}_n, \mathbf{p}_n)$ définie sur $\mathcal{M}$ et respectant la relation $x_{n+1} = (\mathbf{q}_n + \mathbf{p}_{n+1}, \mathbf{p}_{n+1})$.

Remarque : On pourrait prendre, plus généralement, $\mathbb{Z}^d$ comme espace des vitesses. Dans ce cas on peut voir que pour déterminer la position $\mathbf{q}_{n+1}$ de la marche il suffit de connaître les positions

$\mathbf{q}_{n-1}$ et $\mathbf{q}_n$, car $\mathbf{p}_n = \mathbf{q}_n - \mathbf{q}_{n-1}$. On a donc affaire à une marche aléatoire d'ordre 2. Et on peut dire que les marches aléatoires cinématiques sont des cas particuliers de marches aléatoires d'ordre 2. Mais du point de vu du physicien il est intéressant de voir cela comme des marches aléatoires dans l'espace des phases position-vitesse.

Pour une marche aléatoire cinématique, le noyau de transition peut alors s'écrire :

$$p(x, \bar{x}) = T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) \delta_{\mathbf{q}, \mathbf{q}+\bar{\mathbf{p}}} \quad (2.1)$$

où pour tout $\mathbf{q} \in \mathbb{Z}^d$, $T_{\mathbf{q}} : \mathcal{P} \times \mathcal{P} \rightarrow [0, 1]$ et :
 $\forall \mathbf{q} \in \mathbb{Z}^d, \forall \mathbf{p} \in \mathcal{P}$,

$$\sum_{\bar{\mathbf{p}}} T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) = 1. \quad (2.2)$$

On notera qu'une marche aléatoire cinématique est entièrement définie par la donnée de $T_{\mathbf{q}}$ pour tout $\mathbf{q} \in \mathbb{Z}^d$. On utilisera donc l'appellation "noyau de transition" à la fois pour parler de p et pour parler de T . Cela ne pose pas de problème d'ambiguïté.

On dit que la marche aléatoire cinématique est *invariante par translation* si le noyau de transition $T_{\mathbf{q}}$ ne dépend pas de la variable spatiale $\mathbf{q}$.

Exemples

1. La *marche aléatoire cinématique symétrique* est définie par le noyau de transition $T_{\mathbf{q}}$ tel que $T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) = 1/(2d)$, pour tout $\mathbf{q} \in \mathbb{Z}^d$ et pour tout $\mathbf{p}, \bar{\mathbf{p}} \in \mathcal{P}$. Cela revient à tirer une nouvelle vitesse à chaque pas, uniformément et indépendamment. La séquence $(\mathbf{q}_n)_{n \in \mathbb{N}}$ a alors la même loi qu'une marche aléatoire simple symétrique sur $\mathbb{Z}^d$.
2. La *marche aléatoire non rebroussante* est un exemple de marche avec mémoire sur $\mathbb{Z}^d$ cette marche peut se définir naturellement comme une marche aléatoire cinématique. Elle est définie par :

$$T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) = \begin{cases} 0 & \text{si } \bar{\mathbf{p}} = -\mathbf{p} \\ \frac{1}{2d-1} & \text{sinon} \end{cases}$$

pour tout $\mathbf{q} \in \mathbb{Z}^d$.

soit $p : \mathcal{M} \times \mathcal{M} \rightarrow [0, 1]$ le noyau de transition d'une marche aléatoire cinématique, on note R l'opération de renversement du temps $R(\mathbf{q}, \mathbf{p}) = (\mathbf{q} - \mathbf{p}, -\mathbf{p})$, on dit qu'une marche aléatoire cinématique est *réversible* si et seulement si

$$\forall x, y \in \mathcal{M}, p(x, y) = p(Ry, Rx).$$

2.1.3 Marche aléatoire cinématique en milieu aléatoire

Avec le formalisme précédent, le modèle des miroirs peut s'exprimer simplement comme une marche aléatoire caractérisée par un noyau de transition $T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}})$. Mais ce noyau de transition est lui même aléatoire. Ceci est caractéristique d'une marche aléatoire en environnement aléatoire. On utilisera indifféremment l'appellation *marche aléatoire en environnement aléatoire* (MAEA) ou l'appellation *marche aléatoire en milieu aléatoire* (MAMA) plus élégante en français.

On propose ici un formalisme général des marches aléatoires cinématiques en milieu aléatoire. Le but est d'exprimer le modèle des miroirs dans ce formalisme puis ensuite de décrire des modèles plus généraux. L'idée est que le modèle des miroirs est en pratique une marche déterministe en environnement aléatoire et que ce modèle ressemble fortement à une marche aléatoire auto-évitante. Par la généralisation on construira un modèle ressemblant à une marche aléatoire faiblement auto-évitante. On précisera ces analogies en temps utile.

On pose :

$$\mathcal{T} := \{T : \mathcal{P} \rightarrow \mathcal{P} | \forall \mathbf{p} \in \mathcal{P}, \sum_{\mathbf{p}' \in \mathcal{P}} T(\mathbf{p}, \mathbf{p}') = 1\}$$

L'ensemble $\mathcal{T}$ est l'ensemble des valeurs possibles d'un noyau de transitions en un point donné. On définit un milieu aléatoire par un ensemble de variables aléatoires indexées par $\mathbb{Z}^d$ et à valeurs dans $\mathcal{T}$. Intuitivement on place sur chaque point de $\mathbb{Z}^d$ une règle (aléatoire) donnant la nouvelle vitesse en fonction de l'ancienne. On dit que l'environnement est invariant par translation si les variables sont iid. Dans ce manuscrit, on ne considérera que des environnements invariants par translation. Une fois l'environnement aléatoire généré, on peut lancer une marche aléatoire cinématique dans cet environnement.

2.1.4 (Re-)Définition du modèle des miroirs

On redéfinit ici le modèle des miroirs, qui est une marche *déterministe* en milieu aléatoire, ce qui est un cas particulier de MAMA. Intuitivement on peut résumer le modèle des miroirs à quatre propriétés fondamentales :

- En tout point, la nouvelle vitesse de la particule est déterminée par une règle *déterministe*, choisie aléatoirement à l'instant initial.
- Cette règle est une permutation de l'ensemble des vitesses, la dynamique est donc injective, on peut alors calculer toutes les positions passées et futures à partir de la position présente.
- La dynamique est réversible.
- Sous les trois conditions précédentes, l'environnement est choisi uniformément.

On va donc déterminer l'ensemble des règles de changement de vitesse possibles en tout point du modèle des miroirs. Soit $S_{\mathcal{P}}$ l'ensemble des permutations de $\mathcal{P}$, pour tout $\pi \in S_{\mathcal{P}}$ on lui associe une règle T^π définie par :

$$T^\pi(\mathbf{p}, \mathbf{p}') = \begin{cases} 1 & \text{si } \mathbf{p}' = \pi(\mathbf{p}) \\ 0 & \text{sinon} \end{cases}$$

De plus, les règles du modèle des miroirs sont réversibles, cela se traduit par une condition de réversibilité sur les permutations correspondantes, cette condition est la suivante :

$$\pi^{-1}(-\mathbf{p}) = -\pi(\mathbf{p}). \quad (2.3)$$

On ajoute aussi que la marche ne peut pas faire demi-tour :

$$\pi(\mathbf{p}) \neq -\mathbf{p}. \quad (2.4)$$

On définit donc de l'ensemble Π suivant des permutations possibles du modèle des miroirs :

$$\Pi = \{\pi \in S_{\mathcal{P}} | \forall \mathbf{p} \in \mathcal{P}, \pi(\mathbf{p}) \neq -\mathbf{p}, \pi(-\pi(\mathbf{p})) = -\mathbf{p}\}$$

une variante de cet ensemble a déjà été utilisée au chapitre 1 afin de définir le modèle des miroirs avec les conventions du chapitre 1.

Pour définir le modèle des miroirs il suffit alors de choisir uniformément une permutation $\pi_{\mathbf{q}}$ pour chaque position $\mathbf{q}$ et de considérer l'application T correspondante. Plus précisément, soit :

$$\mathcal{T}_{mir} := \{T^\pi | \pi \in \Pi\} \quad (2.5)$$

l'ensemble des règles possibles correspondantes au modèle des miroirs. Alors ce modèle est défini en posant le noyau de transition T tel que les variables $T_{\mathbf{q}}$ soient iid et de loi uniforme sur $\mathcal{T}_{mir}$. Ce modèle est satisfaisant pour un physicien, mais il présente des difficultés mathématiques dont nous aurons un aperçu à la fin du chapitre 3. Mathématiquement il est intéressant d'étudier un modèle plus simple dans un premier temps.

2.1.5 Un autre modèle naturel : le modèle des permutations

Le modèle des miroirs présente une dynamique réversible, ce qui du point de vu physique est très significatif. Mais d'un point de vue mathématique ce n'est pas le modèle le plus général ni

le plus simple. Après avoir pris du recul sur le modèle des miroirs il devient évident que pour le simplifier il suffit de faire sauter la condition de réversibilité. On obtient alors un autre modèle déterministe où chaque règle de modification des vitesses est encore basée sur une permutation mais où toutes les permutations possibles sont acceptées. On appelle donc ce modèle le *modèle des permutations*. Le noyau de transition T de ce modèle est alors défini tel que les variables $T_{\mathbf{q}}$ soient iid et de loi uniforme sur T_{perm} . Avec :

$$\mathcal{T}_{perm} := \{T^\sigma | \sigma \in \mathcal{S}_P\}. \quad (2.6)$$

Pour différencier les deux modèles on utilisera la lettre π pour désigner les permutations utilisées dans le modèle des miroirs et la lettre σ pour désigner les permutations dans le modèle du même nom.

Remarque sur le nom de ce modèle : Le nom "modèle des permutations" est apparu assez naturellement et il est clair. Cependant, il manque de spécificité et il pourrait être utile sur le plan de la communication d'avoir un nom plus original. Si un jour quelqu'un étudie à nouveau ce modèle il ou elle pourrait l'appeler *Le modèle des échangeurs*. L'idée vient de la non-réversibilité, dans un échangeur auto-routier vous êtes sur une voie à sens unique, donc le sens dans lequel vous parcourez le chemin est déterminant. Mais il s'agit ici d'échangeurs particuliers, dans lesquels vous n'avez pas le choix de la route que vous prendrez. Il est sans doute possible d'inventer encore d'autres noms pour ce modèle, je souhaite de l'inspiration à ceux qui voudront travailler dessus.

Ce modèle étant plus simple mathématiquement c'est sur celui-ci que nous allons présenter notre démarche d'approche perturbative du coefficient de diffusion. Mais pour l'instant, discutons à quel point ce modèle est proche d'un autre modèle très important en probabilité : la marche aléatoire auto-évitante.

2.1.6 Lien entre les modèles étudiés et les marches auto-évitantes

En explorant quelque peu les deux modèles présentés précédemment on repère que les trajectoires sont soit des trajectoires ouvertes soit des boucles auto-évitantes. En effet, la trajectoire est déterministe dans un environnement aléatoire. Une fois l'environnement déterminé on peut donc calculer toutes les positions futures de la particule. Et donc si la particule repasse deux fois au même point $(\mathbf{q}, \mathbf{p})$ de l'espace des phases alors elle va refaire exactement le même trajet et ce une infinité de fois. On peut donc voir que la particule peut être bloquée dans une boucle auto-évitante, ou alors dans une trajectoire ouverte auto-évitante. On peut donc considérer que le modèle des permutations est une forme de marche aléatoire auto-évitante sur l'espace des phases $\mathcal{M}$. Il y a une bijection simple entre M et l'ensemble des arêtes orientées de $\mathbb{Z}^d$. La voici :

$$(\mathbf{q}, \mathbf{p}) \mapsto (\mathbf{q} - \mathbf{p}, \mathbf{q})$$

le point $(\mathbf{q}, \mathbf{p})$ correspond à l'arête orientée que la particule a utilisée pour arriver au point $(\mathbf{q}, \mathbf{p})$. Ainsi on peut dire que le modèle des permutations est une marche aléatoire auto-évitante sur les arêtes orientées de $\mathbb{Z}^d$.

Pour le modèle des miroirs la condition de réversibilité associée à l'absence de demi-tour a pour conséquence qu'une arête empruntée dans un sens ne pourra pas être empruntée dans l'autre sens. C'est facile à prouver. Ainsi on peut dire que le modèle des miroirs est une marche aléatoire auto-évitante sur les arêtes non-orientées de $\mathbb{Z}^d$.

Dans les deux cas, le comportement de la marche est particulier au niveau de l'origine, mais pour les autres positions l'analogie est parfaite. C'est pourquoi nous avons utilisé la *lace expansion* de Gordon Slade pour étudier ces modèles. On notera que nos modèles sont des processus aléatoires, contrairement à la marche aléatoire auto-évitante classique. Ils sont plus proches de la marche myope, aussi nommée *true self-avoiding random walk* [1, 4].

2.1.7 Versions faiblement corrélées des modèles, marche faiblement auto-évitante

Ces modèles sont difficiles à étudier, on ne sait jamais comment commencer les preuves, c'est pourquoi nous avons cherché à les simplifier encore plus et pour ça nous nous sommes inspirés de l'analogie avec la marche aléatoire auto-évitante. Une simplification connue de cette marche est la *marche aléatoire faiblement auto-évitante* [34] qui est un modèle intermédiaire entre la marche auto-évitante et la marche aléatoire simple, avec un paramètre ε qui permet une transition continue entre les deux modèles. Nous avons donc cherché par analogie à créer une version, disons, "faiblement aléatoire" de nos modèles. Et finalement nous avons trouvé une manière satisfaisante de le faire en utilisant une sorte de combinaison linéaire de nos modèles déterministes avec des marches aléatoires markoviennes. Voici comment nous procédons.

Pour le modèle des permutations on constate d'abord que si l'on découvre une permutation pour la première fois alors elle nous envoie uniformément sur n'importe quel point voisin. Donc, si l'on supprimait la mémoire au fur et à mesure alors la particule suivrait une *marche aléatoire cinématique symétrique* définie à la sous-section 2.1.2. On dira alors que la *marche aléatoire cinématique symétrique* est la marche aléatoire associée au modèle des permutations. L'idée est alors de combiner les deux processus ainsi : on pose $\varepsilon \in [0, 1]$ et à chaque pas on tire une pièce déséquilibrée. Si la pièce fait face (avec une probabilité ε) alors on fait un pas en suivant la règle donnée par la permutation qui est au point où on est. Mais si la pièce fait pile (avec une probabilité $1 - \varepsilon$) alors on ignore l'environnement et on fait un pas de marche aléatoire cinématique symétrique. Ce processus est une marche aléatoire cinématique en milieu aléatoire dont le noyau de transition est :

$$T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) = \varepsilon \delta_{\bar{\mathbf{p}}, \sigma_{\mathbf{q}}(\mathbf{p})} + \frac{1 - \varepsilon}{2d} \quad (2.7)$$

où la permutation $\sigma_{\mathbf{q}}$ est une permutation aléatoire tirée uniformément sur $\mathcal{S}_{\mathcal{P}}$. On notera que $\delta_{\mathbf{p}', \sigma_{\mathbf{q}}(\mathbf{p})}$ est le noyau de transition du modèle des permutations et que $\frac{1}{2d}$ est celui de la marche aléatoire cinématique symétrique. On appelle le modèle ainsi défini le modèle des permutations faiblement aléatoire. Quand il n'y a pas d'ambiguïté on dira juste "modèle des permutations" pour faire court.

Pour le modèle des miroirs la marche aléatoire associée est la marche aléatoire non-rebrousante vue aussi à la sous-section 2.1.2. On étudiera donc la MAMA associée au modèle des miroirs qui, pour $\varepsilon \in [0, 1]$ fixé a le noyau de transition suivant :

$$T_{\mathbf{q}}(\mathbf{p}, \bar{\mathbf{p}}) = \varepsilon \delta_{\bar{\mathbf{p}}, \pi_{\mathbf{q}}(\mathbf{p})} + (1 - \varepsilon) \frac{1 - \delta_{-\mathbf{p}, \bar{\mathbf{p}}}}{2d - 1} \quad (2.8)$$

où la permutation $\pi_{\mathbf{q}}$ est une permutation aléatoire tirée uniformément sur Π . Nous n'aurons malheureusement que peu de temps pour aborder ce modèle mais on en dira quelques mots au chapitre 3. En attendant nous allons étudier le modèle des permutations faiblement aléatoire et développer une approche perturbative du coefficient de diffusion qui est la contribution principale de cette thèse. On ne traitera que le modèle des permutations mais le modèle des miroirs est très probablement analysable selon la même approche.

Remarque importante : Pour $\varepsilon < 1$ les deux modèles satisfont les conditions de [7], par conséquent ces modèles sont diffusifs. Par contre, dans cet article, Bálint Tóth utilise une approche radicalement différente de la nôtre et ne donne aucune piste pour le calcul de la valeur du coefficient de diffusion. On retiendra seulement que le coefficient de diffusion existe.

2.2 Approche perturbative du coefficient de diffusion pour le modèle des permutations dans $\mathbb{Z}^d$

2.2.1 Étude du déplacement carré moyen

On se place dans le modèle des permutations dans lequel on cherche à calculer le déplacement carré moyen. On lance une particule depuis l'origine, on mesure sa position $\mathbf{q}$ au temps n et on cherche l'espérance de $\mathbf{q}^2$.

La quantité qui nous intéresse est le déplacement carré moyen au temps n . Celui-ci dépend de σ et on va s'intéresser simplement à son espérance selon la loi $\mathbb{Q}$ qui définit la variable σ . Si le système est diffusif alors il existe un coefficient de diffusion $\kappa \in \mathbb{R}_+^*$ tel que le déplacement carré moyen au temps n est équivalent à $n\kappa$ quand n tend vers $+\infty$. On pose donc : $p_n(x, y, \sigma, \varepsilon)$ la probabilité que la particule lancée en x au temps 0, dans l'environnement σ , atteigne y au temps n . Pour embellir les notations on pose aussi que si $x = (\mathbf{q}, \mathbf{p})$ alors $\|x\| := \|\mathbf{q}\|$ et on pose $O := (\mathbf{e}_1, \mathbf{e}_1)$. On verra plus tard pourquoi on a choisi $(\mathbf{e}_1, \mathbf{e}_1)$ comme origine et pas $(0, \mathbf{e}_1)$ ou autre chose. On étudie alors la quantité suivante :

$$\begin{aligned} K(n) &:= \frac{1}{n} \mathbb{E}_{\mathbb{Q}} \left[\sum_{x \in \mathcal{M}} \|x\|^2 p_n(O, x, \sigma, \varepsilon) \right] \\ &= \frac{1}{n} \sum_{x \in \mathcal{M}} \|x\|^2 \mathbb{E}_{\mathbb{Q}}[p_n(O, x, \sigma, \varepsilon)] \end{aligned} \quad (2.9)$$

On cherche à prouver qu'il existe $\kappa \in \mathbb{R}$ tel que $K(n) \rightarrow \kappa$ quand $n \rightarrow +\infty$.

Remarque : L'échange de la somme et de l'espérance est possible parce que la somme sur $x \in \mathcal{M}$ est ici de fait une somme finie, étant donné que tous les termes sont nuls pour $\|x\| > n$. Ce sera souvent le cas dans la suite et on ne le précisera plus.

Le noyau de transition de la chaîne pour un environnement σ donné est :

$$p((\mathbf{q}, \mathbf{p}), (\mathbf{q}', \mathbf{p}'), \sigma, \varepsilon) = \delta_{\mathbf{q}', \mathbf{q} + \mathbf{p}'} \left(\varepsilon \delta_{\mathbf{p}', \sigma_q(\mathbf{p})} + \frac{1 - \varepsilon}{2d} \right) \quad (2.10)$$

On a alors, pour $x_0, x_n \in \mathcal{M}$:

$$p_n(x_0, x_n, \sigma, \varepsilon) := \sum_{x_1, \dots, x_{n-1} \in \mathcal{M}} \prod_{k=0}^{n-1} p(x_k, x_{k+1}, \sigma, \varepsilon). \quad (2.11)$$

Pour $x, y \in \mathcal{M}$ on définit $p(x, y)$ comme le noyau de transition pour $\varepsilon = 0$, ce noyau ne dépend alors plus de σ , donc :

$$p((\mathbf{q}, \mathbf{p}), (\mathbf{q}', \mathbf{p}')) := \frac{\delta_{\mathbf{q}', \mathbf{q} + \mathbf{p}'}}{2d} \quad (2.12)$$

On constate que l'on peut écrire :

$$p(x, y, \sigma, \varepsilon) = p(x, y) + \varepsilon \hat{p}(x, y, \sigma). \quad (2.13)$$

Avec :

$$\hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma) := \delta_{\bar{\mathbf{q}}, \mathbf{q} + \bar{\mathbf{p}}} \left(\delta_{\bar{\mathbf{p}}, \sigma_q(\mathbf{p})} - \frac{1}{2d} \right) \quad (2.14)$$

Dans (2.13) le terme $\hat{p}$ doit être vu comme une perturbation. On a alors 3 propriétés générales que l'on résume en une seule.

Proposition 2.1.

$$\sum_{x \in \mathcal{M}} \hat{p}(x, y, \sigma) = \sum_{y \in \mathcal{M}} \hat{p}(x, y, \sigma) = \mathbb{E}_{\mathbb{Q}}[\hat{p}(x, y, \sigma)] = 0 \quad (2.15)$$

La preuve est très simple en écrivant simplement le calcul. Le facteur $\delta_{\bar{\mathbf{q}}, \mathbf{q} + \bar{\mathbf{p}}}$ dans (2.14) permet, dans le cadre du modèle des permutations, de réécrire plus spécifiquement que :

$$\sum_{\mathbf{p} \in \mathcal{P}} \hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma) = \sum_{\bar{\mathbf{p}} \in \mathcal{P}} \hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma) = 0 \quad (2.16)$$

Pour alléger les calculs on va abréger les sommes du style $\sum_{x_1, \dots, x_n \in \mathcal{M}}$ en $\sum_{\underline{x}}$. Cette notation signifie que l'on somme sur toutes les variables x qui apparaissent dans le membre de droite de l'égalité mais pas dans le membre de gauche. On trouve donc que pour $x_0, x_n \in \mathcal{M}$:

$$\begin{aligned} p_n(x_0, x_n, \sigma, \varepsilon) &= \sum_{\underline{x}} \prod_{k=0}^{n-1} p(x_k, x_{k+1}, \sigma, \varepsilon). \\ &= \sum_{\underline{x}} \prod_{k=0}^{n-1} \varepsilon \hat{p}(x_k, x_{k+1}, \sigma) + p(x_k, x_{k+1}). \\ &= \sum_{\underline{x}} \sum_{m=0}^n \sum_{\substack{P \subset \llbracket 0, n-1 \rrbracket \\ \|p\|=m}} \prod_{k \in P} \varepsilon \hat{p}(x_k, x_{k+1}, \sigma) \prod_{k \in P^c} p(x_k, x_{k+1}). \end{aligned} \quad (2.17)$$

On rappelle la définition de la notation $\llbracket a, b \rrbracket$ pour les intervalles d'entiers :

$$\forall a, b \in \mathbb{R}, \llbracket a, b \rrbracket := [a, b] \cap \mathbb{Z}. \quad (2.18)$$

Nous utiliserons beaucoup cette notation.

Si $P = \{k_1, \dots, k_m\} \subset \llbracket 0, n-1 \rrbracket$ alors on pose $\forall i \in \llbracket 0, m \rrbracket, y_i := x_{k_i}$ et $\bar{y}_i := x_{k_i+1}$. On pose aussi :

$$p_n(x_0, x_n) := \sum_{\underline{x}} \prod_{k=0}^{n-1} p(x_k, x_{k+1}) \quad (2.19)$$

qui est la probabilité que la marche aléatoire simple cinématique atteigne x_n en n pas en partant de x_0 . On renomme aussi les points de départ et d'arrivée : $\bar{y}_0 := x_0$ et $y_{m+1} := x_n$. On trouve dans ce cas, après quelques lignes de calcul :

$$p_n(\bar{y}_0, y_{m+1}, \sigma, \varepsilon) = \sum_{m=0}^n \varepsilon^m \sum_{\underline{y}} \sum_{\substack{n_0, \dots, n_m \\ n_0 + \dots + n_m = n-m}} \prod_{i=1}^m \hat{p}(y_i, \bar{y}_i, \sigma) \prod_{i=0}^m p_{n_i}(\bar{y}_i, y_{i+1}). \quad (2.20)$$

La somme $\sum_{\underline{y}}$ est à la fois sur les indices y et $\bar{y}$. On verra plus tard que $p_n(x, y)$ est facile à calculer à partir des résultats standards sur la marche aléatoire simple dans $\mathbb{Z}^d$. On souhaite aussi une notation abrégée pour la somme sur les n_i , mais il faut garder l'idée que cette somme est contrainte. Voyons donc $\underline{n}$ comme un multi-indice, en l'occurrence un vecteur de $\mathbb{N}^{m+1}$. On considère que l'on peut sommer sur $\underline{n} \in \mathcal{S}_n$ où $\mathcal{S}_n$ est la sphère de rayon n pour la norme 1. Cette convention de sommation permet de garder la contrainte $n_0 + \dots + n_m = n - m$ sous la forme très concise $\underline{n} \in \mathcal{S}_n$. On omet volontairement de préciser la dimension de l'espace dans lequel vit la sphère $\mathcal{S}_n$ par souci de concision. Cela revient à ne pas préciser combien d'indices sont sommés.

On s'intéresse à l'espérance de p_n , qui vaut :

$$\mathbb{E}_{\mathbb{Q}}[p_n(\bar{y}_0, y_{m+1}, \sigma, \varepsilon)] = \sum_{m=0}^n \varepsilon^m \sum_{\underline{y}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}(y_i, \bar{y}_i, \sigma) \right] \prod_{i=0}^m p_{n_i}(\bar{y}_i, y_{i+1}). \quad (2.21)$$

On note que le terme $m = 0$ vaut $p_n(\bar{y}_0, y_{m+1})$ et que le terme $m = 1$ est nul d'après la proposition 2.1. Calculons rapidement le terme $m = 0$.

On définit d'abord $g_n(\mathbf{q})$ comme étant la probabilité qu'une marche aléatoire simple partant de zéro sur $\mathbb{Z}^d$ arrive au point $\mathbf{q}$ au temps n . Ici, il ne s'agit pas d'une marche cinématique mais d'une marche aléatoire classique, cette quantité est donc bien connue.

On a alors l'égalité :

$$p_n((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}})) = \delta_n \delta_{\mathbf{q}, \bar{\mathbf{q}}} \delta_{\mathbf{p}, \bar{\mathbf{p}}} + \frac{1 - \delta_n}{2d} g_{n-1}(\bar{\mathbf{q}} - \bar{\mathbf{p}} - \mathbf{q}). \quad (2.22)$$

L'idée de cette équation est que si $n = 0$ alors seul le chemin trivial de zéro pas contribue, c'est le premier terme. Et si $n > 0$ alors on peut décomposer le chemin en deux parties. D'abord on va du point $\mathbf{q}$ au point $\bar{\mathbf{q}} - \bar{\mathbf{p}}$ en $n - 1$ pas, la vitesse n'important pas, la probabilité de faire cela est $g_{n-1}(\bar{\mathbf{q}} - \bar{\mathbf{p}} - \mathbf{q})$. Puis on fait un dernier pas pour atteindre le point $\bar{\mathbf{q}}$, avec probabilité $1/2d$. La vitesse à l'arrivée est alors $\bar{\mathbf{p}}$. On a donc le second terme.

On a donc, pour $n \geq 1$:

$$\frac{1}{n} \sum_{y_{m+1} \in \mathcal{M}} \|y_{m+1}\| p_n(O, y_{m+1}) = \frac{1}{n} \sum_{\mathbf{q} \in \mathbb{Z}^d} \sum_{\mathbf{p} \in \mathcal{P}} \|\mathbf{q}\| p_n((\mathbf{e}_1, \mathbf{e}_1), (\mathbf{q}, \mathbf{p})) \quad (2.23)$$

$$= \frac{1}{n} \sum_{\mathbf{q} \in \mathbb{Z}^d} \|\mathbf{q}^2\| \sum_{\mathbf{p} \in \mathcal{P}} \delta_n \delta_{\mathbf{e}_1, \mathbf{q}} \delta_{\mathbf{e}_1, \mathbf{p}} + \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_1 - \mathbf{p}) \quad (2.24)$$

$$= \frac{1}{n} \sum_{\mathbf{q} \in \mathbb{Z}^d} \|\mathbf{q}^2\| \sum_{\mathbf{p} \in \mathcal{P}} \frac{1}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_1 - \mathbf{p}) \quad (2.25)$$

$$= \frac{1}{n} \sum_{\mathbf{q} \in \mathbb{Z}^d} \|\mathbf{q}^2\| g_n(\mathbf{q} - \mathbf{e}_1) \quad (2.26)$$

$$= \frac{n+1}{n} \quad (2.27)$$

$$(2.28)$$

On a donc :

$$K(n) = \frac{n+1}{n} + \frac{1}{n} \sum_{m=2}^n \varepsilon^m \sum_{\underline{y}} \|y_{m+1}\|^2 \sum_{\underline{n} \in \mathcal{S}_n} p_{n_0}(O, y_1) \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}(y_i, \bar{y}_i, \sigma) \right] \prod_{i=1}^m p_{n_i}(\bar{y}_i, y_{i+1}). \quad (2.29)$$

Pour $m \geq 2$ on définit une quantité importante :

$$\mathcal{B}_{m,n}(y_1, \bar{y}_m) := \sum_{\underline{y}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}(y_i, \bar{y}_i, \sigma) \right] \prod_{i=1}^{m-1} p_{n_i}(\bar{y}_i, y_{i+1}) \quad (2.30)$$

Dans la suite, sauf précision contraire, on supposera toujours que $m \geq 2$. Avec cette définition on peut écrire :

$$\mathbb{E}_{\mathbb{Q}}[p_n(O, y_{m+1}, \sigma, \varepsilon)] = \sum_{m=0}^n \varepsilon^m \sum_{\underline{y}} \sum_{\underline{n} \in \mathcal{S}_{n-m}} p_{n_0}(O, y_1) \mathcal{B}_{m,n_1}(y_1, \bar{y}_m) p_{n_2}(\bar{y}_m, y_{m+1}). \quad (2.31)$$

On note que la convention de somme sur $\underline{y}$ fonctionne encore quand on ne somme que sur deux variables. La variable y_{m+1} n'est pas sommée car elle apparaît des deux côtés du signe égal. On a alors :

$$K(n) = \frac{n+1}{n} + \frac{1}{n} \sum_{m=2}^n \varepsilon^m \sum_{\underline{y}} \|y_{m+1}\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} p_{n_0}(O, y_1) \mathcal{B}_{m,n_1}(y_1, \bar{y}_m) p_{n_2}(\bar{y}_m, y_{m+1}). \quad (2.32)$$

On va séparer les sommes sur $\mathcal{M}$ en sommes sur $\mathbb{Z}^d$ et sur $\mathcal{P}$, c'est un peu lourd au début mais après on pourra introduire des notations matricielles confortables. On remplace y_i par $(\mathbf{q}_i, \mathbf{p}_i)$ et $\bar{y}_i$ par $(\bar{\mathbf{q}}_i, \bar{\mathbf{p}}_i)$. On a alors :

$$\mathcal{B}_{m,n}(y_1, \bar{y}_m) = \sum_{\underline{\mathbf{q}}} \sum_{\underline{\mathbf{p}}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\bar{\mathbf{q}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \prod_{i=1}^{m-1} p_{n_i}((\bar{\mathbf{q}}_i, \bar{\mathbf{p}}_i), (\mathbf{q}_{i+1}, \mathbf{p}_{i+1})). \quad (2.33)$$

À cause du facteur $\delta_{\mathbf{q}+\bar{\mathbf{p}}, \bar{\mathbf{q}}}$ dans $\hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma)$ on a toujours $\bar{\mathbf{q}} = \mathbf{q} + \bar{\mathbf{p}}$. D'où :

$$\mathcal{B}_{m,n}(y_1, \bar{y}_m) = \sum_{\underline{\mathbf{q}}} \sum_{\underline{\mathbf{p}}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \prod_{i=1}^{m-1} p_{n_i}((\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), (\mathbf{q}_{i+1}, \mathbf{p}_{i+1})). \quad (2.34)$$

Introduisons les notations matricielles qui embellissent les calculs.

2.2.2 Notations matricielles

Remarquons d'abord que $p_n((\bar{\mathbf{q}}, \bar{\mathbf{p}}), (\mathbf{q}, \mathbf{p}))$ est invariant par translation simultanée de $\mathbf{q}$ et $\bar{\mathbf{q}}$. Pour tout $n \in \mathbb{N}$ on définit donc la matrice $P_n(\mathbf{q})$ de dimension $(2d) \times (2d)$ telle que :

$$P_n^{ij}(\mathbf{q}) := p_n((\mathbf{e}_i, \mathbf{e}_i), (\mathbf{q}, \mathbf{e}_j)). \quad (2.35)$$

Ainsi avec $\mathbf{p}_i = \mathbf{e}_{j_i}$ et $\bar{\mathbf{p}}_i = \mathbf{e}_{\bar{j}_i}$:

$$p_n(\bar{y}_i, y_{i+1}) = p((\bar{\mathbf{q}}_i, \bar{\mathbf{p}}_i), (\mathbf{q}_{i+1}, \mathbf{p}_{i+1})) \quad (2.36)$$

$$= P^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.37)$$

Cette définition est choisie en fonction des termes qui apparaissent dans nos équations, pour rendre celles-ci les plus jolies possibles.

De la même manière on définit :

$$\mathcal{B}_{m,n}^{ij}(\mathbf{q}) := \mathcal{B}_{m,n}((0, \mathbf{e}_i), (\mathbf{q} + \mathbf{e}_j, \mathbf{e}_j)). \quad (2.38)$$

Ainsi :

$$\mathcal{B}_{m,n_1}(y_1, \bar{y}_m) = \mathcal{B}_{m,n_1}((\mathbf{q}_1, \mathbf{p}_1), (\bar{\mathbf{q}}_m, \bar{\mathbf{p}}_m)) \quad (2.39)$$

$$= \mathcal{B}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_m - \mathbf{q}_1) \quad (2.40)$$

C'est possible car on a aussi l'invariance par translation pour $\mathcal{B}$. On a alors :

$$K(n) = \frac{n+1}{n} + \frac{1}{n} \sum_{m=0}^n \varepsilon^m \sum_{\underline{\mathbf{q}}} \sum_{\underline{j}} \|\mathbf{q}_{m+1}\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} P_{n_0}^{1j_1}(\mathbf{q}_1) \mathcal{B}_{m,n_1}^{j_1 \bar{j}_m}(\mathbf{q}_m - \mathbf{q}_1) P_{n_2}^{\bar{j}_m j_{m+1}}(\mathbf{q}_{m+1} - \mathbf{q}_m). \quad (2.41)$$

Les indices j sont sommés sur $\mathcal{P}_{ind} := \llbracket -d, d \rrbracket \setminus \{0\}$ qui est l'ensemble des indices possibles des vecteurs de $\mathcal{P}$. Par conséquent, pour $m \geq 2$ on peut définir la matrice :

$$X_m(n) := \frac{1}{n} \sum_{\underline{\mathbf{q}}} \|\mathbf{q}_{m+1}\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_m - \mathbf{q}_1) P_{n_2}(\mathbf{q}_{m+1} - \mathbf{q}_m) \quad (2.42)$$

et avec un petit changement de variable :

$$X_m(n) = \frac{1}{n} \sum_{\underline{\mathbf{q}}} \|\mathbf{q}_1 + \mathbf{q}_2 + \mathbf{q}_3\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \quad (2.43)$$

Avec cette matrice on peut réécrire $K(n)$ simplement :

$$K(n) = \frac{n+1}{n} + \sum_{m=0}^n \varepsilon^m \sum_j X_m^{1j}(n) \quad (2.44)$$

En posant par convention $K_0(n) = (n+1)/n$ et $K_1(n) = 0$ on a alors une belle décomposition :

$$K(n) = \sum_{m=0}^n \varepsilon^m K_m(n) \quad (2.45)$$

avec

$$K_m(n) = \sum_j X_m^{1j}(n) \quad (2.46)$$

On aimerait prouver que la suite $K(n)$ est convergente, mais on va se contenter de prouver le théorème suivant :

Théorème 2.1. $\forall d \geq 3, \forall m \in \mathbb{N}, \exists \kappa_m \in \mathbb{R}$ tel que :

$$\kappa_m = \lim_{n \rightarrow +\infty} K_m(n) \quad (2.47)$$

Avant d'entamer la preuve à proprement parler de ce théorème on a plusieurs notions à introduire.

On avait :

$$K_m(n) = \sum_j X_m^{1j}(n) \quad (2.48)$$

À ce moment, il est utile de constater que la valeur de $K_m(n)$ est invariante par rotation du point de départ, autrement dit, $\forall i \in \mathcal{P}_{ind}$:

$$K_m(n) = \sum_{j \in \mathcal{P}_{ind}} X_m^{ij}(n).$$

D'où :

$$K_m(n) = \frac{1}{2d} \sum_{i,j} X_m^{ij}(n), \quad (2.49)$$

$$= \frac{1}{2d} \vec{1}^T X_m(n) \vec{1} \quad (2.50)$$

Où $\vec{1}$ est le vecteur de dimension $2d$ qui ne contient que des 1 et $\vec{1}^T$ son transposé. À partir de cette expression on peut progresser dans le calcul de $K_m(n)$.

$$\begin{aligned} K_m(n) &= \frac{1}{2d} \vec{1}^T X_m(n) \vec{1} \\ &= \frac{1}{n(2d)} \sum_{\underline{\mathbf{q}}} \|\mathbf{q}_1 + \mathbf{q}_2 + \mathbf{q}_3\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} \vec{1}^T P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \vec{1} \end{aligned} \quad (2.51)$$

Pour aller plus loin dans ce calcul nous avons besoin de l'expression de $P_n^{ij}(\mathbf{q})$ puis d'une propriété importante sur $\mathcal{B}_{m,n}^{ij}(\mathbf{q})$. D'après la définition de la matrice P à l'équation (2.35) on a :

$$P_n^{ij}(\mathbf{q}) := p_n((\mathbf{e}_i, \mathbf{e}_i), (\mathbf{q}, \mathbf{e}_j)). \quad (2.52)$$

Et donc, d'après (2.22), on a :

$$P_n^{ij}(\mathbf{q}) = \delta_n \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_j, \mathbf{q}} + \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j). \quad (2.53)$$

On en déduit facilement que :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} P_n^{ij}(\mathbf{q}) = \delta_n \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1 - \delta_n}{2d}. \quad (2.54)$$

Puis :

$$\sum_{i \in \mathcal{P}_{ind}} \sum_{\mathbf{q} \in \mathbb{Z}^d} P_n^{ij}(\mathbf{q}) = \sum_{j \in \mathcal{P}_{ind}} \sum_{\mathbf{q} \in \mathbb{Z}^d} P_n^{ij}(\mathbf{q}) = 1, \quad (2.55)$$

En définissant C la matrice qui ne contient que des 1 : $\forall i, j \in \mathcal{P}_{ind}, C^{ij} := 1$ on peut réécrire (2.54) ainsi :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} P_n(\mathbf{q}) = \delta_n I + \frac{1 - \delta_n}{2d} C. \quad (2.56)$$

Par ailleurs, les deux égalités de (2.55) s'écrivent matriciellement ainsi :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} \vec{1}^T P_n(\mathbf{q}) = \vec{1}^T \quad (2.57)$$

et

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} P_n(\mathbf{q}) \vec{1} = \vec{1} \quad (2.58)$$

Voici à présent la propriété importante de la matrice $\mathcal{B}$:

Proposition 2.2. $\forall m \in \mathbb{N}, \forall n \in \mathbb{N}, \forall \mathbf{q} \in \mathbb{Z}^d, \forall i \in \mathcal{P}_{ind},$

$$\sum_j \mathcal{B}_{m,n}^{ij}(\mathbf{q}) = \sum_j \mathcal{B}_{m,n}^{ji}(\mathbf{q}) = 0.$$

Ou matriciellement :

$$\mathcal{B}_{m,n}(\mathbf{q}) \vec{1} = \vec{0} \quad (2.59)$$

et

$$\vec{1}^T \mathcal{B}_{m,n}(\mathbf{q}) = \vec{0}^T \quad (2.60)$$

Preuve : À partir des définitions (2.30) et (2.38) on a :

$$\mathcal{B}_{m,n}^{i_1 \bar{j}_m}(\mathbf{q}_m - \mathbf{q}_1) = \sum_{\underline{q}} \sum_{\underline{j}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\bar{\mathbf{q}}_i, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.61)$$

Ainsi, (2.16) permet de conclure. En fait la propriété vient directement de $\hat{p}$ et s'étend directement à $\mathcal{B}_{m,n}$.

□

On reprend l'expression :

$$K_m(n) = \frac{1}{n(2d)} \sum_{\underline{q}} \|\mathbf{q}_1 + \mathbf{q}_2 + \mathbf{q}_3\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} \vec{1}^T P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \vec{1} \quad (2.62)$$

et on écrit : $\|\mathbf{q}_1 + \mathbf{q}_2 + \mathbf{q}_3\|^2 = \|\mathbf{q}_1\|^2 + \|\mathbf{q}_2\|^2 + \|\mathbf{q}_3\|^2 + 2\mathbf{q}_1 \cdot \mathbf{q}_2 + 2\mathbf{q}_1 \cdot \mathbf{q}_3 + 2\mathbf{q}_2 \cdot \mathbf{q}_3$ ce qui nous permet d'écrire $K_m(n)$ comme une somme de 6 termes. Le premier terme, celui correspondant à $\|\mathbf{q}_1\|^2$ est :

$$\frac{1}{n(2d)} \sum_{\underline{q}} \|\mathbf{q}_1\|^2 \sum_{\underline{n} \in \mathcal{S}_{n-m}} \vec{1}^T P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \vec{1}.$$

Une application directe de l'équation (2.58) puis de la propriété (2.2) nous permet de montrer que ce terme est nul car $\|\mathbf{q}_1\|^2$ ne dépend pas de $\mathbf{q}_3$, un raisonnement semblable en appliquant (2.57) nous permet de démontrer que le terme en $\|\mathbf{q}_3\|^2$ est nul aussi car $\|\mathbf{q}_3\|^2$ ne dépend pas de $\mathbf{q}_1$. Finalement, seul le terme en $\mathbf{q}_1 \cdot \mathbf{q}_3$ est non nul. Donc :

$$K_m(n) = \frac{1}{nd} \sum_{\underline{\mathbf{q}}} \mathbf{q}_1 \cdot \mathbf{q}_3 \sum_{\underline{n} \in \mathcal{S}_{n-m}} \vec{1}^T P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \vec{1}.$$

On notera que c'est ici que le fait de moyenner sur toutes les vitesses possibles au niveau du point de départ est bien utile, sinon les termes qui dépendent de $\mathbf{q}_3$ mais pas de $\mathbf{q}_1$ ne se simplifieraient pas.

On calcule maintenant :

$$\begin{aligned} \sum_{\mathbf{q} \in \mathbb{Z}^d} \mathbf{q} P_n^{ij}(\mathbf{q}) &= \sum_{\mathbf{q} \in \mathbb{Z}^d} \mathbf{q} \left(\delta_n \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_j} + \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) \right) \\ &= \delta_n \delta_{\mathbf{e}_i, \mathbf{e}_j} \mathbf{e}_i + \sum_{\mathbf{q} \in \mathbb{Z}^d} (\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q}) \\ &= \delta_n \delta_{\mathbf{e}_i, \mathbf{e}_j} \mathbf{e}_i + \frac{1 - \delta_n}{2d} (\mathbf{e}_i + \mathbf{e}_j) \end{aligned} \tag{2.63}$$

D'où :

$$\sum_{i \in \mathcal{A}} \sum_{\mathbf{q} \in \mathbb{Z}^d} \mathbf{q} P_n^{ij}(\mathbf{q}) = \mathbf{e}_j$$

et

$$\sum_{j \in \mathcal{A}} \sum_{\mathbf{q} \in \mathbb{Z}^d} \mathbf{q} P_n^{ij}(\mathbf{q}) = \mathbf{e}_i.$$

Donc :

$$\begin{aligned} K_m(n) &= \frac{1}{nd} \sum_{\underline{\mathbf{q}}} \mathbf{q}_1 \cdot \mathbf{q}_3 \sum_{\underline{n} \in \mathcal{S}_{n-m}} \vec{1}^T P_{n_0}(\mathbf{q}_1) \mathcal{B}_{m,n_1}(\mathbf{q}_2) P_{n_2}(\mathbf{q}_3) \vec{1}. \\ &= \frac{1}{nd} \sum_{\underline{\mathbf{q}}} \sum_{\underline{n} \in \mathcal{S}_{n-m}} \sum_{\underline{j}} \mathbf{q}_1 P_{n_0}^{j_0 j_1}(\mathbf{q}_1) \mathcal{B}_{m,n_1}^{j_1 j_2}(\mathbf{q}_2) \cdot \mathbf{q}_3 P_{n_2}^{j_2 j_3}(\mathbf{q}_3) \\ &= \frac{1}{nd} \sum_{\mathbf{q}_2 \in \mathbb{Z}^d} \sum_{\substack{n_0, n_1, n_2 \\ n_0 + n_1 + n_2 = n - m}} \sum_{\underline{j}} \mathcal{B}_{m,n_1}^{j_1 j_2}(\mathbf{q}_2) \mathbf{e}_{j_1} \cdot \mathbf{e}_{j_2} \\ &= \frac{1}{d} \sum_{\mathbf{q}_2 \in \mathbb{Z}^d} \sum_{n_1=0}^{n-m} \frac{n - m - n_1}{n} \sum_{j_1, j_2} \mathcal{B}_{m,n_1}^{j_1 j_2}(\mathbf{q}_2) \mathbf{e}_{j_1} \cdot \mathbf{e}_{j_2} \\ &= 2 \sum_{\mathbf{q}_2 \in \mathbb{Z}^d} \sum_{n_1=0}^{n-m} \frac{n - m - n_1}{n} (\mathcal{B}_{m,n_1}^{1,1}(\mathbf{q}_2) - \mathcal{B}_{m,n_1}^{1,-1}(\mathbf{q}_2)) \end{aligned} \tag{2.64}$$

Étudions à présent la convergence de cette expression.

2.2.3 Première approche de la convergence de $K_m(n)$

On rappelle que l'on veut prouver que $\lim_{n \rightarrow +\infty} K_m(n)$ existe. Pour ça on va commencer par séparer les sommes sur $\mathbb{N}$ et les sommes sur $\mathbb{Z}^d$ qui sont mélangées dans $\mathcal{B}$. Pour cela on définit :

$$\mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_{\underline{j}} \sum_{\underline{n} \in \mathcal{S}_n} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \tag{2.65}$$

avec par convention, ici : $\mathbf{q}_1 := 0$. On rappelle que la dépendance en n est cachée dans le fait que la somme sur $\underline{n}$ contient la contrainte $\sum_i n_i = n$. On voit grâce à (2.61) que :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} \mathcal{B}_{m,n}(\mathbf{q}) = \sum_{\underline{q}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m). \quad (2.66)$$

On cherche donc à prouver que :

$$\kappa_m = \lim_{n \rightarrow +\infty} 2 \sum_{n_1=0}^{n-m} \frac{n-m-n_1}{n} \sum_{\underline{q}} (\mathcal{U}_{m,n_1}^{1,1}(\mathbf{q}_2, \dots, \mathbf{q}_m) - \mathcal{U}_{m,n_1}^{1,-1}(\mathbf{q}_2, \dots, \mathbf{q}_m)) \quad (2.67)$$

Ce qui nous donnera une manière de calculer κ_m au passage. Il s'agit d'une forme de somme de Césàro, dans le sens suivant. Si $(v_n)_n$ est une suite réelle et $(S_n)_n$ la suite de ses sommes partielles alors :

$$\sigma_n := \frac{1}{n} \sum_{k=0}^n S_k$$

est la n -ième somme de Césàro de $(v_n)_n$. Et il est facile de prouver que si la suite $(S_n)_n$ converge alors la suite $(\sigma_n)_n$ converge vers la même limite. Par ailleurs on voit bien que :

$$\sigma_n = \sum_{k=0}^n \frac{n+1-k}{n} v_k$$

ce qui ressemble fortement à (2.67). Il est donc suffisant de prouver que la série de terme général :

$$u_n := \sum_{\underline{q}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.68)$$

est convergente car il est facile de prouver que la convergence de $(u_n)_n$ implique l'existence de la limite dans (2.67) tout comme la convergence de $(S_n)_n$ implique celle de $(\sigma_n)_n$ dans l'exemple ci-dessus. Pour prouver la convergence de $(u_n)_n$ on va prouver, à l'aide du théorème de convergence dominée, que :

$$\sum_{n=0}^{+\infty} \sum_{\underline{q}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{\underline{q}} \sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.69)$$

Pour cela on a trois étapes dont seul la première est simple.

- Premièrement, on montrera que la somme $\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ est bien définie.
- Deuxièmement on montrera que cette somme est intégrable selon $(\mathbf{q}_2, \dots, \mathbf{q}_m)$.
- Enfin on trouvera une fonction de domination afin de prouver que l'égalité (2.69) est vraie.

On note que le premier et le troisième points impliquent ensemble le second, mais on prouvera quand même le second à part car cela rend la preuve plus intuitive.

Commençons par le premier point. On définit d'abord la fonction de Green de la marche aléatoire simple (non-cinématique) sur $\mathbb{Z}^d$:

$$G(\mathbf{q}) := \sum_{n=0}^{+\infty} g_n(\mathbf{q}). \quad (2.70)$$

Puis la fonction de Green de la marche aléatoire simple cinématique :

$$\mathbf{G}(\mathbf{q}) := \sum_{n=0}^{+\infty} P_n(\mathbf{q}) \quad (2.71)$$

On trouve facilement que :

$$\mathbf{G}^{ij}(\mathbf{q}) = \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) \quad (2.72)$$

On a alors la proposition suivante :

Proposition 2.3. *Pour tout $d \geq 3$ et pour tout entier naturel m et pour tout $(\mathbf{q}_2, \dots, \mathbf{q}_m) \in (\mathbb{Z}^d)^{m-1}$, la série de terme général $\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ est absolument convergente et sa somme est :*

$$\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_j \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} \mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.73)$$

Preuve : En déplaçant simplement la somme sur $\underline{n} \in \mathcal{S}_n$ dans (2.65) on trouve que

$$\mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.74)$$

On rappelle que si $d \geq 3$,

$$\sum_{n=0}^{+\infty} P_n^{ij}(\mathbf{q}) = \mathbf{G}^{ij}(\mathbf{q}) < +\infty \quad (2.75)$$

On en déduit la convergence du produit de Cauchy à termes positifs qui suit :

$$\sum_{n=0}^{+\infty} \sum_{\substack{n_1, \dots, n_{m-1} \\ n_1 + \dots + n_{m-1} = n-m}} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) = \prod_{i=1}^{m-1} \mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.76)$$

Et par conséquent :

$$\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} \mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.77)$$

Puisque la somme sur $\underline{j}$ est une somme finie et que le facteur en $\mathbb{E}_{\mathbb{Q}}$ ne dépend pas de n .

□

2.2.4 Partitions sans singleton

Attaquons le second point, l'intégrabilité de $\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ selon $\underline{\mathbf{q}}$. Voici la démarche complète de la preuve. D'abord nous allons couper $\mathcal{U}_{m,n}$ en plusieurs morceaux $\mathcal{U}_{W,n}$ indexés par des partitions sans singleton de l'intervalle d'entiers $\llbracket 1, m \rrbracket$. On dénomme ces partitions par la lettre W de manière générique. Cette démarche est inspirée des diagrammes de la *lace expansion*. On aura alors la formule :

$$\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_W \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.78)$$

Où la somme sur W est une somme finie sur l'ensemble des partitions sans singleton de $\llbracket 1, m \rrbracket$. Ensuite on va classier ces partitions en deux catégories, les *connectées* et les *non-connectées*. Puis on résoudra le problème pour les partitions *connectées*. Enfin on va ramener le problème des partitions *non-connectées* à celui des partitions *connectées* ce qui clôturera la preuve du théorème 2.1.

Pour faire la décomposition de $\mathcal{U}_{m,n}$ en termes $\mathcal{U}_{W,n}$ comme on vient de le suggérer on se concentre sur $\hat{p}$ dans l'expression de $\mathcal{U}$. On rappelle (2.14) :

$$\hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma) := \delta_{\bar{\mathbf{q}}, \mathbf{q} + \bar{\mathbf{p}}} \left(\delta_{\bar{\mathbf{p}}, \sigma_q(\mathbf{p})} - \frac{1}{2d} \right)$$

On voit que $\hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma)$ est une variable aléatoire qui ne dépend de l'environnement σ qu'à travers la permutation aléatoire $\sigma_{\mathbf{q}}$. Par conséquent, si $\mathbf{q}_1 \neq \mathbf{q}_2$ alors $\hat{p}((\mathbf{q}_1, \mathbf{p}_1), (\bar{\mathbf{q}}_1, \bar{\mathbf{p}}_1), \sigma)$ et $\hat{p}((\mathbf{q}_2, \mathbf{p}_2), (\bar{\mathbf{q}}_2, \bar{\mathbf{p}}_2), \sigma)$ sont indépendantes. Cela entraîne que s'il existe un $\mathbf{q}_i$ différent de tous les autres alors on trouve facilement que :

$$\mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] = 0$$

car

$$\mathbb{E}_{\mathbb{Q}} [\hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma)] = 0 \quad (2.79)$$

On voit donc (revoir (2.74) si besoin) que $\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ est nul sauf si

$$\forall i \in \llbracket 1, m \rrbracket, \exists j \in \llbracket 1, m \rrbracket, \mathbf{q}_i = \mathbf{q}_j.$$

Autrement dit, $\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ est nul si et seulement si il existe une partition sans singleton W de $\llbracket 1, m \rrbracket$ telle que $W = (W^1, \dots, W^n)$ où les W^i sont des parties de $\llbracket 1, m \rrbracket$ et telle que $\forall i, j \in \llbracket 1, m \rrbracket$ et $\forall k \in \llbracket 1, n \rrbracket$:

$$i, j \in W^k \Leftrightarrow \mathbf{q}_i = \mathbf{q}_j. \quad (2.80)$$

Avec par convention, $\mathbf{q}_1 = 0$. Pour une partition donnée W on définit H_W comme étant l'ensemble des $(\mathbf{q}_2, \dots, \mathbf{q}_m)$ qui respectent la condition (2.80). On définit alors :

$$\mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \mathbf{1}[(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W] \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m). \quad (2.81)$$

Remarque importante : Nous aurions pu définir H_W en incluant $\mathbf{q}_1$ dans la définition, nous avons préféré l'exclure pour avoir de notations plus concises. Mais il faudra toujours penser que si un $\mathbf{q}_1$ apparaît dans une équation alors on a $\mathbf{q}_1 = 0$ par convention. Sauf mention explicite du contraire.

Proposition 2.4. *Les H_W sont deux à deux disjoints.*

Preuve : Soient W_1 et W_2 deux partitions sans singleton différentes de $\llbracket 1, m \rrbracket$. Il existe donc i et $j \in \llbracket 1, m \rrbracket$ qui sont dans la même partie de W_1 et dans deux parties différentes de W_2 . On a alors :

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_{W_1} \Rightarrow \mathbf{q}_i = \mathbf{q}_j$$

et

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_{W_2} \Rightarrow \mathbf{q}_i \neq \mathbf{q}_j$$

avec toujours par convention $\mathbf{q}_1 = 0$. On trouve alors que :

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_{W_1} \Rightarrow (\mathbf{q}_2, \dots, \mathbf{q}_m) \notin H_{W_2}.$$

□

Corolaire 2.1.

$$\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_W \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.82)$$

où la somme est sur l'ensemble des partitions sans singleton.

Preuve : On a vu que s'il n'existe pas de partition sans singleton W telle que $(\mathbf{q}_2, \dots, \mathbf{q}_m)$ respecte la condition (2.80), alors $\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = 0$. Par définition de H_W on en déduit que

$$\mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \neq 0 \Rightarrow \exists W, (\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W$$

la partition W étant sans singleton. Les H_W étant disjoints on conclut.

□

La somme dans le corolaire ci-dessus est une somme finie. On trouve donc que l'intégrabilité de la fonction

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$$

se ramène à l'intégrabilité des fonctions

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$$

pour l'ensemble des W .

D'après la proposition 2.3 on voit que l'intégrabilité de $\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ selon $\underline{\mathbf{q}}$ dépend fortement de la décroissance de $\mathbf{G}(\mathbf{q})$ quand $\|\mathbf{q}\| \rightarrow +\infty$. On rappelle que

$$\mathbf{G}^{ij}(\mathbf{q}) = \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) \quad (2.83)$$

$\mathbf{G}^{ij}(\mathbf{q})$ a donc la même décroissance en $+\infty$ que $G(\mathbf{q})$. Il est bien connu [19] qu'en dimension $d \geq 3$ alors :

$$G(\mathbf{q}) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^{d-2}}\right) \quad (2.84)$$

Mais cette décroissance ne suffit pas pour prouver l'intégrabilité de $\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$, on va donc rentrer un peu plus dans le détail. En appliquant l'équation (2.16) à la formule démontrée dans la proposition 2.3 on voit que l'on peut remplacer $\mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i)$ dans :

$$\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} \mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.85)$$

par $\mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) + A$ à condition que A soit indépendante de $\bar{j}_i$ ou de j_{i+1} . En appliquant 2 fois cette transformation on remplace $\mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i)$ par $\tilde{\mathbf{G}}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i)$ avec

$$\tilde{\mathbf{G}}^{ij}(\mathbf{q}) := \mathbf{G}^{ij}(\mathbf{q}) - \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_i) - \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_j) + \frac{1}{2d} G(\mathbf{q}). \quad (2.86)$$

En définissant pour toute fonction $f : \mathbb{Z}^d \rightarrow \mathbb{R}$

$$\nabla_i f(\mathbf{q}) := f(\mathbf{q} + \mathbf{e}_i) - f(\mathbf{q}) \quad (2.87)$$

on a :

$$\tilde{\mathbf{G}}^{ij}(\mathbf{q}) := \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1}{2d} \nabla_{-i} \nabla_{-j} G(\mathbf{q}). \quad (2.88)$$

Montrons que l'on a amélioré la vitesse de décroissance :

Lemme 2.1. *En dimension $d \geq 3$, $\forall i, j \in \mathcal{P}_{ind}$:*

$$\tilde{\mathbf{G}}^{ij}(\mathbf{q}) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right). \quad (2.89)$$

Preuve : Soient $i, j \in \mathcal{P}_{ind}$, il suffit de montrer que

$$\nabla_i \nabla_j G(\mathbf{q}) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right)$$

On sait [19] qu'il existe une constante, que l'on nomme A_0 , telle que :

$$G(\mathbf{q}) = \frac{A_0}{\|\mathbf{q}\|^{d-2}} + O\left(\frac{1}{\|\mathbf{q}\|^d}\right).$$

Notons que

$$\nabla_i O\left(\frac{1}{\|\mathbf{q}\|^d}\right) = O\left(\frac{1}{\|\mathbf{q}\|^d}\right). \quad (2.90)$$

Calculons :

$$\begin{aligned}
\nabla_i \frac{1}{\|\mathbf{q}\|^n} &= \frac{1}{\|\mathbf{q} + \mathbf{e}_i\|^n} - \frac{1}{\|\mathbf{q}\|^n}, \\
&= \frac{1}{\|\mathbf{q}\|^n} \left(\frac{\|\mathbf{q}\|^n}{\|\mathbf{q} + \mathbf{e}_i\|^n} - 1 \right), \\
&= \frac{1}{\|\mathbf{q}\|^n} \left(\left(\frac{\|\mathbf{q} + \mathbf{e}_i\| - \nabla_i \|\mathbf{q}\|}{\|\mathbf{q} + \mathbf{e}_i\|} \right)^n - 1 \right), \\
&= \frac{1}{\|\mathbf{q}\|^n} \left(\frac{-n \nabla_i \|\mathbf{q}\|}{\|\mathbf{q} + \mathbf{e}_i\|} + O\left(\frac{1}{\|\mathbf{q}\|^2}\right) \right), \\
&= \frac{-n \nabla_i \|\mathbf{q}\|}{\|\mathbf{q}\|^{n+1}} + O\left(\frac{1}{\|\mathbf{q}\|^{n+2}}\right).
\end{aligned} \tag{2.91}$$

Or :

$$\begin{aligned}
\|\mathbf{q} + \mathbf{e}_i\| &= \sqrt{(\mathbf{q} + \mathbf{e}_i) \cdot (\mathbf{q} + \mathbf{e}_i)} \\
&= \|\mathbf{q}\| \sqrt{1 + \frac{2\mathbf{q} \cdot \mathbf{e}_i + 1}{\|\mathbf{q}\|^2}} \\
&= \|\mathbf{q}\| \left(1 + \frac{\mathbf{q} \cdot \mathbf{e}_i}{\|\mathbf{q}\|^2} + O\left(\frac{1}{\|\mathbf{q}\|^2}\right) \right) \\
&= \|\mathbf{q}\| + \frac{\mathbf{q} \cdot \mathbf{e}_i}{\|\mathbf{q}\|} + O\left(\frac{1}{\|\mathbf{q}\|}\right)
\end{aligned}$$

d'où

$$\nabla_i \|\mathbf{q}\| = \frac{\mathbf{q} \cdot \mathbf{e}_i}{\|\mathbf{q}\|} + O\left(\frac{1}{\|\mathbf{q}\|}\right) \tag{2.92}$$

Il suit de cela, grâce à (2.91), que :

$$\begin{aligned}
\nabla_j \left(\frac{\nabla_i \|\mathbf{q}\|}{\|\mathbf{q}\|^{n+1}} \right) &= \frac{\nabla_i \|\mathbf{q} + \mathbf{e}_j\|}{\|\mathbf{q} + \mathbf{e}_j\|^{n+1}} - \frac{\nabla_i \|\mathbf{q}\|}{\|\mathbf{q}\|^{n+1}} \\
&= \frac{(\mathbf{q} + \mathbf{e}_j) \cdot \mathbf{e}_i}{\|\mathbf{q} + \mathbf{e}_j\|^{n+2}} - \frac{\mathbf{q} \cdot \mathbf{e}_i}{\|\mathbf{q}\|^{n+2}} + O\left(\frac{1}{\|\mathbf{q}\|^{n+2}}\right) \\
&= \mathbf{q} \cdot \mathbf{e}_i \frac{1}{\|\mathbf{q}\|^{n+2}} + O\left(\frac{1}{\|\mathbf{q}\|^{n+2}}\right) \\
&= \mathbf{q} \cdot \mathbf{e}_i \left(\frac{-(n+2) \nabla_j \|\mathbf{q}\|}{\|\mathbf{q}\|^{n+3}} + O\left(\frac{1}{\|\mathbf{q}\|^{n+3}}\right) \right) + O\left(\frac{1}{\|\mathbf{q}\|^{n+2}}\right) \\
&= O\left(\frac{1}{\|\mathbf{q}\|^{n+2}}\right)
\end{aligned} \tag{2.93}$$

En appliquant (2.90) puis (2.91) puis (2.93) on calcule finalement :

$$\begin{aligned}
\nabla_j \nabla_i G(\mathbf{q}) &= \nabla_j \nabla_i \frac{A_0}{\|\mathbf{q}\|^{d-2}} + O\left(\frac{1}{\|\mathbf{q}\|^d}\right) \\
&= \nabla_j \frac{-n \nabla_i \|\mathbf{q}\|}{\|\mathbf{q}\|^{d-1}} + O\left(\frac{1}{\|\mathbf{q}\|^d}\right) \\
&= O\left(\frac{1}{\|\mathbf{q}\|^d}\right)
\end{aligned} \tag{2.94}$$

□

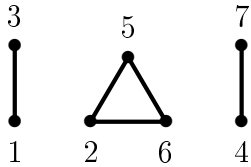
Les manipulations sur l'espérance du produit des $\hat{p}$ nous ont donc permis de passer d'une décroissance en $1/\|\mathbf{q}\|^{d-2}$ à une décroissance en $1/\|\mathbf{q}\|^d$. C'est une amélioration sensible qui va faire la différence. On pourrait se poser la question de la validité éventuelle de la version adaptée de ce lemme en dimension inférieure à 3, mais on ne répondra pas à cette question dans cette thèse.

2.2.5 Représentation graphique des partitions sans singleton.

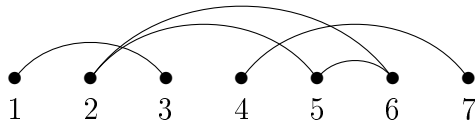
On va introduire une manière visuelle de présenter les partitions sans singleton. On introduira ensuite les notions de partition *connectées* et de partition *non-connectées*. Soit W une partition sans singleton, on définit son graphe associé dont l'ensemble des sommets est $\llbracket 1, m \rrbracket$ et deux sommets sont reliés si et seulement si ils appartiennent à la même sous-partie de la partition W . Par exemple la partition de $\llbracket 1, 7 \rrbracket$ suivante :

$$\{\{1, 3\}, \{2, 5, 6\}, \{4, 7\}\}$$

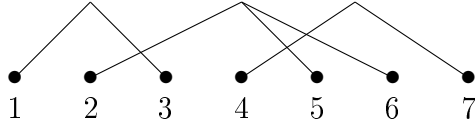
est associée au graphe :



On constate, et c'est vrai de manière générale, que le graphe obtenu est un ensemble de cliques disjointes. Cette représentation n'est pas très parlante, on va utiliser une représentation en ligne comme le fait Gordon Slade dans sa présentation de la *lace expansion* [9]. Cela donne le graphe suivant :



Si les cliques deviennent grandes, la représentation précédente, bien que jolie, devient lourde. On adopte donc une autre représentation des cliques :



Les traits dans cette dernière représentation ne représentent plus des arêtes au sens de la théorie des graphes, nous allons donc utiliser d'autres mots. En lieu et place des cliques nous avons à présent des liens, on peut donc avoir n points reliés tous ensemble par un seul lien. La partie du lien qui le relie spécifiquement à un point est nommée une *patte*, c'est un peu l'équivalent d'une demi-arête. On utilisera aussi le terme de *double patte* pour designer deux pattes d'un même lien reliées à deux points successifs. Voici une illustration à la figure 3.1 :

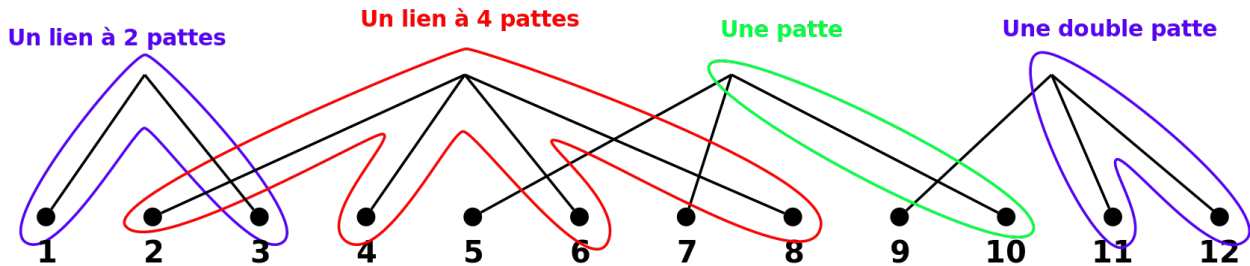


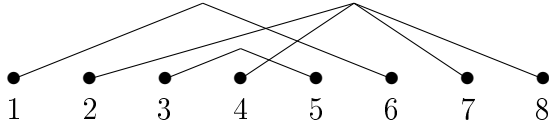
FIGURE 3.1

Introduisons à présent la notion de graphe connecté.

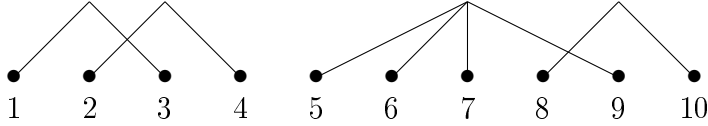
Définition 2.1. On dit qu'un graphe à n sommets numéroté de 1 à n est connecté si et seulement si :

$$\forall i \in \llbracket 2, n-1 \rrbracket, \exists j, k \in \llbracket 1, n \rrbracket, j < i < k \text{ et } j \text{ est relié à } k.$$

Autrement dit, dans la représentation précédente, pour tout point qui n'est pas à une extrémité, il existe un lien qui passe "au-dessus" de ce point. Jusque là nous n'avons vu que des graphes connectés. Voici encore un exemple de graphe connecté :



Et un graphe non-connecté :



Proposition 2.5. Soit $m \in \mathbb{N}$, il existe une bijection entre l'ensemble des partitions sans singleton de $\llbracket 1, m \rrbracket$ et l'ensemble des graphes à m sommets numérotés de 1 à m qui sont des réunions disjointes de cliques sans sommet isolé.

Preuve : C'est construit pour fonctionner. Deux sommets du graphe sont reliés si et seulement si ils sont dans la même partie de la partition correspondante.

□

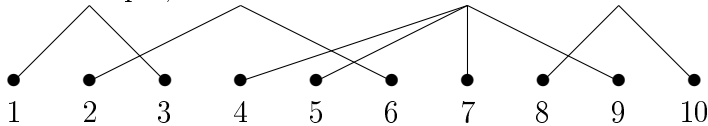
Du fait de cette bijection on dira indifféremment que W est une partition ou que c'est un graphe. Un graphe est non-connecté si et seulement si la partition correspondante peut être scindée en deux sous-partitions, une partition de $\llbracket 1, i \rrbracket$ et une partition de $\llbracket i + 1, n \rrbracket$ pour un certain i . Dans ce cas on dira aussi que la partition est non-connectée. On y reviendra quand on traitera les partitions non-connectées.

En s'inspirant de la *lace expansion* de Gordon Slade on définit une notion de *lace* qui sera utile dans la preuve du théorème 2.2. Les laces sont des graphes connectés minimaux.

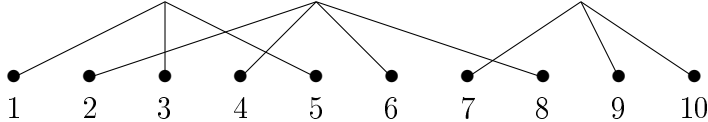
Pour un graphe donné de taille m on peut retirer un lien, l'opération consiste à supprimer tous les sommets associés à ce lien puis à renuméroter les sommets restants de 1 à $m' < m$ sans changer leur ordre.

Définition 2.2. Une *lace* est un graphe connecté minimal, c'est à dire qu'un graphe connecté est une lace si et seulement si en retirant n'importe quel lien sauf le premier ou le dernier on obtient un graphe non-connecté.

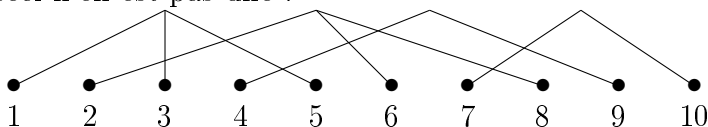
Par exemple, ceci est une lace :



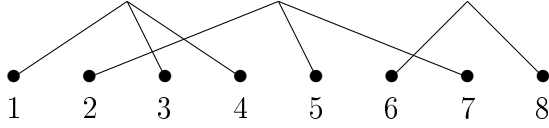
Une autre :



Et ceci n'en est pas une :



Car on peut retirer le lien (4,9) et obtenir un le graphe suivant, toujours connecté :



Les extrémités d'un lien sont le sommet le plus à gauche et le sommet le plus à droite de ce lien.

Proposition 2.6. *Soit W une lace de taille m . Soit l_1 un lien de W . Soit $x_1 \in \llbracket 1, m \rrbracket$ une extrémité de l_1 telle que $x_1 \neq 1$ et $x_1 \neq m$. Alors il existe un unique lien l_2 dont les extrémités x_2 et x'_2 sont telles que :*

$$x_2 < x_1 < x'_2.$$

Preuve : Prouvons l'existence et l'unicité de l_2 .

Existence : Comme W est connecté il existe un lien qui passe au dessus de x_1 . C'est l_2 .

Unicité : Supposons par l'absurde qu'il existe deux liens l_2 et l_3 d'extrémités respectives x_2, x'_2 et x_3, x'_3 . Telles que $x_2 < x_1 < x'_2$ et $x_3 < x_1 < x'_3$. Soit x'_1 la seconde extrémité de l_1 . Supposons sans perte de généralité que $x_1 < x'_1$. Supposons sans perte de généralité que $x_2 < x_3$. Alors si $x'_3 < x'_1$ on peut retirer l_3 et W sera toujours connecté, ce qui est absurde. Et si $x'_1 < x'_3$ alors on peut retirer l_1 et W sera toujours connecté, ce qui est absurde. Cela conclut à l'unicité de l_2 .

□

2.2.6 Majoration de $\sum_{n=0}^{+\infty} \mathcal{U}_{W,n}$

Revenons à l'intégrabilité de

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m).$$

En remplaçant $\mathbf{G}$ par $\tilde{\mathbf{G}}$ dans (2.85) on trouve

$$\sum_{n=0}^{+\infty} \mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} \tilde{\mathbf{G}}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.95)$$

Donc

$$\sum_{n=0}^{+\infty} \mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \mathbf{1}[(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W] \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \prod_{i=1}^{m-1} \tilde{\mathbf{G}}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.96)$$

On n'a plus besoin du terme en espérance de $\hat{p}$. Majorons le par 1. On trouve alors qu'il suffit de prouver l'intégrabilité de

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \prod_{i=1}^{m-1} \tilde{\mathbf{G}}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.97)$$

sur H_W (et pas sur $(\mathbb{Z}^d)^{m-1}$). On peut même définir la fonction suivante :

$$\psi : \begin{cases} \mathbb{Z}^d & \rightarrow \mathbb{R}_+ \\ \mathbf{q} & \mapsto \sup_{(i,j) \in \mathcal{P}_{ind}} |\tilde{\mathbf{G}}^{ij}(\mathbf{q})| \end{cases} \quad (2.98)$$

Il suffit à présent de prouver que

$$\sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H} \prod_{i=1}^{m-1} \psi(\mathbf{q}_{i+1} - \mathbf{q}_i) < +\infty. \quad (2.99)$$

En anticipant un peu on définit alors pour tout entier α la fonction réelle :

$$\varphi_\alpha = \begin{cases} \mathbb{R}_+ & \rightarrow \mathbb{R}_+ \\ x & \mapsto \frac{e^d \ln^\alpha(x+e)}{(x+e)^d} \end{cases} \quad (2.100)$$

où $e = \exp(1)$, cette fonction est définie de manière à être la plus simple possible tout en respectant les contraintes $\varphi_\alpha(0) = 1$ et $\varphi_\alpha(x) \underset{x \rightarrow +\infty}{=} O\left(\frac{\ln^\alpha(x)}{x^d}\right)$. Ainsi, on voit facilement grâce au lemme 2.1 que pour $d \geq 3$:

$$\forall \alpha \in \mathbb{N}, \exists M \in \mathbb{R}, \psi(\mathbf{q}) < M\varphi_\alpha(\|\mathbf{q}\|), \quad (2.101)$$

car

$$\psi(\mathbf{q}) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right).$$

et ψ est définie sur $\mathbb{Z}^d$ donc bornée sur tout compact. On utilise donc les φ_α pour majorer ψ . Et on voit que quel que soit $\alpha \in \mathbb{N}$,

$$\sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{q}_{i+1} - \mathbf{q}_i\|) < +\infty \Rightarrow \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H} \prod_{i=1}^{m-1} \psi(\mathbf{q}_{i+1} - \mathbf{q}_i) < +\infty. \quad (2.102)$$

Une simple étude de fonction montre que φ_0 est une fonction décroissante de $\mathbb{R}_+$ dans $[0, 1]$ et que si $\alpha > 0$ alors $\varphi_\alpha(0) = 1$, φ_α est croissante sur $[0, e^{\alpha/d} - e]$ et φ_α est décroissante sur $[e^{\alpha/d} - e, +\infty[$. De plus :

$$\max_{x \in \mathbb{R}_+} \varphi_\alpha(x) = \varphi_\alpha(e^{\alpha/d} - e) = e^d \left(\frac{\alpha}{ed}\right)^\alpha. \quad (2.103)$$

On a enfin :

$$\forall \alpha, \beta \in \mathbb{N}, \alpha \leq \beta \Rightarrow \forall x \in \mathbb{R}_+, \varphi_\alpha(x) \leq \varphi_\beta(x). \quad (2.104)$$

On définit la norme p classique :

$$\|f\|_p := \left(\sum_{\mathbf{q} \in \mathbb{Z}^d} |f(\mathbf{q})|^p \right)^{1/p},$$

ainsi que l'espace :

$$\mathcal{L}^p(\mathbb{Z}^d) := \left\{ f : \mathbb{Z}^d \rightarrow \mathbb{R} \left| \sum_{\mathbf{q} \in \mathbb{Z}^d} |f(\mathbf{q})|^p < +\infty \right. \right\}.$$

On notera que $\forall p, q \in [1, +\infty], p \leq q \Rightarrow \mathcal{L}^p(\mathbb{Z}^d) \subset \mathcal{L}^q(\mathbb{Z}^d)$.

Proposition 2.7. $\forall p \in]1, +\infty[$, la fonction

$$\mathbf{q} \mapsto \varphi_\alpha(\|\mathbf{q}\|)$$

appartient à $\mathcal{L}^p(\mathbb{Z}^d)$.

Preuve : Comme :

$$\begin{aligned} \varphi_\alpha(\|\mathbf{q}\|) &\underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^\alpha(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right), \\ &\underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{(\|\mathbf{q}\|^d)^p}\right). \end{aligned} \quad (2.105)$$

une simple comparaison série-intégrale suffit à conclure. □

Corolaire 2.2. Si $d \geq 3$, $\forall p \in]1, +\infty[$,

$$\psi \in \mathcal{L}^p(\mathbb{Z}^d).$$

Preuve : Par la majoration $\psi(\mathbf{q}) \leq M\varphi_\alpha(\|\mathbf{q}\|)$, établie en (2.101).

□

2.2.7 Convergence des graphes connectés

Définition 2.3. On dira qu'un graphe W à m sommets correspondant à une partition sans singleton converge si et seulement si :

$\forall \alpha_1, \dots, \alpha_{m-1} \in \mathbb{N}$,

$$\sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \prod_{i=1}^{m-1} \varphi_{\alpha_i}(\|\mathbf{q}_{i+1} - \mathbf{q}_i\|) < +\infty. \quad (2.106)$$

On a alors le théorème suivant :

Théorème 2.2. *Tout les graphes connectés convergent.*

Preuve : Pour simplifier les notations, pour tout graphe W à m sommets et pour tout $\alpha \in \mathbb{N}$, on définit la fonction suivante à valeurs dans $\mathbb{R}_+ \cup \{+\infty\}$:

$$F_\alpha(W) := \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{q}_{i+1} - \mathbf{q}_i\|). \quad (2.107)$$

On a alors que W converge si et seulement si $\forall \alpha \in \mathbb{N}, F_\alpha(W) < +\infty$.

On établit un premier lemme :

Lemme 2.2. $\forall \alpha, \beta \in \mathbb{N}, \forall \mathbf{q} \in \mathbb{Z}^d, \exists M \in \mathbb{R}_+$,

$$\sum_{\mathbf{q}' \in \mathbb{Z}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) \leq M \varphi_{\alpha+\beta+1}(\|\mathbf{q}\|)$$

Preuve : Dans toute la preuve M est un réel assez grand qui peut changer de valeur d'une ligne à l'autre. On voit d'abord que comme $\varphi_\alpha, \varphi_\beta \in \mathcal{L}^2$ alors $\sum_{\mathbf{q}' \in \mathbb{Z}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|)$ est bien défini pour tout $\mathbf{q} \in \mathbb{Z}$. Il reste à prouver que :

$$\sum_{\mathbf{q}' \in \mathbb{Z}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^{\alpha+\beta+1}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right).$$

On commence par une comparaison série-intégrale :

$$\sum_{\mathbf{q}' \in \mathbb{Z}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) \leq M \int_{\mathbb{R}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}'.$$

On va utiliser le changement de variable $\|\mathbf{q}\|\mathbf{q}'' = \mathbf{q}'$ et la notation $u_{\mathbf{q}} := \mathbf{q}/\|\mathbf{q}\|$ pour le vecteur unitaire orienté selon $\mathbf{q}$.

$$\begin{aligned} \int_{\mathbb{R}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' &= \int \frac{e^d \ln^\alpha(\|\mathbf{q}'\| + e)}{(\|\mathbf{q}'\| + e)^d} \frac{e^d \ln^\beta(\|\mathbf{q} - \mathbf{q}'\| + e)}{(\|\mathbf{q} - \mathbf{q}'\| + e)^d} d\mathbf{q}' \\ &= \frac{e^{2d}}{\|\mathbf{q}\|^d} \int \frac{\ln^\alpha(\|\mathbf{q}\|\|\mathbf{q}''\| + e)}{(\|\mathbf{q}''\| + \frac{e}{\|\mathbf{q}\|})^d} \frac{\ln^\beta(\|\mathbf{q}\|\|u_{\mathbf{q}} - \mathbf{q}''\| + e)}{(\|u_{\mathbf{q}} - \mathbf{q}''\| + \frac{e}{\|\mathbf{q}\|})^d} d\mathbf{q}'' \end{aligned} \quad (2.108)$$

On va majorer cette intégrale indépendamment au voisinage de zéro, de $u_{\mathbf{q}}$ et de l'infini. Soit $B(0, 2)$ la boule de centre 0 et de rayon 2. On note que $\mathbf{q}'' \in B(0, 2) \Leftrightarrow \mathbf{q}' \in B(0, 2\|\mathbf{q}\|)$. On a les

inégalités : $\|\mathbf{q}''\| + 1 \geq \|u_{\mathbf{q}} - \mathbf{q}''\| \geq \|\mathbf{q}''\| - 1$. Et évidemment $\|\mathbf{q}''\| + 1 \geq \|\mathbf{q}''\| \geq \|\mathbf{q}''\| - 1$ cela nous permet de majorer la quantité suivante :

$$\begin{aligned}
\mathcal{A}_{\mathbb{R}^d \setminus B(0,2)} &= \int_{\mathbb{R}^d \setminus B(0,2\|\mathbf{q}\|)} \varphi_{\alpha}(\|\mathbf{q}'\|) \varphi_{\beta}(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' \\
&= \frac{e^{2d}}{\|\mathbf{q}\|^d} \int \frac{\ln^{\alpha}(\|\mathbf{q}\| \|\mathbf{q}''\| + e)}{(\|\mathbf{q}''\| + \frac{e}{\|\mathbf{q}\|})^d} \frac{\ln^{\beta}(\|\mathbf{q}\| \|u_{\mathbf{q}} - \mathbf{q}''\| + e)}{(\|u_{\mathbf{q}} - \mathbf{q}''\| + \frac{e}{\|\mathbf{q}\|})^d} d\mathbf{q}'' \\
&\leq \frac{e^{2d}}{\|\mathbf{q}\|^d} \int_{\mathbb{R}^d \setminus B(0,2)} \frac{\ln^{\alpha+\beta}(\|\mathbf{q}\|(\|\mathbf{q}''\| + 1) + e)}{(\|\mathbf{q}''\| - 1 + \frac{e}{\|\mathbf{q}\|})^{2d}} d\mathbf{q}'' \\
&\leq \frac{e^{2d}}{\|\mathbf{q}\|^d} \int_{\mathbb{R}^d \setminus B(0,2)} \frac{\sum_{k=0}^{\alpha+\beta} \binom{\alpha+\beta}{k} \ln^k(\|\mathbf{q}\|) \ln^{\alpha+\beta-k}(\|\mathbf{q}''\| + 1 + e/\|\mathbf{q}\|)}{(\|\mathbf{q}''\| - 1 + \frac{e}{\|\mathbf{q}\|})^{2d}} d\mathbf{q}''
\end{aligned} \tag{2.109}$$

Inversons la somme (finie) et l'intégrale puis majorons les intégrales. On suppose $\|\mathbf{q}\| > e$ car on cherche une borne pour $\|\mathbf{q}\| \rightarrow +\infty$.

$$\begin{aligned}
\mathcal{A}_{\mathbb{R}^d \setminus B(0,2)} &\leq \frac{e^{2d}}{\|\mathbf{q}\|^d} \sum_{k=0}^{\alpha+\beta} \binom{\alpha+\beta}{k} \ln^k(\|\mathbf{q}\|) \int_{\mathbb{R}^d \setminus B(0,2)} \frac{\ln^{\alpha+\beta-k}(\|\mathbf{q}''\| + 1 + e/\|\mathbf{q}\|)}{(\|\mathbf{q}''\| - 1 + \frac{e}{\|\mathbf{q}\|})^{2d}} d\mathbf{q}'' \\
&\leq \frac{e^{2d}}{\|\mathbf{q}\|^d} \sum_{k=0}^{\alpha+\beta} \binom{\alpha+\beta}{k} \ln^k(\|\mathbf{q}\|) \int_{\mathbb{R}^d \setminus B(0,2)} \frac{\ln^{\alpha+\beta-k}(\|\mathbf{q}''\| + 1 + e)}{(\|\mathbf{q}''\| - 1)^{2d}} d\mathbf{q}'' \\
&\leq \frac{\ln^{\alpha+\beta}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d} e^{2d} \int_{\mathbb{R}^d \setminus B(0,2)} \frac{\ln^{\alpha+\beta}(\|\mathbf{q}''\| + 1 + e)}{(\|\mathbf{q}''\| - 1)^{2d}} d\mathbf{q}'' \sum_{k=0}^{\alpha+\beta} \binom{\alpha+\beta}{k} \\
&\leq \frac{\ln^{\alpha+\beta}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d} M
\end{aligned} \tag{2.110}$$

Avec M qui ne dépend pas de $\mathbf{q}$. Ça marche parce que l'intégrale converge. On a donc prouvé que :

$$\mathcal{A}_{\mathbb{R}^d \setminus B(0,2)} \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^{\alpha+\beta}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right). \tag{2.111}$$

Les dernières majorations sont extrêmement grossières, mais elle fonctionnent pour cette preuve.

Étudions maintenant l'intégrale pour $\mathbf{q}' \in B(0, 2\|\mathbf{q}\|)$. On la coupe en deux, l'ensemble V_0 des points plus proches de 0 que de $\mathbf{q}$ et l'ensemble $V_{\mathbf{q}}$ des points plus proches de $\mathbf{q}$ que de 0. Le V fait référence au concept de cellule de Voronoï, on trouve une illustration pour $d = 2$ à la figure 3.2.

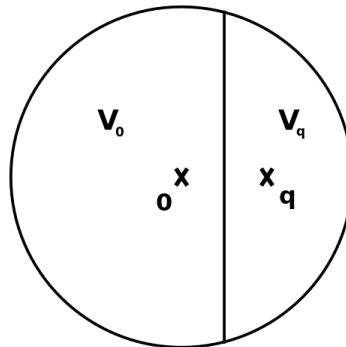


FIGURE 3.2

Pour $\mathbf{q}' \in V_0$ on a $3\|\mathbf{q}\| \geq \|\mathbf{q} - \mathbf{q}'\| \geq \|\mathbf{q}\|/2$. On majore les termes ainsi :

$$\begin{aligned}\varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) &:= \frac{e^d \ln^\beta(\|\mathbf{q} - \mathbf{q}'\| + e)}{(\|\mathbf{q} - \mathbf{q}'\| + e)^d} \\ &\leq \frac{e^d \ln^\beta(3\|\mathbf{q}\| + e)}{(2^{-1}\|\mathbf{q}\| + e)^d}\end{aligned}\quad (2.112)$$

car $\mathbf{q}'$ est plus proche de 0 que de $\mathbf{q}$. Et

$$\begin{aligned}\varphi_\alpha(\|\mathbf{q}'\|) &:= \frac{e^d \ln^\alpha(\|\mathbf{q}'\| + e)}{(\|\mathbf{q}'\| + e)^d} \\ &\leq \frac{e^d \ln^\alpha(3\|\mathbf{q}\| + e)}{(\|\mathbf{q}'\| + e)^d}\end{aligned}\quad (2.113)$$

D'où (avec S_d l'hypersurface d'une hypersphère de rayon 1) :

$$\begin{aligned}\int_{V_0} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' &\leq \int_{V_0} \frac{e^d \ln^\alpha(3\|\mathbf{q}\| + e)}{(\|\mathbf{q}'\| + e)^d} \frac{e^d \ln^\beta(3\|\mathbf{q}\| + e)}{2^{-d}\|\mathbf{q}\|^d} d\mathbf{q}' \\ &\leq \frac{2^d e^{2d} \ln^{\alpha+\beta}(3\|\mathbf{q}\| + e)}{\|\mathbf{q}\|^d} \int_{V_0} \frac{1}{(\|\mathbf{q}'\| + e)^d} d\mathbf{q}' \\ &\leq \frac{2^d e^{2d} \ln^{\alpha+\beta}(3\|\mathbf{q}\| + e)}{\|\mathbf{q}\|^d} \int_{B(0, 2\|\mathbf{q}\|)} \frac{1}{(\|\mathbf{q}'\| + e)^d} d\mathbf{q}' \\ &\leq \frac{2^d e^{2d} S_d \ln^{\alpha+\beta}(3\|\mathbf{q}\| + e)}{\|\mathbf{q}\|^d} \int_0^{2\|\mathbf{q}\|} \frac{r^{d-1}}{(r + e)^d} dr \\ &\leq \frac{2^d e^{2d} S_d \ln^{\alpha+\beta}(3\|\mathbf{q}\| + e)}{\|\mathbf{q}\|^d} \left(1 + \int_1^{2\|\mathbf{q}\|} \frac{1}{r} dr\right) \\ &\leq \frac{2^d e^{2d} S_d \ln^{\alpha+\beta}(3\|\mathbf{q}\| + e)}{\|\mathbf{q}\|^d} (1 + \ln(2\|\mathbf{q}\|))\end{aligned}\quad (2.114)$$

D'où :

$$\int_{V_0} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^{\alpha+\beta+1}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right) \quad (2.115)$$

On peut faire un raisonnement similaire sur $V_{\mathbf{q}}$ et ainsi prouver que :

$$\int_{V_{\mathbf{q}}} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^{\alpha+\beta+1}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right) \quad (2.116)$$

Quand on rassemble les trois intégrales, sur $\mathbb{R}^d \setminus B(0, 2\|\mathbf{q}\|)$, V_0 et $V_{\mathbf{q}}$, on trouve que :

$$\int_{\mathbb{R}^d} \varphi_\alpha(\|\mathbf{q}'\|) \varphi_\beta(\|\mathbf{q} - \mathbf{q}'\|) d\mathbf{q}' \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{\ln^{\alpha+\beta+1}(\|\mathbf{q}\|)}{\|\mathbf{q}\|^d}\right) \quad (2.117)$$

ce qui termine la preuve. □

On rappelle que la notation $(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W$ signifie que pour chaque partie de la partition W , tous les $\mathbf{q}_i$ dont les indices sont dans cette partie sont tous égaux, avec $\mathbf{q}_1 := 0$. Autrement dit : $W = \{W_j : j \in \llbracket 1, n \rrbracket\}$ et $\forall j \in \llbracket 1, n \rrbracket$ et $\forall i, i' \in \llbracket 1, m \rrbracket$,

$$i \in W_j, i' \in W_j \Rightarrow \mathbf{q}_i = \mathbf{q}_{i'}.$$

On va donc considérer des $\mathbf{q}_i$ particuliers, on définit simplement un représentant de classe pour chaque W_j :

$$\forall j \in \llbracket 1, n \rrbracket, i_j := \min\{i \in W_j\}$$

Et on pose :

$$\mathbf{a}_j := \mathbf{q}_{i_j}$$

À l'inverse, on note j_i l'unique $j \in \llbracket 1, n \rrbracket$ tel que $i \in W_j$, j_i est en fait le numéro de la classe de W qui contient i . ainsi :

$$\mathbf{q}_i = \mathbf{a}_{j_i}$$

Pour résumer on associe à tout $j \in \llbracket 1, n \rrbracket$ un entier $i_j \in \llbracket 1, m \rrbracket$ qui est le numéro du représentant de classe de la classe W_j et on associe à tout $i \in \llbracket 1, m \rrbracket$ un entier $j_i \in \llbracket 1, n \rrbracket$ qui est le numéro de la classe qui contient i .

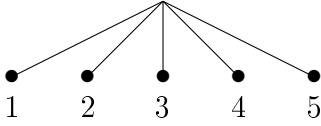
On a donc pour tout entier α :

$$\begin{aligned} F_\alpha(W) &:= \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{q}_{i+1} - \mathbf{q}_i\|), \\ &= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_n} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \end{aligned} \quad (2.118)$$

$$= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_n} \prod_{j=1}^n \prod_{i=i_j}^{i_{j+1}-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|). \quad (2.119)$$

avec les conventions : $i_1 = 1$, $\mathbf{q}_1 = 0$ et $i_{n+1} := m$.

Montrons à présent que toutes les laces convergent, on procède par récurrence. L'initialisation se fera sur les laces à 2 liens, traitons à part les laces à un seul lien. Comme celle-ci :

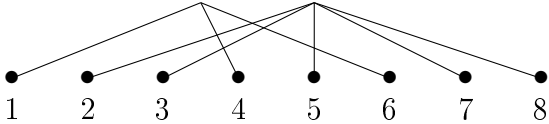


Si W est une telle lace alors on a :

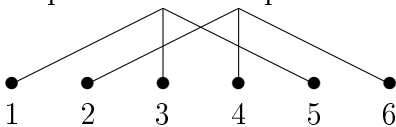
$$\begin{aligned} F_\alpha(W) &= \prod_{i=1}^{m-1} \varphi_\alpha(0), \\ &< +\infty. \end{aligned} \quad (2.120)$$

Donc W converge.

Si W est une lace à 2 liens comme celle-ci :



alors remarquons d'abord que si deux points côte à côte sont reliés par le même lien alors un facteur $\varphi_\alpha(0)$ apparaît et celui-ci ne change rien à la convergence. On peut donc le retirer, ce qui revient à retirer l'un des deux points. On peut faire ça si le lien concerné reliait au moins trois points. On dira alors qu'un tel lien a une double-patte et que l'opération ci-dessus revient à fusionner la double-patte. Par exemple la lace ci-dessus converge si et seulement si la lace ci-dessous converge :



Il suffit donc de montrer la convergence des laces à 2 liens de taille $m = 2k$ de la forme : $W =$

$\{W_1, W_2\}$ avec $W_1 = \{2i - 1 | i \in \llbracket 1, k \rrbracket\}$ et $W_2 = \{2i | i \in \llbracket 1, k \rrbracket\}$. Avec $k \geq 2$. Si W est ainsi alors :

$$F_\alpha(W) = \sum_{\mathbf{a}_{i_2}} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \quad (2.121)$$

$$= \sum_{\mathbf{a}_2} (\varphi_\alpha(\|\mathbf{a}_2 - \mathbf{a}_1\|))^{m-1} < +\infty \quad (2.122)$$

car $m \geq 4$ et $\varphi_\alpha \in \mathcal{L}^2$. On vient donc de démontrer que toutes les laces de taille 2 convergent.

Soit W une lace de taille $m > 2$ avec n liens, sans perte de généralité on peut supposer que cette lace n'a pas de double patte. On crée la lace W' obtenue en supprimant le lien n de W puis en fusionnant les doubles pattes créées. On va prouver que si W' converge alors W converge. On définit m' ainsi :

$$\max\{i \in \llbracket 1, m-1 \rrbracket | \mathbf{q}_i = \mathbf{a}_{n-2}\} = m' - 1.$$

La partition W' est alors une partition sans singleton de $\llbracket 1, m' \rrbracket$. On a alors :

$$\begin{aligned} F_\alpha(W) &= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_n} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \\ &= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_{n-1}} \prod_{i=1}^{m'-2} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \end{aligned} \quad (2.123)$$

On utilise le lemme suivant que l'on prouvera juste après :

Lemme 2.3.

$$\sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \leq M \varphi_{2\alpha+1}(\|\mathbf{a}_{n-1} - \mathbf{a}_{n-2}\|) \quad (2.124)$$

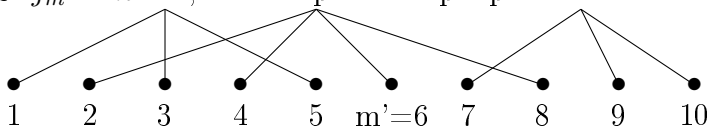
Grâce à ce lemme on prouve :

$$\begin{aligned} F_\alpha(W) &= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_{n-1}} \prod_{i=1}^{m'-2} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \\ &\leq M \sum_{\mathbf{a}_2, \dots, \mathbf{a}_{n-1}} \prod_{i=1}^{m'-1} \varphi_{2\alpha+1}(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \\ &\leq M F_{2\alpha+1}(W'), \\ &< +\infty. \end{aligned} \quad (2.125)$$

On a donc prouvé que la convergence de W' implique la convergence de W . Donc, par récurrence, toutes les laces convergent.

Preuve du lemme 2.3 Par définition de m' , $m' - 1$ est l'extrémité droite du lien $n - 2$, donc $j_{m'-1} = n - 2$. Et comme W est une lace on voit facilement que $j_{m'} = n - 1$ ou $j_{m'} = n$. Traitons les deux cas séparément.

Si $j_{m'} = n - 1$, comme par exemple pour cette lace-ci :



alors on a :

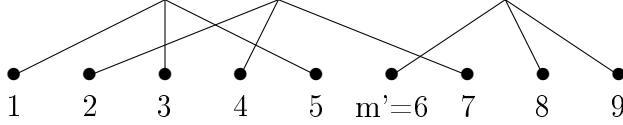
$$\begin{aligned}
\sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) &= \varphi_\alpha(\|\mathbf{a}_{j_{m'}} - \mathbf{a}_{j_{m'-1}}\|) \sum_{\mathbf{a}_n} \prod_{i=m'}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \\
&= \varphi_\alpha(\|\mathbf{a}_{n-1} - \mathbf{a}_{n-2}\|) \sum_{\mathbf{a}_n} \prod_{i=m'}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \\
&\leq M \varphi_\alpha(\|\mathbf{a}_{n-1} - \mathbf{a}_{n-2}\|)
\end{aligned} \tag{2.126}$$

car la somme sur $\mathbf{a}_n$ est convergente. En effet, on peut voir facilement qu'ici $m' \leq m-3$ et donc que :

$$\sum_{\mathbf{a}_n} \prod_{i=m'}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \leq M \sum_{\mathbf{a}_n} (\varphi_\alpha(\|\mathbf{a}_n - \mathbf{a}_{n-1}\|))^2 \tag{2.127}$$

or, $\varphi_\alpha \in \mathcal{L}^2$, donc la somme sur $\mathbf{a}_n$ est convergente.

Si $j_{m'} = n$, comme par exemple pour cette lace-ci :



alors on a :

$$\begin{aligned}
\sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) &\leq M \sum_{\mathbf{a}_n} \prod_{i=m'-1}^{m'} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|) \\
&\leq M \sum_{\mathbf{a}_n} \varphi_\alpha(\|\mathbf{a}_n - \mathbf{a}_{n-2}\|) \varphi_\alpha(\|\mathbf{a}_{n-1} - \mathbf{a}_n\|) \\
&\leq M \varphi_{2\alpha+1}(\|\mathbf{a}_{n-1} - \mathbf{a}_{n-2}\|)
\end{aligned} \tag{2.128}$$

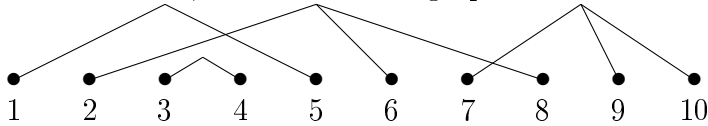
d'après le lemme 2.2.

Cela termine la preuve du lemme 2.3.

□

La dernière étape du théorème consiste à prouver que tous les graphes connectés convergent. Soit W un graphe connecté non minimal (donc pas une lace) il existe donc W' connecté obtenu en enlevant un lien de W . On va prouver que la convergence de W' implique la convergence de W . Par conséquent, la convergence des laces implique la convergence de tous les graphes connectés. On suppose donc par récurrence que W' converge.

Soit donc W un graphe connecté non minimal, on enlève un lien de manière à obtenir W' un graphe connecté. On suppose sans perte de généralité que W ne contient pas de double patte, mais il est possible que W' en contienne. Il existe un cas particulier où le lien supprimé est un lien à 2 pattes côte à côte, comme dans ce graphe :



Dans ce cas on appelle m' l'extrémité gauche du lien en question, son extrémité droite est $m' + 1$. Le numéro de ce lien est $n' = j_{m'} = j_{m'+1}$. Alors :

$$\begin{aligned}
F_\alpha(W) &= \sum_{\mathbf{a}_2, \dots, \mathbf{a}_n} \varphi_\alpha(\|\mathbf{a}_{n'} - \mathbf{a}_{j_{m'-1}}\|) \varphi_\alpha(0) \varphi_\alpha(\|\mathbf{a}_{j_{m'+2}} - \mathbf{a}_{n'}\|) \prod_{\substack{i=1 \\ i \notin [m'-1, m'+1]}}^{m-1} \varphi_\alpha(\|\mathbf{a}_{j_{i+1}} - \mathbf{a}_{j_i}\|), \\
&\leq M F_{2\alpha+1}(W')
\end{aligned} \tag{2.129}$$

La majoration finale étant à nouveau une application directe du lemme 2.2. Donc :

$$\forall \alpha \in \mathbb{N}, F_\alpha(W') < +\infty \Rightarrow \forall \alpha \in \mathbb{N}, F_\alpha(W) < +\infty$$

Passons au cas général. Dans le cas général, rajouter un lien sans double patte consiste à remplacer un terme

$$\prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|) \quad (2.130)$$

par

$$\sum_{\mathbf{q}'} \prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|). \quad (2.131)$$

où r est le nombre de pattes du lien que l'on rajoute.

Pour chaque i , on a deux cas : $\mathbf{q}_i = \mathbf{q}'_i$ et $\mathbf{q}_i \neq \mathbf{q}'_i$. Dans le premier cas on peut majorer $\varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|)$ par une constante, donc :

$$\varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) \leq M \varphi_\alpha(0) = M \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|) \quad (2.132)$$

et dans le second cas on utilise l'inégalité simple suivante :

$$\|\mathbf{q}_i - \mathbf{q}'_i\| \leq \|\mathbf{q}' - \mathbf{q}_i\| + \|\mathbf{q}' - \mathbf{q}'_i\|$$

d'où :

$$\|\mathbf{q}' - \mathbf{q}_i\| \geq \|\mathbf{q}_i - \mathbf{q}'_i\|/2 \text{ ou } \|\mathbf{q}' - \mathbf{q}'_i\| \geq \|\mathbf{q}_i - \mathbf{q}'_i\|/2$$

On a aussi l'inégalité suivante :

$$\exists M \in \mathbb{R}_+^*, \forall x, x' \in \mathbb{R}_+, x \geq x' \Rightarrow \varphi_\alpha(x) \leq M \varphi_\alpha(x') \quad (2.133)$$

qui vient du fait que φ_α est décroissante sur l'intervalle $[e^{\alpha/d} - e, +\infty[$. En fait, on peut même prendre explicitement $M = e^d \left(\frac{\alpha}{ed}\right)^\alpha$. Mais on a aussi une inégalité dans l'autre sens :

$$\exists M \in \mathbb{R}_+^*, \forall x \in \mathbb{R}_+, \varphi_\alpha(x/2) \leq M \varphi_\alpha(x) \quad (2.134)$$

on prouve rapidement cette équation ainsi :

$$\begin{aligned} \varphi_\alpha(x/2) &= \frac{e^d \ln^\alpha(x/2 + e)}{(x/2 + e)^d}, \\ &\leq \frac{e^d \ln^\alpha(x + e)}{(x/2 + e)^d}, \\ &\leq 2^d \frac{e^d \ln^\alpha(x + e)}{(x + 2e)^d}, \\ &\leq 2^d \varphi_\alpha(x). \end{aligned} \quad (2.135)$$

On trouve donc en utilisant (2.133) puis (2.134) :

$$\begin{aligned} \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) &\leq M \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|/2) \\ &\leq M \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|) \end{aligned} \quad (2.136)$$

En combinant (2.132) et (2.136) on trouve que :

$$\forall \mathbf{q}_i, \mathbf{q}'_i, \mathbf{q}' \in \mathbb{Z}^d, \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) \leq M \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|) \quad (2.137)$$

On trouve donc, en appliquant cela ainsi que le lemme 2.2 :

$$\begin{aligned} \sum_{\mathbf{q}'} \prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) &\leq M \prod_{i=2}^r \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|) \sum_{\mathbf{q}'} \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_1\|) \varphi_\alpha(\|\mathbf{q}_1 - \mathbf{q}'\|) \\ &\leq M \prod_{i=1}^r \varphi_{2\alpha+1}(\|\mathbf{q}_i - \mathbf{q}'_i\|) \end{aligned} \quad (2.138)$$

$$(2.139)$$

Donc, si on retire un lien sans double patte à un graphe W pour obtenir un graphe W' on peut trouver $F_\alpha(W')$ à partir de $F_\alpha(W)$ en remplaçant un terme de la forme :

$$\sum_{\mathbf{q}'} \prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) \quad (2.140)$$

par

$$\prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'_i\|). \quad (2.141)$$

Et on vient de prouver que :

$$\sum_{\mathbf{q}'} \prod_{i=1}^r \varphi_\alpha(\|\mathbf{q}' - \mathbf{q}'_i\|) \varphi_\alpha(\|\mathbf{q}_i - \mathbf{q}'\|) \leq M \prod_{i=1}^r \varphi_{2\alpha+1}(\|\mathbf{q}_i - \mathbf{q}'_i\|). \quad (2.142)$$

Donc,

$$\forall \alpha \in \mathbb{N}, F_\alpha(W') < +\infty \Rightarrow \forall \alpha \in \mathbb{N}, F_\alpha(W) < +\infty.$$

On a donc prouvé que la convergence des laces implique la convergence de tous les graphes connectés. Ce qui termine la preuve du théorème 2.2. □

On rappelle alors que, $\forall \alpha \in \mathbb{N}$,

$$\begin{aligned} \sum_{\underline{q}} \left| \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \right| &\leq \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \prod_{i=1}^{m-1} \psi(\mathbf{q}_{i+1} - \mathbf{q}_i) \\ &\leq \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \prod_{i=1}^{m-1} \varphi_\alpha(\|\mathbf{q}_{i+1} - \mathbf{q}_i\|) \end{aligned} \quad (2.143)$$

Par conséquent, une application directe du théorème 2.2 permet de conclure que pour tout graphe connecté W , la fonction :

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$$

est intégrable sur H_W .

Traisons à présent la convergence de la série.

2.2.8 Convergence dominée de la série $\mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ pour W connecté.

On a donc prouvé que $\sum_{\underline{q}} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ existe, sortons maintenant notre cher théorème de convergence dominée !

On cherche à prouver que

$$\lim_{N \rightarrow +\infty} \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{n=0}^N \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$$

On va prouver la convergence pour chaque terme ij de la matrice $\mathcal{U}_{W,n}$. Pour i et j donnés on cherche donc une fonction intégrable $F : (\mathbb{Z}^d)^{m-1} \rightarrow \mathbb{R}_+$ telle que

$$\forall N \in \mathbb{N}, \left| \sum_{n=0}^N \mathcal{U}_{W,N}^{ij}(\mathbf{q}_2, \dots, \mathbf{q}_m) \right| \leq F(\mathbf{q}_2, \dots, \mathbf{q}_m).$$

Comme candidat on a :

$$\sum_{r=0}^{+\infty} |\mathcal{U}_{W,r}^{ij}(\mathbf{q}_2, \dots, \mathbf{q}_m)|$$

Dont il suffirait de prouver l'intégrabilité. On rappelle l'équation (2.74) :

$$\mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \sum_{\underline{j}} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right] \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.144)$$

On va d'abord créer une notation plus compacte pour le terme $\mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right]$. On note que pour deux indices i, i' , $\hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma)$ et $\hat{p}((\mathbf{q}_{i'}, \mathbf{e}_{j_{i'}}), (\mathbf{q}_{i'} + \mathbf{e}_{\bar{j}_{i'}}, \mathbf{e}_{\bar{j}_{i'}}), \sigma)$ sont indépendants si et seulement si $\mathbf{q}_i \neq \mathbf{q}_{i'}$, et cette information est donnée par la partition W . À part cela, les valeurs des $\mathbf{q}_i$ n'ont aucune influence sur la valeur de l'espérance, seuls les j_i sont importants. Avec la convention $\mathbf{q}_1 = 0$, on peut alors définir la quantité : $L_W(j_1, \bar{j}_1, \dots, j_m, \bar{j}_m)$ telle que :

$$\forall(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W, L_W(j_1, \bar{j}_1, \dots, j_m, \bar{j}_m) = \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right]. \quad (2.145)$$

Que l'on abrégera même en $L_W(j_1, \dots, \bar{j}_m)$. Comme :

$$\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \mathbf{1}[(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W] \mathcal{U}_{m,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.146)$$

On a alors :

$$\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \mathbf{1}[(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W] \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.147)$$

Proposition 2.8. $\forall i \in \llbracket 1, m \rrbracket$,

$$\sum_{j_i} L_W(j_1, \dots, \bar{j}_m) = \sum_{\bar{j}_i} L_W(j_1, \dots, \bar{j}_m) = 0.$$

Preuve : Cela découle très simplement de l'équation (2.16) vue en début de section et de la définition de L_W . □

On rappelle la formule (2.53) :

$$P_n^{ij}(\mathbf{q}) = \delta_n \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j). \quad (2.148)$$

D'où

$$P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) = \delta_{n_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{e}_{j_{i+1}}} + \frac{1 - \delta_{n_i}}{2d} g_{n_i-1}(\mathbf{q}_{i+1} - \mathbf{q}_i - \mathbf{e}_{\bar{j}_i} - \mathbf{e}_{j_{i+1}}). \quad (2.149)$$

On voudrait, grâce à la proposition 2.8 refaire le coup de la double dérivée discrète sur g , comme on l'avait fait en remplaçant G par $\nabla_i \nabla_j G$ à l'équation (2.88), mais on a un problème de parité. En effet :

$$\vec{\nabla}_i g_n(\mathbf{q}) = \begin{cases} g_n(\mathbf{q}) & \text{si } n \text{ et } \mathbf{q} \text{ ont la même parité} \\ g_n(\mathbf{q} + \mathbf{e}_i S) & \text{sinon} \end{cases}$$

Ce qui n'est pas très intéressant. Mais on va faire quelque chose de semblable.

On voit grâce à la proposition 2.8 que l'on peut remplacer chaque facteur de la forme $P_n^{ij}(\mathbf{q})$ dans (2.147) par $P_n^{ij}(\mathbf{q}) + A$ pour n'importe quel A indépendant de i ou de j sans changer la valeur de $\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m)$. Utilisons cette propriété pour accélérer la vitesse de décroissance de ces facteurs quand $\mathbf{q} \rightarrow \infty$. Pour cela on ignore le terme $\delta_n \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j}$ qui est local, on met de côté

le facteur $\frac{1-\delta_n}{2d}$, puis on retranche la moyenne de $g(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j)$ sur l'ensemble des $\mathbf{e}_i$. Or cette moyenne est :

$$\frac{1}{2d} \sum_{i \in \mathcal{P}_{ind}} g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) = g_n(\mathbf{q} - \mathbf{e}_j). \quad (2.150)$$

Cette dernière égalité découle simplement de la définition de $g_n(\mathbf{q})$ qui est la probabilité que la marche aléatoire simple aille de 0 à $\mathbf{q}$ en n pas. On vient donc de faire la substitution :

$$g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) \rightarrow g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) - g_n(\mathbf{q} - \mathbf{e}_j). \quad (2.151)$$

On retranche alors la moyenne sur $\mathbf{e}_j$, ce qui correspond à la substitution :

$$g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) - g_n(\mathbf{q} - \mathbf{e}_j) \rightarrow g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) - g_n(\mathbf{q} - \mathbf{e}_j) - g_n(\mathbf{q} - \mathbf{e}_i) + g_{n+1}(\mathbf{q}). \quad (2.152)$$

On définit :

$$\mathcal{A}_n^{ij}(\mathbf{q}) := g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - g_{n+1}(\mathbf{q} + \mathbf{e}_i) - g_{n+1}(\mathbf{q} + \mathbf{e}_j) + g_{n+2}(\mathbf{q})$$

On a alors démontré que le remplacement :

$$g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j) \rightarrow \mathcal{A}_{n-1}(\mathbf{q}). \quad (2.153)$$

Dans (2.147) ne change pas la valeur de $\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m)$.

Par ailleurs, le facteur $L_W(j_1, \dots, \bar{j}_m)$ est invariant par renumérotation des $\mathbf{e}_i$, en particulier :

$$L_W(j_1, \dots, \bar{j}_m) = L_W(-j_1, \dots, -\bar{j}_m). \quad (2.154)$$

En faisant les changements de variable $j_i \rightarrow -j_i, \bar{j}_i \rightarrow -\bar{j}_i$, pour $(0, \mathbf{q}_2, \dots, \mathbf{q}_m)$

$$\begin{aligned} \sum_{n=0}^{+\infty} |\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m)| &= \sum_{n=0}^{+\infty} \sum_{\underline{j}} |L_W(j_1, \dots, \bar{j}_m)| \sum_{\underline{n} \in \mathcal{S}_n} \\ &\quad \prod_{i=1}^{m-1} \left| \delta_{n_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\bar{j}_i, j_{i+1}} + \frac{1 - \delta_{n_i}}{2d} \mathcal{A}_{n_i-1}^{\bar{j}_i, j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \right| \\ &\leq \sum_{n=0}^{+\infty} \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} \sum_{\underline{j}} \left| \delta_{n_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\bar{j}_i, j_{i+1}} + \frac{1 - \delta_{n_i}}{2d} \mathcal{A}_{n_i-1}^{\bar{j}_i, j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \right| \\ &\leq \sum_{\underline{j}} \prod_{i=1}^{m-1} \sum_{n_i=0}^{+\infty} \left| \delta_{n_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\bar{j}_i, j_{i+1}} + \frac{1 - \delta_{n_i}}{2d} \mathcal{A}_{n_i-1}^{\bar{j}_i, j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \right| \end{aligned} \quad (2.155)$$

Le terme $\delta_{n_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\bar{j}_i, j_{i+1}}$ est local, il ne change donc rien à l'intégrabilité, le facteur $\frac{1}{2d}$ ne change rien non plus et le facteur $1 - \delta_{n_i}$ permet le changement de variable $n_i - 1 \rightarrow n_i$. On trouve donc que la fonction

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \sum_{r=0}^{+\infty} |\mathcal{U}_{W,r}^{\mathbf{p}_1, \bar{\mathbf{p}}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m)|$$

est intégrable sur H_W si et seulement si quelques soient les j et les $\bar{j}$,

$$(\mathbf{q}_2, \dots, \mathbf{q}_m) \mapsto \prod_{i=1}^{m-1} \sum_{n_i=0}^{+\infty} \left| \mathcal{A}_{n_i}^{\bar{j}_i, j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \right|$$

est intégrable sur H_W .

On va d'abord prouver la proposition suivante :

Proposition 2.9. $\forall i, j \in \mathcal{P}_{ind}$,

$$\sum_{n=0}^{+\infty} |\mathcal{A}_n^{ij}(\mathbf{q})| \Big|_{\mathbf{q} \rightarrow +\infty} = O\left(\frac{1}{\|\mathbf{q}\|^d}\right) \quad (2.156)$$

Puis conclure à l'aide du théorème 2.2.

Preuve : On rappelle que

$$\mathcal{A}_n^{ij}(\mathbf{q}) := g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - g_{n+1}(\mathbf{q} + \mathbf{e}_i) - g_{n+1}(\mathbf{q} + \mathbf{e}_j) + g_{n+2}(\mathbf{q}).$$

On constate d'abord que $\mathcal{A}_n^{ij}(\mathbf{q}) \neq 0$ si et seulement si n et $\mathbf{q}$ ont la même parité, seuls ces cas-là sont donc à traiter. On approxime alors la valeur de $g_n(\mathbf{q})$ par le théorème central limite local. On pose :

$$\bar{g}_n(\mathbf{q}) := \frac{1}{(2\pi n)^{d/2}} e^{-\frac{\mathbf{q}^2}{2n}} \quad (2.157)$$

On sait alors par le théorème central limite local que le terme d'erreur :

$$|g_n(\mathbf{q}) + g_{n+1}(\mathbf{q}) - 2\bar{g}_n(\mathbf{q})|$$

est petit, dans un sens que l'on précisera en temps utile. Or, si l'on suppose que n et $\mathbf{q}$ ont la même parité alors

$$|g_n(\mathbf{q}) + g_{n+1}(\mathbf{q}) - 2\bar{g}_n(\mathbf{q})| = |g_n(\mathbf{q}) - 2\bar{g}_n(\mathbf{q})|.$$

Le facteur 2 vient donc du fait que la marche aléatoire simple est paire. On écrit donc :

$$\begin{aligned} |\mathcal{A}_n^{ij}(\mathbf{q})| &= |g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - g_{n+1}(\mathbf{q} + \mathbf{e}_i) - g_{n+1}(\mathbf{q} + \mathbf{e}_j) + g_{n+2}(\mathbf{q})| \\ &\leq 2|\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_i) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_j) + \bar{g}_{n+2}(\mathbf{q})| \\ &\quad + |g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - 2\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j)| + |g_{n+1}(\mathbf{q} + \mathbf{e}_i) - 2\bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_i)| \\ &\quad + |g_{n+1}(\mathbf{q} + \mathbf{e}_j) - 2\bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_j)| + |g_{n+2}(\mathbf{q}) - 2\bar{g}_{n+2}(\mathbf{q})| \end{aligned} \quad (2.158)$$

La preuve du théorème 4.3.1. à la page 83 de [19] prouve que si n et $\mathbf{q}$ ont la même parité alors :

$$\sum_{\substack{n=0 \\ n \text{ et } \mathbf{q} \text{ ont même parité}}}^{+\infty} |g_n(\mathbf{q}) - 2\bar{g}_n(\mathbf{q})| \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right). \quad (2.159)$$

Cela traite d'un coup les quatre derniers termes de la somme ci-dessus et on trouve que :

$$\begin{aligned} \sum_{n=0}^{+\infty} |\mathcal{A}_n^{ij}(\mathbf{q})| &\underset{\mathbf{q} \rightarrow +\infty}{=} \sum_{\substack{n=0 \\ n \text{ et } \mathbf{q} \text{ ont même parité}}}^{+\infty} |\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_i) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_j) + \bar{g}_{n+2}(\mathbf{q})| \\ &\quad + O\left(\frac{1}{\|\mathbf{q}\|^d}\right), \end{aligned} \quad (2.160)$$

où la somme est à nouveau uniquement sur les n ayant la même parité que $\mathbf{q}$. Enfin, on peut majorer brutalement cette somme par la somme sur tous les entiers. On a alors plus de soucis de parité et cela suffira pour terminer la preuve.

Pour finir la preuve il ne reste donc qu'à prouver que :

$$\sum_{n=0}^{+\infty} |\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_i) - \bar{g}_{n+1}(\mathbf{q} + \mathbf{e}_j) + \bar{g}_{n+2}(\mathbf{q})| \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right)$$

On va décomposer le reste de la preuve en 3 lemme :

Lemme 2.4. $\forall k \in \mathbb{N}$,

$$\sum_{n=0}^{+\infty} |\bar{g}_{n+k}(\mathbf{q}) - \bar{g}_n(\mathbf{q})| \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right).$$

Lemme 2.5.

$$\sum_{n=0}^{+\infty} |\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_n(\mathbf{q} + \mathbf{e}_i) - \bar{g}_n(\mathbf{q} + \mathbf{e}_j) + \bar{g}_n(\mathbf{q})| \leq \frac{M}{\|\mathbf{q}\|^d} + \sum_{n=0}^{+\infty} 3\bar{g}_n(\mathbf{q}) \left(e^{\frac{3\|\mathbf{q}\|}{n}} - 1 - \frac{3\|\mathbf{q}\|}{n} \right).$$

Pour un certain $M > 0$.

Lemme 2.6.

$$\sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left(e^{\frac{3\|\mathbf{q}\|}{n}} - 1 - \frac{3\|\mathbf{q}\|}{n} \right) \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right).$$

L'enchainement de ces 3 lemmes permet de conclure la preuve ainsi, avec M prenant une valeur différente à chaque ligne :

$$\begin{aligned} \sum_{n=0}^{+\infty} |\mathcal{A}_n^{ij}(\mathbf{q})| &\leq \frac{M}{\|\mathbf{q}\|^d} + \sum_{n=0}^{+\infty} |g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - g_{n+1}(\mathbf{q} + \mathbf{e}_i) - g_{n+1}(\mathbf{q} + \mathbf{e}_j) + g_{n+2}(\mathbf{q})|, \\ &\leq \frac{M}{\|\mathbf{q}\|^d} + \sum_{n=0}^{+\infty} |g_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - g_n(\mathbf{q} + \mathbf{e}_i) - g_n(\mathbf{q} + \mathbf{e}_j) + g_n(\mathbf{q})|, \\ &\leq \frac{M}{\|\mathbf{q}\|^d} + \sum_{n=0}^{+\infty} 3\bar{g}_n(\mathbf{q}) \left(e^{\frac{3\|\mathbf{q}\|}{n}} - 1 - \frac{3\|\mathbf{q}\|}{n} \right), \\ &\leq \frac{M}{\|\mathbf{q}\|^d} \end{aligned} \tag{2.161}$$

Pour prouver ces lemmes on utilisera le lemme 4.3.2 page 82 de [19] :

Lemme 2.7. *Pour tout $b > 1$,*

$$\sum_{n=1}^{+\infty} n^{-b} e^{-x/n} \underset{x \rightarrow +\infty}{=} \frac{\Gamma(b-1)}{x^{b-1}} + O\left(\frac{1}{x^{b+1}}\right) \tag{2.162}$$

Preuve : On trouvera la preuve dans [19], lemme 4.3.2, page 82.

□

Preuve lemme 2.4 : Soit $k > 0$.

$$\begin{aligned} \sum_{n=0}^{+\infty} |\bar{g}_{n+k}(\mathbf{q}) - \bar{g}_n(\mathbf{q})| &= \sum_{n=0}^{+\infty} \frac{1}{(2\pi)^{d/2}} \left| \frac{1}{(n)^{d/2}} - \frac{1}{(n+k)^{d/2}} \right| e^{-\frac{\mathbf{q}^2}{2n}}, \\ &= \sum_{n=0}^{+\infty} \frac{1}{(2\pi n)^{d/2}} \left| 1 - \left(1 - \frac{k}{n+k}\right)^{d/2} \right| e^{-\frac{\mathbf{q}^2}{2n}}, \\ &\leq \sum_{n=0}^{+\infty} \frac{1}{(2\pi n)^{d/2}} \frac{d}{2} \frac{k}{(n+k)} e^{-\frac{\mathbf{q}^2}{2n}}, \\ &\leq \sum_{n=0}^{+\infty} \frac{kd}{2(2\pi)^{d/2}} \frac{1}{n^{d/2+1}} e^{-\frac{\mathbf{q}^2}{2n}}, \end{aligned} \tag{2.163}$$

On a utilisé le fait que :

$$\forall a \in [1, +\infty[, \forall x \in [0, 1], |1 - (1-x)^a| \leq ax. \tag{2.164}$$

En posant $b = d/2 + 1$ et $x = \mathbf{q}^2/2$, on trouve par le lemme 2.7 que

$$\sum_{n=0}^{+\infty} \frac{kd}{2(2\pi)^{d/2}} \frac{1}{n^{d/2+1}} e^{-\frac{\mathbf{q}^2}{2n}} \underset{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right) \tag{2.165}$$

Ce qui conclut la preuve du lemme 2.4.

□

Preuve lemme 2.5 : Commençons par regarder la fonction :

$$\zeta : x \mapsto e^x - 1 - x. \quad (2.166)$$

On voit facilement que ζ est positive sur $\mathbb{R}$, croissante sur $\mathbb{R}_+$ et que $\forall x \in \mathbb{R}$,

$$\zeta(|x|) \geq \zeta(x). \quad (2.167)$$

On démontre tout cela facilement en montrant que :

$$\zeta(x) = \sum_{n=2}^{+\infty} \frac{x^n}{n!} \quad (2.168)$$

On pose alors $\bar{\mathcal{A}}_n^{ij}(\mathbf{q}) := \bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_n(\mathbf{q} + \mathbf{e}_i) - \bar{g}_n(\mathbf{q} + \mathbf{e}_j) + \bar{g}_n(\mathbf{q})$, ainsi :

$$\begin{aligned} \sum_{n=1}^{+\infty} |\bar{\mathcal{A}}_n^{ij}(\mathbf{q})| &= \sum_{n=2}^{+\infty} |\bar{g}_n(\mathbf{q} + \mathbf{e}_i + \mathbf{e}_j) - \bar{g}_n(\mathbf{q} + \mathbf{e}_i) - \bar{g}_n(\mathbf{q} + \mathbf{e}_j) + \bar{g}_n(\mathbf{q})| \\ &\leq \sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left| e^{-\frac{2\mathbf{q} \cdot (\mathbf{e}_i + \mathbf{e}_j) + (\mathbf{e}_i + \mathbf{e}_j)^2}{2n}} - e^{-\frac{2\mathbf{q} \cdot \mathbf{e}_i + \mathbf{e}_i^2}{2n}} - e^{-\frac{2\mathbf{q} \cdot \mathbf{e}_j + \mathbf{e}_j^2}{2n}} + 1 \right| \\ &\leq \sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left| e^{-\frac{2\mathbf{q} \cdot (\mathbf{e}_i + \mathbf{e}_j) + (\mathbf{e}_i + \mathbf{e}_j)^2}{2n}} - 1 + \frac{2\mathbf{q} \cdot (\mathbf{e}_i + \mathbf{e}_j) + (\mathbf{e}_i + \mathbf{e}_j)^2}{2n} \right. \\ &\quad \left. - e^{-\frac{2\mathbf{q} \cdot \mathbf{e}_i + \mathbf{e}_i^2}{2n}} + 1 - \frac{2\mathbf{q} \cdot \mathbf{e}_i + \mathbf{e}_i^2}{2n} - e^{-\frac{2\mathbf{q} \cdot \mathbf{e}_j + \mathbf{e}_j^2}{2n}} + 1 - \frac{2\mathbf{q} \cdot \mathbf{e}_j + \mathbf{e}_j^2}{2n} - \frac{\mathbf{e}_i \cdot \mathbf{e}_j}{n} \right| \end{aligned} \quad (2.169)$$

On voit apparaitre des fonctions ζ et on applique l'inégalité triangulaire.

$$\begin{aligned} \sum_{n=1}^{+\infty} |\bar{\mathcal{A}}_n^{ij}(\mathbf{q})| &\leq \sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left(\zeta \left(-\frac{2\mathbf{q} \cdot (\mathbf{e}_i + \mathbf{e}_j) + (\mathbf{e}_i + \mathbf{e}_j)^2}{2n} \right) + \zeta \left(-\frac{2\mathbf{q} \cdot \mathbf{e}_i + \mathbf{e}_i^2}{2n} \right) \right. \\ &\quad \left. + \zeta \left(-\frac{2\mathbf{q} \cdot \mathbf{e}_j + \mathbf{e}_j^2}{2n} \right) + \left| \frac{\mathbf{e}_i \cdot \mathbf{e}_j}{n} \right| \right) \\ &\leq \sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left(3\zeta \left(\frac{3\|\mathbf{q}\|}{n} \right) + \left| \frac{\mathbf{e}_i \cdot \mathbf{e}_j}{n} \right| \right) \end{aligned} \quad (2.170)$$

On conclut grâce au lemme 2.7 qui permet la majoration :

$$\sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \left| \frac{\mathbf{e}_i \cdot \mathbf{e}_j}{n} \right| \leq \frac{M}{\|\mathbf{q}\|^d}. \quad (2.171)$$

□

Preuve lemme 2.6 : En reprenant la fonction ζ définie plus haut, on cherche à prouver :

$$\sum_{n=1}^{+\infty} \bar{g}_n(\mathbf{q}) \zeta \left(\frac{3\|\mathbf{q}\|}{n} \right) \underset{\mathbf{q} \rightarrow +\infty}{=} O \left(\frac{1}{\|\mathbf{q}\|^d} \right). \quad (2.172)$$

Omettons le facteur $1/(2\pi)^{d/2}$ car on cherche une majoration à un facteur près.

$$\begin{aligned} \sum_{n=1}^{+\infty} n^{-d/2} e^{-\frac{\|\mathbf{q}\|^2}{2n}} \zeta \left(\frac{3\|\mathbf{q}\|}{n} \right) &= \sum_{n=1}^{+\infty} n^{-d/2} e^{-\frac{\|\mathbf{q}\|^2}{2n}} \sum_{k=2}^{+\infty} \frac{(3\|\mathbf{q}\|/n)^k}{k!} \\ &= \sum_{k=2}^{+\infty} \frac{(3\|\mathbf{q}\|)^k}{k!} \sum_{n=1}^{+\infty} n^{-d/2-k} e^{-\frac{\|\mathbf{q}\|^2}{2n}} \end{aligned} \quad (2.173)$$

On peut intervertir les sommes car ce sont des sommes à termes positifs.

On démontre alors une équation proche du lemme 2.7, avec une majoration plus grossière. Le lemme 2.7 ne fonctionnerait pas car on ne pourrait pas majorer la somme sur k du terme d'erreur. Soient $b > 1$ et $x > 0$:

$$\begin{aligned}
\sum_{n=1}^{+\infty} n^{-b} e^{-x/n} &= \int_0^{+\infty} \lceil t \rceil^{-b} e^{-x/\lceil t \rceil} dt \\
&\leq \int_0^{+\infty} t^{-b} e^{-x/(t+1)} dt \\
&\leq \int_1^2 (u-1)^{-b} e^{-x/(u)} du + \int_2^{+\infty} (u-1)^{-b} e^{-x/(u)} du \\
&\leq \int_1^2 (u/2)^{-b} e^{-x/u} du + \int_2^{+\infty} (u/2)^{-b} e^{-x/u} du \\
&\leq 2^b \int_0^{+\infty} u^{-b} e^{-x/u} du \\
&\leq \frac{2^b}{x^{b-1}} \int_0^{+\infty} y^{b-2} e^{-y} dy \\
&\leq \frac{2^b \Gamma(b-1)}{x^{b-1}}
\end{aligned} \tag{2.174}$$

Grâce à cette majoration on a :

$$\begin{aligned}
\sum_{n=1}^{+\infty} n^{-d/2} e^{-\frac{\|\mathbf{q}\|^2}{2n}} \zeta\left(\frac{3\|\mathbf{q}\|}{n}\right) &\leq \sum_{k=2}^{+\infty} \frac{(3\|\mathbf{q}\|)^k}{k!} \frac{2^{d/2+k} \Gamma(d/2+k-1)}{\|\mathbf{q}\|^{d+2k-2}} \\
&\leq \frac{2^{d/2}}{\|\mathbf{q}\|^d} \sum_{k=0}^{+\infty} \frac{6^{k+2}}{(k+2)!} \frac{\Gamma(d/2+k+1)}{\|\mathbf{q}\|^k} \\
&\leq \frac{2^{d/2}}{\|\mathbf{q}\|^d} \sum_{k=0}^{+\infty} \frac{(d+k)!}{(k+2)!} \frac{6^{k+2}}{\|\mathbf{q}\|^k} \\
&\leq \frac{2^{d/2}}{\|\mathbf{q}\|^d} \sum_{k=0}^{+\infty} (d+k)^d \frac{6^{k+2}}{\|\mathbf{q}\|^k} \\
&\stackrel{\mathbf{q} \rightarrow +\infty}{=} O\left(\frac{1}{\|\mathbf{q}\|^d}\right).
\end{aligned} \tag{2.175}$$

□

La preuve de ce dernier lemme conclut la preuve de la proposition.

□

On a donc prouvé que :

$$\lim_{N \rightarrow +\infty} \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{n=0}^N \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$$

Pour W connectée. Il reste juste le cas des partitions non-connectées.

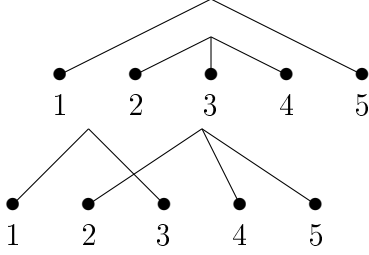
2.2.9 Convergence de la série $\mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$ pour une partition W non-connectée.

On va se ramener à la convergence des partitions connectées. Un graphe est connecté si et seulement si la partition correspondante peut être scindée en deux sous-partitions, une partition de $\llbracket 1, i \rrbracket$ et une partition de $\llbracket i+1, n \rrbracket$ pour un certain i . Dans cette idée-là, on va définir une notion de concaténation des partitions. Pour toute partition W on pose que $\sim_W$ est la relation d'équivalence associée, on a alors la définition suivante.

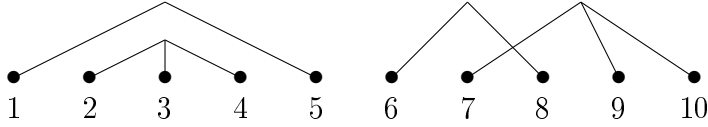
Définition 2.4. Soient $n_1, n_2 \in \mathbb{N}$ et soient W_1 une partition de $\llbracket 1, n_1 \rrbracket$ et W_2 une partition de $\llbracket 1, n_2 \rrbracket$. Alors on définit la concaténation $W = W_1 \star W_2$ comme étant la partition de $\llbracket 1, n_1 + n_2 \rrbracket$ telle que : $\forall i, j \in \llbracket 1, n_1 + n_2 \rrbracket$:

$$i \underset{W}{\sim} j \Leftrightarrow i \underset{W_1}{\sim} j \text{ ou } i + n_1 \underset{W_2}{\sim} j + n_1.$$

Avec la représentation en graphe l'opération est très simple, voici deux graphes :



et leur concaténation :



On note qu'une partition correspond à un graphe non-connecté si et seulement si elle est la concaténation de deux autres partitions.

Cette notion de concaténation va nous permettre de prouver la proposition suivante qui conclut la preuve du théorème 2.1.

Proposition 2.10. Pour toute partition sans singleton W non-connectée,

$$\sum_{n=0}^{+\infty} \sum_{\underline{q}} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) = \sum_{\underline{q}} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.176)$$

De plus, si on définit la matrice

$$\mathcal{V}_W := \sum_{\underline{q}} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.177)$$

alors il existe une matrice $\mathbb{A}$, qui ne dépend que de d telle que pour tout couple (W_1, W_2) de partitions sans singleton :

$$\mathcal{V}_{W_1 \star W_2} = \mathcal{V}_{W_1} \mathbb{A} \mathcal{V}_{W_2} \quad (2.178)$$

Preuve : Soit W une partition sans singleton non-connectée, il existe alors W_1 et W_2 tels que $W = W_1 \star W_2$, étudions $\sum_{\underline{q}} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m)$, dans le cas des partitions non-connectées la notation $\mathcal{B}_{W,n}$ est plus adaptée. On rappelle l'équation (2.66) :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} \mathcal{B}_{m,n}(\mathbf{q}) = \sum_{\underline{q}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_m). \quad (2.179)$$

Par analogie on définit :

$$\mathcal{B}_{W,n}(\mathbf{q}) := \sum_{\mathbf{q}_2, \dots, \mathbf{q}_{m-1}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_{m-1}, \mathbf{q}). \quad (2.180)$$

Ainsi :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} \mathcal{B}_{W,n}(\mathbf{q}) = \sum_{\underline{q}} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m). \quad (2.181)$$

et :

$$\mathcal{B}_{W,n}(\mathbf{q}) := \sum_{\mathbf{q}_2, \dots, \mathbf{q}_{m-1}} \mathcal{U}_{m,n}(\mathbf{q}_2, \dots, \mathbf{q}_{m-1}, \mathbf{q}). \quad (2.182)$$

On rappelle l'équation (2.147) :

$$\mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) := \mathbf{1}[(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W] \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \quad (2.183)$$

Ainsi, on a :

$$\mathcal{B}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_m) = \sum_{\substack{\mathbf{q}_2, \dots, \mathbf{q}_{m-1} \\ (\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W}} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i). \quad (2.184)$$

On a alors un certain i_0 tel que W_1 soit une partition de $\llbracket 1, i_0 \rrbracket$ et W_2 une partition de $\llbracket i_0 + 1, n \rrbracket$. En coupant alors en trois parties les sommes sur les $\mathbf{q}$, sur les $\underline{j}$ et sur les $\underline{n}$, et en notant $H_W + \mathbf{q}$ l'ensemble H_W translaté de $\mathbf{q}$, on trouve :

$$\begin{aligned} \sum_{\mathbf{q}_m} \mathcal{B}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_m) &= \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \\ &= \sum_{\substack{r_1, r_2, r_3 \\ r_1 + r_2 + r_3 = n}} \sum_{\bar{j}_{i_0}, j_{i_0+1}} \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_{i_0}) \in H_{W_1}} \sum_{\bar{j}_1, j_2, \dots, j_{i_0-1}, j_{i_0}} L_{W_1}(j_1, \dots, \bar{j}_{i_0}) \\ &\quad \sum_{\substack{n_1, \dots, n_{i_0} \\ n_1 + \dots + n_{i_0} = r_1}} \prod_{i=1}^{i_0-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \sum_{\substack{\mathbf{q}_{i_0+1} \\ \mathbf{q}_{i_0} \neq \mathbf{q}_{i_0+1}}} P_{r_2}^{\bar{j}_{i_0} j_{i_0+1}}(\mathbf{q}_{i_0+1} - \mathbf{q}_{i_0}) \\ &\quad \sum_{(\mathbf{q}_{i_0+2}, \dots, \mathbf{q}_m) \in (H_{W_2} + \mathbf{q}_{i_0+1})} \sum_{\bar{j}_{i_0+1}, j_{i_0+2}, \dots, \bar{j}_{m-1}, j_m} L_{W_2}(j_{i_0+1}, \dots, \bar{j}_m) \\ &\quad \sum_{\substack{n_{i_0+2}, \dots, n_{m-1} \\ n_{i_0+1} + \dots + n_{m-1} = r_3}} \prod_{i=i_0+1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \end{aligned} \quad (2.185)$$

On notera que $r_2 = i_0 + 1$. Comme n est fixé, toutes les sommes sont des sommes finies, car tous les termes sont nuls dès que $\|\mathbf{q}\| > n$, on a donc pas besoin de précautions pour les manipuler. On voit apparaître l'expression de $\sum_{\mathbf{q}_{i_0}} \mathcal{B}_{W_1, r_1}^{j_1 \bar{j}_{i_0}}(\mathbf{q}_{i_0})$. Puis en faisant les changements de variables $\mathbf{q}_i \rightarrow \mathbf{q}_i - \mathbf{q}_{i_0+1}$ pour tout $i > i_0 + 1$ on trouve l'expression de $\sum_{\mathbf{q}_m} \mathcal{B}_{W_2, r_3}^{j_{i_0+1} \bar{j}_m}(\mathbf{q}_m - \mathbf{q}_{i_0+1})$. On a alors :

$$\begin{aligned} \sum_{\mathbf{q}_m} \mathcal{B}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_m) &= \sum_{\substack{r_1, r_2, r_3 \\ r_1 + r_2 + r_3 = n}} \sum_{\bar{j}_{i_0}, j_{i_0+1}} \sum_{\mathbf{q}_{i_0}} \mathcal{B}_{W_1, r_1}^{j_1 \bar{j}_{i_0}}(\mathbf{q}_{i_0}) \\ &\quad \sum_{\substack{\mathbf{q}_{i_0+1} \\ \mathbf{q}_{i_0} \neq \mathbf{q}_{i_0+1}}} P_{r_2}^{\bar{j}_{i_0} j_{i_0+1}}(\mathbf{q}_{i_0+1} - \mathbf{q}_{i_0}) \sum_{\mathbf{q}_m} \mathcal{B}_{W_2, r_3}^{j_{i_0+1} \bar{j}_m}(\mathbf{q}_m - \mathbf{q}_{i_0+1}). \end{aligned} \quad (2.186)$$

Dans le membre de gauche, on change $\mathbf{q}_m$ en $\mathbf{q}$, c'est une variable muette, et dans le membre de droite on pose : $\mathbf{q}_1 = \mathbf{q}_{i_0}$, $\mathbf{q}_2 = \mathbf{q}_{i_0+1} - \mathbf{q}_{i_0}$ et $\mathbf{q}_3 = \mathbf{q}_m - \mathbf{q}_{i_0+1}$. On trouve alors :

$$\sum_{\mathbf{q}} \mathcal{B}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}) = \sum_{\substack{r_1, r_2, r_3 \\ r_1 + r_2 + r_3 = n}} \sum_{\bar{j}_{i_0}, j_{i_0+1}} \sum_{\mathbf{q}_1} \mathcal{B}_{W_1, r_1}^{j_1 \bar{j}_{i_0}}(\mathbf{q}_1) \sum_{\mathbf{q}_2 \neq 0} P_{r_2}^{\bar{j}_{i_0} j_{i_0+1}}(\mathbf{q}_2) \sum_{\mathbf{q}_3} \mathcal{B}_{W_2, r_3}^{j_{i_0+1} \bar{j}_m}(\mathbf{q}_3). \quad (2.187)$$

La somme sur $\bar{j}_{i_0}, j_{i_0+1}$ est en fait un produit matriciel, en renommant les r en n on a donc finalement :

$$\sum_{\mathbf{q}} \mathcal{B}_{W,n}(\mathbf{q}) = \sum_{\substack{n_1, n_2, n_3 \\ n_1 + n_2 + n_3 = n}} \sum_{\mathbf{q}_1} \mathcal{B}_{W_1, n_1}(\mathbf{q}_1) \sum_{\mathbf{q}_2 \neq 0} P_{n_2}(\mathbf{q}_2) \sum_{\mathbf{q}_3} \mathcal{B}_{W_2, n_3}(\mathbf{q}_3). \quad (2.188)$$

Étudions :

$$\sum_{\mathbf{q}_2 \neq 0} P_{n_2}(\mathbf{q}_2)$$

On avait l'équation 2.53 que l'on rappelle ici :

$$P_n^{ij}(\mathbf{q}) = \delta_n \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1 - \delta_n}{2d} g_{n-1}(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j).$$

On en déduit facilement que :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d \setminus \{0\}} P_n^{ij}(\mathbf{q}) = \delta_n \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1 - \delta_n}{2d} (1 - g_{n-1}(\mathbf{e}_i + \mathbf{e}_j)). \quad (2.189)$$

En définissant $\mathcal{C}_n$ la matrice telle que : $\forall i, j \in \mathcal{A}, \mathcal{C}^{ij} := g_n(\mathbf{e}_i + \mathbf{e}_j)$ on peut réécrire 2.54 ainsi :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} P_n(\mathbf{q}) = \delta_n I + \frac{1 - \delta_n}{2d} (C - \mathcal{C}_{n-1}). \quad (2.190)$$

On note que

$$\mathcal{C}_n = (g_n(2\mathbf{e}_1) - g_n(\mathbf{e}_1 + \mathbf{e}_2))I + (g_n(0) - g_n(\mathbf{e}_1 + \mathbf{e}_2))J + g_n(\mathbf{e}_1 + \mathbf{e}_2)C$$

On déduit facilement de la propriété 2.2 que $C\mathcal{B}_{W,n}(\mathbf{q}) = 0$, donc en posant

$$\begin{aligned} \tilde{\mathcal{C}}_n &:= \mathcal{C}_n - g_n(\mathbf{e}_1 + \mathbf{e}_2)C \\ &= (g_n(2\mathbf{e}_1) - g_n(\mathbf{e}_1 + \mathbf{e}_2))I + (g_n(0) - g_n(\mathbf{e}_1 + \mathbf{e}_2))J \end{aligned} \quad (2.191)$$

on trouve alors que pour tout $W = W_1 \star W_2$ et $\forall, n \in \mathbb{N}, \mathbf{q} \in \mathbb{Z}^d$:

$$\begin{aligned} \sum_{\mathbf{q}} \mathcal{B}_{W,n}(\mathbf{q}) &= \sum_{\substack{n_1, n_2, n_3 \\ n_1 + n_2 + n_3 = n}} \sum_{\mathbf{q}_1} \mathcal{B}_{W_1, n_1}(\mathbf{q}_1) \left(\delta_{n_2} I - \frac{(1 - \delta_{n_2}) \tilde{\mathcal{C}}_{n_2-1}}{2d} \right) \sum_{\mathbf{q}_3} \mathcal{B}_{W_2, n_3}(\mathbf{q}_3) \\ &= \sum_{\substack{n_1, n_3 \\ n_1 + n_3 = n}} \sum_{\mathbf{q}_1} \mathcal{B}_{W_1, n_1}(\mathbf{q}_1) \sum_{\mathbf{q}_3} \mathcal{B}_{W_2, n_3}(\mathbf{q}_3) \\ &\quad - \frac{1}{2d} \sum_{\substack{n_1, n_2, n_3 \\ n_1 + n_2 + n_3 = n-1}} \sum_{\mathbf{q}_1} \mathcal{B}_{W_1, n_1}(\mathbf{q}_1) \tilde{\mathcal{C}}_{n_2} \sum_{\mathbf{q}_3} \mathcal{B}_{W_2, n_3}(\mathbf{q}_3) \end{aligned} \quad (2.192)$$

Sommons sur n , pour cela il est pratique d'utiliser à nouveau la notation $\mathcal{U}$. On avait $W = W_1 \star W_2$, ou W est une partition de $\llbracket 1, m \rrbracket$, W_i est une partition de $\llbracket 1, m_i \rrbracket$ pour $i \in \{1, 2\}$ avec $m = m_1 + m_2$ (m_1 est le i_0 de tout à l'heure et $m_2 = m - m_1$), on sait que :

$$\sum_{\mathbf{q} \in \mathbb{Z}^d} \mathcal{B}_{W,n}(\mathbf{q}) = \sum_{\underline{\mathbf{q}}} \mathcal{U}_{W,n}(\mathbf{q}_2, \dots, \mathbf{q}_m). \quad (2.193)$$

D'où :

$$\begin{aligned} \sum_{\underline{\mathbf{q}}} \sum_{n=0}^N \mathcal{B}_{W,n}(\mathbf{q}) &= \sum_{\underline{\mathbf{q}}} \sum_{n=0}^N \sum_{\underline{n} \in \mathcal{S}_n} \mathcal{U}_{W_1, n_1}(\mathbf{q}_2, \dots, \mathbf{q}_{m_1}) \mathcal{U}_{W_2, n_2}(\bar{\mathbf{q}}_2, \dots, \bar{\mathbf{q}}_{m_2}) \\ &\quad - \frac{1}{2d} \sum_{\underline{\mathbf{q}}} \sum_{n=1}^N \sum_{\underline{n} \in \mathcal{S}_{n-1}} \mathcal{U}_{W_1, n_1}(\mathbf{q}_2, \dots, \mathbf{q}_{m_1}) \tilde{\mathcal{C}}_{n_2} \mathcal{U}_{W_2, n_2}(\bar{\mathbf{q}}_2, \dots, \bar{\mathbf{q}}_{m_2}) \end{aligned} \quad (2.194)$$

On veut prendre la limite quand $N \rightarrow +\infty$, or la série de terme général $\tilde{\mathcal{C}}_n$ est absolument convergente et la série de terme général $\|\mathcal{U}_{W_1, n_1}(\mathbf{q}_2, \dots, \mathbf{q}_{m_1})\|$ est convergente et sa somme est

intégrable. Ici, $\|\cdot\|$ est la norme d'une matrice, donnée par le maximum de la valeur absolue de ses éléments. Cette norme permet d'appliquer le théorème de convergence dominée, c'est tout ce qu'il nous faut. Donc

$$\left(\sum_{n_1=0}^{+\infty} \|\mathcal{U}_{W_1, n_1}(\mathbf{q}_2, \dots, \mathbf{q}_{m_1})\| \right) \left(I + \frac{1}{2d} \sum_{n_2=0}^{+\infty} \|\tilde{\mathcal{C}}_{n_2}\| \right) \left(\sum_{n_3=0}^{+\infty} \|\mathcal{U}_{W_2, n_3}(\bar{\mathbf{q}}_2, \dots, \bar{\mathbf{q}}_{m_2})\| \right)$$

est intégrable et nous permet d'appliquer le théorème de convergence dominée pour prouver que :

$$\begin{aligned} \lim_{n \rightarrow +\infty} \sum_{\mathbf{q}} \sum_{n=0}^N \mathcal{B}_{W, n}(\mathbf{q}) &= \sum_{\mathbf{q}} \sum_{n=0}^{+\infty} \mathcal{B}_{W, n}(\mathbf{q}) \\ &= \left(\sum_{\underline{\mathbf{q}}} \sum_{n_1=0}^{+\infty} \mathcal{U}_{W_1, n_1}(\mathbf{q}_2, \dots, \mathbf{q}_{m_1}) \right) \\ &\quad \left(I - \frac{1}{2d} \sum_{n_2=0}^{+\infty} \tilde{\mathcal{C}}_{n_2} \right) \left(\sum_{\underline{\mathbf{q}}} \sum_{n_2=0}^{+\infty} \mathcal{U}_{W_2, n_2}(\bar{\mathbf{q}}_2, \dots, \bar{\mathbf{q}}_{m_2}) \right) \end{aligned} \quad (2.195)$$

On définit la matrice

$$\mathcal{V}_W := \sum_{\underline{\mathbf{q}}} \sum_{n=0}^{+\infty} \mathcal{U}_{W, n}(\mathbf{q}_2, \dots, \mathbf{q}_m) \quad (2.196)$$

et la matrice

$$\mathbb{A} := I - \frac{1}{2d} \sum_{n=0}^{+\infty} \tilde{\mathcal{C}}_n. \quad (2.197)$$

On a donc :

$$\lim_{n \rightarrow +\infty} \sum_{\mathbf{q}} \sum_{r=0}^n \mathcal{B}_{W, r}(\mathbf{q}) = \mathcal{V}_W$$

et

$$\mathcal{V}_{W_1 \star W_2} = \mathcal{V}_{W_1} \mathbb{A} \mathcal{V}_{W_2} \quad (2.198)$$

Ce qui termine la preuve. □

2.2.10 Retour au calcul des κ_m

On a fini la preuve du théorème 2.1, on va maintenant donner une formule pour calculer les κ_m . On rappelle (2.64) :

$$K_W(n) = 2 \sum_{\mathbf{q}_2 \in \mathbb{Z}^d} \sum_{r=0}^{n-m} \frac{n-m-r}{n} (\mathcal{B}_{W, r}^{1,1}(\mathbf{q}_2) - \mathcal{B}_{W, r}^{1,-1}(\mathbf{q}_2))$$

On note que le terme $(n-m-r)/n$ ne pose aucun problème et donc on a :

$$\lim_{n \rightarrow +\infty} K_W(n) = 2(\mathcal{V}_W^{1,1} - \mathcal{V}_W^{1,-1}) \quad (2.199)$$

on a prouvé que cette limite est bien définie quel que soit W et on peut donc enfin définir :

$$\begin{aligned} \kappa_W &:= \lim_{n \rightarrow +\infty} K_W(n) \\ &= 2(\mathcal{V}_W^{1,1} - \mathcal{V}_W^{1,-1}) \end{aligned} \quad (2.200)$$

D'après l'équation (2.147) :

$$\begin{aligned}
\mathcal{V}_W^{j_1 \bar{j}_m} &= \sum_{\underline{\mathbf{q}}} \sum_{n=0}^{+\infty} \mathcal{U}_{W,n}^{j_1 \bar{j}_m}(\mathbf{q}_2, \dots, \mathbf{q}_m) \\
&= \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \sum_{n=0}^{+\infty} \sum_{\underline{n} \in \mathcal{S}_n} \prod_{i=1}^{m-1} P_{n_i}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \\
&= \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \prod_{i=1}^{m-1} \mathbf{G}^{\bar{j}_i j_{i+1}}(\mathbf{q}_{i+1} - \mathbf{q}_i) \\
&= \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \prod_{i=1}^{m-1} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{e}_{j_{i+1}}} + \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_{\bar{j}_i} - \mathbf{e}_{j_{i+1}})
\end{aligned} \tag{2.201}$$

car :

$$\mathbf{G}^{ij}(\mathbf{q}) = \delta_{\mathbf{e}_i, \mathbf{q}} \delta_{\mathbf{e}_i, \mathbf{e}_j} + \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_i - \mathbf{e}_j). \tag{2.202}$$

On a donc une formule pour calculer κ_W pour de petites valeurs de m . Ce que l'on va faire dès maintenant. Quand W est non-connecté on utilisera la formule :

$$\mathcal{V}_{W_1 \star W_2} = \mathcal{V}_{W_1} \mathbb{A} \mathcal{V}_{W_2}$$

On va voir comment calculer simplement κ_W à partir de cette formule. Premièrement on rappelle que :

$$\tilde{\mathcal{C}}_n := (g_n(2\mathbf{e}_1) - g_n(\mathbf{e}_1 + \mathbf{e}_2))I + (g_n(0) - g_n(\mathbf{e}_1 + \mathbf{e}_2))J$$

donc :

$$\begin{aligned}
\mathbb{A} &:= I - \frac{1}{2d} \sum_{n=0}^{+\infty} \tilde{\mathcal{C}}_n \\
&= I - \frac{(G(2\mathbf{e}_1) - G(\mathbf{e}_1 + \mathbf{e}_2))I + (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))J}{2d}
\end{aligned} \tag{2.203}$$

La seconde étape est de constater que $\mathcal{V}_W^{ij}$ est invariant par rotation de $(\mathbf{e}_i, \mathbf{e}_j)$. On voit donc que :

$$\mathcal{V}_W^{ij} = \begin{cases} \mathcal{V}_W^{1,1} & \text{si } \mathbf{e}_i \cdot \mathbf{e}_j = 1 \\ \mathcal{V}_W^{1,-1} & \text{si } \mathbf{e}_i \cdot \mathbf{e}_j = -1 \\ \mathcal{V}_W^{1,2} & \text{si } \mathbf{e}_i \cdot \mathbf{e}_j = 0 \end{cases}$$

On voit donc que

$$\mathcal{V}_W = a_W I + b_W J + c_W C$$

où I est la matrice identité, $\forall (i, j) \in \mathcal{P}_{ind}, I^{ij} := \delta_{ij}$, J est défini par $\forall (i, j) \in \mathcal{P}_{ind}, J^{ij} := \delta_{i, -j}$ et C est la matrice qui ne contient que des 1, $\forall (i, j) \in \mathcal{P}_{ind}, C^{ij} := 1$. Les notations $a_W := \mathcal{V}^{1,1} - \mathcal{V}^{1,2}$, $b_W := \mathcal{V}^{1,-1} - \mathcal{V}^{1,2}$ et $c_W := \mathcal{V}^{1,2}$ sont des notations que l'on ne réutilisera pas. On a alors :

$$\mathcal{V}_W^{1,1} - \mathcal{V}_W^{1,-1} = a_W - b_W$$

et si $W = W_1 \star W_2$ alors

$$\begin{aligned}
\mathcal{V}_W &= \mathcal{V}_{W_1} \mathbb{A} \mathcal{V}_{W_2} \\
&= (a_{W_1} I + b_{W_1} J + c_{W_1} C) \left(\left(1 - \frac{G(2\mathbf{e}_1) - G(\mathbf{e}_1 + \mathbf{e}_2)}{2d}\right) I - \frac{G(0) - G(\mathbf{e}_1 + \mathbf{e}_2)}{2d} J \right) \\
&\quad (a_{W_2} I + b_{W_2} J + c_{W_2} C)
\end{aligned} \tag{2.204}$$

or,

$$\mathcal{V}_W = a_W I + b_W J + c_W C \quad (2.205)$$

On peut donc facilement trouver l'expression de a_W et de b_W , l'expression de c_W est plus compliquée et n'est pas utile. On trouve alors en quelques lignes de calcul :

$$\begin{aligned} \mathcal{V}_W^{1,1} - \mathcal{V}_W^{1,-1} &= a_W - b_W \\ &= (a_{W_1} - b_{W_1})(a_{W_2} - b_{W_2}) \left(1 + \frac{G(0) - G(2\mathbf{e}_1)}{2d}\right) \\ &= (\mathcal{V}_{W_1}^{1,1} - \mathcal{V}_{W_1}^{1,-1})(\mathcal{V}_{W_2}^{1,1} - \mathcal{V}_{W_2}^{1,-1}) \left(1 + \frac{G(0) - G(2\mathbf{e}_1)}{2d}\right) \end{aligned} \quad (2.206)$$

On pose $\eta := 1 + \frac{G(0) - G(2\mathbf{e}_1)}{2d}$. On a donc :

$$\mathcal{V}_{W_1 \star W_2}^{1,1} - \mathcal{V}_{W_1 \star W_2}^{1,-1} = (\mathcal{V}_{W_1}^{1,1} - \mathcal{V}_{W_1}^{1,-1})\eta(\mathcal{V}_{W_2}^{1,1} - \mathcal{V}_{W_2}^{1,-1}) \quad (2.207)$$

Par conséquent, on trouve par récurrence immédiate que si W est un graphe non-connecté tel que $W = W_1 \star \dots \star W_k$ alors :

$$\kappa_W = 2\eta^{k-1} \prod_{i=1}^k (\mathcal{V}_{W_i}^{1,1} - \mathcal{V}_{W_i}^{1,-1}) = 2\eta^{k-1} \prod_{i=1}^k \frac{\kappa_{W_i}}{2} \quad (2.208)$$

On peut alors conjecturer la formule suivante, non démontrée mais que l'on pourrait utiliser quand même dans des calculs numériques :

$$\kappa = 1 - \frac{2}{\eta} + \frac{2}{\eta} \prod_{W \text{ connectés}} \sum_{k=0}^{+\infty} (\varepsilon^{m(W)} \eta \kappa_W / 2)^k \quad (2.209)$$

$$= 1 - \frac{2}{\eta} + \frac{2}{\eta} \prod_{W \text{ connectés}} \frac{1}{1 - \varepsilon^{m(W)} \eta \kappa_W / 2} \quad (2.210)$$

où $m(W)$ est le m correspondant à W . Rentrons à présent dans le détail du calcul des premier termes de la série.

2.3 Valeurs numériques aux petits ordres

Nous allons consacrer la fin de ce chapitre au calcul des deux premiers termes perturbatifs de la série des κ_m , en l'occurrence κ_2 et κ_3 . Nous rappelons que dans le modèle des permutations, $\kappa_0 = 1$ et $\kappa_1 = 0$. On reprend la formule :

$$\mathcal{V}_W^{j_1, \bar{j}_m} = \sum_{(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W} \sum_{\underline{j}} L_W(j_1, \dots, \bar{j}_m) \prod_{i=1}^{m-1} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{q}_{i+1} - \mathbf{q}_i} \delta_{\mathbf{e}_{\bar{j}_i}, \mathbf{e}_{j_{i+1}}} + \frac{1}{2d} G(\mathbf{q} - \mathbf{e}_{\bar{j}_1} - \mathbf{e}_{j_{m+1}}) \quad (2.211)$$

Ainsi que :

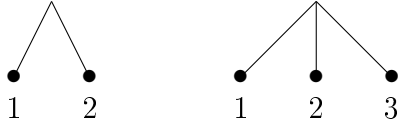
$$\kappa_W = 2(\mathcal{V}_W^{1,1} - \mathcal{V}_W^{1,-1})$$

Quand W est non-connecté on utilisera la formule :

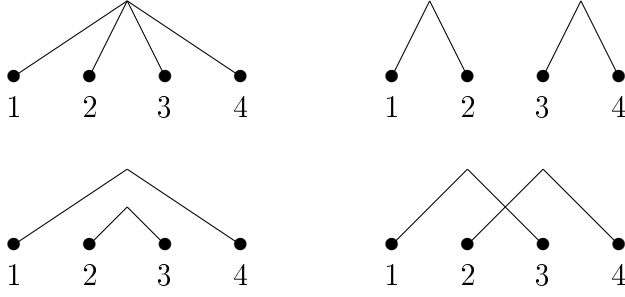
$$\mathcal{V}_{W_1 \star W_2}^{1,1} - \mathcal{V}_{W_1 \star W_2}^{1,-1} = (\mathcal{V}_{W_1}^{1,1} - \mathcal{V}_{W_1}^{1,-1})\eta(\mathcal{V}_{W_2}^{1,1} - \mathcal{V}_{W_2}^{1,-1})$$

Avec $\eta := 1 + \frac{G(0) - G(2\mathbf{e}_1)}{2d}$.

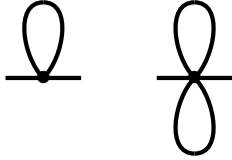
Nous allons lister les partitions sans singletons de $\llbracket 1, m \rrbracket$ pour $m \in \{2, 3, 4\}$. Les voici avec la représentation en graphes pour $m \in \{2, 3\}$:



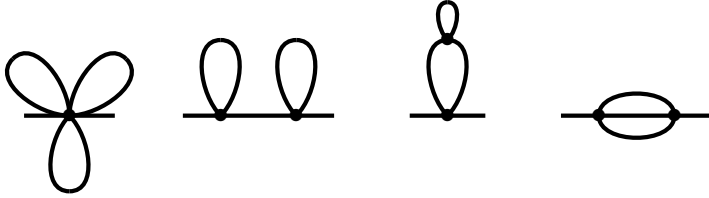
Et pour $m = 4$:



En nous inspirant de la *lace expansion* nous allons utiliser une représentation plus compacte de ces objets. Nous allons représenter la trajectoire parcourue par la particule. Deux points reliés dans l'un des graphes ci-dessus représentent deux instants où la particule passe au même endroit. Par exemple, pour $m = 2$, la seule partition représentée correspond à une trajectoire où la particule passe deux fois au même endroit, et fait donc une boucle entre les deux. On précise qu'elle passe au même endroit mais pas avec la même vitesse. On peut donc représenter cette partition par une trajectoire en forme de boucle. Pour $m = 3$ on aura une double boucle. Voici donc, dans le même ordre que précédemment, les partitions sans singletons, représentées par leurs trajectoires, pour $m \in \{2, 3\}$:



Puis pour $m = 4$:



On est donc prêt à calculer $\kappa_2 = \kappa_{\underline{\emptyset}}$, puis $\kappa_3 = \kappa_{\emptyset}$. Nous allons voir que ces deux calculs sont déjà très longs, nous n'aurons donc malheureusement pas le temps d'aborder le calcul de $\kappa_4 = \kappa_{\infty} + \kappa_{\underline{\emptyset}} + \kappa_{\emptyset} + \kappa_{\ominus}$.

Remarque : La représentation à partir des trajectoires est intéressante car elle est très visuelle, mais à partir de $m = 5$ elle ne fonctionne plus aussi bien. Il faudrait la travailler un peu plus mais ce n'est pas utile ici car on ne dépasse pas $m = 3$.

2.3.1 Calcul de κ_2

Il n'y a qu'une seule partition qui participe à κ_2 , c'est la partition à une seule partie de $\llbracket 1, 2 \rrbracket$ que l'on note $\underline{\emptyset}$. Donc, $\kappa_2 = \kappa_{\underline{\emptyset}}$.

Avant toute chose on va avoir besoin de calculer $L_{\underline{\emptyset}}(j_1, \bar{j}_1, j_2, \bar{j}_2)$. On rappelle la définition de L donnée en (2.145) :

$$\forall(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W, L_W(j_1, \dots, \bar{j}_m) = \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{e}_{j_i}), (\mathbf{q}_i + \mathbf{e}_{\bar{j}_i}, \mathbf{e}_{\bar{j}_i}), \sigma) \right]. \quad (2.212)$$

Pour la suite des calculs il est plus pratique d'écrire des calculs faisant intervenir des vecteurs unitaires $\mathbf{p} \in \mathcal{P}$ plutôt que des indices $j \in \mathcal{P}_{ind}$ comme on l'a beaucoup fait à la section précédente.

On pose donc que :

$$L_W(\mathbf{e}_{j_1}, \dots, \mathbf{e}_{\bar{j}_m}) := L_W(j_1, \dots, \bar{j}_m). \quad (2.213)$$

donc :

$$\forall(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W, L_W(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_m) = \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right]. \quad (2.214)$$

Calculons donc : $L_{\underline{\mathcal{Q}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2)$.

L'ensemble $H_{\underline{\mathcal{Q}}}$ est en fait très simple, $H_{\underline{\mathcal{Q}}} = \{0\}$ car $\mathbf{q}_2 \in H_{\underline{\mathcal{Q}}} \Rightarrow \mathbf{q}_2 = \mathbf{q}_1 = 0$, toujours avec la convention $\mathbf{q}_1 = 0$. Par conséquent :

$$L_{\underline{\mathcal{Q}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) = \mathbb{E}_{\mathbb{Q}} [\hat{p}((0, \mathbf{p}_1), (\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_1), \sigma) \hat{p}((0, \mathbf{p}_2), (\bar{\mathbf{p}}_2, \bar{\mathbf{p}}_2), \sigma)] \quad (2.215)$$

On avait à l'équation (2.14) :

$$\hat{p}((\mathbf{q}, \mathbf{p}), (\bar{\mathbf{q}}, \bar{\mathbf{p}}), \sigma) := \delta_{\bar{\mathbf{q}}, \mathbf{q} + \bar{\mathbf{p}}} \left(\delta_{\bar{\mathbf{p}}, \sigma_q(\mathbf{p})} - \frac{1}{2d} \right) \quad (2.216)$$

donc :

$$\begin{aligned} L_{\underline{\mathcal{Q}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) &= \mathbb{E}_{\mathbb{Q}} \left[\left(\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} - \frac{1}{2d} \right) \left(\delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} - \frac{1}{2d} \right) \right] \\ &= \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}] - \frac{1}{(2d)^2} \end{aligned} \quad (2.217)$$

Car $\forall \mathbf{p}, \bar{\mathbf{p}} \in \mathcal{P}, \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}, \sigma_0(\mathbf{p})}] = \frac{1}{2d}$. C'est la probabilité que la permutation aléatoire σ_0 envoie $\mathbf{p}$ sur $\bar{\mathbf{p}}$. Calculons à présent $\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}]$. C'est la probabilité que la permutation aléatoire σ_0 envoie simultanément $\mathbf{p}_1$ sur $\bar{\mathbf{p}}_1$ et $\mathbf{p}_2$ sur $\bar{\mathbf{p}}_2$. On a :

$$\begin{aligned} \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}] &= \mathbb{Q}(\sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_2) = \bar{\mathbf{p}}_2) \\ &= \mathbb{Q}(\sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1) \mathbb{Q}(\sigma_0(\mathbf{p}_2) = \bar{\mathbf{p}}_2 | \sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1) \end{aligned}$$

On a trois cas possibles :

— Si $\mathbf{p}_1 = \mathbf{p}_2$ et $\bar{\mathbf{p}}_1 = \bar{\mathbf{p}}_2$ alors

$$\mathbb{Q}(\sigma_0(\mathbf{p}_2) = \bar{\mathbf{p}}_2 | \sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1) = 1.$$

On est dans ce cas si et seulement si $\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 2$.

— Si $\mathbf{p}_1 = \mathbf{p}_2$ et $\bar{\mathbf{p}}_1 \neq \bar{\mathbf{p}}_2$ ou si $\mathbf{p}_1 \neq \mathbf{p}_2$ et $\bar{\mathbf{p}}_1 = \bar{\mathbf{p}}_2$ alors

$$\mathbb{Q}(\sigma_0(\mathbf{p}_2) = \bar{\mathbf{p}}_2 | \sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1) = 0.$$

On est dans ce cas si et seulement si $\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 1$.

— Enfin, dans le cas où $\mathbf{p}_1 \neq \mathbf{p}_2$ et $\bar{\mathbf{p}}_1 \neq \bar{\mathbf{p}}_2$, on a :

$$\mathbb{Q}(\sigma_0(\mathbf{p}_2) = \bar{\mathbf{p}}_2 | \sigma_0(\mathbf{p}_1) = \bar{\mathbf{p}}_1) = \frac{1}{2d-1}.$$

On est dans ce cas si et seulement si $\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 0$.

On a donc :

$$\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}] = \begin{cases} \frac{1}{2d} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 2 \\ 0 & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 1 \\ \frac{1}{2d(2d-1)} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 0 \end{cases} \quad (2.218)$$

Et, par conséquent :

$$L_{\underline{d}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) = \begin{cases} \frac{1}{2d} - \frac{1}{(2d)^2} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 2 \\ -\frac{1}{(2d)^2} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 1 \\ \frac{1}{2d(2d-1)} - \frac{1}{(2d)^2} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 0 \end{cases} \quad (2.219)$$

On notera qu'ici, regarder la valeur de $\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2}$ n'est qu'une manière (rapide à écrire) de déterminer dans lequel des trois cas décrits ci-dessus on est. On pourra vérifier en testant les 3 cas possibles que la formule suivante est vraie, et elle est bien pratique :

$$L_{\underline{d}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) = \frac{1}{2d-1} \left(\delta_{\mathbf{p}_1, \mathbf{p}_2} - \frac{1}{2d} \right) \left(\delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} - \frac{1}{2d} \right) \quad (2.220)$$

On établit à présent deux propriétés importante. Premièrement :

Proposition 2.11. *Soit $m \in \mathbb{N} \setminus \{0, 1\}$, W une partition de $\llbracket 1, m \rrbracket$ et $(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) \in \mathcal{P}^{2m}$ alors $\forall i \in \llbracket 1, m \rrbracket$,*

$$\sum_{\mathbf{p} \in \mathcal{P}} L_W(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_{i-1}, \mathbf{p}, \bar{\mathbf{p}}_i, \dots, \bar{\mathbf{p}}_m) = 0$$

et

$$\sum_{\bar{\mathbf{p}} \in \mathcal{P}} L_W(\mathbf{p}_1, \dots, \mathbf{p}_i, \bar{\mathbf{p}}, \mathbf{p}_{i+1}, \dots, \bar{\mathbf{p}}_m) = 0$$

Preuve : On a :

$$\hat{p}((0, \mathbf{p}), (\bar{\mathbf{p}}, \bar{\mathbf{p}}), \sigma) = \delta_{\bar{\mathbf{p}}, \sigma_0(\mathbf{p})} - \frac{1}{2d}$$

D'où, clairement,

$$\sum_{\mathbf{p} \in \mathcal{P}} \hat{p}((0, \mathbf{p}), (\bar{\mathbf{p}}, \bar{\mathbf{p}}), \sigma) = \sum_{\bar{\mathbf{p}} \in \mathcal{P}} \hat{p}((0, \mathbf{p}), (\bar{\mathbf{p}}, \bar{\mathbf{p}}), \sigma) = 0.$$

Sachant (2.214) on déduit facilement que la propriété est vraie. □

Cette propriété est généralisable à une vaste classe de modèles, incluant le modèle des miroirs. La propriété suivante, par contre, est spécifique au modèle des permutations.

Proposition 2.12. *Soit $m \in \mathbb{N} \setminus \{0, 1\}$, W une partition de $\llbracket 1, m \rrbracket$ et $(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) \in \mathcal{P}^{2m}$ alors :*

$$L_W(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) = L_W(\mathbf{p}_1, -\bar{\mathbf{p}}_1, \dots, \mathbf{p}_i, -\bar{\mathbf{p}}_i, \dots, \mathbf{p}_m, -\bar{\mathbf{p}}_m)$$

Preuve : On va prouver plus généralement que pour toute permutation s de $\mathcal{P}$ on a :

$$L_W(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) = L_W(\mathbf{p}_1, s(\bar{\mathbf{p}}_1), \dots, \mathbf{p}_i, s(\bar{\mathbf{p}}_i), \dots, \mathbf{p}_m, s(\bar{\mathbf{p}}_m))$$

Soit $(\mathbf{q}_2, \dots, \mathbf{q}_m) \in H_W$, on a alors :

$$L_W(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) = \mathbb{E}_{\mathbb{Q}} \left[\prod_{i=1}^m \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \quad (2.221)$$

or, $\mathbf{q}_i \neq \mathbf{q}_j$ si et seulement si $\hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma)$ et $\hat{p}((\mathbf{q}_j, \mathbf{p}_j), (\mathbf{q}_j + \bar{\mathbf{p}}_j, \bar{\mathbf{p}}_j), \sigma)$ sont indépendants, donc si la partition W a n sous-parties, on note ça $|W| = n$, on peut écrire $W = (W_1, \dots, W_n)$ et donc :

$$L_W(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) = \mathbb{E}_{\mathbb{Q}} \left[\prod_{j=1}^{|W|} \prod_{i \in W_j} \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \quad (2.222)$$

Or, par définition de H_W , $\mathbf{q}_i \neq \mathbf{q}_{i'}$ si et seulement si i et i' sont dans la même sous partie W_j . Par conséquent, on a :

$$\begin{aligned} L_W(\mathbf{p}_1, \bar{\mathbf{p}}_1, \dots, \mathbf{p}_m, \bar{\mathbf{p}}_m) &= \prod_{j=1}^{|W|} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i \in W_j} \hat{p}((\mathbf{q}_i, \mathbf{p}_i), (\mathbf{q}_i + \bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \\ &= \prod_{j=1}^{|W|} \mathbb{E}_{\mathbb{Q}} \left[\prod_{i \in W_j} \hat{p}((0, \mathbf{p}_i), (\bar{\mathbf{p}}_i, \bar{\mathbf{p}}_i), \sigma) \right] \end{aligned} \quad (2.223)$$

Or, on a vu que :

$$\hat{p}((\mathbf{q}, \mathbf{p}), (\mathbf{q} + \bar{\mathbf{p}}, \bar{\mathbf{p}}), \sigma) = \delta_{\bar{\mathbf{p}}, \sigma_{\mathbf{q}}(\mathbf{p})} - \frac{1}{2d}. \quad (2.224)$$

D'où, en notant $S_{\mathcal{P}}$ le groupe des permutations de $\mathcal{P}$ et avec le changement de variable $\sigma'_0 = s \circ \sigma_0$,

$$\begin{aligned} L_W(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_m) &= \prod_{j=1}^{|W|} \frac{1}{(2d)!} \sum_{\sigma_0 \in S_{\mathcal{P}}} \prod_{i \in W_j} \left(\delta_{\bar{\mathbf{p}}, \sigma_0(\mathbf{p})} - \frac{1}{2d} \right) \\ &= \prod_{j=1}^{|W|} \frac{1}{(2d)!} \sum_{\sigma'_0 \in S_{\mathcal{P}}} \prod_{i \in W_j} \left(\delta_{\bar{\mathbf{p}}, s^{-1} \circ \sigma'_0(\mathbf{p})} - \frac{1}{2d} \right) \\ &= \prod_{j=1}^{|W|} \frac{1}{(2d)!} \sum_{\sigma'_0 \in S_{\mathcal{P}}} \prod_{i \in W_j} \left(\delta_{s(\bar{\mathbf{p}}), \sigma'_0(\mathbf{p})} - \frac{1}{2d} \right) \\ &= L_W(\mathbf{p}_1, s(\bar{\mathbf{p}}_1), \dots, \mathbf{p}_m, s(\bar{\mathbf{p}}_m)) \end{aligned} \quad (2.225)$$

□

On a alors tous les outils pour calculer $\kappa_{\underline{\mathbb{L}}}$. Notons $\mathcal{V}_{\underline{\mathbb{L}}}^{\mathbf{e}_i, \mathbf{e}_j} := \mathcal{V}_{\underline{\mathbb{L}}}^{i, j}$ et commençons par calculer $\mathcal{V}_{\underline{\mathbb{L}}}^{\mathbf{p}_1, \bar{\mathbf{p}}_2}$:

$$\begin{aligned} \mathcal{V}_{\underline{\mathbb{L}}}^{\mathbf{p}_1, \bar{\mathbf{p}}_2} &= \sum_{\mathbf{q}_2 \in H_{\underline{\mathbb{L}}}} \sum_{\bar{\mathbf{p}}_1, \mathbf{p}_2} L_{\underline{\mathbb{L}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) \left(\delta_{\bar{\mathbf{p}}_1, \mathbf{q}_2 - \mathbf{q}_1} \delta_{\bar{\mathbf{p}}_1, \mathbf{p}_2} + \frac{1}{2d} G(\mathbf{q}_2 - \mathbf{q}_1 - \bar{\mathbf{p}}_1 - \mathbf{p}_2) \right) \\ &= \frac{1}{2d} \sum_{\bar{\mathbf{p}}_1, \mathbf{p}_2} L_{\underline{\mathbb{L}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) G(\bar{\mathbf{p}}_1 + \mathbf{p}_2) \end{aligned}$$

car $H_{\underline{\mathbb{L}}} = \{0\}$. En isolant les cas $\bar{\mathbf{p}}_1 = \mathbf{p}_2$ et $\bar{\mathbf{p}}_1 = -\mathbf{p}_2$ puis en appliquant la propriété 2.11 on trouve :

$$\begin{aligned} \mathcal{V}_{\underline{\mathbb{L}}}^{\mathbf{p}_1, \bar{\mathbf{p}}_2} &= \frac{1}{2d} \sum_{\mathbf{p}_2} L_{\underline{\mathbb{L}}}(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_2, \bar{\mathbf{p}}_2) (G(2\mathbf{e}_1) - G(\mathbf{e}_1 + \mathbf{e}_2)) + L_{\underline{\mathbb{L}}}(\mathbf{p}_1, -\mathbf{p}_2, \mathbf{p}_2, \bar{\mathbf{p}}_2) (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2)) \\ &+ \frac{1}{2d} \sum_{\bar{\mathbf{p}}_1, \mathbf{p}_2} L_{\underline{\mathbb{L}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) G(\mathbf{e}_1 + \mathbf{e}_2) \\ &= \frac{1}{2d} \sum_{\mathbf{p}_2} L_{\underline{\mathbb{L}}}(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_2, \bar{\mathbf{p}}_2) (G(2\mathbf{e}_1) - G(\mathbf{e}_1 + \mathbf{e}_2)) + L_{\underline{\mathbb{L}}}(\mathbf{p}_1, -\mathbf{p}_2, \mathbf{p}_2, \bar{\mathbf{p}}_2) (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2)) \end{aligned} \quad (2.226)$$

Calculons $\mathcal{V}_{\underline{\mathbb{L}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{L}}}^{1,-1}$. Grâce à l'équation ci-dessus et à la propriété 2.12 on trouve que :

$$\begin{aligned} L_{\underline{\mathbb{L}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{e}_1) &= L_{\underline{\mathbb{L}}}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, -\mathbf{e}_1) \\ L_{\underline{\mathbb{L}}}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, \mathbf{e}_1) &= L_{\underline{\mathbb{L}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{e}_1) \end{aligned}$$

d'où

$$\mathcal{V}_{\underline{\mathbb{L}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{L}}}^{1,-1} = \frac{1}{2d} \sum_{\mathbf{p}_2} (L_{\underline{\mathbb{L}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{e}_1) - L_{\underline{\mathbb{L}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{e}_1)) (G(2\mathbf{e}_1) - G(0))$$

On rappelle l'équation 2.220 :

$$L_{\underline{\mathbb{Q}}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2) = \frac{1}{2d-1} \left(\delta_{\mathbf{p}_1, \mathbf{p}_2} - \frac{1}{2d} \right) \left(\delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} - \frac{1}{2d} \right)$$

dont on déduit que

$$\begin{aligned} \sum_{\mathbf{p}_2} (L_{\underline{\mathbb{Q}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{e}_1) - L_{\underline{\mathbb{Q}}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{e}_1)) &= \sum_{\mathbf{p}_2} \frac{1}{2d-1} \left(\delta_{\mathbf{e}_1, \mathbf{p}_2} - \frac{1}{2d} \right) (\delta_{\mathbf{p}_2, \mathbf{e}_1} - \delta_{\mathbf{p}_2, -\mathbf{e}_1}) \\ &= \frac{1}{2d-1} \left(1 - \frac{1}{2d} \right) + \frac{1}{2d-1} \left(\frac{1}{2d} \right) \\ &= \frac{1}{2d-1} \end{aligned}$$

Au final, on trouve que

$$\begin{aligned} \kappa_{\underline{\mathbb{Q}}} &= 2(\mathcal{V}_{\underline{\mathbb{Q}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{Q}}}^{1,-1}) \\ &= \frac{G(2\mathbf{e}_1) - G(0)}{d(2d-1)} \end{aligned} \quad (2.227)$$

On obtient du même coup $\kappa_{\underline{\mathbb{Q}}\underline{\mathbb{Q}}}$, car $\underline{\mathbb{Q}}\underline{\mathbb{Q}} = \underline{\mathbb{Q}} \star \underline{\mathbb{Q}}$:

$$\begin{aligned} \kappa_{\underline{\mathbb{Q}}\underline{\mathbb{Q}}} &= 2(\mathcal{V}_{\underline{\mathbb{Q}}\underline{\mathbb{Q}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{Q}}\underline{\mathbb{Q}}}^{1,-1}) \\ &= 2(\mathcal{V}_{\underline{\mathbb{Q}} \star \underline{\mathbb{Q}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{Q}} \star \underline{\mathbb{Q}}}^{1,-1}) \\ &= 2(\mathcal{V}_{\underline{\mathbb{Q}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{Q}}}^{1,-1}) \eta(\mathcal{V}_{\underline{\mathbb{Q}}}^{1,1} - \mathcal{V}_{\underline{\mathbb{Q}}}^{1,-1}) \\ &= \frac{\eta}{2} (\kappa_2)^2 \end{aligned} \quad (2.228)$$

2.3.2 Calcul de κ_3

Il n'y a à nouveau qu'une seule partition qui participe à κ_3 , c'est la partition à une seule partie de $\llbracket 1, 3 \rrbracket$ que l'on note $\mathfrak{g}$. Donc, $\kappa_3 = \kappa_{\mathfrak{g}}$. De plus, on a $H_{\mathfrak{g}} = \{0, 0\}$.

Calculons maintenant $\kappa_{\mathfrak{g}} = 2(\mathcal{V}_{\mathfrak{g}}^{1,1} - \mathcal{V}_{\mathfrak{g}}^{1,-1})$. Commençons par calculer, grâce à la propriété 2.11 :

$$\begin{aligned} \mathcal{V}_{\mathfrak{g}}^{\mathbf{p}_1, \bar{\mathbf{p}}_3} &= \frac{1}{(2d)^2} \sum_{\substack{\mathbf{q}_2, \mathbf{q}_3 \\ (\mathbf{q}_1, \mathbf{q}_2, \mathbf{q}_3) \in H_{\mathfrak{g}}}} \sum_{\bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2, \mathbf{p}_3} L_{\mathfrak{g}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2, \mathbf{p}_3, \bar{\mathbf{p}}_3) (\delta_{\bar{\mathbf{p}}_1, \mathbf{q}_2 - \mathbf{q}_1} \delta_{\bar{\mathbf{p}}_1, \mathbf{p}_2} + G(\mathbf{q}_2 - \mathbf{q}_1 - \bar{\mathbf{p}}_1 - \mathbf{p}_2)) \\ &\quad (\delta_{\bar{\mathbf{p}}_2, \mathbf{q}_3 - \mathbf{q}_2} \delta_{\bar{\mathbf{p}}_2, \mathbf{p}_3} + G(\mathbf{q}_3 - \mathbf{q}_2 - \bar{\mathbf{p}}_2 - \mathbf{p}_3)) \\ &= \frac{1}{(2d)^2} \sum_{\bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2, \mathbf{p}_3} L_{\mathfrak{g}}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2, \mathbf{p}_3, \bar{\mathbf{p}}_3) G(\bar{\mathbf{p}}_1 + \mathbf{p}_2) G(\bar{\mathbf{p}}_2 + \mathbf{p}_3) \\ &= \frac{1}{(2d)^2} \sum_{\mathbf{p}_2, \mathbf{p}_3} L_{\mathfrak{g}}(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \bar{\mathbf{p}}_3) (G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 \\ &\quad + (L_{\mathfrak{g}}(\mathbf{p}_1, -\mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \bar{\mathbf{p}}_3) + L_{\mathfrak{g}}(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, \bar{\mathbf{p}}_3)) \\ &\quad (G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))(G(0) - G(\mathbf{e}_1 + \mathbf{e}_2)) \\ &\quad + L_{\mathfrak{g}}(\mathbf{p}_1, -\mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, \bar{\mathbf{p}}_3) (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 \end{aligned} \quad (2.229)$$

Puis, d'après la propriété 2.12

$$L_{\mathfrak{g}}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) = L_{\mathfrak{g}}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1)$$

et

$$L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) = L_{\emptyset}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1)$$

et

$$L_{\emptyset}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) = L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1)$$

on en déduit que :

$$\begin{aligned} \kappa_{\emptyset} &= 2(\mathcal{V}_{\emptyset}^{1,1} - \mathcal{V}_{\emptyset}^{1,-1}) \\ &= \frac{2}{(2d)^2} \sum_{\mathbf{p}_2, \mathbf{p}_3} (L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) - L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1))(G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 \\ &\quad + (L_{\emptyset}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) - L_{\emptyset}(\mathbf{e}_1, -\mathbf{p}_2, \mathbf{p}_2, -\mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1))(G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 \\ &= \frac{1}{2d^2} ((G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 - (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))^2) \\ &\quad \sum_{\mathbf{p}_2, \mathbf{p}_3} L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) - L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1) \end{aligned} \quad (2.230)$$

Pour calculer cette dernière somme on a besoin de l'expression de $L_{\emptyset}$, celle-ci est plus compliquée que celle de $L_{\underline{\emptyset}}$ mais le principe est le même, allons-y ! D'abord, avec un raisonnement similaire à celui qui a mené à 2.217 on trouve que :

$$\begin{aligned} L_{\emptyset}(\mathbf{p}_1, \bar{\mathbf{p}}_1, \mathbf{p}_2, \bar{\mathbf{p}}_2, \mathbf{p}_3, \bar{\mathbf{p}}_3) &= \mathbb{E}_{\mathbb{Q}} \left[\left(\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} - \frac{1}{2d} \right) \left(\delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} - \frac{1}{2d} \right) \left(\delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)} - \frac{1}{2d} \right) \right] \\ &= \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} \delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)}] - \frac{1}{2d} (\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}] \\ &\quad + \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} \delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)}] + \mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)} \delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)}]) + \frac{2}{(2d)^3}. \end{aligned} \quad (2.231)$$

On rappelle 2.218 :

$$\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)}] = \begin{cases} \frac{1}{2d} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 2 \\ 0 & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 1 \\ \frac{1}{2d(2d-1)} & \text{si } \delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} = 0 \end{cases} \quad (2.232)$$

Et en reprenant le raisonnement utilisé pour $L_{\underline{\emptyset}}$ on trouve que la valeur de $\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} \delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)}]$ dépend des valeurs de $\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2}$, $\delta_{\mathbf{p}_1, \mathbf{p}_3} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_3}$ et $\delta_{\mathbf{p}_2, \mathbf{p}_3} + \delta_{\bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3}$, non ordonnées. On définit donc une notation très temporaire :

$$A(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3, \bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3) = (\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2}, \delta_{\mathbf{p}_1, \mathbf{p}_3} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_3}, \delta_{\mathbf{p}_2, \mathbf{p}_3} + \delta_{\bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3})$$

On note que $A(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3, \bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3)$ n'a que 7 valeurs possibles (sans ordre) dont voici la liste :

$$\{(0, 0, 0), (2, 0, 0), (2, 2, 2), (1, 0, 0), (1, 1, 0), (1, 1, 1), (2, 1, 1)\}$$

On trouve donc que :

$$\mathbb{E}_{\mathbb{Q}} [\delta_{\bar{\mathbf{p}}_1, \sigma_0(\mathbf{p}_1)} \delta_{\bar{\mathbf{p}}_2, \sigma_0(\mathbf{p}_2)} \delta_{\bar{\mathbf{p}}_3, \sigma_0(\mathbf{p}_3)}] = \begin{cases} \frac{1}{2d(2d-1)(2d-2)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (0, 0, 0) \\ \frac{1}{2d(2d-1)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 0, 0) \\ \frac{1}{2d} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 2, 2) \\ 0 & \text{sinon} \end{cases}$$

Et par conséquent on peut calculer $L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3)$ pour chacune des 7 valeurs de A , on trouve :

$$L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = \begin{cases} \frac{4}{(2d)^3(2d-1)(2d-2)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (0, 0, 0) \\ \frac{-2}{(2d)^3(2d-1)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 0, 0) \\ \frac{2d-2}{(2d)^3(2d-1)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 0, 0) \\ \frac{2d-2}{(2d)^3(2d-1)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 1, 0) \\ \frac{2}{(2d)^3} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 1, 1) \\ \frac{-2d+1}{(2d)^3} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 1, 1) \\ \frac{(2d-1)(2d-2)}{(2d)^3} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 2, 2) \end{cases} \quad (2.233)$$

Si on regarde les rapports des différentes valeurs possibles de $L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3)$ on trouve que le rapport de deux valeurs est toujours le produit de $1-d$ et $1-2d$ ou leurs inverses. On peut le montrer en réécrivant le tableau ci-dessus différemment :

$$4d^3 L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = \begin{cases} \frac{1}{(1-2d)(1-d)} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (0, 0, 0) \\ \frac{1}{1-2d} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 0, 0) \\ \frac{1-d}{1-2d} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 0, 0) \\ \frac{1-d}{1-2d} & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 1, 0) \\ 1 & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (1, 1, 1) \\ 1-2d & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 1, 1) \\ (1-d)(1-2d) & \text{si } A(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = (2, 2, 2) \end{cases} \quad (2.234)$$

On intuite à partir de là la formule :

$$L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = \frac{1}{4d^3(2d-1)(d-1)} (1-d\delta_{\mathbf{p}_1, \mathbf{p}_2})(1-d\delta_{\mathbf{p}_1, \mathbf{p}_3})(1-d\delta_{\mathbf{p}_2, \mathbf{p}_3}) \left(1 - \left(\frac{d}{d-1}\right)^2 \delta_{\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3}\right) \\ (1-d\delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2})(1-d\delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_3})(1-d\delta_{\bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3}) \left(1 - \left(\frac{d}{d-1}\right)^2 \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3}\right) \quad (2.235)$$

où $\delta_{\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3} := \delta_{\mathbf{p}_1, \mathbf{p}_2} \delta_{\mathbf{p}_1, \mathbf{p}_3}$. On démontre cette formule en la testant sur les 7 cas possibles. On peut alors développer cette formule pour en obtenir une autre encore plus compacte que voici, on ne détaille pas les calculs qui ne sont pas mystérieux.

$$L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = \frac{1}{4d^3(2d-1)(d-1)} (1-d(\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\mathbf{p}_1, \mathbf{p}_3} + \delta_{\mathbf{p}_2, \mathbf{p}_3}) + 2d^2 \delta_{\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3}) \\ (1-d(\delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2} + \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_3} + \delta_{\bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3}) + 2d^2 \delta_{\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3}) \quad (2.236)$$

On note que l'on peut démontrer cette dernière formule directement en la testant sur les 7 cas possibles et ainsi ignorer complètement la formule (2.235). Mais on a choisi de montrer les deux formules qui ont chacune leurs avantages.

Cette dernière formule combinée à (2.230) nous permet de calculer κ_3 . On va faire le calcul en deux parties, premièrement on va calculer rapidement

$$X_1 := (G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 - (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))^2$$

puis

$$X_2 := \sum_{\mathbf{p}_2, \mathbf{p}_3} L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) - L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1).$$

On leur a donné des noms juste pour plus de commodité dans l'écriture des calculs. On peut réécrire (2.230) ainsi :

$$\kappa_3 = \frac{1}{2d^2} X_1 X_2 \quad (2.237)$$

Pour calculer X_1 on va exprimer $G(\mathbf{e}_1 + \mathbf{e}_2)$ en fonction de $G(2\mathbf{e}_1)$ et $G(0)$ c'est facile car c'est la fonction de Green de la marche aléatoire simple sur $\mathbb{Z}^d$, on a donc la formule suivante :

$$G(\mathbf{q}) = \sum_{i \in \mathcal{P}_{ind}} \frac{1}{2d} G(\mathbf{q} + \mathbf{e}_i). \quad (2.238)$$

On l'applique pour $\mathbf{q} = 0$:

$$G(0) = 1 + G(\mathbf{e}_1), \quad (2.239)$$

puis pour $\mathbf{q} = \mathbf{e}_1$:

$$G(\mathbf{e}_1) = \frac{1}{2d} [G(0) + G(2\mathbf{e}_1) + (2d - 2)G(\mathbf{e}_1 + \mathbf{e}_2)] \quad (2.240)$$

Et on en déduit que :

$$G(\mathbf{e}_1 + \mathbf{e}_2) = \frac{1}{2d - 2} [(2d - 1)G(0) - 2d - G(2\mathbf{e}_1)] \quad (2.241)$$

Par conséquent :

$$\begin{aligned} X_1 &:= (G(2\mathbf{e}_2) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 - (G(0) - G(\mathbf{e}_1 + \mathbf{e}_2))^2 \\ &= G(2\mathbf{e}_2)^2 - G(0)^2 + 2G(\mathbf{e}_1 + \mathbf{e}_2)(G(0) - G(2\mathbf{e}_1)) \\ &= \frac{1}{d - 1} ((d - 1)G(2\mathbf{e}_1)^2 - (d - 1)G(0)^2 + ((2d - 1)G(0) - 2d - G(2\mathbf{e}_1))(G(0) - G(2\mathbf{e}_1))) \\ &= \frac{1}{d - 1} (dG(2\mathbf{e}_1)^2 + dG(0)^2 - 2dG(2\mathbf{e}_1)G(0) + 2dG(2\mathbf{e}_1) - 2dG(0)) \\ &= \frac{d}{d - 1} (G(2\mathbf{e}_1) - G(0)) (G(2\mathbf{e}_1) - G(0) + 2) \end{aligned} \quad (2.242)$$

On garde ça de côté et on passe au calcul de X_2 grâce à la formule (2.244). On pose :

$$Y(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3) := 1 - d(\delta_{\mathbf{p}_1, \mathbf{p}_2} + \delta_{\mathbf{p}_1, \mathbf{p}_3} + \delta_{\mathbf{p}_2, \mathbf{p}_3}) + 2d^2 \delta_{\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3}. \quad (2.243)$$

Ainsi, (2.244) peut s'écrire :

$$L_{\emptyset}(\mathbf{p}_1, \dots, \bar{\mathbf{p}}_3) = \frac{1}{4d^3(2d - 1)(d - 1)} Y(\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3) Y(\bar{\mathbf{p}}_1, \bar{\mathbf{p}}_2, \bar{\mathbf{p}}_3) \quad (2.244)$$

Et par conséquent :

$$\begin{aligned} X_2 &:= \sum_{\mathbf{p}_2, \mathbf{p}_3} L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, \mathbf{e}_1) - L_{\emptyset}(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_2, \mathbf{p}_3, \mathbf{p}_3, -\mathbf{e}_1) \\ &= \sum_{\mathbf{p}_2, \mathbf{p}_3} \frac{1}{4d^3(2d - 1)(d - 1)} Y(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3) (Y(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3) - Y(-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3)). \end{aligned} \quad (2.245)$$

Or :

$$Y(\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3) - Y(-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3) = d(\delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{-\mathbf{e}_1, \mathbf{p}_3} - \delta_{\mathbf{e}_1, \mathbf{p}_2} - \delta_{\mathbf{e}_1, \mathbf{p}_3}) + 2d^2(\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - \delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3}). \quad (2.246)$$

Donc,

$$\begin{aligned}
X_2 &= \frac{d}{4d^3(2d-1)(d-1)} \sum_{\mathbf{p}_2, \mathbf{p}_3} (1 - d(\delta_{\mathbf{e}_1, \mathbf{p}_2} + \delta_{\mathbf{e}_1, \mathbf{p}_3} + \delta_{\mathbf{p}_2, \mathbf{p}_3}) + 2d^2\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3}) \\
&\quad (\delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{-\mathbf{e}_1, \mathbf{p}_3} - \delta_{\mathbf{e}_1, \mathbf{p}_2} - \delta_{\mathbf{e}_1, \mathbf{p}_3} + 2d(\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - \delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3})). \\
&= \frac{1}{4d^2(2d-1)(d-1)} \sum_{\mathbf{p}_2, \mathbf{p}_3} \delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{-\mathbf{e}_1, \mathbf{p}_3} - \delta_{\mathbf{e}_1, \mathbf{p}_2} - \delta_{\mathbf{e}_1, \mathbf{p}_3} + 2d(\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - \delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3}) \\
&\quad + d(-\delta_{\mathbf{e}_1, \mathbf{p}_2}\delta_{-\mathbf{e}_1, \mathbf{p}_3} + \delta_{\mathbf{e}_1, \mathbf{p}_2} + \delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - 2d\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - \delta_{\mathbf{e}_1, \mathbf{p}_3}\delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} + \delta_{\mathbf{e}_1, \mathbf{p}_3} \\
&\quad - 2d\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - 2\delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} + 2\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} - 2d\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} + 2d\delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3}) \\
&\quad + 2d^2\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3}(-2 + 2d) \\
&= \frac{1}{4d^2(2d-1)(d-1)} \sum_{\mathbf{p}_2, \mathbf{p}_3} (d-1)(\delta_{\mathbf{e}_1, \mathbf{p}_2} + \delta_{\mathbf{e}_1, \mathbf{p}_3}) + \delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{-\mathbf{e}_1, \mathbf{p}_3} \\
&\quad - d(\delta_{\mathbf{e}_1, \mathbf{p}_2}\delta_{-\mathbf{e}_1, \mathbf{p}_3} + \delta_{\mathbf{e}_1, \mathbf{p}_3}\delta_{-\mathbf{e}_1, \mathbf{p}_2}) + (2d + d - 2d^2 + d - 2d^2 + 2d - 2d^2 - 4d^2 + 4d^3)\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} \\
&\quad + (-2d - 2d + 2d^2)\delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} \tag{2.247}
\end{aligned}$$

On peut calculer ça facilement car :

$$\sum_{\mathbf{p}_2, \mathbf{p}_3} \delta_{\mathbf{e}_1, \mathbf{p}_2} = 2d \text{ et } \sum_{\mathbf{p}_2, \mathbf{p}_3} \delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} = \sum_{\mathbf{p}_2, \mathbf{p}_3} \delta_{\mathbf{e}_1, \mathbf{p}_2} \delta_{-\mathbf{e}_1, \mathbf{p}_3} = 1. \tag{2.248}$$

D'où :

$$\begin{aligned}
X_2 &= \frac{1}{4d^2(2d-1)(d-1)} \sum_{\mathbf{p}_2, \mathbf{p}_3} (d-1)(\delta_{\mathbf{e}_1, \mathbf{p}_2} + \delta_{\mathbf{e}_1, \mathbf{p}_3}) + \delta_{-\mathbf{e}_1, \mathbf{p}_2} + \delta_{-\mathbf{e}_1, \mathbf{p}_3} \\
&\quad - d(\delta_{\mathbf{e}_1, \mathbf{p}_2}\delta_{-\mathbf{e}_1, \mathbf{p}_3} + \delta_{\mathbf{e}_1, \mathbf{p}_3}\delta_{-\mathbf{e}_1, \mathbf{p}_2}) + (6d - 10d^2 + 4d^3)\delta_{\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} \\
&\quad + (-4d + 2d^2)\delta_{-\mathbf{e}_1, \mathbf{p}_2, \mathbf{p}_3} \\
&= \frac{1}{4d^2(2d-1)(d-1)} (4d(d-1) + 4d - 2d + 6d - 10d^2 + 4d^3 - 4d + 2d^2) \\
&= \frac{4d^3 - 4d^2}{4d^2(2d-1)(d-1)} \\
&= \frac{1}{2d-1} \tag{2.249}
\end{aligned}$$

Un résultat si simple pour un calcul si long pose question, il est possible que certaines simplifications astucieuses permettent un calcul beaucoup plus simple de κ_3 et on pourrait espérer une généralisation possible à κ_m pour $m > 3$.

On reprend donc la formule (2.237) :

$$\kappa_3 = \frac{1}{2d^2} X_1 X_2 \tag{2.250}$$

D'où :

$$\kappa_3 = \frac{1}{d(2d-1)(2d-2)} (G(2\mathbf{e}_1) - G(0)) (G(2\mathbf{e}_1) - G(0) + 2) \tag{2.251}$$

On remarque d'ailleurs que :

$$\kappa_3 = \kappa_2 \frac{G(2\mathbf{e}_1) - G(0) + 2}{2d-2} \tag{2.252}$$

2.3.3 Résumé des résultats analytiques

On pourrait continuer à calculer d'autres valeurs de κ_m mais les calculs deviennent de plus en plus compliqués au fur et à mesure, on décide donc de s'arrêter là. On souhaite obtenir une

approximation de la valeur de κ en fonction de d et ε . On reprend donc la formule (2.45) :

$$K(n) = \sum_{m=0}^n \varepsilon^m K_m(n) \quad (2.253)$$

Le théorème 2.1 montre que chaque terme de la somme converge vers $\varepsilon^m \kappa_m$ quand $n \rightarrow +\infty$, on conjecture donc la formule suivante :

Conjecture 2.1. $\forall d \geq 3, \forall \varepsilon \in [0, 1]$

$$\kappa = \sum_{m=0}^{+\infty} \varepsilon^m \kappa_m.$$

Puis on tronque cette formule pour obtenir :

$$\kappa \simeq \kappa_0 + \varepsilon^2 \kappa_2 + \varepsilon^3 \kappa_3. \quad (2.254)$$

On a pas de garantie théorique que cette formule soit une bonne approximation de κ , c'est pourquoi nous allons la tester numériquement. Faisons donc le point sur les résultats quantitatifs donnés par cette formule, on étudiera $d = 3$, $d = 4$ et $d = 5$. On doit d'abord calculer $G(2\mathbf{e}_1)$ et $G(0)$ pour $d = 3$ à $d = 5$, pour cela on utilise la formule :

$$G(\mathbf{q}) = \int_{\mathbf{k} \in [-\pi, \pi]} \frac{e^{-i\mathbf{k} \cdot \mathbf{q}}}{1 - \frac{1}{d} \sum_{j=1}^d \cos(\mathbf{k} \cdot \mathbf{e}_j)} d\mathbf{k}. \quad (2.255)$$

Cette formule se démontre par transformation de Fourier. On trouvera une idée de la preuve au chapitre 1 de [8].

Pour le cas particulier du calcul de $G(0)$ pour $d = 3$, qui est une intégrale de Watson, on utilise la très belle formule :

$$G(0) = \frac{\sqrt{6}}{32\pi^3} \Gamma\left(\frac{1}{24}\right) \Gamma\left(\frac{5}{24}\right) \Gamma\left(\frac{7}{24}\right) \Gamma\left(\frac{11}{24}\right). \quad (2.256)$$

On notera que cette formule a été démontrée à l'origine en 1977 par Glasser et Zucker mais qu'ils ont fait une erreur dans l'article d'origine. Ce n'est pas très grave mais malheureusement je n'ai pas pu trouver de version corrigée de cet article. Cette erreur est clarifiée et corrigée à la page 10 de [33] par Zucker lui-même. Je préfère citer cette source plutôt que l'article d'origine car j'ai perdu un certain temps à comprendre le problème et je souhaite éviter cette perte de temps à d'autres.

À partir de ces formules, le logiciel Mathématica nous donne les approximations suivantes :

d	$G(2\mathbf{e}_1)$	$G(0)$
3	0.2573358863	1.516386059151978
4	0.0659641	1.23947
5	0.0275044	1.15631

On trouve alors :

d	κ_2	κ_3	κ_{00}
3	-0.0839367	-0,0155482168	0,0042618887
4	-0.041910925	-0,0057731887	0,0010070936
5	-0,02508452	-0,0027316936	0,0003501306

On garde ces résultats et on va à présent s'attaquer à l'approche numérique du modèle des permutations. On va obtenir ainsi des approximations numériques de κ à comparer aux valeurs ci-dessus.

Chapitre 3

Approche numérique

Le but de ce chapitre est de présenter une approche numérique des modèles étudiés dans cette thèse. On présentera d'abord les algorithmes utilisés puis les résultats des simulations numériques effectuées. L'approche numérique est complémentaire de l'approche analytique et permet d'évaluer la pertinence de cette dernière. Elle permet aussi de guider des recherches futures et de donner des valeurs approchées de quantités très difficiles à calculer. C'est une très bonne manière d'aborder de nouveaux problèmes. C'est d'ailleurs par là que nous avons commencé l'étude du modèle des miroirs. Sans plus attendre, voici les algorithmes utilisés.

3.1 Simulation des marches aléatoires en milieu aléatoire

Tous les modèles étudiés dans cette thèse sont des marches aléatoires en milieu aléatoire (MAMA). Certaines sont des marches déterministes en milieu aléatoire, ce qui peut être vu comme une sous-catégorie des MAMAs. Simuler de tels modèles peut se faire à partir d'une forme d'algorithme générale. Le but de cette section et de la suivante est de donner cette trame générale permettant d'écrire un algorithme de simulation pour n'importe quelle MAMA sur un graphe discret muni d'une relation d'ordre totale sur ses sommets (cette relation pouvant être complètement arbitraire, par exemple l'ordre lexicographique, c'est juste utile algorithmiquement).

3.1.1 Algorithme général

On souhaite simuler des réalisations d'une MAMA. Une première démarche extrêmement naïve consiste à prendre au pied de la lettre la définition du modèle et à tirer aléatoirement un environnement complet, puis à lancer une marche aléatoire dans l'environnement ainsi créé. Cette démarche est extrêmement coûteuse en ressources informatiques car il faut créer et stocker en mémoire tout un environnement alors que la plupart des positions ne seront même pas visitées.

Une manière plus intelligente de procéder consiste à créer l'environnement au fur et à mesure de son exploration. Voici la structure de l'algorithme permettant cela :

1. La particule est à la position $\mathbf{q}$ avec la vitesse $\mathbf{p}$.
2. Est-ce que le noyau de transition $T_{\mathbf{q}}$ existe déjà ? Si non, tirer aléatoirement $T_{\mathbf{q}}$ et le stocker en mémoire.
3. Tirer aléatoirement la nouvelle vitesse $\mathbf{p}$ à partir de l'ancienne et du noyau $T_{\mathbf{q}}$.
4. Faire un pas : $\mathbf{q} \leftarrow \mathbf{q} + \mathbf{p}$.
5. Revenir au point 1.

Cet algorithme nécessite une structure de données pour stocker l'environnement au fur et à mesure. Choisissons une telle structure.

3.1.2 Structures de données

On cherche une structure offrant deux possibilités. On veut pouvoir trouver rapidement si l'environnement $T_{\mathbf{q}}$ a déjà été créé à une position $\mathbf{q}$ donnée et on veut pouvoir stocker rapidement en mémoire une nouvelle partie de l'environnement à chaque fois que l'on visite une position pour la première fois. En terme algorithmique on cherche donc une structure efficace en recherche et en insertion.

Précisons quelques points techniques. En fonction des résultats recherchés, on va vouloir simuler une marche dans un environnement donné et/ou pendant un temps donné. Par exemple, pour obtenir la probabilité de traverser dans le tore au chapitre 1, on fixe l'espace : le tore de taille N contenant N^d points puis on lance la marche pour un temps indéterminé que l'on notera n mais qui n'est pas connu à l'avance, on l'arrête quand elle touche le bord droit ou gauche. À l'opposé, pour avoir le déplacement carré moyen, on veut lancer la marche sans limites spatiales et la laisser faire un nombre de pas n déterminé à l'avance pour mesurer sa position au bout de n pas.

Pour ces deux situations nous cherchons à optimiser nos algorithmes. En algorithmique, les quantités à optimiser sont la complexité en espace et la complexité en temps, qui sont des notions très classiques. Schématiquement, la complexité en espace est le nombre de cases mémoires nécessaires pour faire tourner le programme et la complexité en temps est le temps de calcul. Sur le plan théorique, ces quantités sont définies à une constante près, c'est pourquoi on utilise la notation O , dite "grand O". Par exemple, si on dit que la complexité en temps est $O(n^2)$, cela signifie que le temps de calcul augmente proportionnellement au carré du nombre de pas réalisés par la marche (ce nombre de pas est noté n). En pratique on cherchera un bon compromis entre la complexité en temps et la complexité en espace.

Une première idée de structure consiste à créer un très grand tableau en mémoire vive comme si on voulait stocker tout l'environnement et à le remplir au fur et à mesure. C'est la solution la plus simple en pratique et celle que j'ai utilisée pour les premières simulations du chapitre 1. Dans ce cas, la quantité de mémoire utilisée est égale à la taille du tableau N^d multipliée par le nombre d'octets nécessaires pour stocker un point de l'environnement. Pour le modèle des miroirs ou celui des permutations, on a besoin de stocker $2d$ entiers à chaque fois, c'est ce qu'il faut pour coder une permutation de l'espace des vitesses $\mathcal{P}$. On rappelle que les miroirs sont des cas particuliers de permutation de $\mathcal{P}$. La complexité en espace totale est alors égale à $2dN^d$ octets car on peut stocker les entiers en question sur un seul octet, sauf si $2d > 256$, ce qui n'est jamais le cas. Le nombre d'opérations à chaque pas ne dépend alors ni de n ni de N , la complexité en temps est donc $O(n)$ où n est le nombre de pas de la marche. Il est utile de voir que dans cette situation on a $E[n] = O(Nd)$.

On voit donc que cette solution est très gourmande en mémoire mais elle convient en dimension 2 ou 3. Par exemple, pour $d = 3$, avec 6 gigaoctets de mémoire vive, mon ordinateur peut simuler jusqu'à $N = 1000$, ce qui est suffisant. Mais dans ce cas la majeure partie de la mémoire affectée n'est pas utilisée, ce qui est dommage. Pour le déplacement carré moyen c'est vraiment gênant car on ne connaît pas à l'avance l'espace nécessaire. On peut bricoler des solutions mais mieux vaut changer de structure de données. Explorons rapidement les solutions (trop) simples.

Une première idée pour diminuer la complexité en espace est de remplacer le tableau par une liste des positions visitées. On a alors la complexité en espace égale à $O(n)$, ce qui est beaucoup mieux. Mais le problème est que pour savoir si une position a déjà été visitée il faut parcourir toute la liste. Chaque pas de la marche prend alors un temps $O(n)$ et faire n pas nous prend alors un temps $O(n^2)$, ce qui est énorme. On peut ramener le temps de recherche à $O(\log_2(n))$ en ordonnant la liste mais c'est alors le temps d'insertion qui devient $O(n)$, cela ne fait que déplacer le problème.

Heureusement il existe une solution adaptée pour traiter ce problème, une structure de données nommée "arbre équilibré". Cette structure permet des recherches et des insertions en temps $O(\log_2(n))$, ce qui nous permet d'obtenir une complexité en temps de $O(n \log_2(n))$, ce qui est très acceptable. Dans nos algorithmes nous utilisons la première version connue d'arbre équilibré : l'arbre AVL. Cet arbre tire son nom des noms respectifs de ses deux inventeurs, Georgii Adelson-Velsky et

Evgenii Landis, qui l'ont publié en 1962, voir [25]. On ne présentera pas en détail l'arbre AVL ici car c'est une structure classique en informatique théorique, on pourra en trouver une description dans l'ouvrage de Donald E. Knuth [21] qui consacre un chapitre aux arbres équilibrés, dont l'arbre AVL. L'arbre AVL est une structure de données qui permet de rendre très facilement accessibles de grandes quantités de données à condition d'avoir une relation d'ordre totale sur ces données. Cette relation d'ordre peut être tout à fait arbitraire. Et comme nous travaillons sur des données indexées par des positions dans $\mathbb{Z}^d$, nous avons choisi l'ordre lexicographique sur $\mathbb{Z}^d$. Cela règle notre problème de complexité et nous avons donc un algorithme suffisamment performant pour obtenir les résultats recherchés.

On notera qu'il existe une autre structure de données qui paraît très adaptée au problème : la table de hachage [20, 21]. Le principe de la table de hachage est d'avoir un tableau d'une taille raisonnable (nous allons préciser cela) et d'avoir une fonction nommée *fonction de hachage* qui à chaque donnée associe une position dans le tableau. Le point important est que cette fonction n'est pas injective car le tableau est plus petit que l'ensemble des données potentielles. Du coup, il est très facile de ranger une donnée, il suffit d'appliquer la fonction de hachage pour savoir où la ranger et il est aussi très simple de savoir si une donnée est déjà stockée, il suffit de regarder dans la case où elle devrait être, grâce à la fonction de hachage. Mais le prix à payer est que deux données peuvent être envoyées dans la même case mémoire, on appelle cela une collision. En cas de collision on s'arrange pour stocker les deux données au même endroit, typiquement en stockant dans le tableau la première donnée ainsi qu'un pointeur vers la seconde, en gros en commençant une liste chaînée. Si on fait ça une fois ça ne pose pas de problème, mais si les collisions sont nombreuses alors l'accès aux données est fortement ralenti et on perd tout l'intérêt de la table de hachage. L'idée est donc de trouver une fonction de hachage intelligente et adaptée au problème donné.

Pour le problème qui nous intéresse, on propose d'utiliser comme table de hachage un tableau de dimension d et de côté N . Ce tableau contient N^d cases et il faudra prendre N assez grand pour que N^d soit plus grand que le nombre total de données à stocker durant l'exécution du programme. Typiquement, si on veut faire n pas de la MAMA étudiée il suffit de prendre $N^d > 3n$ (valeur très arbitraire), cela limite raisonnablement les collisions. Si la place mémoire n'est pas trop limitée, on peut prendre N^d plus grand que ça. Reste le choix de la fonction de hachage avec l'idée que deux données effectivement rencontrées lors de l'exécution doivent le plus rarement possible entrer en collision. On rappelle que les données sont des permutations à des positions données et que deux permutations sont toujours à deux positions différentes de $\mathbb{Z}^d$. Pour limiter les collisions on propose simplement de prendre comme fonction de hachage les coordonnées modulo N de la position de la permutation. Ainsi on aura une collision si deux permutations sont à deux positions différentes dont les coordonnées sont égales modulo N , elle seront donc forcément distantes d'au moins une distance N donc la probabilité de collision est a priori assez faible. Il n'y apparemment pas de phénomène particulier qui causerait un nombre anormal de collision. On peut donc dire que notre fonction de hachage est la projection par l'opération de modulo de $\mathbb{Z}^d$ dans $(\mathbb{Z}/N\mathbb{Z})^d$. On suppose que dans notre cas, la table de hachage peut être légèrement plus performante que l'arbre AVL. On en attend en pratique un algorithme plus rapide que pour l'arbre AVL mais légèrement plus volumineux en place mémoire, ce qui n'est pas un problème. Malheureusement nous n'avons pas eu le temps de tester cette structure de données. Pour bien comprendre les tables de hachage on trouvera un bon tutoriel dans [20], les articles Wikipédia en Français et en Anglais sont aussi de bonne qualité. Pour une référence plus académique on pourra consulter [21].

Avant de donner les résultats des simulations on va encore discuter deux petits points techniques : les nombres pseudo-aléatoires et la parallélisation.

3.1.3 Nombres pseudo-aléatoires

Les modèles étudiés sont aléatoires, cependant aucun ordinateur ne peut générer un nombre aléatoire. On a alors recours à des générateurs d'entiers pseudo-aléatoires. Un tel générateur est un

algorithme qui génère une suite d'entiers (déterministe) qui ressemble beaucoup à une suite d'entiers aléatoires. Cette ressemblance est quantifiée et les générateurs de nombres pseudo-aléatoires sont les fruits d'une théorie mathématique avancée. La suite de nombres générée est généralement périodique. Si le nombre d'utilisations du générateur dans un même programme est supérieur à sa période, alors le caractère aléatoire est perdu. Cela peut arriver avec des générateurs simples qui ont une période de 2^{32} .

Ces générateurs doivent être initialisés en leur fournissant un entier, généralement codé sur 32 bit, qui permet de déterminer à quel endroit de la période on va commencer à faire fonctionner notre générateur. Cet entier est appelé une graine. Il est important de ne pas donner deux fois la même graine au générateur sous peine de travailler deux fois avec la même séquence de nombres pseudo-aléatoires.

Dans tous les algorithmes, sauf mention contraire explicite, nous avons utilisé le générateur *Mersenne Twister 19937*, très utilisé en physique statistique, qui a une période de $2^{19937} - 1$. Cet algorithme est de très bonne qualité. Son défaut majeur est d'être prévisible ce qui le rend caduc en cryptographie. Mais pour la simulation il est idéal. Dans tout le reste de la discussion, on considérera donc que l'on utilise des nombres réellement aléatoires.

3.1.4 Parallélisation

Une fois l'algorithme conçu et le code écrit, le travail de simulation consiste à réaliser de très nombreuses fois la même expérience. Par exemple, pour le déplacement carré moyen, pour chaque valeur de d et de ε , on lance plus ou moins 100 000 réalisations différentes de la marche, et ce pour différentes valeurs de n , où n est le nombre de pas avant la mesure du déplacement carré moyen. Le temps de calcul est donc assez élevé et il est très utile de paralléliser le calcul. Pour cela, on lance des copies du programme sur tous les processeurs disponibles grâce à la fonction *fork* du langage C. Il faut juste s'assurer que les différentes réalisations sont bien indépendantes, il serait stupide d'utiliser 20 processeurs pour calculer exactement 20 fois la même chose. Mais il serait aussi problématique qu'il y ait des corrélations entre les différentes réalisations. Pour éviter cela, il faut que chaque processus ait une graine différente, on procède de la manière suivante. On utilise une graine commune, entrée à la main et généralement tirée au dé (tirer la graine au dé est une forme de perfectionnisme peu utile mais peu coûteux) et on utilise cette graine pour générer des graines secondaires à l'aide du générateur aléatoire par défaut du langage C. Ces graines secondaires seront utilisées par les différentes copies du processus principal pour initialiser le générateur *Mersenne Twister 19937*. Ainsi, on considère que les nombres pseudo-aléatoires utilisés par les divers processus sont indépendants.

Sur le calculateur du LPSM nous avons pu utiliser jusqu'à 24 cœurs en parallèle, ce qui nous a permis de gagner un facteur 10 par rapport à un ordinateur de bureau à 2 cœurs. Cette parallélisation est assez modeste comparée à des simulations tournant sur des milliers de cœurs ou plus mais elle nous a été bien utile.

3.2 Résultats et comparaisons avec l'approche perturbative

3.2.1 Déplacement carré moyen dans le modèle des permutations

Le déplacement carré moyen est la quantité la plus importante. Ce sont ces simulations qui nous permettent de juger l'utilité de l'approche perturbative présentée au chapitre 2. On va le mesurer d'abord pour le modèle des permutations afin de comparer aux résultats analytiques, ensuite on le mesurera pour le modèle des miroirs car il y a un phénomène de discontinuité intéressant lié au critère de réversibilité. Pour chacun de ces modèles, on va simuler la marche pour diverses valeurs de d et de ε .

Pour chaque couple (d, ε) , on simule la marche pendant n pas et on mesure la distance à l'origine au bout de ces n pas, n est le temps d'excursion. On répète l'opération m fois avec typiquement

$m = 10^6$. Pour chaque réalisation i on obtient un estimateur statistique de la quantité $K(n)$ définie à l'équation (2.9). Cet estimateur est :

$$\hat{K}(n, i) := \frac{\mathbf{q}(n, i)^2}{n} \quad (3.1)$$

où $\mathbf{q}(n, i)$ est la position de la particule au temps n dans la simulation numéro i . On rappelle que pour d et ε fixés :

$$K(n) = \frac{1}{n} \sum_{x \in \mathcal{M}} \|x\|^2 \mathbb{E}_{\mathbb{Q}}[p_n(O, x, \sigma, \varepsilon)] \quad (3.2)$$

Par conséquent, $\hat{K}(n, i)$ est un estimateur de $K(n)$ et $\tilde{K}(n, m) := \frac{1}{m} \sum_{i=1}^m \hat{K}(n, i)$ est un meilleur estimateur de la même quantité. Ces estimateurs sont sans biais, on a donc :

$$K(n) = \mathbb{E}[\hat{K}(n, i)] \quad (3.3)$$

Et on a la convergence en loi :

$$\frac{1}{m} \sum_{i=1}^m \hat{K}(n, i) \xrightarrow{m \rightarrow +\infty} \mathcal{N}\left(K(n), \frac{\text{Var}[\hat{K}(n, 1)]}{m}\right) \quad (3.4)$$

Ce qui nous permet, pour n assez grand, d'obtenir $K(n)$ avec l'intervalle de confiance à 95% suivant :

$$K(n) \in \left[\frac{1}{m} \sum_{i=1}^m \hat{K}(n, i) - 1.96 \sqrt{\frac{\text{Var}[\hat{K}(n, 1)]}{m}}, \frac{1}{m} \sum_{i=1}^m \hat{K}(n, i) + 1.96 \sqrt{\frac{\text{Var}[\hat{K}(n, 1)]}{m}} \right]. \quad (3.5)$$

Cette expression fait intervenir la variance de $\hat{K}(n, 1)$ qui est une quantité inconnue. On va donc utiliser un estimateur sans biais de cette variance que voici :

$$\hat{\theta}(n, m) = \frac{1}{m-1} \left(\sum_{i=1}^m \hat{K}(n, i)^2 - \frac{1}{m} \left(\sum_{i=1}^m \hat{K}(n, i) \right)^2 \right). \quad (3.6)$$

On peut ainsi tracer notre estimateur $\tilde{K}(n, m)$ et son intervalle de confiance pour différentes valeurs du temps d'excursion n . Par exemple pour $d = 3$ et $\varepsilon = 0.8$ on obtient le graphe 3.1. Pour chaque point on a le nombre m de tests qui vaut entre 1 000 000 et 2 000 000.

On voit alors que $K(n)$ semble converger quand $n \rightarrow +\infty$ car les différences de valeurs en fonction de n sont de l'ordre de grandeur de l'intervalle de confiance à 95%. Les variations de $\tilde{K}(n, m)$ en fonction de n sont donc l'ordre de grandeur des fluctuations statistiques attendues si $K(n)$ était convergent. Ce n'est pas surprenant, cela correspond au résultat de Bálint Tóth [7]. Nous n'avons pas de résultat théorique sur la vitesse de convergence donc on en est réduit à faire un choix arbitraire "à la main" pour évaluer κ . On explique les variations de l'estimateur $\tilde{K}(n, m)$ par les variations fondamentales de $K(n)$ et par les variations aléatoires de l'estimateur. On estime donc à la main à partir de quelle valeur de n les variations fondamentales de $K(n)$ sont petites devant les variations aléatoires de l'estimateur $\tilde{K}(n, m)$. En prenant ainsi n assez grand on considère que $\tilde{K}(n, m)$ est un bon estimateur de κ , et l'intervalle de confiance à 95% est donné par la formule (3.5). On se contentera de cette approximation satisfaisante.

Ainsi, on peut obtenir la valeur estimée de κ en fonction de d et de ε . Voici les résultats obtenus :

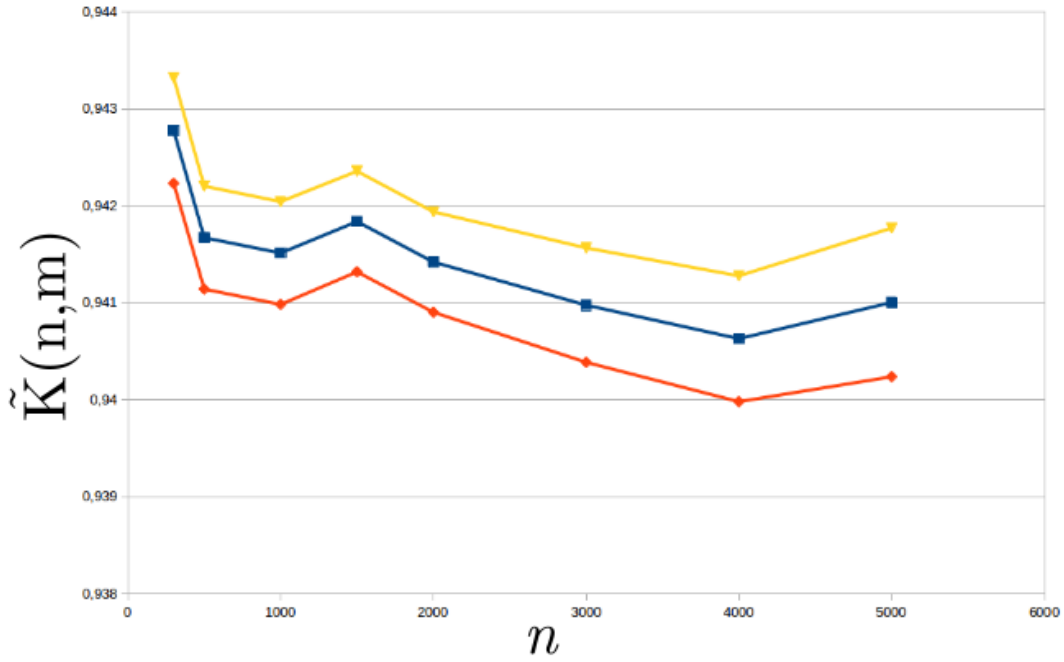


FIGURE 3.1 – L'estimateur $\tilde{K}(n, m)$ en fonction de n pour $d = 3$ et $\varepsilon = 0.8$. Les courbes jaune et orange représentent $\tilde{K}(n, m)$ plus ou moins l'estimateur de son écart type : $\sqrt{\hat{\theta}(n, m)}$.

$\varepsilon \backslash d$	3	4
0	1	1
0.1	0.9990 ± 0.0006	.
0.2	0.9967 ± 0.0006	.
0.3	0.9925 ± 0.0009	.
0.4	0.9863 ± 0.0009	.
0.5	0.9775 ± 0.0009	0.9895 ± 0.0009
0.6	0.9680 ± 0.0009	.
0.7	0.9565 ± 0.0009	0.9803 ± 0.0009
0.8	0.9410 ± 0.0009	0.9750 ± 0.0009
0.9	0.9225 ± 0.0009	0.9691 ± 0.0009
0.95	0.9110 ± 0.0009	0.9660 ± 0.0009
0.97	0.9050 ± 0.0009	.
0.99	0.8970 ± 0.0009	.
1	0.8890 ± 0.0009	0.9625 ± 0.0009

On notera que la valeur pour $\varepsilon = 0$ est une valeur théorique bien connue, c'est le coefficient de diffusion de la marche aléatoire simple.

On peut alors comparer ces valeurs à l'approximation :

$$\kappa \simeq \kappa_0 + \varepsilon^2 \kappa_2 + \varepsilon^3 \kappa_3. \quad (3.7)$$

Voici les comparaisons sous forme de graphiques, à la figure 3.2 pour $d = 3$ et à la figure 3.3 pour $d = 4$.

On constate que l'adéquation entre les deux courbes est excellente en dimension 3, un peu moins bonne en dimension 4, notamment car la plage de valeur de ε n'est pas la même. On attendrait une meilleure adéquation pour ε proche de 0. En tous cas, ces résultats sont encourageants quant à la qualité de l'approche perturbative du chapitre 2. Cela vaudrait le coup à l'avenir d'étoffer

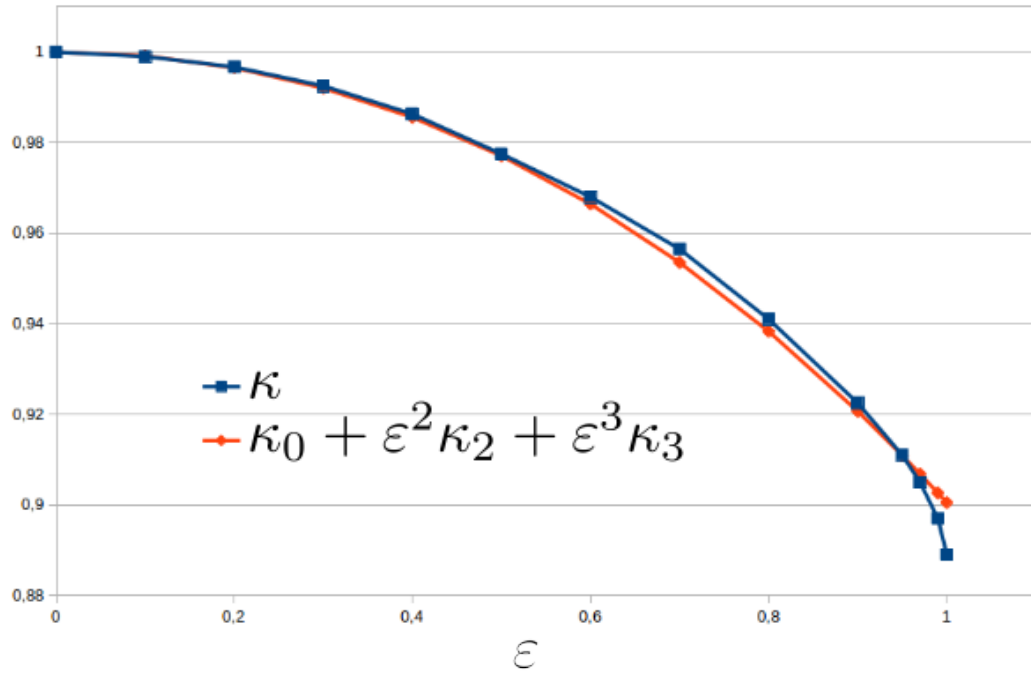


FIGURE 3.2 – Valeurs numériques de κ (en bleu) et $\kappa_0 + \varepsilon^2 \kappa_2 + \varepsilon^3 \kappa_3$ en fonction de ε pour $d = 3$. L'intervalle de confiance à 95% sur la valeur de κ est moins large que les points bleus.

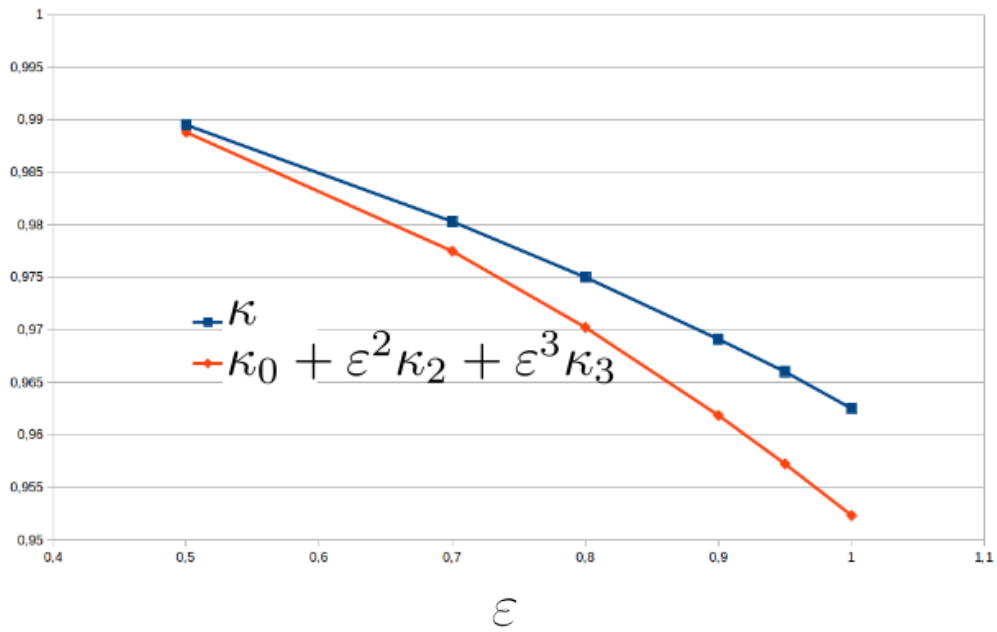


FIGURE 3.3 – Valeurs numériques de κ (en bleu) et $\kappa_0 + \varepsilon^2 \kappa_2 + \varepsilon^3 \kappa_3$ en fonction de ε pour $d = 4$. L'intervalle de confiance à 95% sur la valeur de κ est moins large que les points bleus.

les résultats analytiques en calculant κ_4 et d'améliorer la précision des résultats numériques (en augmentant le temps de calcul et le nombre de processeurs) pour mieux tester cette adéquation. En attendant nous allons voir les résultats numériques sur le modèle des miroirs.

3.2.2 Modèle des Miroirs

On a fait les mêmes simulations pour le modèle des miroirs. Malheureusement nous n'avons pas eu le temps de rédiger les calculs de κ_2 et κ_3 pour ce modèle, mais ces quantités existent et sont tout à fait calculables. On se contente donc des résultats numériques qui, comme on va le voir, ont un intérêt en soit.

Premièrement, on note que pour $\varepsilon = 0$ le modèle des miroirs coïncide avec la marche aléatoire non-rebrousante vue au chapitre 2. On trouvera une très bonne étude de cette marche dans [24]. Un point important est que le coefficient de diffusion de cette marche est :

$$\kappa_0 = \frac{d}{d-1}.$$

Voici donc le tableau des résultats numériques, seulement en dimension 3 :

$\varepsilon \backslash d$	3
0.5	1.496 ± 0.002
0.7	1.489 ± 0.002
0.9	1.474 ± 0.002
0.95	1.468 ± 0.002
0.97	1.4645 ± 0.002
0.99	1.4625 ± 0.002
1	1.538 ± 0.002

Et les voici sous forme de graphe à la figure 3.4.

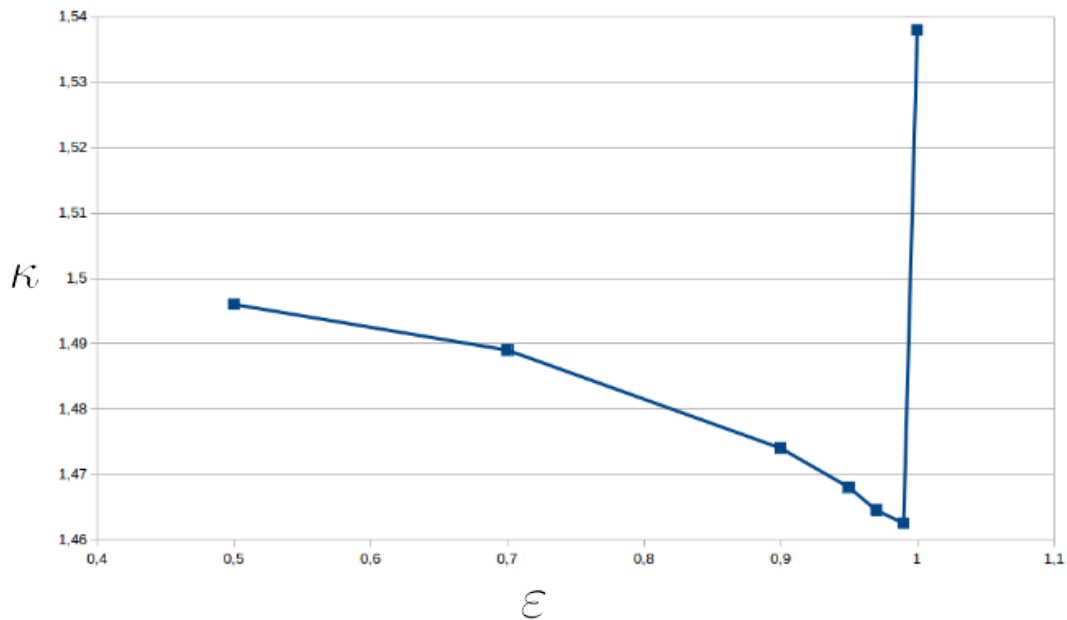


FIGURE 3.4 – Valeurs numériques de κ (en bleu) en fonction de ε pour $d = 3$ dans le modèle de miroirs. L'intervalle de confiance à 95% sur la valeur de κ est à peine plus large que les points bleus.

Bien entendu, la limite singulière en $\varepsilon = 1$ saute aux yeux. On va en donner immédiatement une interprétation heuristique dans la section qui vient.

3.3 Interprétation heuristique de la limite singulière dans le modèle des miroirs

Nous avons vu que la fonction $\kappa : \varepsilon \mapsto \kappa(\varepsilon)$ qui donne le coefficient de diffusion κ en fonction de ε semble discontinue en $\varepsilon = 1$. Ce résultat numérique surprenant nous a donné du fil à retordre. En effet, nous avons repéré depuis longtemps que le coefficient de diffusion du modèle des miroirs était légèrement supérieur au coefficient de diffusion de la marche aléatoire auto-évitante (Cf. Figure 1.6 du chapitre 1), ce qui nous paraissait logique et nous a, entre autres, poussés à chercher une approche perturbative pour approcher ce coefficient analytiquement. Mais tous nos calculs donnaient des valeurs de κ_2 et κ_3 négatives, ce qui ne collait pas avec les résultats numériques. En introduisant le modèle des permutations le problème a subitement disparu. Le coefficient de diffusion était alors inférieur à celui de la marche aléatoire simple, ce qui était cohérent avec les valeurs de κ_2 et κ_3 qui étaient aussi négatives dans ce modèle. Le problème avait disparu mais il n'était toujours pas expliqué. C'est seulement après avoir introduit le paramètre ε et avoir vu les résultats ci-dessus qu'un début d'explication est apparu. Ensuite nous avons trouvé l'explication heuristique suivante qui montre que le phénomène est lié à la réversibilité. Le phénomène est encore à creuser.

Pour comprendre le phénomène, il faut se pencher sur les trajectoires de la particule dans le cas où $\varepsilon = 1$, dans ce cas le modèle est déterministe en milieu aléatoire. La particule est donc dans une forme de "tunnel" qu'elle suit indéfiniment. Dans le cas du modèle des miroirs, on peut voir numériquement que pour $d \geq 3$ la plupart des points $(\mathbf{q}, \mathbf{p}) \in \mathcal{M}$ sont sur des trajectoires ouvertes. Considérons donc que l'on est sur une trajectoire ouverte. Une telle trajectoire va passer par une infinité de miroirs et certain de ces miroirs vont être visités plusieurs fois, c'est là que le problème apparaît. En effet, maintenant, si ε est très légèrement inférieur à 1, disons $\varepsilon = 1 - 10^{-6}$ pour fixer les idées, alors à chaque pas il y a une probabilité 10^{-6} pour que la marche prenne une direction autre que celle dictée par le miroir. Et dans ce cas, si la trajectoire passe au moins deux fois en ce point, alors la particule peut reprendre la trajectoire mais en sens inverse. C'est possible dans le modèle des miroirs qui est réversible et dans lequel la même trajectoire peut être parcourue dans les deux sens, dans le modèle des permutations ça ne marche pas et on n'observe pas de discontinuité de $\kappa(\varepsilon)$ au voisinage de $\varepsilon = 1$. Dans ce cas où la particule repart en sens inverse elle va alors faire environ 1 000 000 de pas "en arrière" sur sa trajectoire précédente, ce qui la ramène fortement vers son point de départ et a un effet macroscopique sur sa distance à l'origine et donc sur son coefficient de diffusion.

Nous allons détailler le phénomène. Posons un certain $\varepsilon < 1$, a priori proche de 1 et posons $N = 1/(1-\varepsilon)$ qui est a priori grand (par exemple on peut imaginer $N = 1\,000\,000$, ce qui correspond à $\varepsilon = 1 - 10^{-6}$). En moyenne, la particule fait un pas de marche aléatoire non-rebroussante (qui ignore les miroirs) tous les N pas, on dira qu'elle fait une "erreur" tous les N pas. On appelle p la probabilité qu'à un instant donné la particule soit sur un miroir touché 2 fois par la trajectoire (la deuxième fois peut-être dans le passé ou dans le futur), on appelle un tel point une "auto-intersection". On néglige les miroirs vus plus de 2 fois. L'histoire est alors la suivante.

- D'abord la particule fait n_1 pas sur sa trajectoire délimitée par les miroirs, n_1 est une variable aléatoire de loi exponentielle et d'espérance N .
- Ensuite elle fait une "erreur".
- Avec probabilité p elle était sur une auto-intersection.
- Alors avec probabilité $1/(2d-1)$ elle reprend la même trajectoire mais dans le sens inverse, on appelle ça une "erreur fatale", qui va lui faire perdre beaucoup de temps.
- Dans ce cas elle repart en arrière pendant n_2 pas, n_2 est une variable aléatoire de loi exponentielle et d'espérance N .
- Le nombre de pas fait en arrière sur l'ancienne trajectoire est $\min(n_1, n_2)$ qui est une variable aléatoire de loi exponentielle et d'espérance $N/2$.
- À ce moment la particule sort de sa trajectoire fermée pour aller sur une nouvelle trajectoire, et elle est à la position où elle était n_3 pas plus tôt avec $n_3 = 2 \min(n_1, n_2)$. Elle c'est donc

comme si elle avait fait du sur place pendant n_3 pas, en fait elle a fait un aller-retour avec $n_3/2$ pas à l'aller et autant au retour. n_3 est donc le "temps perdu", et $\mathbb{E}[n_3] = N$.

- Finalement, à chaque pas, la particule a donc une probabilité $p/((2d-1)N)$ de perdre en moyenne un temps N . Donc sur une trajectoire de longueur $n \gg N$ le temps perdu est en moyenne $np/(2d-1)$. Et cette quantité ne dépend pas de ε

Donc, il semblerait que si la particule fait n pas dans le modèle des miroirs avec $\varepsilon < 1$ et $1-\varepsilon \ll 1$ alors elle s'éloigne de l'origine autant qu'en faisant $n(1-p/(2d-1))$ pas dans le modèle des miroirs avec $\varepsilon = 1$. Le déplacement carré moyen est alors $\kappa(1)n(1-p/(2d-1))$ où $\kappa(1)$ est le coefficient de diffusion pour le modèle des miroirs avec $\varepsilon = 1$. Par conséquent, on peut conjecturer que :

$$\lim_{\varepsilon \rightarrow 1^-} \kappa(\varepsilon) = \kappa(1) \left(1 - \frac{p}{2d-1}\right) \quad (3.8)$$

Il ne reste plus qu'à évaluer p numériquement. Une quantité très simple à extraire des simulations est le nombre total de miroirs visités. Pour $d = 3$ on trouve que la quantité *Nombre de miroirs visités/Nombre de pas* tend vers 0.88 pour un nombre de pas assez grand. Si on suppose que le nombre de miroirs visités 3 fois ou plus est négligeable alors on trouve que 24% du temps la particule est au contact d'un miroir visité 2 fois. Plus précisément 12% du temps elle est au contact d'un miroir déjà visité par le passé et 12% du temps elle est au contact d'un miroir qu'elle reverra dans le futur. On en déduit que $p = 0.24$. Par conséquent, pour $d = 3$ on pourrait prévoir que :

$$\begin{aligned} \lim_{\varepsilon \rightarrow 1^-} \kappa(\varepsilon) &= \kappa(1) \left(1 - \frac{p}{2d-1}\right) \\ &= 1.538 \left(1 - \frac{0.24}{5}\right), \\ &= 1.464. \end{aligned} \quad (3.9)$$

Par simulation numérique direct on a $\kappa(0.99) = 1.4625 \pm 0.002$. Aux incertitudes près cette prédiction est excellente, il est rare qu'une heuristique fonctionne aussi bien pour le modèle des miroirs.

Conclusion

Résumé des travaux effectués

Dans cette thèse nous avons présenté les différents sens du mot *diffusion* en physique statistique, de la distance carré moyenne à la loi de Fick en passant par le mouvement Brownien. Au chapitre 1 nous avons défini mathématiquement la loi de Fick pour le modèle des miroirs dans le tore $\llbracket 1, N \rrbracket \times (\mathbb{Z}/N\mathbb{Z})^{d-1}$. Puis nous avons donné des éléments de preuve concrets pour aller vers la démonstration de cette loi pour $D \geq 3$, simulations numériques à l'appui. Pour cela nous avons donc dû généraliser le modèle des miroirs en dimension supérieure à 2. Pour tenter de terminer cette preuve nous avons développé, au chapitre 2, une approche perturbative, que nous avons testée sur un problème plus facile.

Pour trouver un problème plus simple nous avons dû nous éloigner de la physique en construisant une variante non réversible du modèle des miroirs : le modèle des permutations. Nous avons donc utilisé notre approche perturbative pour étudier le coefficient de diffusion κ dans le modèle des permutations pour $d \geq 3$. Nous avons ainsi pu trouver une série bien définie pour approcher κ .

Pour prouver que cette série est bien définie nous avons généralisé l'idée de *lace expansion* de Gordon Slade développée pour traiter l'épineux problème de la marche aléatoire auto-évitante. Il reste à clarifier sous quelles conditions cette série converge vers κ . Mais nous avons tout de même pu calculer ses premiers termes κ_0 , $\varepsilon^2 \kappa_2$ et $\varepsilon^3 \kappa_3$.

Nous avons aussi conçu des algorithmes non triviaux pour obtenir une valeur approchée de κ par simulation numérique. La comparaison des résultats numériques et des résultats perturbatifs montre que notre approche perturbative est pertinente pour le modèle des permutations. Pour le modèle des miroirs les résultats numériques montrent une discontinuité du coefficient de diffusion à la transition entre le déterministe et le modèle faiblement aléatoire. Cette discontinuité est absente dans le modèle des permutations. Nous avons donné une explication heuristique de cette discontinuité.

Problèmes ouverts

Ce travail ouvre de nombreuses voies qui ne demandent qu'à être explorées. La plus évidente est celle de la généralisation de la démarche perturbative à une classe plus large de modèles. La preuve sur le modèle des permutations est une forme de cas d'école qui permet de bien comprendre les idées mais la démarche est aussi faisable sur le modèle des miroirs. Un article qui présente la version généralisée de cette approche perturbative est en cours de rédaction. Ce sera l'occasion de préciser les succès et les limites de la méthode.

Deux améliorations seront à apporter à cette méthode. La première, calculatoire, est le calcul des termes suivants du développement, et pour cela une automatisation du calcul est à envisager. En effet il semble que le nombre de lignes de calcul augmente drastiquement en passant du calcul de κ_m au calcul de κ_{m+1} il y a donc un travail mathématique et/ou un travail informatique à faire pour aborder ces calculs. La seconde amélioration est purement théorique : il faut prouver que la série $\sum \kappa_m$ converge, et cette preuve s'annonce ardue mais pas infaisable. La *lace expansion* a sans doute des briques supplémentaires à apporter, le reste de la preuve reposant sur la créativité et la persévérance de celui ou celle qui s'y attaquera.

L'approche numérique n'est pas non plus épuisée et je regrette de ne pas l'avoir poussée plus loin par manque de temps. La structure générale des algorithmes permet d'explorer de nombreux modèles, notamment des modèles réversibles qui permettront de tester la pertinence de l'heuristique présentée à la section 3.3. Ce travail pourrait éventuellement faire l'objet d'un stage de recherche.

Un autre point à remarquer est que les expressions de κ_2 et κ_3 font intervenir des différences entre des valeurs de la fonction de Green de la marche aléatoire. Et il est donc probable que l'on puisse réutiliser la démarche en dimension deux, où la fonction de Green n'existe pas mais où nous pourrions travailler avec le noyau de potentiel (*potential kernel*) présenté à la section 4.4 de [19]. Ainsi, la démarche pourrait se révéler fructueuse en dimension 2.

Bibliographie

- [1] Illés Horváth, Bálint Tóth, Bálint Vető, *Diffusive limits for “true” (or myopic) self-avoiding random walks and self-repellent Brownian polymers in $d \geq 3$* , (2010)
- [2] D. J. Gates*, Lattice Wind-Tree Models. I. Absence of Diffusion, *Journal of Mathematical Physics* 13, 1005 (1972); doi : 10.1063/1.1666080
- [3] D. J. Gates*, Lattice Wind-Tree Models. II. Analytic Property, *Journal of Mathematical Physics* 13, 1315 (1972); doi : 10.1063/1.1666138
- [4] D. Amit, G. Parisi, L. Peliti : *Asymptotic behavior of the ‘true’ self-avoiding walk*, *Phys. Rev. B*, 27 : 1635–1645 (1983)
- [5] L.A. Bunimovich, Ya. G. Sinai, Ya. G. *Statistical properties of the Lorentz gas with periodic configuration of scatterers*. *Commun. Math. Phys.* 78, (1980) 479-497
- [6] X. P. Kong and E. G. D. Cohen, *Anomalous diffusion in a lattice-gas wind-tree model*, *Phys. Rev. B* (1989)
- [7] Bálint Tóth, *Persistent Random Walks in Random Environment*, *Probability Theory and Related Fields* (1986)
- [8] Gordon Slade, *The Lace Expansion and its Applications*, Lecture notes for the XXXIV th Saint-Flour International Probability School, July 2004
- [9] N. Madras and G. Slade, *The Self-Avoiding Walk*, Birkhäuser, Boston, (1993)
- [10] J. Bricmont and A. Kupiainen, *Random Walks in Asymmetric Random Environments*, *Commun. Math. Phys.* 142, 345-420 (1991)
- [11] Gordon Slade, *Self-Avoiding Walk*, *The Mathematical Intelligencer*, vol. 16 NO 1 1994
- [12] Raffaele Esposito and Mario Pulvirenti, *From Particles to Fluids*, *Handbook of Mathematical Fluid Dynamics*, (2004)
- [13] Th. W. Ruijgrok and E.G.D. Cohen, *Deterministic Lattice Gas Models*, *Physics Letters A* (1988)
- [14] P. and T. Ehrenfest, *Begriffliche Grundlagen der statistischen Auffassungin der Mechanik* *Encykl. d. Math. Wissensch.* IV 2 II, Heft 6, 90 S (1912) (in German, translated in :) *The conceptual foundations of the statistical approach in mechanics*, (trans. Moravicsik, M. J.), 10-13 Cornell Univer-sity Press, Itacha NY, (1959).
- [15] T. Bodineau, I. Gallagher, L. Saint-Raymond, *The Brownian motion as the limit of a deterministic system of hard-spheres*, *Inventiones mathematicae*, 1-61 (2016).
- [16] Jean-Marc Lévy-Leblond, *De la matière, relativiste, quantique, interactive*, (2006) Édition du Seuil ISBN : 2-02-084836-8
- [17] Daniel Bernoulli, *Hydrodynamica, sive de viribus et motibus fluidorum commentarii* (1738).
- [18] H. Spohn, *Large Scale Dynamics of Interacting Particles*, ISBN 978-3-642-84371-6 (1991)
- [19] G. Lawler and V. Limic *Random Walk : A Modern Introduction*
 Pour les citations de ce livre nous nous référons au brouillon en accès libre à l’adresse suivante :
www.math.uchicago.edu/~lawler/srwbook.pdf

- [20] Mathieu Nebra (Mateo21), *Apprenez à programmer en C!*, Les tables de hachage, (2013)
[https : //openclassrooms.com/fr/courses/19980-apprenez-a-programmer-en-c/19978-les-tables-de-hachage](https://openclassrooms.com/fr/courses/19980-apprenez-a-programmer-en-c/19978-les-tables-de-hachage)
- [21] Donald E. Knuth, *The art of computer programming Volume 3*, 2nd édition (1998)
- [22] Y. Chiffaudel et R. Lefevre, *The mirrors Model : Macroscopic diffusion without noise or chaos*, JPhys A, arxiv.org/abs/1506.04727, (2016)
- [23] Cédric Villani, *Théorème Vivant*, ISBN : 978-2-246-79882-8, (2012)
- [24] Robert Fitzner, Remco van der Hofstad *Non-backtracking Random Walk* J Stat Phys (2013) 150 :264-284
- [25] G. Adelson-Velskii et E. M. Landis, *An Algorithm for the Organization of Information*. Soviet Mathematics Doklady, 3 :1259–1263, 1962.
- [26] V. V . Anshelevich, K. M. Khanin and Ya. G. Sinai *Symmetric Random Walks in Random Environments* Commun.Math.Phys.85,449-470(1982)
- [27] G. R. Grimmett and H. Kesten, *Percolation theory at Saint-Flour*, Springer (2012)
- [28] A. S. Kraemer, D. Sanders *Zero Density of Open Paths in the Lorentz Mirror Model for Arbitrary Mirror Probability* J. Stat. Phys. 156 (5), (2014), 908-916
- [29] G. Lawler *Weak Convergence of a Random Walk in a Random Environment* Commun. Math. Phys. 87, 81-87 (1982)
- [30] R. Lefevre *Macroscopic diffusion from a Hamilton-like dynamics*, Journal of Statistical Physics, Volume 151 (5),(2013) 861-869
- [31] R. Lefevre, *Fick's law in a random lattice Lorentz gas*, [http ://arxiv.org/abs/1404.5694](http://arxiv.org/abs/1404.5694), Arch. Rat. Mech. and Anal. : Volume 216, Issue 3 (2015), Page 983-1008
- [32] T. W. Ruijgrok, E. G. D. Cohen, *Deterministic lattice gas models* Phys. Lett. A 133 (1988) 415
- [33] I.J. Zucker, *70+ Years of the Watson Integrals*, J Stat Phys (2011)
- [34] David Brydges, Gordon Slade, *Renormalisation group analysis of weakly self-avoiding walk in dimensions four and higher*, (2010)
- [35] G. Basile, A. Nota, F. Pezzotti, M. Pulvirenti *Derivation of the Fick's law for the Lorentz model in a low density regime* Communications in Mathematical Physics 04/2014 ; 336(3)
- [36] L. A. Bunimovich and S. E. Troubetzkoy, Journal of Statistical Physics, Vol. 67, Nos. 1/2, 1992
- [37] G. Casati, C. Mejia-Monasterio and T. Prosen Phys. Rev. Lett. 101, (2008) 016601
- [38] G. Kozma, V. Sidoravicius, arXiv :1311.7437
- [39] I. Horváth, B.Tóth, B. Veto, Prob. Th. and Rel. F. 153 (2012) 691-726