



Statistique de l'interférence quantique et circuits quantiques aléatoires

Ludovic Arnaud

► To cite this version:

Ludovic Arnaud. Statistique de l'interférence quantique et circuits quantiques aléatoires. Physique mathématique [math-ph]. Université Paul Sabatier - Toulouse III, 2009. Français. NNT: . tel-00599990

HAL Id: tel-00599990

<https://theses.hal.science/tel-00599990>

Submitted on 12 Jun 2011

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.



THÈSE

En vue de l'obtention du

DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par *l'Université Toulouse III - Paul Sabatier*

Discipline ou spécialité : *Physique*

Présentée et soutenue par *Ludovic Arnaud*

Le *17 decembre 2009*

Titre : *Statistique de l'interférence quantique et circuits quantiques aléatoires*

JURY

*Pablo Arrighi
Mohamed Aziz Bouchene
Daniel Braun
Nicolas Cerf
Stéphane Nonnenmacher
Denis Ullmo*

*Rapporteur
Président du jury
Directeur de thèse
Examineur
Examineur
Rapporteur*

Ecole doctorale : *Science de la matière*

Unité de recherche : *Laboratoire de Physique Théorique*

Directeur(s) de Thèse : *Daniel Braun*

Rapporteurs : *Pablo Arrighi et Denis Ullmo*

A la mémoire de Jonathan « Jojo Jo » Ribeiro.

A mes parents et à mon petit frère.

Remerciements

Voici le chapitre zéro de cette thèse. Il contient une liste non exhaustive bien entendu des remerciements et autres reconnaissances relatives à son aboutissement.

Je tiens à remercier Didier Poilblanc, premièrement pour avoir accepté que mon stage de Master se déroule au sein du Laboratoire de Physique Théorique quand il était directeur, et deuxièmement pour son soutien à l'époque où je cherchais un financement. Sans lui, je n'aurais jamais pu avoir la chance de faire ma thèse dans ce laboratoire et de ce fait je lui en suis très reconnaissant. Pour les mêmes raisons, je tiens à remercier André Neveu pour son soutien au ministère à l'époque où la bourse se faisait attendre. Je remercie Clément Sire, qui a succédé à Didier en tant que directeur du LPT, ceci depuis quelques années maintenant. Sa capacité de meneurs d'hommes, alliée à sa sympathie naturelle ainsi que la manière qu'il a de parler de science en faisant en sorte que son interlocuteur se sente aussi compétant que lui, n'ont d'égal que ses très bons goûts musicaux.

Je tiens à remercier Daniel Braun, mon directeur de thèse. Travailler avec lui a vraiment été un plaisir. Je le remercie pour sa gentillesse, sa disponibilité et pour toutes les choses que j'ai apprises grâce à lui. Encore Merci.

Je remercie l'ensemble de personnes qui ont gentiment acceptées de faire partie de mon jury lors de la soutenance : Nicolas Cerf, Mohamed Aziz Bouchène et Stéphane Nonnenmacher. Je remercie d'autant plus Pablo Arrighi et Denis Ullmo qui ont aussi assuré la tâche de rapporteurs.

De manière globale, je tiens à remercier l'ensemble des personnes qui gravitent au sein de l'IRSAMC et qui m'ont permis d'évoluer dans des conditions optimales. A commencer par Sylvia (ainsi que Gisèle) pour son aide souvent précieuse. Je remercie aussi ma voisine de bureau Sandrine, certes pour son aide et ses conseils de nature essentiellement informatique, mais aussi pour les nombreuses pauses café/thé partagées qui était bien sympathiques. Je tiens aussi à remercier très chaleureusement Fabienne Alary, pour sa gentillesse et sa serviabilité (et aussi pour ses très bons goûts musicaux). En ce qui concerne les personnes du LPT, je tiens à remercier les *très fortement corrélés fermions* que sont Fabien, Matthieu et Sylvain, pour leur aide en ce qui concerne l'utilisation du cluster de calcul.

Je tiens à remercier Hélène Carrère, Michel Bonnet, Sébastien Lachaize, Thomas Belhadj, Cyril Jaudet, ainsi que tous les moniteurs et les vacataires avec qui j'ai eu le plaisir d'enseigner au département de physique de l'INSA.

Je me dois de remercier l'ensemble de mes professeurs, les *bons* mais aussi les *mauvais* (qui m'ont montré l'exemple à ne pas suivre). Parmi les très bons, j'ai une pensée toute particulière pour mes professeurs de la période collège-lycée. Je tiens donc à remercier Mme Choulet, pour ne jamais m'avoir mis de 20/20 et pour m'avoir orienté vers plus de rigueur, nécessaire à tout travail scientifique. Je remercie Mr Andrieux pour sa gentillesse et pour tous les livres qu'il m'a prêtés. Je n'oublie pas non plus les conseils majeurs donnés par Mme Baudau. Pour finir,

j'adresse mes remerciements au *premier* de tous ces professeurs, Mr Ortéga. Je le remercie pour ne pas avoir soufflé la petite flamme qui brûlait en moi (mais plutôt pour l'avoir attisée) à une époque où mes seules matières de réussite était la musique et le dessin.

Je voudrais aussi remercier deux personnes que je n'ai pas vue depuis longtemps, mais qui ont chacune eu beaucoup d'influence. La première est Christophe Naudi. Pédagogue de première, il ne s'ait pas contenté de m'apprendre le solfège rythmique et la coordination nécessaire au maniement d'une batterie. Il m'a tout simplement appris à apprendre, de la manière la plus naturelle possible et pour cela je le remercie. Je tiens aussi à remercier Richard Garcia pour sa générosité. Je n'oublie pas la quantité de livres intéressants que j'ai découvert grâce à lui, livres qui m'ont grandement influencés.

Je voudrais remercier l'ensemble des thésards et des post-docs avec qui j'ai eu le plaisir d'interagir pendant ces trois ans. Plus ou moins par ordre chronologique de rencontre : Julien Sopik, Guillaume Roux, Gaspard Bousquet (merci pour le T-shirt), Andréas Abendschein, Manuela Capello, Arnaud Ralko (et son célèbre *au putaing !*), Julien Baglio, Benoit Roubert, John Martin, Ignacio García-Mata, Lorand Horvath, Luca Delfini, Sylvain Vidal, David Schwandt, Vincent Démery, Mickael Pasek, Sébastien Weber, Marie Barthélémy et Nicolas Thiré. Parmi toute ces personnes, certains sont devenus des amis proches. Je pense en particulier à mes compagnons de route qu'étaient Fabien "ftroussel" Trouselet et Thomas "Pépoune Garcia" Portet... Ah oui j'allais oublier quelqu'un, Clément : bon c'est l'histoire de deux super potes, passionnés de Physique, qui ont bien conscience que faire une thèse de doctorat n'est pas la chose la plus certaine de l'Univers observable. Sauf qu'ils en font chacun une au même endroit. Pour le meilleur et pour le pire jusqu'à ce que le postdoc les sépare. Au passage, j'en profite pour m'excuser auprès de l'ensemble des personnes qui furent déranger par nos discussions souvent très agitées. Je ne serais pas surpris d'apprendre qu'un chercheur du LCAR tout en a bas est au courant que j'ai oublié un facteur 2 dans un calcul ou que Clément n'aime pas le punk californien. Enfin... merci man.

Je voudrais maintenant profiter de l'occasion pour remercier mes amis les plus proche qui ont tous eu une certaine influence : Xavier, Jean-Christophe, David (et pour ce livre pas mal qu'il m'a conseillé un jour ;)), Kéké, Julien, Laure, Eric et Aurélie (pour leur soutien constant), Avédis, Romain et Guihem. Et Ghyslain bien sûr mon colocataire-chanteur-guitariste-bassiste-maman-cuisinier. Que dire mec, de ces trois ans de cohabitation ? Je suis désolé pour les miettes de pain, vraiment. Métaphoriquement parlant, je serais tenté de rapprocher tout ces bons moments, à celui qui consiste à jouer en live le passage à 1 : 55 de *final plaisant pain*... Je te remercie pour tout ça.

Pour finir, je voudrais remercier tous les membres de ma famille et bien entendu mes parents, pour leur soutien sans limite depuis le tout début. Je veux que vous sachiez a quel point je vous respecte tout les deux pour ce que vous êtes et à quel point je vous aime. Cette thèse vous est dédiée.

Table des matières

Remerciements	iv
Avant-propos	xiii
1. L'interférence comme une ressource de l'information quantique	1
I. L'interférence en physique	1
II. Théorie de l'information quantique	2
II.1. Bits quantiques	2
II.2. Mesure de l'information quantique et lien avec l'information classique . .	4
II.3. Calculs quantiques et algorithmes quantiques	4
III. Ressources de l'information quantique	7
III.1. L'intrication	7
III.2. L'interférence quantique	9
III.3. Nuances sur les capacités des ressources quantiques	10
IV. Une mesure quantitative de l'interférence quantique	10
IV.1. Expression générale	11
IV.2. Cas d'une propagation purement unitaire	12
IV.3. Cas particulier pour un seul qubit	12
V. Résultats	13
V.1. Interférence dans l'algorithme de Shor	13
V.2. Interférence dans l'algorithmes de Grover	13
V.3. Processus n'impliquant pas de l'interférence	14
VI. Conclusion partielle	14
2. Statistique de l'interférence dans les algorithmes quantiques	15
I. Motivation	15
II. Distribution de l'interférence dans des ensembles circulaires	16
II.1. L'ensemble circulaire unitaire CUE	16
II.2. Distribution calculée numériquement	17
II.3. Calculs analytiques des distributions dans le cas $N=2$	18
II.4. Calculs analytiques des deux premiers moments	19
III. Distribution de l'interférence dans des circuits aléatoires	21
III.1. Ensembles de circuits aléatoires	21
III.2. Résultats numériques	22
IV. Conclusion partielle	24

3. Statistique de l'interférence en présence de décohérence	25
I. Introduction	25
II. Statistique de l'interférence dans un système quantique couplé à un spin	25
II.1. Propagateur pour un application complètement positive	26
II.2. Interférence dans un système quantique couplé à un seul spin	26
II.3. Résultats numériques	27
II.4. Résultats analytiques	29
2.4.a. Interférence moyenne	29
2.4.b. Second moment de la distribution d'interférence	32
III. Interférence dans un système quantique couplé à plusieurs spins	37
IV. Conclusion partielle	39
4. L'information quantique au service de la production de matrices aléatoires	41
I. Généralités et motivations	41
I.1. Opérateurs pseudo-aléatoires et circuits quantiques aléatoires	41
I.2. Théorie des k-design	42
I.3. L'ensemble de circuit unitaire UCE	43
II. Convergence de UCE vers CUE	44
II.1. Distribution des différences entre phases propres voisines	44
III. Efficacité de la convergence	46
III.1. Distribution relative au éléments de matrice	46
III.2. Moments de la distribution du carré des éléments de matrice	51
III.3. Corrélations entre les éléments de matrice	55
IV. Conclusion partielle	57
Conclusion générale et perspectives	58
Appendices	60
A. Méthode d'intégration invariante	63
I. Introduction	63
II. Propriétés et écriture diagrammatique	63
II.1. Notations et propriétés générales	63
II.2. Écriture diagrammatique	64
III. Calculs des intégrales	65
III.1. Rotation	66
III.2. Unitarité	67
IV. Relations entre intégrales et formules explicites	68
IV.1. La <i>fan relation</i>	68
IV.2. La <i>double fan relation</i>	69
IV.3. Formules explicites	69
V. Conclusion	71
B. Détails sur le calcul du second moment	73
I. Le terme A	73
II. Le terme B	76
C. Valeurs des intégrales nécessaires pour le calcul du second moment	77

Publications	79
Bibliographie	79
Index	84

Avant-propos

On peut considérer que l'un des derniers paradigmes proposés en Physique est celui engendré par le mariage entre la *théorie de l'information* et la *mécanique quantique*. La discipline qui en résulte est maintenant connue sous le nom d'information quantique, le domaine du calcul quantique étant une de ses variantes principales. En effet, bien qu'étant encore toute jeune, cette discipline laisse entrevoir de nombreuses perspectives intéressantes : la téléportation quantique bien qu'issue d'idées théoriques récentes, a déjà reçu plusieurs vérifications expérimentales [Zei97]. Dans la même veine, des dispositifs commerciaux utilisant la cryptographie quantique, commencent à être testés pour assurer la sécurité de transactions bancaires. Les idées de l'information quantique sont aussi utilisées dans d'autres domaines de la Physique. Par exemple, le problème de la conservation de l'information dans les trous noirs, semble se résoudre dans le cadre de l'information quantique [SL04]. Certains physiciens vont même jusqu'à évoquer la reformulation de l'ensemble des lois de la Physique dans un formalisme où le concept d'information est celui qui serait le plus fondamental [Ste98], sans oublier les nouvelles questions que pose le domaine aux mathématiques pures.

Cependant, parmi toutes les attentes et tous les espoirs de cette jeune discipline, la construction d'un ordinateur quantique reste la plus populaire au vu de sa rapidité de calcul potentielle comparée à son homologue classique.

Il est donc tout à fait naturel de chercher à comprendre pourquoi dans certaines mesures, la nature quantique de notre monde semble résoudre certains problèmes computationnels de manière plus efficace qu'elle le ferait classiquement. Alors dans ce cas, où faut-il chercher ? L'information quantique étant le mariage de deux disciplines, la logique première est d'identifier quels sont les concepts propres de la théorie. Parmi ces concepts, le plus étudié est celui de *l'intrication*. Il est maintenant admis dans la communauté que l'intrication est une ressource essentielle de l'information quantique [JL03]. L'autre concept important, typiquement quantique, mais qui est beaucoup moins étudié cependant est *l'interférence quantique*. L'interférence quantique est la capacité qu'a un processus quantique à favoriser des états classiquement défavorisés et *vice-versa*. Dans ce sens l'interférence est la propriété quantique qui modifie le flot d'information classique. L'étude d'un des aspects de l'interférence est le sujet central de cette thèse.

Dans le premier chapitre, nous rappellerons quelques propriétés de l'interférence dans le cadre général de la Physique puis dans le cadre plus ciblé de l'information quantique. Nous définirons de manière succincte l'information quantique et le calcul quantique. Nous introduirons une mesure de l'interférence adaptée à l'étude des algorithmes quantiques et passerons en revue plusieurs résultats la concernant. Dans le deuxième chapitre, nous nous intéresserons au comportement statistique de l'interférence lors de processus unitaires, puis dans le troisième chapitre nous nous intéresserons à la manière dont est modifié ce comportement en présence de décohérence. Les réponses obtenues dans ces deux chapitres amèneront bien sûr d'autres

questions, en particulier des questions concernant un des outils utilisés, celui de la théorie des matrices aléatoires. Ainsi dans le quatrième chapitre nous nous intéresserons au liens existant entre les circuits quantiques aléatoires et les matrices aléatoires, et dans quels sens l'équivalence entre ces deux types d'entités peut se révéler utile pour produire de manière efficace des ensembles de matrices aléatoires.

L'interférence comme une ressource de l'information quantique

I. L'interférence en physique

La première équation qui apparaîtra dans cette thèse est sûrement aussi simple quelle est peu intuitive à première vue :

$$\text{lumière} + \text{lumière} = \text{nuit}$$

Elle porte l'essence même du concept d'interférence tel qui l'est apparu au XVIIIème siècle avec les premières expériences d'optiques qui ont mises en évidence le caractère ondulatoire de la lumière. L'idée principale contenue dans cette équation est celle-ci : Les intensités lumineuse ne s'ajoutes pas simplement de manière algébrique. Cette remarque permet d'en déduire que la lumière ne peut pas être entièrement décrite par son intensité et qu'un autre concept indépendant, la phase, doit être pris en compte.

L'interférence en optique a permis la mise au point de technique telle l'holographie. Récemment encore, une expérience d'optique utilisant des impulsions laser femtosecondes a permis d'implémenter le calcul d'une somme de Gauss et de ce fait de factoriser des entiers grâce au phénomène d'interférence [BCSG08].

Dans le cadre de la mécanique quantique, le concept d'interférence est intimement lié à celui de superposition d'états représentés par des amplitude de probabilité. En fait, il permet d'affirmer qu'un état quantique dans une superposition cohérente est physiquement discernable d'un mélange statistique. L'exemple qui suit permet de bien appréhender cela : soient trois sources (par exemple des canons qui émettent des électrons) d'états quantiques bien définis et distincts : dans une certaine base orthogonale (reliée à la mesure dans le sens où chaque état de la base peut être parfaitement distingué des autres) ces états s'écrivent respectivement.

$$|\psi+\rangle = \frac{1}{\sqrt{2}}(|a\rangle + |b\rangle) \tag{1.1}$$

$$|\psi-\rangle = \frac{1}{\sqrt{2}}(|a\rangle - |b\rangle) \tag{1.2}$$

$$|\psi_s\rangle = |a\rangle \text{ ou } |b\rangle \text{ avec une probabilité } \frac{1}{2} \tag{1.3}$$

Les deux premiers canons produisent donc chacun des superpositions cohérentes tandis que le troisième produit un mélange statistique. Pour ces trois type d'états, la probabilité de faire

une mesure de l'état $|a\rangle$ est $\frac{1}{2}$. De ce fait une mesure projective directe ne permet pas de les distinguer. Comment faire si les étiquettes indiquant la nature des états produits par les canons ont été égarées ? La capacité qu'ont les états quantiques à interférer peut être utilisée, ceci en appliquant une même transformation unitaire sur chacun de ces états. Physiquement cela peut se traduire par l'évolution temporelle du système sous un certain hamiltonien H pendant un temps t et dans ce cas la transformation unitaire est $U = \exp(-i\hbar Ht)$. Pour notre exemple considérons une transformation unitaire précise, celle donnée par la matrice :

$$\frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$$

Appliquée sur les états précédemment définis, juste avant la mesure, cette transformation permet de discriminer les états et de remettre la bonne étiquette sur chaque canon : les probabilités de mesure de la composante $|a\rangle$ pour les états $|\psi+\rangle$, $|\psi-\rangle$ et $|\psi_s\rangle$ seront respectivement 0,1 et $\frac{1}{2}$. Cette expérience est l'analogue d'une expérience d'optique de type fentes de Young. Ici le rôle joué par l'état $|\psi+\rangle$ correspond au chemin d'interférence constructive qui se traduit par une frange brillante sur l'écran. L'état $|\psi-\rangle$ correspond au chemin d'interférence destructive engendrant une frange sombre. L'état $|\psi_s\rangle$ correspond au cas de deux sources incohérentes qui n'amènent aucun motif d'interférence. Quant à l'opérateur d'évolution U , il traduit le rôle de la propagation des ondes lumineuses et de leur superposition cohérente sur l'écran.

II. Théorie de l'information quantique

Dans la théorie classique de l'information, il existe une dualité entre d'une part les concepts mathématiques de la théorie et d'autre part leur représentation physique. Suivant la formulation de Landauer l' *information est physique*. Même si cela peut paraître trivial, le concept de bit n'émerge que parce qu'il est possible de le représenter concrètement par une boîte contenant soit une boule ou non, par un chat qu'il soit vivant ou mort ou plus sérieusement par un condensateur chargé ou non. De la même manière, un processus informationnel, une communication ou un calcul par exemple, se matérialise sous la forme d'un processus physique mettant en jeu des quantités bien définies telles que l'énergie ou l'entropie. Ainsi le monde physique impose ses propres contraintes. Les meilleurs exemples de telles contraintes sont la limitation par c de la vitesse de transport de l'information, ainsi que le principe de Landauer qui énonce que l'effacement d'un bit d'information à un coup entropique valant $k_B \ln 2$.

Les idées de l'information quantique peuvent être introduites de la même manière que les principes de la mécanique quantique par les trois points suivants :

- 1 L'information quantique est contenue dans un état quantique.
- 2 La mesure modifie l'information quantique et donne des résultats probabilistes.
- 3 Un processus informationnel se traduit par une évolution unitaire.

II.1. Bits quantiques

Le premier point permet d'introduire le concept de bit quantique, plus communément appelé *quantum bit* ou *qubit* : pour ce faire il suffit de démarrer avec la notion de bit de l'information classique, c'est à dire avec un système à 2 états (nommés communément 0 et 1) puis de le traiter de manière quantique. Son état physique $|q\rangle$ est maintenant décrit par un vecteur dans un espace de Hilbert $\mathcal{H}_2$ de dimension 2. Dans le cas général ce vecteur s'écrit comme une combinaison linéaire de 2 états orthogonaux, pouvant donc être parfaitement distingués l'un de

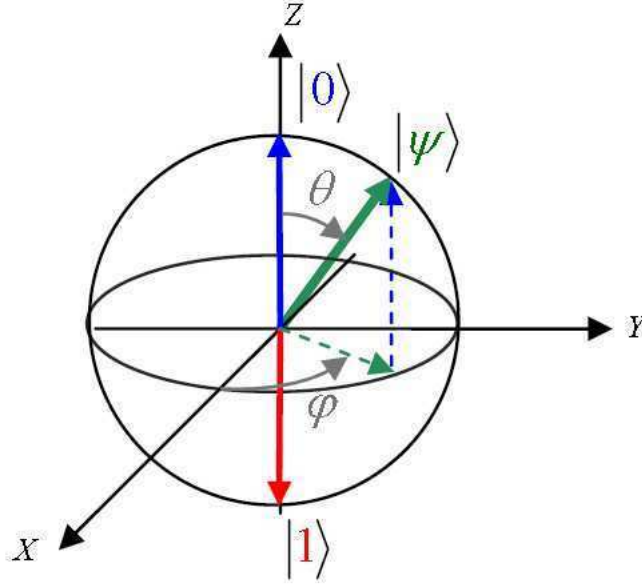


FIGURE 1.1.: Représentation géométrique de l'état d'un bit quantique sur la sphère de Bloch (Source : Wikimedia Commons).

l'autre lors d'un processus de mesure approprié. Ces états de base sont notés $|0\rangle$ et $|1\rangle$ et ainsi l'état d'un qubit s'écrit :

$$|q\rangle = a_0|0\rangle + a_1|1\rangle \text{ avec } a_0, a_1 \in \mathbb{C} \text{ tels que } |a_0|^2 + |a_1|^2 = 1 \quad (1.4)$$

Rappelons qu'en mécanique quantique, pour toute transformation $U(1)$ effectuée sur le vecteur d'état, la Physique reste équivalente. Ainsi pour une phase globale γ les état $|q\rangle$ et $e^{i\gamma}|q\rangle$ représente le même état physique. De ce fait il est commun de réécrire l'état (1.4) sous la forme

$$|q\rangle = \cos\left(\frac{\theta}{2}\right)|0\rangle + e^{i\phi}\sin\left(\frac{\theta}{2}\right)|1\rangle \quad (1.5)$$

où $\theta \in [0, \pi[$ et $\phi \in [0, 2\pi[$ est la *phase relative*. Cette paramétrisation incluant la normalisation et l'invariance $U(1)$ est ce que l'on appelle communément la paramétrisation de Bloch menant à la visualisation géométrique de l'état d'un qubit par la *sphère de Bloch* représentée sur la figure (1.1). Sur cette sphère deux états orthogonaux entre eux sont représentés par des points antipodaux (comme par exemple les états $|0\rangle$ et $|1\rangle$), c'est à dire représenté par des vecteurs colinéaires dans l'espace tridimensionnel $\mathbf{R}^3$ qui porte la sphère de Bloch. L'état d'un système global contenant n_q sous-systèmes à 2 états, c'est à dire n_q qubits, vit dans un espace de Hilbert de dimension 2^{n_q} qui est le produit tensoriel des n_q espaces de Hilbert de dimensions 2 correspondant à chaque qubit individuel. Cet ensemble de qubits constitue un *registre quantique* ou une *mémoire quantique*. De la même manière que l'état d'un qubit individuel, l'état le plus général d'un registre quantique s'écrit comme une combinaison linéaire à 2^{n_q} composantes :

$$|q_{n_q}\rangle = \sum a_{i_1 i_2 \dots i_{n_q}} |i_1 i_2 \dots i_{n_q}\rangle \text{ avec les } a_{i_1 \dots i_{n_q}} \in \mathbb{C} \text{ tels que } \sum |a_{i_1 \dots i_{n_q}}|^2 = 1 \quad (1.6)$$

Dans cette expression les sommes portent sur l'ensemble des indices $i_k = \{0, 1\}$. Les vecteurs $|i_1 i_2 \dots i_{n_q}\rangle$ forment une base très utilisée en information quantique nommée la *base computationnelle*. Ces vecteurs sont obtenus comme produits tensoriels ordonnés des vecteurs de base

de chaque qubits et la notation doit se comprendre comme $|i_1 i_2 \dots i_{n_q}\rangle = |i_1\rangle \otimes |i_2\rangle \otimes \dots \otimes |i_{n_q}\rangle$. Par exemple l'état d'un registre à 2 qubits s'écrit avec les 4 états de base $|00\rangle$, $|01\rangle$, $|10\rangle$ et $|11\rangle$. L'écriture binaire permet de mieux comprendre l'appellation *computationnelle* pour cette base puisqu'en définitive elle permet d'écrire l'état d'un registre quantique à n_q qubits, comme une superposition cohérente de l'ensemble des nombres entiers entre 0 et $2^{n_q} - 1$. Notons que l'obtention d'une généralisation de la paramétrisation de Bloch s'amenuise avec le nombre de qubits. Il existe cependant des travaux pour 2 et 3 qubits basés sur la fibration de Hopf des sphères S^5 et S^7 [MD01].

II.2. Mesure de l'information quantique et lien avec l'information classique

Le deuxième point concerne la mesure de l'information quantique. Imaginons qu'un expérimentateur possède un moyen de fabriquer une multitude de qubits dans le même état (1.4). En mesurant l'état d'un de ces qubits, par exemple dans la base $\{|0\rangle, |1\rangle\}$, l'expérimentateur récupère l'information « mesure dans l'état $|0\rangle$ » avec une probabilité $p_0 = |a_0|^2$ et l'information « mesure dans l'état $|1\rangle$ » avec une probabilité $p_1 = |a_1|^2$, l'état du qubit après la mesure devenant certain (respectivement $|0\rangle$ et $|1\rangle$). En d'autres termes, pour une mesure unique l'observateur récupère un bit aléatoire, c'est à dire de l'information classique. En réalisant une série de mesures équivalentes à la précédente, l'observateur peut reconstruire les probabilités p_0 et p_1 et ainsi récupérer la valeur des modules de α et β . Tout ceci n'est que transposition de la mesure quantique avec la langage de l'information quantique.

Une chose intéressante que peut faire l'expérimentateur, c'est de faire la mesure dans une autre base que la base computationnelle (donc ici dans la base $\{|0\rangle, |1\rangle\}$). C'est typiquement le genre de procédure qu'il doit utiliser s'il veut récupérer l'information de la phase relative qui est perdue lors de la procédure précédemment décrite. Par exemple une mesure dans la base orthonormale $\{|+\rangle, |-\rangle\}$ où les vecteurs sont définis par

$$|+\rangle = \frac{|0\rangle + |1\rangle}{2} \text{ et } |-\rangle = \frac{|0\rangle - |1\rangle}{2}, \quad (1.7)$$

permet à l'observateur d'avoir accès aux probabilités

$$\begin{aligned} p_+ &= \frac{|a_0 + a_1|^2}{2} = \frac{1}{2}(|a_0|^2 + |a_1|^2 + \text{Re}(a_0^* a_1)) = \frac{1}{2}(1 + \cos(\theta/2) \sin(\phi)) \\ p_- &= \frac{|a_0 - a_1|^2}{2} = \frac{1}{2}(1 - \cos(\theta/2) \sin(\phi)) \end{aligned}$$

qui permettent de remonter à la phase relative. Précisons que dans les expressions précédentes la détermination de ϕ est dépendante de l'angle θ , c'est à dire des modules de α (ou de β). Cette dissymétrie dans la mesure des modules et des phases qui semble donner au module un statut supérieur, n'est qu'apparente et résulte de l'arbitraire de la paramétrisation (1.5).

II.3. Calculs quantiques et algorithmes quantiques

Le troisième point permet d'introduire le concept de calcul quantique : l'évolution temporelle d'un registre quantique constitue un calcul quantique. En mécanique quantique, l'évolution temporelle d'un système étant représentée par l'action d'un opérateur unitaire, il suit que

n'importe qu'elle transformation unitaire est à considérer comme un calcul quantique. Ainsi de manière formelle, si $|\psi_i\rangle$ représente l'état initial d'un registre quantique - c'est à dire l'état d'entrée - alors le résultat d'un calcul particulier U sur ce registre est l'état final - l'état de sortie - $|\psi_f\rangle = U|\psi_i\rangle$

Un des résultats les plus intéressants dans le domaine de l'information quantique est la possibilité de construire des *algorithmes quantiques*. Qu'est-ce que cela signifie ? Nous avons défini un algorithme quantique comme un opérateur unitaire agissant sur un registre quantique donc en tout généralité cet opérateur transforme l'état de tous les qubits du registre. Bien sûr rien n'empêche de considérer des opérateurs unitaires « plus petits », c'est à dire qui agissent sur un nombre restreint de qubits sans toucher aux autres. Le cas limite étant celui où seul un qubit subit la transformation unitaire et où les $n_q - 1$ autres ne sont pas modifiés. Le résultat important est que si on considère un petit ensemble $\mathcal{G} = \{G_k\}$ d'opérateurs unitaires agissant sur 1 ou 2 qubits, opérateurs qualifiés de locaux, alors on peut approximer (avec une précision arbitraire) n'importe quel algorithme U comme une séquence S_k de n_g opérateurs choisis dans $\mathcal{G}$:

$$U = \prod_{S_k}^{n_g} G_k. \quad (1.8)$$

Chaque G_k constitue une *porte quantique*, analogue des portes logiques de l'informatique usuelle et un ensemble $\mathcal{G}$ donné est un *ensemble universel de portes* permettant de produire n'importe quels algorithmes quantiques [DiV95, Bar95, SW95]. Ainsi tout opérateur unitaire $U(2)$, comme par exemple l'ensemble des opérateurs de Pauli, constitue une porte quantique à un qubit, et tout opérateur $U(4)$ constitue une porte à deux qubits. Parmi les portes les plus couramment utilisées citons la porte de Hadamard

$$H = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix},$$

agissant sur 1 qubit et la porte de NON-contrôlée (*CNOT*) faisant agir un opérateur de Pauli σ_x sur un qubit *cible* en fonction de l'état d'un qubit de *contrôle*. L'appellation provient du fait que l'opérateur σ_x joue l'analogue quantique de la porte NON (puisque'elle inverse entre eux les vecteurs $|0\rangle$ et $|1\rangle$) et que son action est liée à l'état du qubit de contrôle. A titre d'exemple pour 2 qubits, le premier étant choisi comme contrôle et le deuxième comme cible, la matrice correspondant à la porte CNOT s'écrit

$$CNOT_{\{1,2\}} = |0\rangle\langle 0| \otimes \mathbf{1}_2 + |1\rangle\langle 1| \otimes \sigma_x = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix},$$

où $\mathbf{1}_2$ est l'opérateur identité de taille 2×2 . De manière générale ces deux opérations peuvent être définies pour des registres quantiques contenant un nombre arbitraire de qubits. Une porte de Hadamard agissant sur le qubit i et une porte CNOT agissant sur le qubit cible t contrôlé par le qubit c sont respectivement définies par

$$H_{\{i\}} = \mathbf{1}_2 \otimes \mathbf{1}_2 \otimes \dots H \dots \otimes \mathbf{1}_2 \otimes \mathbf{1}_2, \quad (1.9)$$

$$CNOT_{\{c,t\}} = \mathbf{1}_2 \otimes \dots |0\rangle\langle 0| \dots \otimes \mathbf{1}_2 \otimes \mathbf{1}_2 + \mathbf{1}_2 \otimes \mathbf{1}_2 \otimes \dots |1\rangle\langle 1| \dots \otimes \sigma_x \otimes \dots \otimes \mathbf{1}_2, \quad (1.10)$$

. où H est en i -ème position, les projecteurs $|j\rangle\langle j|$ sont en c -ième position et σ_x est en t -ième position. Notons que la porte CNOT peut être généralisée à 3 qubits par la porte de Toffoli, utilisant non pas un seul qubit de contrôle mais deux.

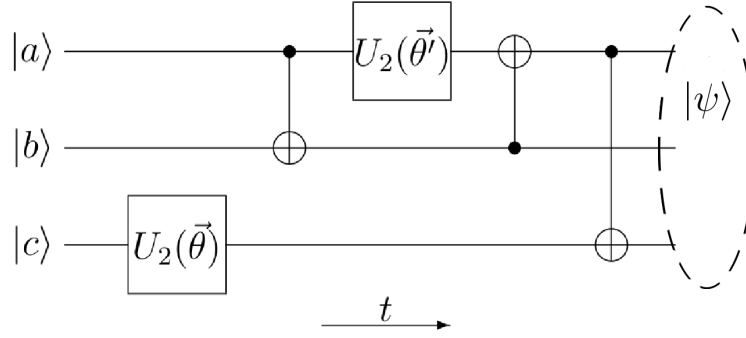


FIGURE 1.2.: Circuit quantique constitué de 2 portes unitaires $U(2)$ et de 3 portes $CNOT$. Le qubit de contrôle (cible) est attachés aux points noirs $\bullet$ (respectivement $\oplus$). La bulle pointillée dénommée $|\psi\rangle$ illustre le fait que le registre est dans un état non-séparable (cf. prochaine section).

A la manière de ce qui se fait en logique classique, on peut aussi dessiner des circuits quantiques comme sur la figure (1.2) C'est cette vision en terme de porte quantique qui permet de cerner la puissance computationnelle qu'aurait un système physique capable d'implémenter un algorithme quantique, système qui porte aujourd'hui le nom d'ordinateur quantique. En effet en 1994, Peter Shor [Sho94] démontra qu'un algorithme quantique pouvait factoriser un entier N avec un nombre d'opérations élémentaires n_g exponentiellement plus petit que le meilleur algorithme classique connu¹ (bien que pour l'instant, rien n'interdise l'existence d'un algorithme classique aussi efficace que l'algorithme de Shor). En 1996 Lov Grover explicita un algorithme quantique capable de retrouver un motif dans une base de données de taille N un nombre d'opérations élémentaires se comportent comme $\mathcal{O}(\sqrt{N})$ alors que le meilleur algorithme classique en utilise $\mathcal{O}(N)$. Une accélération exponentielle comparée au meilleur algorithme classique est aussi prédite pour le *shifted character problem* [GSV00], le problème du sous-groupe caché [Kup03] et pour la résolution de systèmes d'équations linéaires [WHL09]. On sait aussi qu'un marcheur quantique peut traverser un graphe exponentiellement plus vite que n'importe quel marcheur aléatoire classique, ce qui permet des solutions efficaces pour d'autres problèmes [MCD02]. Récemment, Aharonov *et.al* ont proposé un algorithme quantique qui approxime efficacement le polynôme de Jones comme aucune racine primitive de l'identité [ALM06]. Tout ces exemples précis tendent à montrer que dans certaine mesure, la manipulation de l'information en vertu des règles de la mécanique quantique est plus efficace que sa manipulation classique. Cependant il n'y a à l'heure actuelle aucun résultat ni aucun théorème expliquant comment cela se manifeste exactement et qu'elles en sont les causes. De ce fait, de nombreuses questions émergent quand on tente d'analyser la *supériorité* de la manipulation quantique pour le traitement de l'information et ces questions guident de nombreux travaux. Qu'est ce qui rend supérieur la manipulation quantique de l'information ? Ou en d'autres termes, qu'elles sont les ressources de l'information quantique ? C'est ce que nous allons tenter de voir dans la prochaine partie.

1. L'algorithme de Shor peut factoriser un entier N en $\mathcal{O}(\log(N)^3)$ opérations alors qu'on ne connaît aucun algorithme classique capable de résoudre ce problème en $\mathcal{O}(\log(N)^k)$ opérations, $\forall k$. Classiquement il est nécessaire d'effectuer approximativement $\mathcal{O}(\exp((\log N)^{1/3}(\log \log N)^{2/3}))$ opérations.

III. Ressources de l'information quantique

S'il y a effectivement des ressources spécifiques à l'information quantique, elles sont à chercher dans ce qui différencie la mécanique quantique de la mécanique classique. On peut considérer que deux concepts plus ou moins indépendants émergent à la suite d'une telle recherche [BD00]. Ces deux concepts sont l'*intrication* et l'*interférence*.

III.1. L'intrication

La notion d'intrication et d'état intriqué est évoquée dès 1935 par Erwin Schrödinger [Sch35] puis dans le papier d'Einstein, de Podolski et de Rosen [EPR35] pour mettre l'accent sur certaines bizarreries de la mécanique quantique et la nécessité de trouver un moyen d'améliorer cette théorie. En effet, les états intriqués exhibent des propriétés n'ayant pas d'analogue classique, en particulier l'apparition de corrélations non-triviales lors de la mesure d'un système de deux particules. Cependant, l'étude de ce genre d'état dans les années soixante par Bell [Bel64] a permis de faire un pas en avant en montrant que les états intriqués pouvaient justement servir à trancher entre la mécanique quantique et une théorie à variables cachées susceptible de la remplacer. La vérification expérimentale par Aspect *et.al* [ADR82] ainsi, que toutes les autres expériences effectuées depuis, ont montré que les prédictions de la mécanique quantique restent en accord avec l'expérience contrairement aux prédictions d'une quelconque théorie à variables cachées.

La notion d'intrication est liée à la manière dont on construit l'espace des états d'un système composite, à partir des espaces des états de chacune de ses parties. Soient N systèmes quantiques s_i dont les états vivent dans des espaces de Hilbert $\mathcal{H}_{(i)}$, alors l'état du système global S vit dans un espace de Hilbert qui est le produit tensoriel des espaces de Hilbert de chaque sous système :

$$\mathcal{H}_S = \bigotimes_{i=1}^N \mathcal{H}_{(i)}$$

Il en suit que l'état le plus général du système global ne peut pas s'écrire comme le produit tensoriel des états individuels des sous-systèmes, tout simplement parce que ces états individuels n'existent pas à proprement parler, ils sont enchevêtrés. A titre d'exemples et aussi pour revenir à une formulation plus proche de l'information quantique, considérons un système à deux qubits. L'état le plus général de ce système est donné par (1.6) qui dans ce cas se simplifie en :

$$|q_2\rangle = a_{00}|00\rangle + a_{01}|01\rangle + a_{10}|10\rangle + a_{11}|11\rangle$$

Ici, l'état du système est entièrement déterminé par la valeur des composantes a_{ij} . Par exemple l'état $\frac{1}{\sqrt{2}}(|00\rangle + |01\rangle)$ peut aussi s'écrire $|0\rangle \otimes \frac{(|0\rangle + |1\rangle)}{\sqrt{2}}$, ce qui permet d'affirmer que le *premier* qubit est dans l'état $|0\rangle$ alors que de manière indépendante le *second* est dans l'état $\frac{(|0\rangle + |1\rangle)}{\sqrt{2}}$. Mais cette propriété n'existe pas dans le cas de l'état $\frac{1}{\sqrt{2}}(|01\rangle + |10\rangle)$. Il ne peut pas s'écrire comme le produit tensoriel d'états relatifs aux deux qubits. Dans ce cas précis, la description de l'état d'un qubit individuel par un état pur, indépendamment des autres n'est pas possible. On peut au mieux recourir à une formulation en terme de matrice densité. Seul le système dans son ensemble possède une description par un état pur. C'est un état intriqué.

Par définition, on dit qu'un état quantique associé à un système composé est *séparable* s'il peut s'écrire comme le produit tensoriel d'états quantiques de ses parties. Dans le cas contraire on dit qu'il est intriqué. Quels sont les indices montrant que l'intrication semble constituer une ressource de l'information quantique ?

Premièrement, dès que l'on considère un système à plusieurs qubits, l'espace de Hilbert se trouve occupé par une proportion d'états intriqués exponentiellement plus grande que la proportion d'états séparables [ZHSL99]. Ainsi un processus d'information quantique met obligatoirement en jeu de l'intrication. Ceci est d'autant plus vrai que le nombre de qubits est grand.

Deuxièmement, la notion d'intrication donne au concept de qubit un rôle plus subtil et plus profond que celui qui est réservé au bit classique. En effet, même si l'information quantique est représentée physiquement par n_q objets physiques, sa description se répartit à proprement parler sur 2^{n_q} entités individuelles, par exemple sur les états de la base computationnelle. Bien sûr, pendant la préparation de l'état initial ou lors du processus de mesure, ce sont bien n_q objets qui interviennent, et un ensemble de qubits ne permettent évidemment pas de représenter exponentiellement plus d'information que le ferait un même nombre de bits [BF07]. Cependant avant le processus de mesure, lors de l'évolution unitaire, l'intrication retire momentanément leur individualité aux n_q objets physiques, et seul l'état du registre tout entier vivant dans un espace de Hilbert de dimension exponentielle 2^{n_q} reste parfaitement descriptible. Comme c'est dans cet espace de Hilbert que se déroule tout processus informationnel on pourrait être tenté de penser que c'est *ici* que l'intrication rencontre l'efficacité quantique. Cependant il convient d'être prudent avec ce genre de raisonnement qualitatif : considérons l'action d'une porte quantique agissant pour simplifier sur le premier qubits d'un registre à n_q qubits, dont l'état est donné par la relation (1.6). En écrivant la porte $G^{(1)} = U \otimes \mathbf{1} \otimes \dots \otimes \mathbf{1}$, l'état du registre après son action est

$$\begin{aligned} |q'_{n_q}\rangle &= G^{(1)}|q_{n_q}\rangle \\ |q'_{n_q}\rangle &= \sum a'_{i_1 i_2 \dots i_{n_q}} |i_1 i_2 \dots i_{n_q}\rangle \\ \text{où} \quad a'_{i_1 i_2 \dots i_{n_q}} &= \sum_{j_1} U_{i_1 j_1} a_{j_1 i_2 \dots i_{n_q}} \end{aligned} \tag{1.11}$$

Quantiquement ce processus consiste, par définition, en une seule étape qui est l'action de l'opérateur unitaire U sur le premier qubit. D'un point de vu naïf, on pourrait penser que pour simuler ce processus sur un ordinateur classique il soit nécessaire de recourir à un nombre exponentiel d'opérations arithmétiques élémentaires puisqu'il faut effectuer la mise à jour (1.11) des $2^{n_q} - 1$ amplitude $a'_{i_1 i_2 \dots i_{n_q}}$. Alors que d'un autre côté, si on suppose que l'état du registre reste séparable à tout instant du calcul, c'est dire que les amplitudes prennent la forme $a_{i_1 i_2 \dots i_{n_q}} = a_{i_1} b_{i_2} \dots z_{i_{n_q}}$, simuler l'action de la porte (1.11) reste efficace en n'utilisant qu'un nombre polynomial d'opérations arithmétiques élémentaires (opérations permettant la mise à jour de l'amplitude a_{i_1} associée au premier qubit).

De récents travaux [JL03] ont obtenu certains résultats concernant la possibilité ou non de simuler efficacement un calcul quantique sur un ordinateur classique en vertu de la manière dont se comporte l'intrication lors du calcul. Ils montrent que si l'intrication reste tout au long du calcul *p-blocked*, c'est à dire qu'elle reste confinée sur un nombre de qubits $p < n_q$, alors le calcul reste simulable classiquement de manière efficace. Le cas de la simulation d'un calcul quantique avec des états mixtes reste par contre un problème encore ouvert même si ce papier donne quelques indications à son sujet.

À l'heure actuelle, l'intrication est sûrement l'une des propriétés les plus étudiées en information quantique et ceci depuis maintenant presque vingt ans. Et il reste sûrement encore beaucoup de choses intéressantes à apprendre à son sujet. Son étude s'est même étendue au delà de l'information quantique (problème à N -corps quantiques, matière condensée, physique des hautes énergies) et constitue désormais un champ à part entière. Parmi les nombreux travaux

qui la concernent, celui de la manière de quantifier son existence dans les systèmes quantiques s'est vite développé. Cependant, comme l'intrication est une notion très générale qui devient de nature multiple dès que le nombre de partitions du système quantique dépassent deux, les manières de la quantifier peuvent être diverses, choisi en fonction du problème étudié et il existe de ce fait de nombreuses mesures quantitatives utilisables. Citons par exemple l'entropie d'intrication, l'entropie linéaire, la négativité, la concurrence, l'entropie géométrique, etc. Toutes ces mesures ont la propriété d'être invariantes sous l'action de transformations locales (c'est à dire des opérations unitaires appliquées sur un seul qubit).

III.2. L'interférence quantique

Dans la première partie, nous avons évoqué le concept d'interférence tel qu'il apparaît en physique et plus particulièrement en physique quantique. Que se passe-t-il quand on examine le concept d'interférence avec le point de vue de l'information quantique et pourquoi pourrait-elle en constituer une ressource ?

Pensons à l'exemple typique qui permet de différencier un phénomène classique d'un phénomène quantique mettant en jeux des interférences. Il s'agit habituellement de deux chemins A et B , favorisés classiquement qui interfèrent destructivement au profit de deux autres chemins C et D moins probables classiquement. Ou alors on peut penser à la vision de Feynman dans laquelle un système physique explore la totalité des chemins possibles et dans laquelle le principe de moindre action émerge dans la limite classique. Ces deux exemples exhibent la capacité que possède l'interférence quantique pour modifier le flux classique d'information. C'est dans ce sens qu'il faut visualiser la manière dont l'interférence pourrait être nécessaire pour engendrer une accélération exponentielle dans les algorithmes quantiques. Historiquement, c'est d'ailleurs au travers du concept d'interférence, qu'ont été découverts les premiers algorithmes quantiques, en particulier celui de Deutsch². Ainsi l'étude de l'interférence dans le cadre de l'information quantique est un sujet intéressant d'autant que le nombre d'études à son égard dans la communauté reste faible, en comparaison par exemple des études concernant l'intrication.

Il convient ensuite d'énumérer les propriétés qui caractérisent le concept d'interférence et l'expérience familière des doubles fentes d'Young est à garder en tête pour mieux les appréhender :

Propriété 1 : L'interférence ne doit pas être vue comme une propriété d'un état quantique mais comme une propriété de son évolution. C'est la propagation qui est interférente. Cette idée est contenue dans l'expression *interférence entre deux chemins quantiques* souvent utilisée en mécanique quantique. Ce ne sont pas les états qui interfèrent entre eux. Ce sont les *chemins* qu'emprunte le système quantique lors de son évolution. Il est intéressant de noter cette propriété de complémentarité que possède l'interférence par rapport à l'intrication : l'intrication est une propriété des états quantiques tandis que l'interférence est une propriété de leur évolution.

Propriété 2 : L'interférence est reliée à la notion de cohérence, dans le sens où la première notion s'efface sans la seconde. Cependant une propagation cohérente n'engendre pas nécessairement un motif d'interférence comme c'est le cas si l'expérimentateur met l'écran avant les deux fentes, bien que la propagation soit cohérente entre la source et l'écran. En effet il faut au moins que deux amplitudes de probabilités soient superposées pour observer un motif d'interférence. De

2. Algorithme de Deutsch permet d'affirmer si une fonction binaire à un bit reste soit constante ($f(0) = 0, f(1) = 0$ ou $f(0) = 1, f(1) = 1$) soit équilibrée ($f(0) = 0, f(1) = 1$ ou $f(0) = 1, f(1) = 0$) en n'envoyant qu'une seule requête à la fonction f . Pour résoudre le problème classiquement, il faut envoyer 2 requêtes à la fonction f . Cet algorithme peut-être généralisé dans le cas d'une fonction binaire à plusieurs bits où le nombre de requêtes augmente polynomialement avec le nombre de bits alors que classiquement ce nombre augmente exponentiellement [DJ92].

manière générale une propagation est d'autant plus interférente qu'elle superpose avec des poids égaux un maximum d'amplitude de probabilité.

Propriété 3 : Pour des raisons similaires, l'interférence quantique est une notion qui dépend de la base dans le sens où le nombre d'amplitudes dans une superposition cohérente dépend lui aussi de la base. Ainsi, dans une expérience d'interférence, l'apparition d'un motif d'interférence dépend du point de vue choisi par l'expérimentateur pour faire la mesure. Mathématiquement cela veut dire que l'interférence dépend de la base de représentation des états quantiques. Pour bien cerner cela, utilisons encore l'analogie de l'expérience des doubles-fentes. Une source ponctuelle placée en $x = x_0$ peut être représentée comme une source monochromatique de couleur $\mathbf{k}_0$, d'amplitude $\Psi(x)$. Après les fentes, le motif apparaît sur l'écran, c'est à dire quand on mesure $|\Psi'(x)|^2$, l'intensité lumineuse dans l'espace des positions. Mais si l'expérimentateur choisi de mesurer $|\hat{\Psi}(\mathbf{k})|^2$, l'intensité dans l'espace de vecteur d'onde, il n'observera dans ce cas aucun motif d'interférence mais juste la distribution initiale parfaitement localisée en $\mathbf{k}_0$.

III.3. Nuances sur les capacités des ressources quantiques

A ce stade il convient d'être plus précis sur la manière d'envisager le lien qui est supposé exister entre l'accélération associée aux calculateurs quantiques et la présence ou non d'intrication ou d'interférence. Ainsi on voit que deux points de vue différents peuvent être envisagés :

1) Soit ce sont bel et bien l'intrication et l'interférence qui engendrent l'accélération exponentielle. Ils en sont la cause.

2) Soit l'accélération exponentielle est d'une nature plus subtile, l'intrication et l'interférence ne sont juste que des produits secondaires de cette accélération.

Pour l'instant, les méthodes d'investigations mise en œuvre, ne sont pas en mesure de pouvoir choisir entre ces deux points de vue.

Dans [JL03] il est évoqué qu'il est peut-être inapproprié de penser que l'intrication est une ressource clef tout simplement parce que c'est une propriété qui dépend de la description mathématique de la mécanique quantique, en l'occurrence la formulation en terme d'amplitude de probabilité. Les auteurs rappellent qu'une infinité de descriptions mathématiques équivalentes $\mathcal{D}$ de la mécanique quantique sont envisageables et que pour chaque description, il doit exister une propriété donnée $Prop(\mathcal{D})$ dépendant de la description, qui différencie une classe particulière d'état, cette classe n'ayant pour ainsi dire aucune particularité absolue, si ce n'est l'incapacité d'être efficacement représenté classiquement. Par exemple quand $\mathcal{D}$ est la description par des amplitudes de probabilité, $Prop(\mathcal{D})$ est l'intrication. Les auteurs conjecturent alors que dans le cadre de l'information quantique, l'accélération exponentielle des algorithmes quantiques doit être reliée au comportement de l'ensemble des $Prop(\mathcal{D})$ qui ne peuvent pas être efficacement représentées classiquement.

IV. Une mesure quantitative de l'interférence quantique

L'introduction d'une mesure quantitative de l'interférence est nécessaire si on veut pouvoir analyser l'implication du phénomène dans le cadre des processus informationnels. Une telle mesure est proposée dans [BG06]. Notons que, comme pour l'intrication, la généralité du concept d'interférence fait que la mesure issue de ces travaux n'est certainement pas unique et rien n'empêcherait d'en définir d'autres. Cependant c'est la mesure introduite dans cet article qui sera utilisée dans les chapitre 2 et 3 de cette thèse. Signalons qu'au cours de cette thèse et

ceci pour des raisons de simplicité, nous utiliserons souvent le terme *interférence*, pour désigner l'interférence quantique elle-même, mais aussi pour désigner sa mesure.

IV.1. Expression générale

Dans [BG06], l'expression générale d'une mesure d'interférence est obtenue en recherchant tout d'abord une mesure de la cohérence pour un propagateur P qui transforme une matrice densité d'un registre quantique selon $\rho' = P\rho$. P est un super-opérateur dont les composantes dans la base computationnelle sont telles que

$$\rho'_{ij} = \sum_{k,l} P_{ij,kl} \rho_{kl}. \quad (1.12)$$

Cette équation est la manière la plus générale de décrire un algorithme quantique dans lequel peuvent intervenir diverses contributions non-unitaires telles que la décohérence produite par un environnement externe. Un ordinateur quantique peut être vue comme une boîte noire envoyant une matrice densité initiale sur une distribution de probabilité finale, le processus de lecture faisant partie de l'algorithme. Si on ne regarde la propagation que d'un seul état, il est impossible de dire si la propagation est cohérente ou non en n'ayant accès qu'aux probabilités finales $p'_i = |\rho'_{ii}|^2$. Le cas extrême où la distribution finale peut être obtenue en propageant la distribution initiale par une matrice stochastique³ convenablement choisie correspond à une propagation incohérente puisque la distribution finale ne dépend plus des phases initiales. Ainsi pour pouvoir quantifier la cohérence, il est naturel de regarder la dépendance de la distribution finale envers les phases initiales.

Pour un état pur initial $\rho = |\psi\rangle\langle\psi|$, avec $|\psi\rangle = \sum_{j=1}^N a_j |j\rangle$, les amplitudes a_j avec les phases φ_j , $a_j = |a_j|e^{i\varphi_j}$, mènent à des probabilités finales

$$p'_i = |\rho'_{ii}| = \sum_{j,k} P_{ii,jk} e^{i(\varphi_j - \varphi_k)} |a_j a_k|. \quad (1.13)$$

On peut définir S la matrice réelle de « sensibilité de phase » de composantes $S_{il} = \partial p'_i / \partial \varphi_l$, ainsi que la matrice définie semi-positive SS^T telle que $SS^T = 0$ ssi $\partial p'_i / \partial \varphi_l = 0 \forall i, l = (1, \dots, N)$. La mesure de cohérence recherchée $\mathcal{C}(P)$ peut être obtenue en moyennant sur toutes les phases initiales, la trace de la matrice SS^T

$$\mathcal{C}(P) = \frac{1}{(2\pi)^N} \int_0^{2\pi} d\varphi_1 \dots d\varphi_N \text{Tr}(SS^T). \quad (1.14)$$

On peut montrer que cette expression se simplifie comme

$$\mathcal{C}(P) = 2 \sum_{i,k} \sum_{l \neq k} |P_{ii,lk} a_k a_l|^2. \quad (1.15)$$

Cette mesure dépend encore des amplitudes de l'état initial. Pour pouvoir acquérir le statut d'une mesure de l'interférence, les auteurs de [BG06] préconisent de considérer $\mathcal{C}(P)$ pour des états équipartitionnés sur la base computationnelle, c'est à dire des états tels que $|a_i| = 1/\sqrt{N}$ pour tous les $i = (1, \dots, N)$. Ainsi, après multiplication par un préfacteur $N^2/2$, on obtient la mesure d'interférence

$$\mathcal{I}(P) = \sum_{i,k,l} |P_{ii,kl}|^2 - \sum_{i,k} |P_{ii,kk}|^2. \quad (1.16)$$

3. Matrice stochastique : Matrice positive qui propage des probabilités selon $p' = Mp$. Ses composantes M_{ij} sont comprises entre 0 et 1 et sont telles que $\sum_i M_{ij} = 1$.

De cette façon, la mesure ne dépend plus d'aucun état particulier mais seulement du propagateur P ce qui est attendu d'une mesure de l'interférence en vertu de la propriété 1. Par la forme de son expression, on démontre facilement que $0 \leq \mathcal{I}(P)$. Si l'ensemble des valeurs propres de P sont plus petites que 1 (comme dans le cas d'une application quantique dissipative) alors $\mathcal{I} \leq N^2$. Notons aussi que pour une propagation classique les composantes de P s'écrivent $P_{ij,kl} = M_{ik}\delta_{ij}\delta_{kl}$ et qu'il est facile de montrer que dans ce cas $\mathcal{I}(P)$ s'annule exactement.

IV.2. Cas d'une propagation purement unitaire

Dans le cas d'une propagation unitaire, les composantes du propagateur P s'écrivent $P_{ij,kl} = U_{ik}^*U_{jl}$ et l'expression (1.16) se simplifie en

$$\mathcal{I} = \mathcal{I}(U) = N - \sum_{ij} |U_{ij}|^4. \quad (1.17)$$

Cette expression est parfaitement en accord avec les propriétés associées à l'interférence : Tout d'abord, elle ne dépend exclusivement que des composantes de l'opérateur d'évolution U (propriété 1). Ensuite il est facile de montrer que $0 \leq \mathcal{I} \leq N - 1$. La borne inférieure est obtenue dans le cas d'opérateurs U ayant des colonnes parfaitement localisées (une composante de module 1 et les autres de valeur nulle) comme pour l'opération d'identité, pour des opérations de permutation des états de base ou des rotations de phases locales. La borne maximale est obtenue dans le cas d'opérateurs U ayant des colonnes équipartitionnées, c'est à dire avec des colonnes remplies de composantes de module $\frac{1}{\sqrt{N}}$. En fait la somme contenue dans l'expression (1.17) est l'*inverse participation ratio* qui permet de caractériser le degré de localisation dans un vecteur. De manière générale, la mesure d'interférence est maximisée pour des évolutions réalisant le passage de la base computationnelle à sa base conjuguée, comme le fait par exemple la *transformée de Fourier quantique*⁴. Toutes ces remarques sont en accord avec la propriété 2. Finalement $\mathcal{I}$ n'est manifestement pas un invariant et il importe de préciser dans quelle base l'expression est écrite. Dans toute cette thèse, les expressions (1.16) et (1.17) seront écrites dans la base computationnelle (propriété 3). Ajoutons aussi qu'il n'existe pas *a priori* de relation simple entre d'une part $\mathcal{I}(U_1U_2)$ et d'autre part $\mathcal{I}(U_1)$ et $\mathcal{I}(U_2)$.

IV.3. Cas particulier pour un seul qubit

Pour avoir une vision plus intuitive de la manière dont ce comporte l'interférence, considérons le cas à un seul qubit (c'est à dire pour $N = 2$). La transformation de Hadamard (1.9) réalise une propagation maximalisant la mesure $\mathcal{I}$. Ceci est très satisfaisant, car c'est effectivement la transformation de Hadamard, qui réalise le changement de base (1.7) permettant la mesure expérimentale de la phase relative ϕ transformation que l'on retrouve dans un interféromètre (par exemple dans un interféromètre de Mach-Zender elle est réalisée par les lames séparatrices). Dans [BG06], la transformation de Hadamard est vue comme la transformation élémentaire permettant de définir une unité d'interférence, le *i-bit* par l'expression $\log_2(\mathcal{I} + 1)$. Le cas $N = 2$ permet de cerner une autre propriété de l'interférence : comme $H^2 = \mathbf{1}$, deux transformations

4. La transformée de Fourier quantique est la transformation unitaire définie sur une base orthonormale $|0\rangle, |1\rangle \dots |N-1\rangle$ par l'application

$$|j\rangle \rightarrow \frac{1}{\sqrt{N}} \sum_{k=0}^{N-1} e^{2\pi i j k} |k\rangle$$

Voir [NC00]. Dans le cas $N = 2$, la transformée de Fourier quantique se résume à la transformation de Hadamard.

de Hadamard consécutives sur le même qubit apportent une quantité d'interférence qui est nulle. Ceci est une propriété inhérente à la définition et il convient de différencier deux types d'interférence : L'interférence *localisée* qui est celle correspondant à une portion d'algorithme et l'interférence *accumulée*, de nature globale, qui est celle de l'algorithme tout entier.

V. Résultats

Dans cette partie, nous allons faire une revue des différents travaux utilisant la mesure introduite dans [BG06].

V.1. Interférence dans l'algorithme de Shor

L'algorithme de factorisation de Shor U_S [Sho94] permet de factoriser un entier R en un nombre polynomial d'opérations. L'état initial est réparti sur deux registres respectivement de taille $2L$ et L avec $L = \lceil \log_2 R \rceil + 1$ où $\lceil \cdot \rceil$ dénote la partie entière. Le nombre de qubits nécessaires est donc $n_q = 3L$. L'algorithme peut être décomposé en trois étapes : La première étape est l'application indépendamment sur les $2L$ premiers qubits, d'une transformation de Hadamard. Ceci génère une grande quantité d'interférence valant $2^{3L} - 2^L$. Notons que comme les transformations de Hadamard sont des transformations locales, aucune intrication n'est créée à ce stade. La deuxième partie est une séquence de $O(n^3)$ opérations, séquence pouvant être vue comme une matrice de permutation (chaque ligne et chaque colonne n'ayant qu'une seule composante non-nulle). Contrairement à la première, cette étape engendre de l'intrication mais aucune interférence. La dernière étape est une transformée de Fourier quantique appliquée sur le premier registre. Cette transformation génère aussi bien de l'intrication que de l'interférence. Dans [BG06] les auteurs préconisent de ne pas prendre en compte l'interférence trivialement produite dans la première étape de l'algorithme. Ils ne considèrent donc que la version dite *non-générique* de l'algorithme, c'est à dire l'algorithme sans les opérations de Hadamard initiales. Ils démontrent ainsi que l'interférence dans cette version non-générique de l'algorithme, exclusivement due à la transformée de Fourier quantique, est produite en quantité exponentielle en fonction du nombre de qubits (typiquement $2^{3L} - 2^L$).

V.2. Interférence dans l'algorithmes de Grover

L'algorithme de recherche de Grover U_G [Gro97] permet de trouver un élément donné dans une liste non structurée de N éléments en un nombre d'opérations $O(\sqrt{N})$ contre $O(N)$ classiquement. Il fonctionne sur un ensemble de n_q qubits tel que $N = 2^{n_q}$ et il peut aussi être décomposé en plusieurs étapes : Comme pour l'algorithme de factorisation, la première étape consiste en l'application d'opérations de Hadamard indépendantes sur l'ensemble des qubits et les auteurs préfèrent ne pas la considérer. Les autres étapes, qui constituent la version non-générique de l'algorithme, sont une succession d'opération qui engendrent aussi bien de l'intrication que de l'interférence. Cependant il est montré dans [BG06] que l'interférence produite dans cette partie est constante et indépendante du nombre de qubits et vaut approximativement 3 i-bits. Ceci est en contraste flagrant avec ce qui se passe dans l'algorithme de Shor. Ainsi les auteurs conjecturent que la différence d'efficacité entre ces deux algorithmes est liée à la quantité d'interférence produite dans leur partie non-générique.

V.3. Processus n'impliquant pas de l'interférence

Plusieurs travaux se basant sur la mesure (1.16), ont montré que l'interférence n'intervenait pas dans certains processus.

Dans [LBB07] il est montré que l'interférence n'intervient pas dans la propagation d'un état quantique le long d'une chaîne de spin.

Dans [BG08] il est montré que si l'on considère l'algorithme de Shor soumis à des imperfections, la quantité exponentielle d'interférence qu'il produit normalement sans imperfection est détruite.

Dans [RB08], les auteurs se sont intéressés à une classe de cloneur quantique particulier. Rappelons qu'en information quantique il est impossible de trouver un opérateur unitaire permettant de transformer un état inconnu de la forme $|\varphi_1\rangle|\varphi_2\rangle$ en un état de la forme $|\varphi_1\rangle|\varphi_1\rangle$, c'est à dire réaliser une copie parfaite de l'état du premier registre dans le second. C'est le théorème du non-clonage quantique [WZ82]. Cependant on peut considérer des cloneurs quantiques, c'est à dire des algorithmes U_c réalisant une copie avec une certaine fidélité, sachant qu'une fidélité parfaite de 1 est impossible à atteindre. Il est montré dans [RB08] que le maximum de fidélité admissible, valant $5/6$, peut être atteint par la classe de cloneurs considérés, sans produire d'interférence.

VI. Conclusion partielle

Dans ce premier chapitre, nous avons évoqué la possibilité qu'à la manipulation quantique de l'information pour résoudre certains problèmes computationnels de manière plus efficace que classiquement. Cette efficacité n'étant pas bien comprise à l'heure actuelle, la question de savoir quelle en est la cause se pose. N'ayant pas d'analogie classique, il est légitime de regarder si les phénomènes d'intrication et d'interférence quantique jouent un rôle à ce niveau. En particulier nous avons considéré la mesure d'interférence introduite dans [BG06]. Cette mesure se comporte de manière différente quand on considère les algorithmes de Shor et Grover. Elle montre aussi que certains processus ne produisent pas d'interférence pour remplir leurs rôles. De ce fait, en considérant des algorithmes ou des processus précis, on se rend compte que l'interférence est une quantité qui peut être présente dans certains et pas dans d'autres. Ainsi se posent les questions suivantes : si on choisit un algorithme quantique au hasard, quelles sont les chances pour qu'il contienne une interférence nulle ? Qu'elles sont les chances pour qu'il contienne une interférence maximale ? Est-ce que certaines valeurs d'interférence sont plus probables que d'autres ou est-ce que toutes les valeurs permises sont équiprobables ? Répondre à ces questions est intéressant car cela permettrait de savoir si l'interférence est un concept qui caractérise vraiment un algorithme quantique, et s'il convient de penser que c'est bel et bien une ressource computationnelle. Nous allons tenter de répondre à ces questions dans le prochain chapitre.

Statistique de l'interférence dans les algorithmes quantiques

- *Le principe du bigle moi, dit Nicolas,
 que Monsieur connaît sans doute,
 repose sur la production d'interférences
 par deux sources animées d'un mouvement
 oscillatoire rigoureusement synchrone.*
 - *J'ignorais, dit Colin, que cela mît en œuvre
 des éléments de physique aussi avancée*

Boris Vian, l'écume des jours.

I. Motivation

Le but de ce chapitre est d'avoir une vision globale de la manière dont la mesure d'interférence introduite au premier chapitre se comporte vis à vis de l'ensemble des algorithmes quantiques. Il convient de préciser qu'ici, le terme algorithme quantique est à prendre au sens le plus large du terme et nous ne faisons aucunement référence à des algorithmes précis, résolvant des problèmes précis. En d'autre terme nous appelons algorithme tout opérateur unitaire agissant sur un registre de n_q qubits et l'ensemble des algorithmes quantiques auxquels nous faisons référence n'est rien d'autre que le groupe $U(N)$ où $N = 2^{n_q}$. Comme nous l'avons signalé dans le précédent chapitre, la mesure

$$\mathcal{I}(U) = N - \sum_{ij} |U_{ij}|^4. \quad (2.1)$$

s'annule pour l'opération d'identité ou pour une permutation des états de bases alors qu'elle est maximale ($\mathcal{I} = 1$) pour une transformation de Hadamard. Ces exemples constituent des cas limites et il peut être intéressant de considérer des cas intermédiaires, pour des opérations unitaires quelconques.

Pour un seul qubit, il n'est pas difficile d'obtenir une *cartographie* de la mesure (2.1) : Un élément du groupe $U(2)$ peut être paramétré par la matrice unitaire

$$U_2 = e^{i\alpha} \begin{pmatrix} \sqrt{1-\xi}e^{i\psi} & \sqrt{\xi}e^{i\chi} \\ -\sqrt{\xi}e^{-i\chi} & \sqrt{1-\xi}e^{-i\psi} \end{pmatrix} \quad (2.2)$$

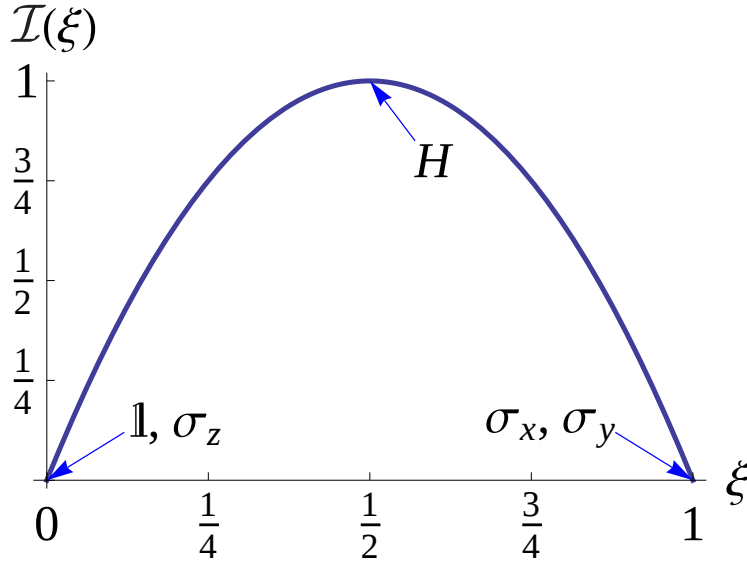


FIGURE 2.1.: Interférence calculée pour un opérateur unitaire de taille 2. Le paramètre ξ contrôle l'équipartition dans la matrice et les flèches indiquent où sont situés par rapport à ce paramètre, les opérateurs unitaires usuels à un qubit.

où α, ψ, χ choisis dans $[0, 2\pi[$ et ξ choisi dans $[0, 1[$, sont des paramètres réels. En injectant cette matrice dans (2.1) on obtient la relation $\mathcal{I}(U_2) = \mathcal{I}(\xi) = 4(\xi - \xi^2)$ représentée sur la figure (2.1). Si on fait abstraction des phases qui n'interviennent pas dans l'expression, on peut associer à chaque point de l'axe horizontal ξ un type d'opérateur. Par exemple le point $\xi = 0$ correspond à la matrice identité (ou à la matrice de Pauli σ_z) et le point $\xi = 1$ correspond aux matrices réalisant une permutation des états de base telle les matrices de Pauli σ_x ou σ_y . Pour tous ces opérateurs, on observe bien une interférence nulle. Le maximum d'interférence $\mathcal{I} = 1$ se situe en $\xi = 1/2$, point qui correspond bien à l'opération de Hadamard. Ainsi l'interférence ne dépend que d'un nombre réduit de paramètres, en l'occurrence de ξ qui contrôle l'équipartition dans la matrice U_2 .

Obtenir le même genre d'information pour des opérateurs plus grand serait intéressant. Cependant, dès que le nombre de qubits augmente, cela devient très difficile même pour seulement 2 qubits. En effet, il est nécessaire pour ceci de manipuler un espace de paramètres continus qui croît exponentiellement avec le nombre de qubits et même les méthodes numériques ne permettent pas d'aller très loin, sans parler de la difficulté pour représenter les résultats. Ainsi la méthode d'investigation utilisée dans ce chapitre sera une analyse statistique. L'idée est la suivante : Choisir des algorithmes quantiques au hasard, les injecter dans la mesure (2.1) et obtenir des informations sur la statistique de l'interférence.

II. Distribution de l'interférence dans des ensembles circulaires

II.1. L'ensemble circulaire unitaire CUE

Cette stratégie pour obtenir des informations statistiques sur l'interférence dépend bien sûr de la manière dont les algorithmes quantiques vont être distribués. Ainsi il convient dès le

départ de spécifier un ensemble d'algorithmes quantiques dont les éléments sont distribués de manière précise. Sans connaissance particulière sur les algorithmes quantiques et de manière à n'en favoriser aucun, il est naturel de choisir l'ensemble des matrices unitaires uniformément distribuées par rapport à la mesure de Haar¹ de $U(N)$. Cet ensemble est connu comme étant l'*ensemble unitaire circulaire* (circular unitary ensemble ou CUE) dans la classification des ensembles circulaires de Dyson. Comme la mesure d'interférence (2.1) ne dépend que du module des éléments de matrice U_{ij} il peut être intéressant de considérer un ensemble de matrices unitaires réelles, c'est à dire un ensemble de matrices orthogonales elles aussi uniformément distribuée par rapport à la mesure de Haar du groupe orthogonal $O(N)$. Si la mesure de Haar est invariante à droite et à gauche, c'est à dire si on a $d\mu(O) = d\mu(V_1 O V_2)$ quelles que soient V_1 et V_2 appartenant à $O(N)$, cet ensemble, moins répandu, est l'*ensemble orthogonal de Haar* (Haar orthogonal ensemble ou HOE) [PZK98].

II.2. Distribution calculée numériquement

Numériquement il est possible de calculer la distribution de la mesure d'interférence $P_{CUE,N}(\mathcal{I})$ pour des matrices tirées de CUE, produites en grand nombre, en utilisant la paramétrisation de Hurwitz du groupe $U(N)$ [Hur97, PZK98] qui permet de construire chaque matrice unitaire U de taille N à partir du produit de $\frac{N(N-1)}{2}$ matrices $E^{(i,j)}(\phi, \psi, \chi)$ de transformation unitaire élémentaire agissant dans un sous-espace de dimension deux. Les éléments de matrice non-nuls de la transformation unitaire élémentaire s'écrivent

$$\begin{aligned} E_{kk}^{(i,j)} &= 1, \forall k = 1, \dots, N; k \neq i, j \\ E_{ii}^{(i,j)} &= \cos(\phi) e^{i\psi} \\ E_{ij}^{(i,j)} &= \sin(\phi) e^{i\chi} \\ E_{ji}^{(i,j)} &= -\sin(\phi) e^{-i\chi} \\ E_{jj}^{(i,j)} &= \cos(\phi) e^{-i\psi} \end{aligned} \tag{2.3}$$

A partir de cette transformation, on peut construire les $N - 1$ matrices de « rotation »

$$\begin{aligned} E_1 &= E^{(N-1,N)}(\phi_{01}, \psi_{01}, \chi_1) \\ E_2 &= E^{(N-2,N-1)}(\phi_{12}, \psi_{12}, 0) E^{(N-1,N)}(\phi_{02}, \psi_{02}, \chi_2) \\ &\vdots \\ E_{N-1} &= E^{(1,2)}(\phi_{N-2,N-1}, \psi_{N-2,N-1}, 0) E^{(2,3)}(\phi_{N-3,N-1}, \psi_{N-3,N-1}, 0) \dots \\ &\dots E^{(N-1,N)}(\phi_{0,N-1}, \psi_{0,N-1}, \chi_{N-1}), \end{aligned} \tag{2.4}$$

et finalement, la matrice unitaire est formée par

$$U = e^{i\alpha} E_1 E_2 \dots E_{N-1}, \tag{2.5}$$

où les angles α, ϕ_{rs} et ψ_{rs} sont choisis uniformément sur l'intervalle $[0, 2\pi]$, tandis que ϕ_{rs} est tel que $\phi_{rs} = \arcsin(\xi_{rs}^{1/(2r+2)})$ avec ξ_{rs} choisit uniformément sur $[0, 1]$.

La figure (2.2) montre les distributions obtenues pour N entre 2 et 8. Quand N croît, la distribution devient de plus en plus piquée, centrée autour d'une valeur proche de N . La figure (2.2) montre aussi les distributions $P_{HOE,N}(\mathcal{I})$ pour HOE. Numériquement cet ensemble est

1. La mesure de Haar $d\mu_N(U)$ de $U(N)$ est la mesure invariante associée au groupe unitaire et peut être vue comme la généralisation du concept d'élément de volume pour un groupe de Lie. Ainsi pour V appartenant à $U(N)$ l'invariance signifie qu'on a les relations $d\mu_N(U) = d\mu_N(VU) = d\mu_N(UV)$.

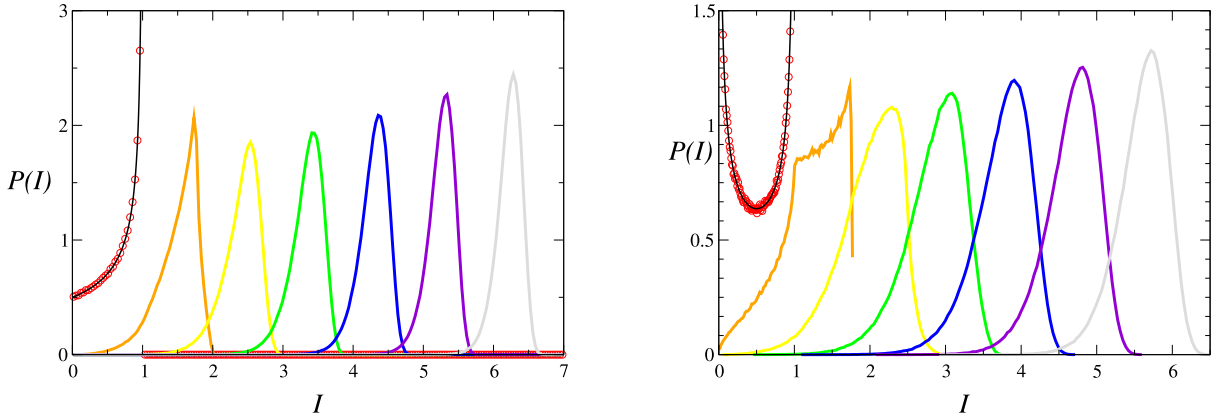


FIGURE 2.2.: Distribution de l'interférence dans CUE (gauche) et HOE (droite) pour $N = 2$ à $N = 8$ par pas de 1 (courbes de la gauche vers la droite). Les courbes pleines pour $N = 2$ représentent les résultats analytiques exprimés par les équations (2.6) et (2.7).

construit en diagonalisant des matrices appartenant à l'ensemble orthogonal gaussien² (GOE) [Meh91, PZK98]. Un fait remarquable est l'allure symétrique de la distribution pour $N = 2$. Pour $N > 4$, la distribution devient mono-nodale, et de plus en plus piquée quand N augmente, comme dans le cas unitaire.

II.3. Calculs analytiques des distributions dans le cas $N=2$

Dans le cas $N = 2$, la distribution $P_{CUE,2}(\mathcal{I})$ associée à CUE peut-être facilement calculée analytiquement en utilisant la paramétrisation (4.3). Pour obtenir des matrices uniformément distribuées par rapport à la mesure de Haar de $U(2)$ il suffit de choisir les paramètres α, ψ, χ et ξ uniformément distribués sur leurs ensembles de définition. Tout comme précédemment, il suit que $\mathcal{I} = \mathcal{I}(\xi) = 4(\xi - \xi^2)$ n'est fonction que de la variable ξ uniformément distribuée entre 0 et 1 et $P(\mathcal{I})$ peut être obtenue à partir de $P(\xi)$.

$$P(\mathcal{I}) = \int_0^1 \tilde{P}(\xi) \delta(\mathcal{I} - \mathcal{I}(\xi)) d\xi = \int_0^1 \tilde{P}(\xi) \sum_{\xi_0} \frac{\delta(\xi - \xi_0)}{|\mathcal{I}'(\xi_0)|} d\xi,$$

où les ξ_0 sont les racines de $\mathcal{I} - \mathcal{I}(\xi)$ et où le prime indique la dérivée par rapport à ξ . Dans notre cas, les deux racines ξ_0 pour $\mathcal{I}$ sont $\xi_{\pm} = (1 \pm \sqrt{1 - \mathcal{I}})/2$ et la dérivée est $\mathcal{I}'(\xi) = 4(1 - 2\xi)$.

2. Les ensembles gaussiens de matrices aléatoires GUE, GOE et GSE sont respectivement les ensembles de matrices hermitiennes, symétriques réelles et quaternions réelles, dont les éléments sont distribués par rapport à la loi normale. Les sigles rappellent que les mesures de Haar associées à chacun de ces ensembles sont respectivement invariantes sous des transformations unitaires, orthogonales et symplectiques [Meh91]

On a donc,

$$\begin{aligned}
P(\mathcal{I}) &= \int_0^1 \tilde{P}(\xi) \left(\frac{\delta(\xi - \xi_+)}{|\mathcal{I}'(\xi_+)|} + \frac{\delta(\xi - \xi_-)}{|\mathcal{I}'(\xi_-)|} \right) d\xi \\
&= \int_{-\infty}^{+\infty} \tilde{P}(\xi) \left(\frac{\delta(\xi - \xi_+)}{|4(1 - 2\xi_+)|} + \frac{\delta(\xi - \xi_-)}{|4(1 - 2\xi_-)|} \right) d\xi \\
&= \frac{1}{|4\sqrt{1 - \mathcal{I}}|} + \frac{1}{|-4\sqrt{1 - \mathcal{I}}|} \\
&= \frac{1}{2\sqrt{1 - \mathcal{I}}} = P_{CUE,2}(\mathcal{I}).
\end{aligned} \tag{2.6}$$

ce qui est en bon accord avec le résultat numérique représenté sur la figure (2.2).

De la même manière que dans le cas unitaire, il est facile de calculer analytiquement dans le cas $N = 2$, la distribution $P_{HOE,2}(\mathcal{I})$ associée à HOE en considérant la paramétrisation d'une matrice de rotation 2×2 par un angle uniformément distribué entre 0 et π . L'analogie du développement (2.6) donne

$$P_{HOE,2}(\mathcal{I}) = \frac{1}{\pi\sqrt{\mathcal{I}(1 - \mathcal{I})}}. \tag{2.7}$$

représenté aussi sur la figure (2.2). Le calcul des distributions $P_{HOE,2}(\mathcal{I})$ et $P_{HOE,2}(\mathcal{I})$ reste *a priori* inaccessible analytiquement pour $N > 2$. Cependant, il est possible d'obtenir les expressions exactes des deux premiers moments de ces distributions pour tout N .

II.4. Calculs analytiques des deux premiers moments

Ces calculs se basent sur une méthode d'intégration invariante introduite dans [AL03]. C'est une méthode clef dans les résultats analytiques présents dans cette thèse et elle sera utilisée intensivement dans le prochain chapitre. Pour ces raisons, l'annexe A lui est entièrement réservée. Cette méthode permet entre autre d'obtenir des expressions complètes pour des valeurs moyennes sur $U(N)$ de la forme

$$\begin{aligned}
Z_{U,N}(m_1, m_2, m_3) &\equiv \langle |U_{i_1 j_1}|^{2m_1} |U_{i_1 j_2}|^{2m_2} |U_{i_2 j_2}|^{2m_3} \rangle_{U,N} \\
&= \int d\mu_N(U) |U_{i_1 j_1}|^{2m_1} |U_{i_1 j_2}|^{2m_2} |U_{i_2 j_2}|^{2m_3} \\
&= \frac{m_1! m_2! m_3! (N-2)! (N-1)!}{(N+m_1-2)! (N+m_3-2)!} \\
&\quad \times \frac{(N+m_1+m_3-2)!}{(N+m_1+m_2+m_3-1)!},
\end{aligned} \tag{2.8}$$

où $d\mu_N(U)$ est normalisé par $\int d\mu_N(U) = 1$, et i_1, i_2, j_1, j_2 sont des indices arbitraires des éléments de matrices d'un élément quelconque U de $U(N)$. C'est justement le genre d'intégrales que l'on obtient si lors du calcul des deux premiers moments de $P(\mathcal{I})$. Pour la valeur moyenne on a

$$\langle \mathcal{I} \rangle_{U,N} = N - \sum_{ij} \langle |U_{ij}|^4 \rangle_{U,N} = N - N^2 Z_{U,N}(2, 0, 0) = N \left(1 - \frac{2}{N+1} \right) \tag{2.9}$$

$$\xrightarrow{N \rightarrow \infty} N - 2, \tag{2.10}$$

De la même manière pour le deuxième moment on a

$$\begin{aligned}
 \left\langle \left(\sum_{i,k} |U_{i,k}|^4 \right)^2 \right\rangle_{U,N} &= (N(N-1))^2 Z_{U,N}(2, 0, 2) \\
 &+ 2N^2(N-1)Z(2, 2, 0) \\
 &+ N^2 Z_{U,N}(4, 0, 0) \\
 &= 4 \frac{N^2 + 2N - 1}{(N+1)(N+3)}.
 \end{aligned} \tag{2.11}$$

Ainsi l'écart quadratique moyen de la distribution d'interférence pour CUE vaut

$$\sigma_{U,N} = \frac{2}{N+1} \sqrt{\frac{N-1}{N+3}}, \tag{2.12}$$

qui décroît comme $\sim 2/N$ pour N grand.

La méthode d'intégration invariante précédemment évoquée peut être généralisée pour l'ensemble HOE [Bra06] est la relation correspondante à (2.8) s'écrit :

$$\begin{aligned}
 Z_{O,N}(m_1, m_2, m_3) &\equiv \langle (O_{i_1 j_1})^{m_1} (O_{i_1 j_2})^{m_2} (O_{i_2 j_2})^{m_3} \rangle_{O,N} \\
 &= \int d\mu_N(O) (O_{i_1 j_1})^{m_1} (O_{i_1 j_2})^{m_2} (O_{i_2 j_2})^{m_3} \\
 &= \frac{2^{2-N} \Gamma(\frac{1+m_1}{2}) \Gamma(\frac{1+m_2}{2}) \Gamma(\frac{1+m_3}{2})}{\pi \Gamma(\frac{N+m_1-1}{2}) \Gamma(\frac{N+m_3-1}{2})} \\
 &\quad \times \frac{\Gamma(N-1) \Gamma(\frac{N+m_1+m_3-1}{2})}{\Gamma(\frac{N+m_1+m_2+m_3}{2})}
 \end{aligned}$$

où $d\mu_N(O)$ est normalisé par $\int d\mu_N(O) = 1$, i_1, i_2, j_1, j_2 sont des indices arbitraires des éléments de matrices d'un élément quelconque O de $O(N)$, m_1, m_2, m_3 sont tous pairs et Γ est la fonction gamma de Euler. Ainsi la valeur moyenne de l'interférence dans HOE est donnée par

$$\begin{aligned}
 \langle \mathcal{I} \rangle_{O,N} &= N - N^2 Z_{O,N}(4, 0, 0) = N \left(1 - \frac{3}{N+2} \right) \\
 &\xrightarrow{N \rightarrow \infty} N - 3,
 \end{aligned} \tag{2.13}$$

Ainsi dans la limite asymptotique $N \rightarrow \infty$, un algorithme quantique réel de taille N tiré parmi HOE contient en moyenne légèrement moins d'interférence qu'un algorithme quantique de même taille tiré parmi CUE. Cependant la taille de l'espace de Hilbert doit être doublé pour pouvoir exprimer un algorithme complexe arbitraire sous une forme réelle [Aha03]. Cela signifie qu'un algorithme écrit sous une forme réelle nécessite approximativement le double d'interférence de sa version complexe.

A partir du second moment

$$\begin{aligned}
 \left\langle \left(\sum_{i,k} (O_{ik})^4 \right)^2 \right\rangle_{O,N} &= (N(N-1))^2 Z_{O,N}(4, 0, 4) \\
 &+ 2N^2(N-1)Z_{O,N}(4, 4, 0) \\
 &+ N^2 Z_{O,N}(8, 0, 0) \\
 &= \frac{3N(-4 + 3N(N+5))}{(N+1)(N+2)(N+6)},
 \end{aligned} \tag{2.14}$$

on obtient la variance de l'interférence pour HOE :

$$\sigma_{O,N}^2(\mathcal{I}) = \frac{24N(N-1)}{(N+2)^2(N^2+7N+6)}. \quad (2.15)$$

Ainsi, l'écart-type $\sigma_{O,N}(\mathcal{I})$ décroît comme $\sim 2\sqrt{6}/N$ pour N grand.

III. Distribution de l'interférence dans des circuits aléatoires

III.1. Ensembles de circuits aléatoires

Pour aller plus loin dans notre investigation sur l'interférence, il est légitime de considérer des algorithmes quantiques plus réalistes que ceux basés sur les ensembles de matrices aléatoires CUE et HOE. Au chapitre 1, nous avons déjà évoqué le fait que n'importe quel algorithme quantique, c'est à dire n'importe quelle transformation unitaire dans l'espace produit tensoriel, peut être approximé (avec une précision arbitraire) par une séquence de portes quantiques agissant au plus sur 2 qubits [DiV95, Bar95, SW95]. Plus précisément, un ensemble universel de portes quantique peut être construit en considérant une transformation $U(4)$ (telle la porte NON-contrôlée (CNOT) agissant sur deux qubits) en combinaison avec l'ensemble $U(2)$ de transformations à un qubit. Notons aussi que n'importe quel algorithme quantique peut aussi être représentée exclusivement par des matrices unitaires réelles, c'est à dire par des matrices orthogonales à condition de doubler la taille de l'espace de Hilbert, ceci en utilisant un ensemble universel constitué de la transformation de Hadamard et de la transformation de Toffoli [Shi02, Aha03]. Il est donc naturel de vouloir appliquer la mesure (1.17) sur des algorithmes quantiques basés sur cette construction à partir de séquences de portes. Bien sur, il est exclu de considérer des séquences concrètes résolvant des problèmes concrets, ceci pour les mêmes raisons précédemment évoquées, et la vision statistique est toujours celle qui nous préoccupe. De ce fait l'idée générale est de choisir des algorithmes quantiques construits comme des séquences *aléatoires* de portes quantiques. Nous avons donc introduit deux ensembles de circuits quantiques aléatoires, l'*ensemble de circuits unitaires* (unitary circuit ensemble ou UCE) et l'*ensemble de circuits orthogonaux* (orthogonal circuit ensemble), construits chacun à la manière des algorithmes quantiques réalistes par des séquences de n_g portes quantiques choisies aléatoirement suivant le protocole suivant.

- Le type de porte (c'est à dire si c'est une porte agissant sur 1, 2 ou 3 qubits) est choisi avec une probabilité p .
- Le(s) qubit(s) sur le(s)quelle agis(sent) la porte est(sont) choisi(s) uniformément parmi l'ensemble des qubits du registre.
- Le protocole est réitéré n_g fois.

L'ensemble UCE est basé sur l'ensemble universel constitué de la porte U_2 à 1 qubit (uniformément distribuée par rapport à la mesure de Haar de $U(2)$) et de la porte *CNOT*. L'ensemble OCE est basé sur l'ensemble universel constitué de la porte de Hadamard à 1 qubit et de la porte de Toffoli à 3 qubits (cf. figure (2.3) pour une instance de CUE constituée de cinq portes agissant sur trois qubits). Ces constructions permettent d'obtenir des ensembles algorithmes quantiques réalistes moins généraux que des matrices aléatoires car ils dépendent des 2 paramètres externes p et n_g . On peut donc examiner numériquement directement la distribution d'interférence $P(\mathcal{I})$ pour ces deux ensembles de circuits aléatoires.

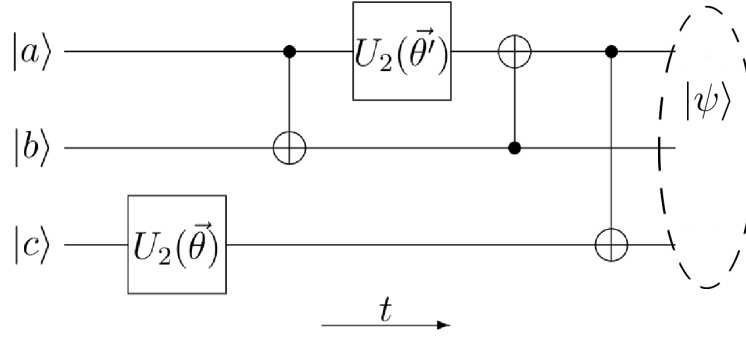


FIGURE 2.3.: Instance possible de CUE pour $n_q = 3$ et n_g . Les deux angles différents θ et θ' signifient que les deux portes aléatoires $U(2)$ sont différentes.

III.2. Résultats numériques

La figure (2.4) montre comment la distribution d'interférence pour UCE (pour $n_q = 4$) évolue entre $n_g = 10$ et $n_g = 100$ d'une distribution plutôt large et uniforme vers une distribution très piquée, qui se trouve être celle associée à CUE. De son côté, la distribution d'interférence $P_{OCE,N}(\mathcal{I})$ pour OCE fluctue beaucoup plus pour un petit nombre de portes, mais s'approche aussi rapidement de celle associée à HOE. A ce stade il est très intéressant de remarquer que quand n_g le nombre de portes dans la séquence est suffisamment grand, nos ensembles de circuits se comportent vis à vis de la distribution d'interférence de la même manière que les ensembles de matrices aléatoires précédemment étudiés. Ainsi les résultats analytiques développés dans la première partie se trouvent être totalement utilisables pour nos ensembles de circuits aléatoire. En fait un ensemble de circuits aléatoires similaires est introduit dans [EWS03], dans lequel une même porte quantique est continuellement itérée, porte qui est construite à partir d'un hamiltonien d'interaction entre proches voisins. D'après les résultats publiés dans [ELL05], on peut s'attendre, au moins dans le pire des cas, à une convergence exponentielle avec le nombre de portes, de CUE (HOE) vers UCE (OCE), respectivement. L'analyse précise de la vitesse de convergence sera l'objet du dernier chapitre de cette thèse. Ainsi il n'y a donc rien d'étonnant d'observer un comportement similaire pour l'interférence entre CUE et UCE (respectivement HOE et OCE). Pour examiner quantitativement l'équivalence en terme d'interférence et ainsi pouvoir appréhender dans quelle mesure les résultats analytiques pour les ensembles de matrices aléatoires sont applicables pour les ensembles de circuits aléatoires, nous utilisons la quantité

$$F_{\mathcal{I}} = \int_0^\infty \left(\sqrt{P_{UCE}(\mathcal{I})} - \sqrt{P_{CUE}(\mathcal{I})} \right)^2 ds = 2 \left(1 - \int_0^\infty \sqrt{P_{UCE}(\mathcal{I})P_{CUE}(\mathcal{I})} d\mathcal{I} \right) \quad (2.16)$$

et similairement pour OCE et HOE. Notons que les distributions $P_{CUE}(\mathcal{I})$ et $P_{HOE}(\mathcal{I})$ sont calculées numériquement, ceci pour une même dimension $N = 2^n$. La figure (2.5) montre le résultat pour $F_{\mathcal{I}}(n_g)$ dans le cas $n_q = 5$. Les courbes pour les autres valeurs de n_q examinées ($n_q \in \{4, 6, 7, 8\}$) sont très similaires si ce n'est que l'équivalence augmente avec n_q . $F_{\mathcal{I}}(n_g)$ aussi bien pour UCE que pour OCE est très bien fitté par une gaussienne, du moins jusqu'au point où un phénomène de saturation se manifeste. Notons que cette saturation n'est qu'un artefact numérique provenant du nombre fini de réalisations statistiques. On peut s'en convaincre en faisant varier le nombre de réalisations statistiques. Un fit de $\ln F_{\mathcal{I}}(n_g)$ par $a - c(n_q, p)n_g^2$ dans la région $2 \geq F_{\mathcal{I}}(n_g) \geq 0.01$ donne $c \sim 10^{-4}$ pour $4 \leq n_q \leq 6$, avec un maximum pour c autour de $p \sim 0.5$, ceci dans les deux cas OCE et UCE (voir figure (2.6)).

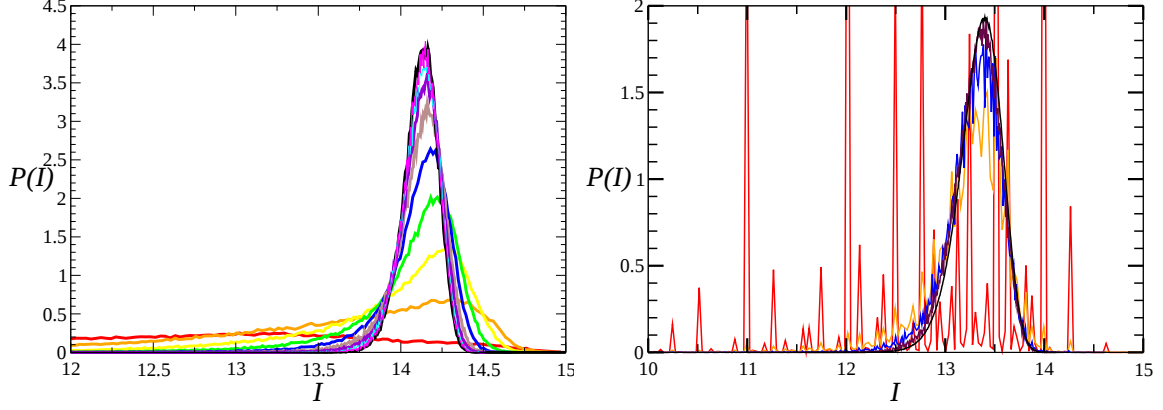


FIGURE 2.4.: Distribution de l'interférence $P_{UCE,N}(\mathcal{I})$ (gauche), pour $n_g = 10, 20, 30, 40, 50, 60, 70, 80, 90, 100$, les maxima augmentant dans cet ordre, et $P_{OCE,N}$ (droit) pour $n_g = 20, 50, 70, 100$ (rouge, orange, bleu, marron respectivement) comparées à $P(\mathcal{I})$ pour les ensembles de matrices aléatoires CUE et HOE (courbes noires). Toutes les courbes sont représentées pour $n_q = 4$ ($N = 16$), $n_r = 10^5$ circuits UCE et OCE, et $n_r = 10^7$ matrices CUE et HOE.

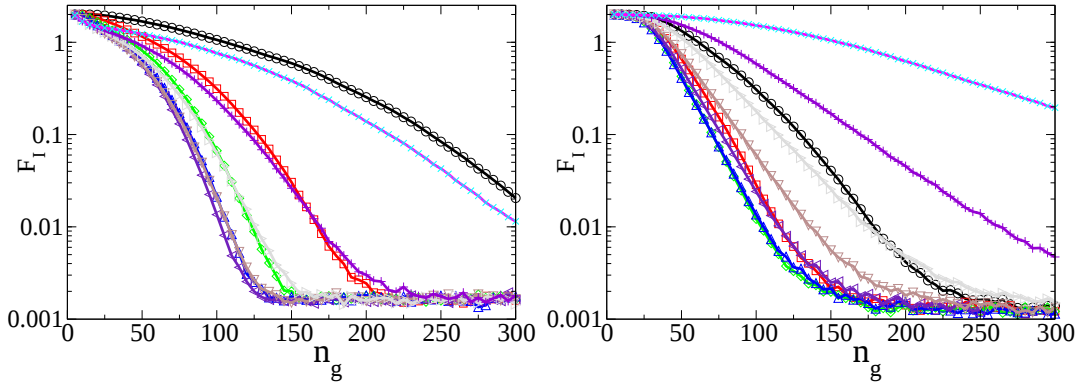


FIGURE 2.5.: Équivalence de $P(\mathcal{I})$ pour UCE (gauche) et OCE (droite) avec la distribution d'interférence de CUE et HOE, respectivement, en fonction de n_g pour $n_q = 5$ qubits ($N = 32$) et différentes valeurs de p . Les différentes valeurs de p sont indiqués suivant les symboles : $p = 0.1$ carrés noirs, $p = 0.2$ carrés rouges, $p = 0.3$ losanges verts, $p = 0.4$ triangle hauts bleus, $p = 0.5$ triangles indigo gauche, $p = 0.6$ triangles bas brun, $p = 0.7$ triangle gris droit, $p = 0.8$ violette et $p = 0.9$ magenta (gauche).

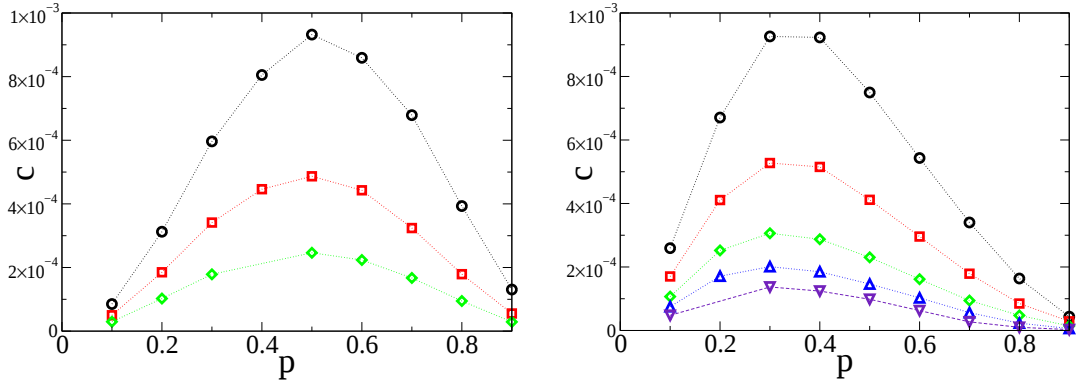


FIGURE 2.6.: Taux de convergence Gaussien c fitté en fonction de la probabilité p , pour différents nombres de qubits (cercles, carré, losange pour $n_q = 4, 5, 6$), respectivement, et de plus des triangles hauts et des triangles bas pour $n_q = 7, 8$ pour UCE.

La question qui reste ouverte concerne la rapidité de convergence des ensembles de circuits aléatoires vers les ensembles de matrices aléatoires. Ceci laisse entrevoir la possibilité de créer des ensembles imitant les ensembles de matrices aléatoires distribuées uniformément, à partir de circuit quantique. Dans nos résultats, une telle efficacité se manifesterait dans le nombre de portes quantiques nécessaires pour atteindre une certaine fidélité $F_{\mathcal{I}}$, qui augmenterait au plus polynomialement avec le nombre de qubits n_q . Ceci nécessiterait que les exposants $c(n_q, p)$ et $b(n_q, p)$ ne décroissent pas plus vite qu'une puissance négative de n_q pour un p donné. Ce genre de questions est posé dans [EWS03, ELL05] mais aucune réponse définitive n'est donnée. Pour vérifier ces affirmations numériquement il est nécessaire de considérer des valeurs bien plus grande de n_q et de sortir du cadre de l'étude de l'interférence (qui n'associe qu'un seul nombre $\mathcal{I}$ par opérateur unitaire). C'est ce qui nous préoccupera dans le dernier chapitre.

IV. Conclusion partielle

Dans ce chapitre, nous avons observé tout d'abord que si on choisit un algorithme quantique agissant sur un nombre de qubits assez grand, le choix se faisant au hasard, sans connaissances particulières sur sa structure interne ou sur le rôle qu'il est sensé rendre, alors il est fort probable que cet algorithme contiennent une quantité d'interférence, mesuré par $\mathcal{I}(U)$ (2.1), qui est très proche de la valeur maximale atteinte par des transformations du type Hadamard. Il est intéressant de remarquer que ce résultat est un analogue du résultat présenté dans [HWW06] où est étudié la quantité d'intrication dans un système bipartite de grande dimension.

Ensuite, si on considère des modèles d'algorithmes quantiques plus réalistes, basés sur des circuits quantiques construits comme des séquences de portes quantiques tirées aléatoirement dans un ensemble universel de portes, alors le résultat précédent concernant l'interférence se trouve totalement applicable dès que le nombre de portes n_g est suffisamment grand.

Finalement, il est suggéré que la construction de circuits aléatoires à partir de séquences aléatoires de portes pourrait constituer un moyen efficace de créer des ensembles de matrices aléatoires distribuées selon CUE. L'étude de l'efficacité de cette construction sera l'objet du dernier chapitre.

Statistique de l'interférence en présence de décohérence

I. Introduction

Dans le chapitre précédent, nous avons démontré qu'un algorithme quantique, de taille suffisamment grande et choisi au hasard, avait de grandes chances de contenir une quantité d'interférence quantique, mesurée à l'aide de (1.17), proche de la valeur maximale admise. Ces résultats ne s'appliquant qu'au cas précis d'une évolution purement unitaire, il est peut-être intéressant de regarder l'effet de la décohérence et de la perte de l'unitarité associée sur cette quantité d'interférence. Pour ce faire, nous allons dans ce chapitre considérer le cas d'un système quantique couplé à un environnement thermique, l'interférence étant mesurée avec (1.16). Comme dans le précédent chapitre, des calculs numériques aussi bien qu'analytiques nous permettront d'obtenir des résultats sur le comportement de l'interférence dans ce type de systèmes.

II. Statistique de l'interférence dans un système quantique couplé à un spin

Dans cette partie nous allons considérer la propagation d'un système quantique de taille fini interagissant avec un environnement lui même consistant en un autre système quantique de taille fini. Le propagateur d'un tel système est une application complètement positive de la matrice densité initiale représentant le système vers sa matrice densité finale [NC00]. Bien qu'un environnement de dimension finie ne constitue pas un vrai bain thermique dans le sens où il n'induit pas un comportement irréversible, l'étude du cas simplifié de dimension finie est motivée en théorie de l'information quantique où on rencontre souvent des qubits auxiliaires qui sont ajoutés à un registre quantique principal. L'autre avantage de ce type de modèle est la possibilité de moyenner sur les degrés de liberté associés à l'environnement, ce qui mène dès l'apparition de corrélations ou d'intrication entre système et environnement, à un mécanisme de décohérence. Ainsi ce type de système est parfaitement adapté à l'étude de l'influence de la décohérence sur l'interférence quantique. Pour ce type de systèmes, le choix de l'état initial de l'environnement est la seule liberté possible et il convient pour notre étude de choisir cet état initial comme un état thermique, état pouvant être vu comme le résultat de l'interaction de l'environnement de taille finie avec son propre bain thermique. En choisissant l'environnement

comme un seul spin initialement à l'équilibre thermique, nous allons voir que l'on peut obtenir l'expression de l'interférence selon la mesure (1.16) et que cette expression permet d'obtenir des informations statistiques dans le cas où le système et l'environnement suivent de manière jointe une évolution unitaire .

II.1. Propagateur pour un application complètement positive

Soit un système bipartite séparé en un sous-système $\mathcal{S}$ (d'espace de Hilbert $\mathcal{H}_S$ de dimension n) et un environnement $\mathcal{E}$ (d'espace de Hilbert $\mathcal{H}_E$ de dimension m). Notons respectivement W et W' les matrices densité initiale et finale de ce système global. Nous choisissons un état initial séparable $W = \sigma \otimes \epsilon$ des matrices densité initiale σ et ϵ du système et de l'environnement, respectivement. L'évolution temporelle du système global $\mathcal{S} + \mathcal{E}$ dans l'espace de Hilbert produit tensoriel $\mathcal{H}_S \otimes \mathcal{H}_E$ de dimension $N = n \times m$, peut être considérée comme purement unitaire, du moment où le système global est assimilable à un système fermé sur l'échelle de temps qui nous intéresse. Cette évolution est entièrement caractérisée par une matrice unitaire U d'après l'équation d'évolution de la matrice densité $W' = UWU^\dagger$. En composante cette équation s'écrit

$$W'_{\alpha_1\alpha_2,\beta_1\beta_2} = \sum_{\gamma_1,\gamma_2,\delta_1,\delta_2}^{n,m} U_{\alpha_1\alpha_2,\gamma_1\gamma_2} W_{\gamma_1\gamma_2,\delta_1\delta_2} U_{\beta_1\beta_2,\delta_1\delta_2}^*$$

où les indices 1 et 2 renvoient respectivement aux états de base du système et de l'environnement. L'état final de la matrice densité réduite du système est obtenue en prenant la trace partielle sur les indices de l'environnement, $\sigma' = \text{Tr}_{\mathcal{E}} W'$, ou de manière explicite, $\sigma'_{\alpha_1\beta_1} = \sum_{\alpha_2}^m W'_{\alpha_1\alpha_2,\beta_1\alpha_2}$. De (3.1) et de la matrice densité initiale pour le système global $W_{\gamma_1\gamma_2} = \rho_{\gamma_1\delta_1} \epsilon_{\gamma_2\delta_2}$ on obtient le propagateur du système $\mathcal{S}$ uniquement,

$$\sigma'_{\alpha_1\beta_1} = \sum_{\gamma_1,\delta_1}^n P_{\alpha_1\beta_1,\gamma_1\delta_1} \sigma_{\gamma_1\delta_1}$$

où les composantes du propagateur sont données par

$$P_{\alpha\beta,\gamma\delta} = \sum_{\mu,\nu,\rho}^m U_{\alpha\mu,\gamma\nu} \epsilon_{\nu\rho} U_{\beta\mu,\delta\rho}^* \quad (3.1)$$

Ce propagateur P est un super-opérateur qui envoie l'opérateur densité initial σ sur l'opérateur densité final σ' . Cette procédure garantit que le propagateur est bien une application complètement positive et traduit bien une évolution physique consistante [NC00]. Comme attendu, P dépend non seulement de U mais aussi de l'état initial de l'environnement ϵ . Pour obtenir des résultats explicites, nous devons considérer des cas particuliers d'environnement. Nous commencerons avec un seul spin, initialement à l'équilibre thermique puis généraliserons au cas à plusieurs spins, également choisis comme initialement à l'équilibre thermique.

II.2. Interférence dans un système quantique couplé à un seul spin

Nous considérons la situation où l'environnement est un seul spin de taille $(d-1)/2$. L'espace de Hilbert de cet environnement est donc de taille $m = d$. Nous nous plaçons dans le cas où les niveaux d'énergie du spin sont équidistants, chaque niveau séparé de ses voisins par une énergie $\hbar\Omega$, comme c'est le cas pour des atomes ou des noyaux placés dans un champ magnétique externe

II. STATISTIQUE DE L'INTERFÉRENCE DANS UN SYSTÈME QUANTIQUE COUPLÉ À UN SPIN

sous effet Zeeman linéaire. Dans sa base propre, les éléments de matrices de l'hamiltonien de spin $H^{(1)}$ s'écrivent

$$H_{\nu\rho}^{(1)} = \hbar\Omega \nu \delta_{\nu\rho} \quad (3.2)$$

où $1 \leq \nu \leq d$ et $\nu - 1$ est le nombre d'excitations du spin. Comme précédemment évoqué, nous choisissons le spin initialement à l'équilibre thermique à la température $T = \frac{1}{k_B\beta}$, tel que sa matrice densité s'écrit comme

$$\epsilon = \frac{e^{-\beta H}}{\text{Tr}(e^{-\beta H})} \rightarrow \epsilon_{\nu\rho} = \frac{1}{Z} e^{-\beta \hbar\Omega \nu} \delta_{\nu\rho} \quad (3.3)$$

avec la fonction de partition définie par

$$Z \equiv Z(x) = \sum_{\nu}^d e^{-\beta \hbar\Omega \nu} = \frac{1 - e^{-dx}}{e^x - 1}, \quad (3.4)$$

et $x = \beta \hbar\Omega$. Le propagateur P se simplifie en,

$$P_{\alpha\beta,\gamma\delta} = \frac{1}{Z} \sum_{\mu,\nu}^d U_{\alpha\mu,\gamma\nu} U_{\beta\mu,\delta\nu}^* e^{-x\nu}. \quad (3.5)$$

En injectant (3.5) dans (1.16), nous obtenons finalement l'expression pour l'interférence contenue dans la propagation du système $\mathcal{S}$,

$$\begin{aligned} \mathcal{I} &= \sum_{\alpha,\gamma \neq \delta}^n |P_{\alpha\alpha,\gamma\delta}|^2 = \frac{1}{Z^2} \sum_{\alpha,\gamma \neq \delta}^n \left| \sum_{\mu,\nu}^d U_{\alpha\mu,\gamma\nu} U_{\alpha\mu,\delta\nu}^* e^{-x\nu} \right|^2 \\ &= \frac{1}{Z^2} \sum_{\alpha,\gamma \neq \delta}^n \sum_{\mu,\nu,\rho,\sigma}^d e^{-x(\nu+\sigma)} U_{\alpha\mu,\gamma\nu} U_{\alpha\mu,\delta\nu}^* U_{\alpha\rho,\gamma\sigma}^* U_{\alpha\rho,\delta\sigma}. \end{aligned}$$

Nous sommes maintenant en mesure d'étudier les propriétés statistiques de $\mathcal{I}$ basée sur la statistique de U . Dans le même ordre d'idée qu'au chapitre 2, sans connaissance particulière sur l'évolution temporelle du système global (ou sur l'ensemble d'algorithmes spécifiques à considérer), il est naturel de choisir U dans CUE, c'est à dire uniformément distribué par rapport à la mesure de Haar dU du groupe unitaire $U(N)$. Ainsi dans le cas limite où la taille de l'environnement m tend vers un, nous nous attendons à retrouver les mêmes résultats que ceux obtenus dans le chapitre 2 pour la statistique de l'interférence lors d'une évolution purement unitaire.

II.3. Résultats numériques

Pour de petites dimensions n et m , on peut obtenir la distribution de l'interférence $P(\mathcal{I})$ numériquement. Pour ce faire nous avons produit de grands ensembles de matrices unitaires aléatoires de taille $N = n \times m$ distribuées selon CUE, en utilisant la paramétrisation de Hurwitz [Hur97, PZK98]. Pour obtenir une statistique convenable, chaque ensembles contient 10^6 matrices pour le calcul des distributions. La figure (3.1) représente $P(\mathcal{I})$ pour des systèmes avec des tailles de $n = 2$ à 4, couplés à des environnements de taille $m = 1$ à 4, ceci à une température inverse choisie arbitrairement à $x = 0.1$. Dans le cas $n = 2$, où le calcul analytique est possible pour $m = 1$ où $P(\mathcal{I}) = \frac{1}{2\sqrt{1-\mathcal{I}}}$ (cf. chapitre 1), les distributions sont très

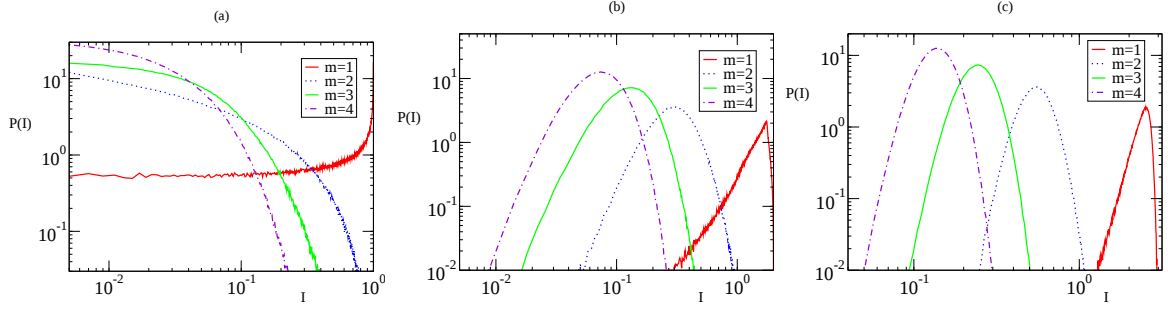


FIGURE 3.1.: Distributions de l'interférence représentées en échelle log-log, calculées numériquement à $x = 0.1$ pour $n = 2, 3, 4$ (graphes (a), (b) et (c) respectivement). Dans chaque graphe m varie entre 1 et 4. Dans chaque cas le nombre de réalisations statistiques vaut $n_r = 10^6$.

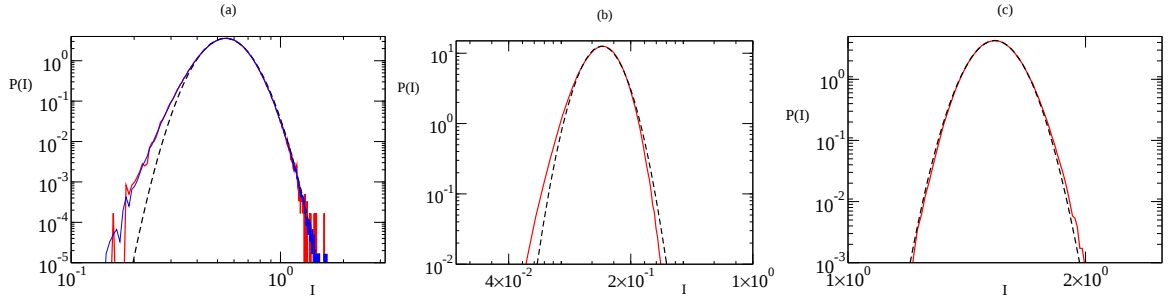


FIGURE 3.2.: Fit des distributions $P(\mathcal{I})$ calculées numériquement (courbes rouges) avec une distribution log-normal (courbes noires pointillées) à $x = 0.1$ pour $(n, m) = (4, 2)$, $(4, 4)$ et $(8, 2)$ (graphes (a), (b) et (c) respectivement). Tous les fits sont réalisés pour les distributions numériques calculées avec $n_r = 10^6$ excepté pour le premier graphe (a) où le fit est réalisé par rapport à la distribution représentée par la courbe bleue pour laquelle $n_r = 10^7$.

larges. Pour des valeurs plus grandes de n , les distributions deviennent de plus en plus piquées ainsi que de plus en plus symétriques par rapport à leur maximum, ceci en échelle log-log. Les queues des distributions décroissent rapidement dans les cas non-unitaires ($m \neq 1$). Dans chaque cas, nous observons que la valeur la plus probable pour $\mathcal{I}$ ainsi que la largeur du pic décroissent avec m . Comme on s'y attend intuitivement, la décohérence due au couplage avec l'environnement détruit l'interférence d'autant plus efficacement que l'environnement est de grande taille. Pour m fixé, la distribution se comporte qualitativement comme au chapitre 2 dans le cas d'une propagation purement unitaire, c'est à dire avec une valeur maximale et une moyenne de l'interférence augmentant avec n et une largeur de la distribution diminuant avec ce même paramètre. Un changement dans la température se traduit essentiellement par un décalage de la distribution. Ceci est dû au comportement de la valeur moyenne de l'interférence avec la température comme nous le verrons dans la prochaine partie (cf. eq.(3.16)) et justifie le fait d'avoir représenté les distributions pour la seule valeur arbitraire $x = 0.1$. La figure (3.2) montre que $P(\mathcal{I})$ est bien fittée par une distribution log-normal,

$$P(\mathcal{I}) = \frac{1}{\mathcal{I}\sqrt{2\pi}\sigma} \exp\left(-\frac{(\log(\mathcal{I}) - \mu)^2}{2\sigma^2}\right). \quad (3.6)$$

en particulier proche du maximum, alors que des déviations apparaissent sur les bords de la distribution. Notons que les bords de la distribution apparaissent coupés sur les données numériques mais que ceci n'est qu'un effet de taille fini du nombre de réalisations utilisées pour les calculs. Cet effet numérique est visible par exemple pour $(n, m) = (4, 2)$ où nous avons poussé le nombre de réalisations, initialement fixé à 10^6 , jusqu'à 10^7 . Dans ce dernier cas, on observe bel et bien que les coupures apparaissent pour des valeurs de $\mathcal{I}$ plus éloignées de la valeur moyenne.

Les distributions obtenues numériquement suggèrent que pour $n > 2$, $P(\mathcal{I})$ est bien caractérisée par ces deux premiers moments. C'est pour cela que la suite de ce chapitre est consacré aux résultats analytiques concernant ces deux premiers moments, permettant de confirmer de manière quantitative les observations numériques faites dans cette partie.

II.4. Résultats analytiques

2.4.a. Interférence moyenne

Pour notre système, l'interférence moyenne $\langle \mathcal{I} \rangle$ s'obtient à partir de (3.6),

$$\langle \mathcal{I} \rangle = \frac{1}{Z^2} \sum_{\alpha, \gamma \neq \delta}^n \sum_{\mu, \nu, \rho, \sigma}^d e^{-x(\nu+\sigma)} \langle U_{\alpha\mu, \gamma\nu} U_{\alpha\mu, \delta\nu}^* U_{\alpha\rho, \gamma\sigma}^* U_{\alpha\rho, \delta\sigma} \rangle \quad (3.7)$$

où $\langle . \rangle \equiv \int dU(.)$ correspond à la moyenne sur le groupe unitaire $U(N)$. Pour des monômes composés d'un nombre relativement faible de facteurs $U_{\alpha\mu, \gamma\nu}$, leurs moyennes peuvent être obtenues par la technique d'intégration invariante introduite dans [AL03, AL04, Bra06], technique qui permet d'obtenir certaines formules utilisées dans le précédent chapitre. Cette technique permet une écriture et une manipulation de ces valeurs moyennes de monômes en termes diagrammatiques (cf. Appendice A). Ainsi en terme de diagramme l'équation précédente s'écrit

$$\mathcal{I} = \frac{1}{Z^2} \sum_{\alpha, \gamma \neq \delta}^n \sum_{\mu, \nu, \rho, \sigma}^d e^{-x(\nu+\sigma)} \begin{array}{c} \alpha\rho \\ \alpha\mu \end{array} \begin{array}{c} \text{---} \\ \text{---} \end{array} \begin{array}{c} \delta\sigma \\ \gamma\sigma \\ \delta\nu \\ \gamma\nu \end{array} \quad (3.8)$$

Nous renvoyons à l'annexe A pour une explication et une dérivation détaillée de l'écriture en terme de diagrammes. Nous résumons cependant ici ses points principaux : par soucis de simplicité, notons les composantes comme il se fait habituellement avec des indices romains (i, j , etc). Chaque indice de ligne (colonne) apparaissant dans les éléments de matrices du monôme moyenné est représenté par des points sur la droite (gauche). Un facteur U_{ij} est représenté par une ligne fine pointillée reliant les points i et j tandis qu'un facteur complexe conjugué U_{kl}^* est représenté par une ligne fine continue entre les points k et l . Quand un élément de matrice apparaît avec une multiplicité t , une seule ligne est dessinée avec la valeur t indiquée à proximité. Un facteur réel du style $|U_{ij}|^2$ est représenté par une ligne continue épaisse, pouvant aussi avoir une certaine multiplicité. Dans [AL03] (cf. Appendice A), il est montré que l'invariance de la mesure de Haar dU sous des transformations unitaires arbitraires, engendre deux propriétés majeures pour ces moyennes de monômes. En terme diagrammatique, ces propriétés expriment que :

- (a) La valeur d'un diagramme ne dépend pas de la valeur spécifique de ces points mais seulement de sa forme, c'est à dire de la manière dont les points sont connectés entre eux. Ceci signifie qu'un diagramme peut être dessiné sans spécifier explicitement la valeur de

ces points. Par exemple,

$$\langle U_{11}U_{11}^*U_{12}U_{12}^* \rangle = \langle U_{24}U_{24}^*U_{26}U_{26}^* \rangle = \text{diagramme} \quad (3.9)$$

- (b) Si pour au moins un point dans le diagramme, le nombre de lignes fines continues connectées à ce point diffère du nombre de lignes pointillées, alors la valeur du diagramme est zéro. Par exemple,

$$\langle U_{11}U_{12}^*U_{23}^*U_{24} \rangle = \text{diagramme} = 0 \quad (3.10)$$

$$\langle U_{11}U_{12}^*U_{21}U_{22}^* \rangle = \text{diagramme} \neq 0. \quad (3.11)$$

Dans l'équation (3.8), les sommes sur tous les indices de lignes et de colonnes font apparaître différentes formes de diagrammes, dépendant de la manière dont certains points coïncident. Cependant de nombreuses configurations sommées correspondent à des diagrammes de valeur nulle (d'après la propriété b). Les seules configurations qui contribuent sont celles dont les points $(\gamma\nu)$ et $(\delta\nu)$ coïncident respectivement sur les points $(\gamma\sigma)$ et $(\delta\sigma)$, c'est à dire les configurations telles que $\nu = \sigma$. De ce fait l'expression pour $\langle \mathcal{I} \rangle$ se réduit à

$$\langle \mathcal{I} \rangle = \frac{1}{Z^2} \sum_{\alpha, \gamma \neq \delta}^n \sum_{\mu, \nu, \rho}^d e^{-2x\nu} \begin{matrix} \alpha\rho \\ \alpha\mu \end{matrix} \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} \quad (3.12)$$

$$= \frac{1}{Z^2} \sum_{\alpha, \gamma \neq \delta}^n \left(\sum_{\nu}^d e^{-2x\nu} \right) \left(\sum_{\mu=\rho}^d \alpha\mu \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} + \sum_{\mu \neq \rho}^d \alpha\rho \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} \right). \quad (3.13)$$

Arrivé à ce stade, seulement deux types de diagrammes se révèlent nécessaires, et puisque d'après la propriété (a) leurs valeurs ne dépendent pas des indices de sommations, on a

$$\langle \mathcal{I} \rangle = \frac{Z(2x)}{Z^2(x)} n^2(n-1) \left(d \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} + d(d-1) \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} \right). \quad (3.14)$$

Les valeurs pour ces deux diagrammes sont calculées dans [AL03],

$$\begin{aligned} \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} &= \frac{1}{N(N+1)} \\ \begin{matrix} \delta\nu \\ \gamma\nu \end{matrix} &= \frac{1}{N(N^2-1)}. \end{aligned}$$

Le préfacteur peut être réécrit comme

$$h(x) \equiv \frac{Z(2x)}{Z^2(x)} = \coth(dx/2) \tanh(x/2), \quad (3.15)$$

et finalement on obtient en fixant $d = m$,

$$\langle \mathcal{I}(n, m, x) \rangle = \coth\left(\frac{mx}{2}\right) \tanh\left(\frac{x}{2}\right) \frac{nm(n-1)^2}{(n^2m^2-1)}. \quad (3.16)$$

II. STATISTIQUE DE L'INTERFÉRENCE DANS UN SYSTÈME QUANTIQUE COUPLÉ À UN SPIN

Voici un des résultats analytiques centraux de cette thèse que nous allons discuter en détail.

Premièrement, on observe que la dépendance en température est entièrement contenue dans le préfacteur $h(x)$. Ces limites pour $x \rightarrow 0$ et $x \rightarrow \infty$ sont respectivement $1/m$ et 1 et on peut vérifier que $dh(x)/dx$ est toujours positif sur $[0, \infty[$, signifiant que le préfacteur augmente de manière monotone sur cet intervalle. Ainsi pour une taille de système donnée, l'interférence moyenne décroît pour une augmentation de la température de l'état initial de l'environnement, comme on s'y attend intuitivement et comme on l'observe numériquement. Seule la dimension de l'environnement $m = d$ intervient dans la dépendance en température. En fait, ceci est vrai pour tous les moments de $P(\mathcal{I})$, cette dépendance en température étant contenue dans le facteur $\exp(-x\nu)$ qui est toujours sommé sur $\nu = 1, \dots, m$.

Dans le cas particulier $m = 1$ (se qui entraîne $n = N$), nous retrouvons comme attendu l'expression (2.9) pour une propagation purement unitaire :

$$\langle \mathcal{I}(n, 1, x) \rangle = \frac{N(N-1)^2}{N^2-1} = \frac{N(N-1)}{N+1} = \langle \mathcal{I}_U(N) \rangle. \quad (3.17)$$

Dans ce cas, aucune intrication ou corrélations avec l'environnement ne peuvent intervenir, l'unique état de l'environnement se factorisant, de manière à ce que la dynamique du système $\mathcal{S}$ reste entièrement unitaire.

Contrairement à ce qu'on pourrait attendre naïvement, le résultat du cas unitaire n'est pas retrouvé à température nulle, $x \rightarrow \infty$. A la place nous trouvons

$$\lim_{x \rightarrow \infty} \langle \mathcal{I}(n, m, x) \rangle = \frac{N(n-1)^2}{N^2-1}, \quad (3.18)$$

et où nous rappelons que $N = n \times m$. Pour $n \gg 1$ et m fixes, nous avons comme comportement asymptotique

$$\langle \mathcal{I}_U(n) \rangle = n - 2 + \mathcal{O}\left(\frac{1}{n}\right) \quad (3.19)$$

$$\langle \mathcal{I}(n, m, x \rightarrow \infty) \rangle = \frac{n-2}{m} + \mathcal{O}\left(\frac{1}{n}\right) \simeq \frac{\langle \mathcal{I}_U(n) \rangle}{m} \quad (3.20)$$

Ainsi pour $n \gg 1$, l'interférence moyenne se comporte comme dans le cas unitaire, linéairement avec le taille du système, mais diminuée approximativement par un facteur m . La raison de cette réduction est due au fait que même si l'environnement est dans son état fondamental (car $T=0$), la dynamique unitaire commune engendre de l'intrication entre les deux sous-systèmes $\mathcal{S}$ et $\mathcal{E}$, de telle façon qu'après l'opération de trace partielle sur l'environnement, il en résulte l'évolution non-unitaire de $\mathcal{S}$. Par conséquent la perte de cohérence se manifeste elle-même par la réduction de l'interférence moyenne. Dans la limite opposée de température infinie, $x \rightarrow 0$, le préfacteur $h(x)$ engendre une réduction supplémentaire par un autre facteur m ,

$$\lim_{x \rightarrow 0} \langle \mathcal{I}(n, m, x) \rangle = \frac{n(n-1)^2}{N^2-1}. \quad (3.21)$$

La réduction additionnelle est aussi visible dans le développement asymptotique pour $n \gg 1$, qui donne dans ce cas

$$\langle \mathcal{I}(n, m, x = 0) \rangle = \frac{n-2}{m^2} + \mathcal{O}\left(\frac{1}{n}\right) \simeq \frac{\langle \mathcal{I}_U(n) \rangle}{m^2}. \quad (3.22)$$

Pour $m \gg 1$ et n fixés on obtient

$$\langle \mathcal{I}(n, m, x \rightarrow \infty) \rangle = \frac{(n-1)^2}{nm} + \mathcal{O}\left(\frac{1}{m^3}\right) \quad (3.23)$$

$$\langle \mathcal{I}(n, m, x = 0) \rangle = \frac{(n-1)^2}{nm^2} + \mathcal{O}\left(\frac{1}{m^3}\right). \quad (3.24)$$

Les relations (3.23) et (3.24) illustrent que pour $n > 1$ fixé, $\langle \mathcal{I} \rangle$ décroît comme $1/m$ ($1/m^2$) pour une température nulle (température infinie). Sur la figure (3.3) est représentée $\langle \mathcal{I}(n, m, x) \rangle$ pour différentes températures en fonction de n et m . On observe que pour une température donnée, $\langle \mathcal{I}(n, m, x) \rangle$ croît avec n , mais décroît avec m . Pour n grand, avec m et x fixés, l'augmentation est essentiellement proportionnelle à n , comme dans le cas unitaire, avec un taux réduit par un facteur $h(x)/m$. Pour m grand, avec n et x fixés, la décroissance de $\langle \mathcal{I}(n, m, x) \rangle$ se comporte grossièrement comme $1/m$ avec un préfacteur $\frac{(e^x-1)^2}{e^{2x}-1} \frac{n(n-1)^2}{n^2}$. Plus généralement, une augmentation dans la dimension de l'environnement fait diminuer l'interférence moyenne comme une loi de puissance. Sur la figure (3.3) cette loi de puissance évolue entre m^{-2} pour $x = 0.001$ jusqu'à m^{-1} pour $x=10$ et n fixé. Cependant, on ne doit pas conclure, qu'un système système quantique couplé à un bain thermique de dimension infinie ne manifeste pas d'effet d'interférence quantique. En effet, il faut bien garder à l'esprit que dans notre modèle, le couplage entre le système et son environnement est de manière générique un couplage fort, dans le sens où l'évolution unitaire globale représenté par l'opérateur U , ne fait aucunement la distinction entre les deux sous-systèmes $\mathcal{S}$ et $\mathcal{E}$. En d'autre terme on ne fait pas de découpage entre des hamiltoniens bien distincts pour le système, pour l'environnement et pour le couplage. Il est donc naturel qu'un tel couplage fort détruise la cohérence et donc l'interférence alors que la situation peut se révéler différente dans le cas d'un couplage faible.

2.4.b. Second moment de la distribution d'interférence

Pour apprécier la largeur de la distribution d'interférence en fonction de m, n et x , nous pouvons calculer le second moment de $P(\mathcal{I})$. En prenant le carré de l'expression (3.6) il suit

$$\mathcal{I}^2 = \frac{1}{Z^4} \sum_{\alpha, \gamma \neq \delta}^n \sum_{a, g \neq d}^n \sum_{\mu, \rho, p, r}^d \sum_{\nu, \sigma, q, s}^d e^{-x(\nu+\sigma+n+s)} U_{\alpha\mu, \gamma\nu} U_{\alpha\mu, \delta\nu}^* U_{\alpha\rho, \gamma\sigma}^* U_{\alpha\rho, \delta\sigma} U_{ap, gq}^* U_{ap, dq} U_{ar, gs} U_{ar, ds}^*$$

Cette expression peut être traitée de la même manière que l'expression obtenue précédemment pour l'interférence moyenne, en moyennant sur le groupe $U(N)$ et en utilisant l'écriture diagrammatique. Le fait que huit facteurs U apparaissent, rend le calcul analytique de $\langle \mathcal{I}^2 \rangle$ plus difficile. Comme nous allons le voir, 19 diagrammes sont nécessaires pour mener le calcul jusqu'au bout. Par soucis de clarté, nous ne développerons dans cette partie, que la ligne principale pour le réaliser, les détails précis étant rassemblés dans les annexes B et C.

La première étape est d'introduire de nouvelles notations et conventions d'écriture pour simplifier et raccourcir l'allure des expressions mathématiques :

- 1 Pour les six indices relatifs à $\mathcal{S}$, on remplace $(\alpha, \gamma, \delta, a, g, d)$ par $(\alpha_1, \alpha_2, \alpha_3, \alpha_4, \alpha_5, \alpha_6)$.
- 2 Nous faisons une chose similaire pour les huit indices relatifs à $\mathcal{E}$ en remplaçant $(\mu, \rho, p, r, \nu, \sigma, q, s)$ par $(\mu_1, \mu_2, \mu_3, \mu_4, \mu_5, \mu_6, \mu_7, \mu_8)$.
- 3 En ce qui concerne les éléments de matrices, nous abandonnons les lettres redondantes α et μ qui n'apportent aucune information. De ce fait, un élément de matrice $U_{\alpha_i \mu_j, \alpha_k \mu_l}$ s'écrit maintenant $U_{ij, kl}$ et $U_{11, 11}$ ne doit alors plus être considéré comme le premier élément de

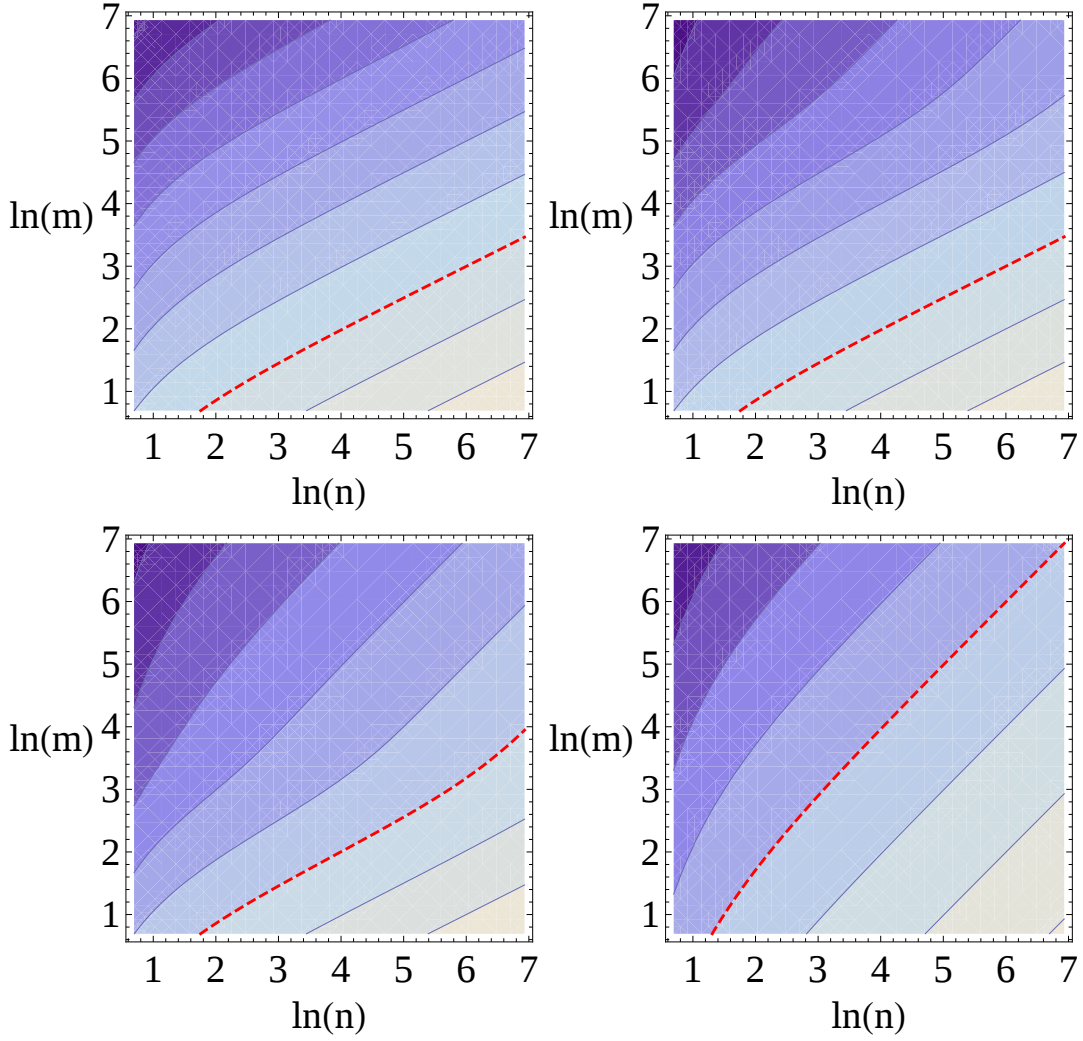


FIGURE 3.3.: Graphe de contour de $\ln(\langle \mathcal{I}(n, m) \rangle)$ pour $x=0.001, 0.01, 0.1$ et 10 (d'en haut à gauche à en bas à droite), pour n et m entre 2 et 1024 . La distance entre les contours vaut 2 et la ligne rouge pointillée indique que le contour $\ln(\langle \mathcal{I}(n, m) \rangle) = 0$. Les valeurs augmentent avec les tons de couleurs allant du sombre au clair.

la matrice U mais comme l'élément de composantes $(\alpha_1 \mu_1, \alpha_1 \mu_1)$. On rappelle que tous les indices α (μ) prennent leur valeur entre 1 et n (m), respectivement.

- 4 Les contraintes $\gamma \neq \delta$ et $g \neq d$ qui s'écrivent maintenant $\alpha_2 \neq \alpha_3$ et $\alpha_5 \neq \alpha_6$ sont implicites et ne sont donc pas notées *a priori*.
- 5 Pour les sommes on introduit la notation $\{\alpha_i, \mu_j\}$ pour dénoter l'ensemble de tous les indices qui apparaissent explicitement dans le sommant comme indices d'éléments de matrice (ou de manière équivalente comme points sur un diagramme), à l'exception des indices qui apparaissent explicitement sous un autre signe somme dans la même expression. Par exemple dans le cas $\sum_{\{\alpha_i, \mu_j\}} \sum_{\mu_5, \mu_7}$, la première somme est sur tous les indices α et μ qui apparaissent dans le sommant, à l'exception des indices μ_5 et μ_7 , considérés séparément.

Avec ces notations, $\langle \mathcal{I}^2 \rangle$ s'écrit

$$\langle \mathcal{I}^2 \rangle = \frac{1}{Z^4} \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} e^{-x(\mu_5 + \mu_6 + \mu_7 + \mu_8)} \langle U_{11,35} U_{11,45}^* U_{12,36}^* U_{12,46} U_{23,57}^* U_{23,67} U_{24,58} U_{24,68}^* \rangle \quad (3.25)$$

$$= \frac{1}{Z^4} \sum_{\{\alpha_i, \mu_j\}}^{(n,d)} e^{-x(\mu_5 + \mu_6 + \mu_7 + \mu_8)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,8 \\ 5,8 \\ 6,7 \\ 5,7 \\ 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{array} \quad (3.26)$$

Insistons à nouveau sur le fait que les indices de U qui apparaissent dans l'expression (B.1) sont des indices d'indices, c'est à dire que par exemple $U_{11,35} \equiv U_{\alpha_1 \mu_1, \alpha_3 \mu_5}$. Comme pour l'équation (3.7), les seules configurations non-nulles se manifestent pour des diagrammes sans branches ouvertes. Elles correspondent aux trois contraintes distinctes suivant pour les indices de sommations : $\mu_5 = \mu_6$ et $\mu_7 = \mu_8$, ou $\alpha_3 = \alpha_5, \alpha_4 = \alpha_6, \mu_5 = \mu_7$, et $\mu_6 = \mu_8$, ou $\alpha_3 = \alpha_6, \alpha_4 = \alpha_5, \mu_5 = \mu_8$, et $\mu_6 = \mu_7$. Ces contraintes donnent les trois sommes suivantes,

$$\begin{aligned} \langle \mathcal{I}^2 \rangle &= \frac{1}{Z^4} \left(\sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\mu_5, \mu_7}^m e^{-2x(\mu_5 + \mu_7)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,7 \\ 5,7 \\ 4,5 \\ 3,5 \end{array} + \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\alpha_3, \alpha_4}^n \sum_{\mu_5 \neq \mu_6}^m e^{-2x(\mu_5 + \mu_6)} \begin{array}{c} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{array} \right. \\ &\quad \left. + \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\alpha_3, \alpha_4}^n \sum_{\mu_5 \neq \mu_6}^m e^{-2x(\mu_5 + \mu_6)} \begin{array}{c} 2,3 \\ 1,2 \\ 2,4 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{array} \right). \quad (3.27) \end{aligned}$$

Les deux derniers termes sont égaux comme on peut le voir en échangeant les indices $\alpha_3 \leftrightarrow \alpha_4$, ce qui simplifie l'expression en

$$\begin{aligned} \langle \mathcal{I}^2 \rangle &= \frac{1}{Z^4} \left(\sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\mu_5 \neq \mu_7}^m e^{-2x(\mu_5 + \mu_7)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,7 \\ 5,7 \\ 4,5 \\ 3,5 \end{array} + 2 \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\alpha_3 \neq \alpha_4}^n \sum_{\mu_5 \neq \mu_6}^m e^{-2x(\mu_5 + \mu_6)} \begin{array}{c} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{array} \right) \\ &\equiv \frac{1}{Z^4} (A + 2B). \quad (3.28) \end{aligned}$$

Les termes A et B dépendent de 19 diagrammes différents qui peuvent être calculés par intégration invariante (les valeurs exactes et les relations utilisées pour calculer ces diagrammes sont répertoriées dans l'appendice C) . Pour A on a

$$A = \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\mu_5, \mu_7}^m e^{-2x(\mu_5 + \mu_7)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,7 \\ 5,7 \\ 4,5 \\ 3,5 \end{array} \quad (3.29)$$

$$= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \left(\sum_{\mu_5 = \mu_7}^m e^{-4x\mu_5} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,5 \\ 5,5 \\ 4,5 \\ 3,5 \end{array} + \sum_{\mu_5 \neq \mu_7}^m e^{-2x(\mu_5 + \mu_7)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6,7 \\ 5,7 \\ 4,5 \\ 3,5 \end{array} \right) \quad (3.30)$$

$$= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \left(\begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} \left(\sum_{\mu_5 = \mu_7}^m e^{-4x\mu_5} \right) + \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ \bar{4} \\ \bar{3} \end{array} \left(\sum_{\mu_5 \neq \mu_7}^m e^{-2x(\mu_5 + \mu_7)} \right) \right) \quad (3.31)$$

$$= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \left(f(x) \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} + g(x) \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \diagup \\ \diagdown \\ \diagup \\ \diagdown \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ \bar{4} \\ \bar{3} \end{array} \right). \quad (3.32)$$

II. STATISTIQUE DE L'INTERFÉRENCE DANS UN SYSTÈME QUANTIQUE COUPLÉ À UN SPIN

A partir de (3.31), alors que sur le côté gauche des diagrammes les indices $\mu_1 \dots \mu_4$ apparaissent encore explicitement, on remarque que les indices sur le côté droit renvoient maintenant seulement à un seul indice α et qu'ils peuvent être barrés. Ceci est encore une simplification d'écriture. Un *point barré* est un point qui ne peut pas se confondre avec un *point normal* même si leurs valeurs correspondantes sont les mêmes. Explicitement, les points notés $\bar{5}$ et $\bar{6}$ s'écrivent respectivement α_5, μ_5 et α_6, μ_7 pour $\mu_5 \neq \mu_7$. Ainsi la restriction $\mu_5 \neq \mu_7$ implique que les deux points supérieurs ne peuvent pas se confondre avec les deux en position inférieure. En d'autre terme, les points normaux et barrés ne peuvent fusionner qu'avec des points du même type. Rappelons aussi que la contrainte $\alpha_5 \neq \alpha_6$ est implicite dans ces expressions. Les fonctions $f(x)$ et $g(x)$ sont définies par

$$f(x) = \left(\frac{1 - e^{-4xd}}{e^{4x} - 1} \right) = Z(4x) \quad (3.33)$$

$$g(x) = e^{-6x} \left(\frac{1 - e^{-2xd}}{1 - e^{-2x}} \right) \left(\frac{1 - e^{-2x(d-1)}}{1 - e^{-4x}} \right) = Z^2(2x) - Z(4x). \quad (3.34)$$

Toutes ces notations, simplifications et définitions nous permettent maintenant de parvenir à l'expression analytique de $\langle \mathcal{I}^2 \rangle$. Le reste du calcul consiste à effectuer une par une les sommes puis à identifier tous les poids des configurations. Le développement explicite de A et B donne finalement

$$\langle \mathcal{I}^2 \rangle = \frac{n}{Z^4} \left[f(x) \left(A_1 + (n-1)A_3 \right) + g(x) \left(A_2 + (n-1)A_4 + n(n-1)B_1 + n(n-1)^2 B_2 \right) \right]. \quad (3.35)$$

Ici le paramètre d dans les fonctions f et g est $d = m$. Les termes A_i et B_i sont définis et calculés explicitement dans l'appendice B. Ils ne dépendent que de n et m .

Dans le cas $m = d = 1$ ($n = N$) tous les préfacteurs $m[i]$ (cf. Appendice B) sont nuls si $i \geq 1$. Avec les mêmes paramètres que pour (3.33,3.34), $f(x) = Z = 1$, et $g(x) = 0$. Ainsi l'expression (3.35) de $\langle \mathcal{I}^2 \rangle$ se simplifie considérablement,

$$\begin{aligned} \langle \mathcal{I}^2 \rangle &= \left(NA_1 + N(N-1)A_3 \right) \\ &= \left(N(N[3] \begin{array}{c} \diagup \\ \diagdown \end{array} + 4N[2] \begin{array}{c} \diagup \\ \diagdown \end{array} + 2N[1] \begin{array}{c} \diagup \\ \diagdown \end{array}) \right. \\ &\quad \left. + N(N-1)(N[3] \begin{array}{c} \diagup \\ \diagdown \end{array} + 4N[2] \begin{array}{c} \diagup \\ \diagdown \end{array} + 2N[1] \begin{array}{c} \diagup \\ \diagdown \end{array}) \right) \\ &= \frac{N(N^3 - 5N + 8) - 4}{(N+1)(N+3)}. \end{aligned}$$

Comme attendu, on récupère l'expression (2.12), obtenue au chapitre précédent, de l'écart quadratique moyen $\sigma_{\mathcal{I}} = \frac{2}{N+1} \sqrt{\frac{N-1}{N+3}}$ dans le cas d'une propagation purement unitaire.

Sur la figure (3.4) est représenté l'écart quadratique moyen de la distribution de $\mathcal{I}$, $\sigma_{\mathcal{I}}(n, m, x) = (\langle \mathcal{I}^2 \rangle - \langle \mathcal{I} \rangle^2)^{1/2}$ pour quatre températures différentes en fonction de n et m . Pour une température donnée, $\sigma_{\mathcal{I}}(n, m, x)$ décroît avec n et m . Le graphe de contour en échelle log-log-log montre que cette décroissance se comporte comme une loi de puissance aussi bien en n qu'en m . Les exposants correspondants peuvent être obtenus en faisant le développement asymptotique de

l'expression de la variance $var(n, m, x) = \sigma_{\mathcal{I}}^2(n, m, x)$ pour $n \gg 1$ ou pour $m \gg 1$, dans les limites de températures nulle ou infinie. Pour m fixé et $n \gg 1$ on obtient

$$var(n, m, x \rightarrow \infty) = \frac{2(m-1)^2}{nm^4} - \frac{4(m^4 - 3m^3 + 3m^2 - 5m + 3)}{m^6 n^2} + \mathcal{O}\left(\frac{1}{n^3}\right) \quad (3.36)$$

$$var(n, m, x = 0) = \frac{2(m^2 - 1)}{nm^6} + \frac{8 - 4m^4}{m^8 n^2} + \mathcal{O}\left(\frac{1}{n^3}\right). \quad (3.37)$$

Ceci doit être comparé au cas unitaire, dans lequel le développement asymptotique s'écrit $var_U(n) = \frac{4}{n^2} + \mathcal{O}\left(\frac{1}{n^3}\right)$, comme on peut le voir facilement à partir de (3.36) et (3.37) en prenant $m = 1$ (voir aussi relation (2.12) du chapitre précédent). On remarque que la variance décroît plus lentement en fonction de n en présence de décohérence, c'est à dire comme $1/n$ au lieu de $1/n^2$ dans le cas unitaire. En d'autres termes, la décohérence tend à ralentir la propriété qu'a la distribution à devenir très piquée et étroite pour n grand. La diminution en loi de puissance de n implique cependant que même dans le cas non-unitaire, la distribution d'interférence devient pour $n \gg 1$ une distribution très piquée et étroite centrée autour de l'interférence moyenne, qui elle-même augmente avec n d'après les équations (3.18) et (3.21)).

Le développement asymptotique de $var(n, m, x)$ en fonction de $m \gg 1$ pour n fixé donne

$$var(n, m, x \rightarrow \infty) = \frac{2(n-1)^2}{n^3 m^2} + \mathcal{O}\left(\frac{1}{m^3}\right) \quad (3.38)$$

$$var(n, m, x = 0) = \frac{(n-1)^2}{n^3 m^4} + \mathcal{O}\left(\frac{1}{m^6}\right). \quad (3.39)$$

Dans ce cas aussi, une augmentation de la taille de l'environnement amincit la distribution d'interférence. Cependant, d'après (3.23,3.24), l'interférence moyenne décroît comme $1/m$ ($1/m^2$) pour $x \rightarrow \infty$ ($x \rightarrow 0$). Il est intéressant de noter que dans ce cas, la largeur relative, c'est à dire l'écart quadratique moyen divisé par la valeur moyenne, est asymptotiquement indépendant de la taille de l'environnement.

Les résultats numériques de la section II.3 sont en très bon accord avec les résultats analytiques, comme on peut le voir sur le tableau (3.1) où sont comparées les moyennes et les écarts quadratiques moyens obtenus numériquement avec leurs valeurs analytiques correspondantes, ceci pour des tailles données.

n	m	$\langle \mathcal{I} \rangle$ (num.)	$\langle \mathcal{I} \rangle$ (ana.)	$\sigma_{\mathcal{I}}$ (num.)	$\sigma_{\mathcal{I}}$ (ana.)
4	2	0,57279	0,57286	0,11728	0,11719
4	4	0,14296	0,14293	0,03260	0,03255
4	8	0,03702	0,03702	0,00864	0,00864
8	2	1,54120	1,54109	0,09022	0,09409
8	4	0,38796	0,38796	0,02670	0,02666

TABLE 3.1.: Comparaisons entre les valeurs numériques et analytiques de $\langle \mathcal{I} \rangle$ et $\sigma_{\mathcal{I}}$. Toutes les valeurs sont arrondies jusqu'au cinquième chiffre après la virgule.

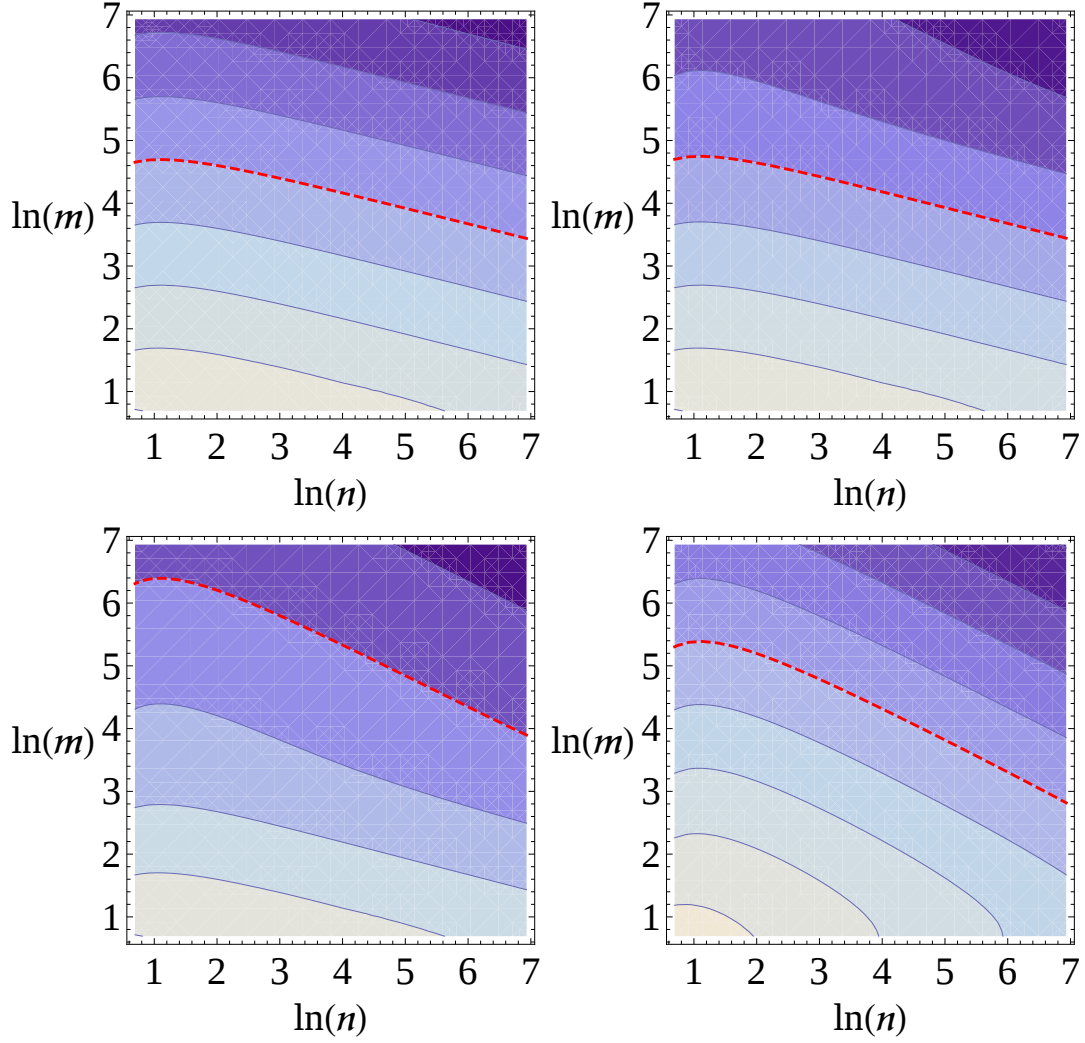


FIGURE 3.4.: Graphe de contour de $\ln(\sigma_{\mathcal{I}}(n, m))$ Pour $x = 0.001, 0.01, 0.1$ et 10 (d'en haut à gauche à en bas à droite), pour n et m entre 2 et 1024 . La distance entre les contours vaut 2 et la ligne rouge pointillée indique que le contour $\ln(\langle \mathcal{I}(n, m) \rangle) = -10$, excepté pour le dernier graphe où la distance entre les contours vaut 1 et la ligne rouge pointillée indique $\ln(\sigma_{\mathcal{I}}(n, m)) = -6$. Les valeurs augmentent avec les tons de couleurs allant de sombres à clairs.

III. Interférence dans un système quantique couplé à plusieurs spins

Dans cette section, nous généralisons les calculs analytiques précédents dans la situation où l'environnement est constitué de s spins indépendants, chacun ayant d niveaux d'énergie espacés de $\hbar\Omega$. Dans ce cas précis, l'environnement a une dimension de $m = d^s$. Le hamiltonien d'un tel système quantique s'écrit

$$H^{(s)} = \sum_{k=1}^s H_k^{(1)}, \quad (3.40)$$

où $H_k^{(1)}$ est le hamiltonien du spin k (cf eq.(3.2)). Les composantes de $H^{(s)}$ dans sa base propre sont

$$H_{\nu\rho}^{(s)} = \hbar\Omega \left(\sum_{k=1}^s \nu_k \right) \delta_{\nu\rho} \quad (3.41)$$

avec comme notation pour les indices $\nu = (\nu_1, \nu_2, \dots, \nu_s)$ et $\rho = (\rho_1, \rho_2, \dots, \rho_s)$

La matrice densité de l'état thermique de ce système factorise comme $\epsilon^{(s)} = \epsilon^{(s)} = \epsilon^{(1)\otimes s}$, ce qui donne en terme de composantes

$$\epsilon_{\nu\rho}^{(s)} = \frac{e^{-x S(\nu)} \delta_{\nu\rho}}{Z^s} \quad (3.42)$$

avec $x = \beta\hbar\Omega$, $S(\nu) = \sum_{k=1}^s \nu_k$, et où Z est la fonction de partition de l'état thermique d'un unique spin, introduite dans la section précédente. Grâce à cette propriété de factorisation de la matrice densité $\epsilon^{(s)}$, généraliser le calcul précédent dans le but d'obtenir les expressions de $\langle \mathcal{I} \rangle$ et $\langle \mathcal{I}^2 \rangle$ pour ce type d'environnement n'est pas difficile. Pour ceci, il suffit de remplacer Z par Z^s dans les équations (3.14), (3.33), et (3.34), en prenant $d = m^s$ au lieu de $d = m$. Ceci est à nouveau une conséquence du fait que la valeur d'un diagramme ne dépend pas de la valeur de ces points. De ce fait, même si les diagrammes portent des valeurs de points issues d'indice composite reflétant la nature composite de l'environnement, leurs valeurs restent les mêmes que dans le cas à un seul spin. Seules leurs multiplicités et les facteurs de dépendance en température sont modifiés. Comme les spins sont choisis comme non interagissant entre eux, la somme sur le facteur thermique donne juste une puissance du facteur thermique à un seul spin, de la même manière que pour l'expression de la fonction de partition. Cela signifie que les fonctions f et g sont modifiées selon

$$\begin{aligned} f(x) &\rightarrow \sum_{\mu} e^{-4xS(\mu)} = \sum_{\mu_1}^d \dots \sum_{\mu_s}^d e^{-4x\mu_1} \dots e^{-4x\mu_s} = Z^s(4x) = f^s(x) \\ g(x) &\rightarrow \sum_{\mu \neq \nu} e^{-2x(S(\mu)+S(\nu))} = \sum_{\mu, \nu} e^{-2x(S(\mu)+S(\nu))} - \sum_{\mu} e^{-4xS(\mu)} = Z^s(2x) - Z^s(4x). \end{aligned}$$

Les expressions pour $\langle \mathcal{I} \rangle$ et $\langle \mathcal{I}^2 \rangle$ deviennent

$$\langle \mathcal{I} \rangle = n^2(n-1)h^s(x) \left(d \left\langle \begin{array}{c} \diagup \quad \diagdown \end{array} \right\rangle + d(d-1) \left\langle \begin{array}{c} \diagup \quad \diagdown \\ \diagdown \quad \diagup \end{array} \right\rangle \right) \quad (3.43)$$

$$= \left(\coth\left(\frac{dx}{2}\right) \tanh\left(\frac{x}{2}\right) \right)^s \left(\frac{nd^s(n-1)^2}{n^2d^{2s}-1} \right), \quad (3.44)$$

$$\begin{aligned} \langle \mathcal{I}^2 \rangle &= \frac{n}{Z^{4s}} \left[f^s(x) \left(A_1 + (n-1)A_3 \right) \right. \\ &\quad \left. + g^s(x) \left(A_2 + (n-1)A_4 + n(n-1)B_1 + n(n-1)^2B_2 \right) \right]. \end{aligned} \quad (3.45)$$

Ici l'argument m dans les termes A_i et B_i des expressions (3.43,3.45) est $m = d^s$. Cela signifie que s spins de taille $(d-1)/2$ agissent de manière très similaire qu'un seul spin de taille $(d^s-1)/2$, sur les deux premiers moments de $P(\mathcal{I})$. La seule différence provient des préfacteurs dépendants de la température $f(x)$, $g(x)$ et $h(x)$. Pour un seul spin de taille $(d^s-1)/2$, le paramètre d dans les équations (3.15,3.33,3.34) est donné par la dimension de l'environnement

$m = d^s$, mais dans les équations (3.43,3.45) nous avons $s = 1$ pour un seul spin. Pour s spins de taille $(d-1)/2$, dans les équations (3.15,3.33,3.34), la dimension reste d , est s est le nombre de spins dans (3.43,3.45). Pour finir notons que dans les deux limites $x \rightarrow 0$ et $x \rightarrow \infty$ les expressions coïncident : à température nulle ou à température infinie, un environnement de s spins de taille $(d-1)/2$ et un environnement d'un seul spin de taille $(d^s-1)/2$, influencent de la même manière le comportement statistique de l'interférence dans un système quantique fini, du moins en ce qui concerne les deux premiers moments de la distribution de l'interférence.

IV. Conclusion partielle

Dans ce chapitre nous avons analysé quantitativement comment l'interférence quantique est affectée par la décohérence. Le modèle étudié est celui d'un système quantique couplé à un environnement, l'ensemble évoluant de manière jointe selon une évolution unitaire. En nous basant sur le modèle statistique dans lequel l'évolution du système global est choisi dans l'ensemble de matrices aléatoires CUE, nous avons montré que l'interférence moyenne augmente approximativement linéairement avec la dimension n de l'espace de Hilbert associé au système, mais décroît comme une puissance de la dimension m de l'espace de Hilbert associé à l'environnement. Cette puissance dépend de la température initiale de l'environnement, la décroissance se comportant comme $1/m^2$ pour $T = 0$ et comme $1/m^3$ pour $T \rightarrow \infty$. De son côté, la largeur de la distribution décroît plus lentement quand la décohérence devient importante, ceci comme $1/\sqrt{n}$ (au lieu de $1/n$ dans le cas unitaire) pour m fixé. De ce fait comme dans le cas sans décohérence, pour $n \gg 1$ et m fixé, la distribution de l'interférence devient un pic très fin centré sur la valeur moyenne. Numériquement, nous avons observé que la distribution d'interférence est bien fittée par une distribution log-normale pour n suffisamment grand. Ceci implique que le nombre de *i-bits* (concept introduit dans [BG06]) est en bonne approximation distribuée selon une loi normale.

L'information quantique au service de la production de matrices aléatoires

Dans ce chapitre nous allons exposer les résultats numériques montrant comment l'ensemble UCE, introduit dans le chapitre 2, reproduit efficacement diverses quantités statistiques propres à l'ensemble CUE, telle la distribution (du logarithme du module carré) des éléments de matrice, les moments jusqu'à l'ordre 16 de la distribution du module carré des éléments de matrice et les corrélations jusqu'à l'ordre 16 entre éléments d'une même colonne. Ces résultats numériques poussés jusqu'à de grandes tailles de matrices (28 qubits pour la distribution, 15 qubits pour les moments et les corrélations) montrent que ces quantités sont reproduites efficacement dans le sens où le nombre de portes nécessaires pour reproduire ces quantités avec une précision arbitraire ϵ n'évolue pas plus vite que $n_q \ln(n_q/\epsilon)$

I. Généralités et motivations

I.1. Opérateurs pseudo-aléatoires et circuits quantiques aléatoires

Les matrices aléatoires unitaires jouent un rôle important dans de nombreuses tâches reliées au traitement quantique de l'information tels que le quantum data hiding [PWM02], la distinction d'état quantique [Sen05], le chiffrement quantique [AS04], le superdense coding d'états quantiques [HHL04] ainsi que l'estimation du bruit [ELL05]. Dans toutes ces applications un ensemble aléatoire de matrices unitaires U de taille $N \times N$ distribuées uniformément par rapport à la mesure de Haar du groupe unitaire de dimension N (l'ensemble unitaire circulaire ou CUE) est nécessaire [Meh91]. En principe, toute matrice unitaire agissant sur les vecteurs d'un espace de Hilbert de dimension $N = 2^{n_q}$ de n_q qubits peut être approximée avec une précision arbitraire en utilisant un ensemble universel de portes quantiques agissant sur un ou deux qubits en même temps [Deu85, DiV95, SW95, BBC95, Bar95]. Cependant, comme une matrice unitaire contient N^2 paramètres réels indépendants, le nombre de portes n_g nécessaires pour la construire augmente exponentiellement avec le nombre de qubits. Plus précisément, $\mathcal{O}(N^2(\ln N)^3)$ portes sont nécessaires pour approximer tous les éléments de matrice de U en utilisant un ensemble universel de porte fixé [NC00]. Ceci rend la construction des ensembles de matrices aléatoires distribuées précisément par rapport à la mesure de Haar du groupe unitaire, hautement inefficace. Une des méthodes explicites pour la construction de matrice CUE, bien qu'inefficace, se base sur la paramétrisation de Hurwitz utilisé au chapitre 2 et 3 (Voir [PZK98]).

Dans un papier fondateur, Emerson *et.al* ont introduit le concept d'opérateurs unitaires pseudo-aléatoires, c'est à dire celui d'opérateurs unitaires aléatoires distribués selon une distribution qui imite une distribution uniforme par rapport à la mesure de Haar de $U(N)$ [EWS03]. La construction de ces opérateurs est motivée par des idées provenant du chaos quantique. En effet, si on considère un système classique intégrable, la statistique des niveaux d'énergie de la version quantifiée du système suit approximativement une statistique poissonnienne. En contre-partie la statistique de niveau d'énergie de la version quantique d'un système chaotique, suit approximativement celles associées à des matrices aléatoires. Dans [EWS03] la construction se base sur un circuit quantique consistant en des rotations $U(2)$ aléatoires sur chaque qubit, suivit par une porte à deux qubits qui implémente une interaction de type Ising entre proches voisins. Ils montrent que ce circuit produit des matrices unitaires dont la distribution des éléments converge exponentiellement avec le nombre de portes quantiques sur la distribution bien connue des éléments de matrice de CUE [Haa00]. Plus tard, Emerson, Livine et Lloyd ont montré que la distribution de probabilité jointe des éléments de matrice d'un produit d'opérateurs unitaires créé par un circuit quantique aléatoire (composé d'un ensemble universel continu ou discret), converge exponentiellement avec le nombre de portes vers celle de CUE [ELL05]. Cependant cette convergence se fait avec un taux qui diminue exponentiellement avec le nombre de qubits, ceci laissant ouverte la question sur l'efficacité de la construction d'opérateur unitaire pseudo-aléatoire avec de telles méthodes. De plus, la distribution $P(U_{ij})$ des éléments de matrice contient seulement une petite quantité d'information comparée à la distribution de probabilité jointe dans son entier. Ainsi il n'est pas clair dans la manière dont les corrélations entre éléments de matrice convergent vers leur valeur CUE équivalentes.

Différentes statistiques ont été étudiées pour différents types d'ensembles de circuits aléatoires. Par exemple dans [Ž07] l'efficacité pour générer de l'intrication bipartite entre deux sous-systèmes a été étudié numériquement pour un circuit quantique composé de portes $U(4)$, chacune étant le produit d'une porte fixe à deux qubits fixe et de deux portes à un seul qubit tirées aléatoirement de manière uniforme par rapport à la mesure de Haar de $U(2)$. Dans ces travaux, une convergence exponentielle est observée avec un taux qui se comporte avec le nombre de qubits comme $n_q \ln n_q$. Oliveira *et.al* ont introduit une technique basée sur les chaînes de Markov pour analyser le même genre de question et ils ont trouvé une borne supérieure de $\mathcal{O}(n_q^3)$ portes quantique nécessaires pour atteindre une précision (absolue) donnée ϵ pour la moyenne de la quantité d'intrication bipartite [ODB07, DOB07]. Finalement la raison de notre intérêt pour ce genre de sujet dans cette thèse provient des travaux du chapitre 1 sur la distribution d'interférence $\mathcal{I}$ entre le modèle de circuit aléatoire UCE (et OCE) et les ensembles CUE (respectivement HOE).

I.2. Théorie des k-design

D'un autre côté l'étude des opérateurs pseudo-aléatoires, est en étroite relation avec la théorie des k -designs. Différentes définitions équivalentes existent et par exemple Dankert *et.al* définissent un k -design unitaire comme un ensemble discret $\mathcal{G} = \{U_{(j)}\}$ de k matrices unitaires de taille N tel que pour n'importe quel polynôme $\mathcal{P}$ d'élément de matrice de degré égal ou inférieur à k , on ait

$$\frac{1}{k} \sum_{j=1}^k \mathcal{P}(U_{(j)}) = \int (dU) \mathcal{P}(U) \quad (4.1)$$

où (dU) est la mesure de Haar de $U(N)$. Ils définissent aussi un k -design ϵ -approximé par

$$(1 - \epsilon) \int (dU) \mathcal{P}(U) \leq \frac{1}{k} \sum_{j=1}^k \mathcal{P}(U_{(j)}) \leq (1 + \epsilon) \int (dU) \mathcal{P}(U) \quad (4.2)$$

[DCEL06]. Récemment Harrow *et.al* ont montré qu'un circuit aléatoire de taille polynomiale en n_q mène à un 2-design ϵ -approximé. En fonction de l'ensemble de portes utilisé, le nombre n_g nécessaire pour atteindre une précision donnée ϵ se comporte comme $\mathcal{O}(n_q(n_q + \log 1/\epsilon))$ ou comme $\mathcal{O}(n_q \ln(n_q/\epsilon))$. Ils conjecturent qu'un circuit aléatoire sur n_q qubits composé de $\text{poly}(n_q, k)$ portes à deux qubits aléatoires choisies dans un ensemble universel est un k -design ϵ -approximé [HL08]. Les résultats précédents, impliquent que des circuits aléatoires peuvent efficacement reproduire les moyennes typiques de CUE des termes du type $|U_{ij}|^2$, $|U_{ij}|^4$ et $|U_{ij}U_{kl}|^2$ (Les termes du type $U_{ij}U_{kl}^*$ moyennés sur CUE s'annulent, cf. Annexe 1). Ces résultats confirment ceux trouvés dans [Z07], dans le sens où les quantités qui y sont étudiées numériquement peuvent être écrites comme des polynômes d'ordre (2,2) en terme de U_{ij} et de U_{kl}^* .

I.3. L'ensemble de circuit unitaire UCE

Rappelons que l'ensemble de circuits unitaires (UCE) introduit dans le chapitre 2, consiste en des algorithmes quantiques qui utilisent deux type de portes (Voir la figure (2.3) du chapitre 2) qui reproduit celle du chapitre 2) : la porte $U(2)$ agissant sur un seul qubit et la porte CNOT agissant sur deux qubits. Chaque algorithme est construit comme une séquence aléatoire de telles portes, dans laquelle la probabilité d'avoir une porte U_2 est p et la probabilité d'avoir une porte CNOT est $1 - p_g$. Le choix des qubits sur lesquels agissent les portes est fait uniformément et indépendamment sur l'ensemble des qubit du circuit. La porte U_2 est choisie uniformément par rapport à la mesure de Haar invariante du groupe $U(2)$. En considérant les quatre angles, α , ψ , χ choisis aléatoirement et uniformément dans l'intervalle $[0, 2\pi[$, et $\varphi = \arcsin(\xi^{1/2})$ où ξ est choisi aléatoirement et uniformément dans l'intervalle $[0, 1[$, toute matrice U_2 est paramétrée par

$$U_2(\boldsymbol{\theta}) = e^{i\alpha} \begin{pmatrix} \cos \varphi e^{i\psi} & \sin \varphi e^{i\chi} \\ -\sin \varphi e^{-i\chi} & \cos \varphi e^{-i\psi} \end{pmatrix} \equiv \begin{pmatrix} c & s \\ -\bar{s} & \bar{c} \end{pmatrix}, \quad (4.3)$$

ici les quatre angles sont notés $\boldsymbol{\theta} = (\alpha, \psi, \chi, \varphi)$ [PZK98]. Notons que l'ensemble des phases α participant à la construction du circuit ne modifie au final que la phase globale de l'algorithme. D'après les résultats de [ELL05] il est clair que dans la limite où le nombre de portes $n_g \rightarrow \infty$ pour un nombre de qubits n_q fixé, UCE converge vers CUE. En suivant la notation de [HL08], l'ensemble UCE, qui est bien constitué de portes $U(4)$, peut être résumé comme

$$\Gamma_{UCE} = \left\{ \left\{ \frac{d\mu(U_2)}{4}, U_2 \otimes \mathbf{1}_2 \right\}, \left\{ \frac{d\mu(U_2)}{4}, \mathbf{1}_2 \otimes U_2 \right\}, \left\{ \frac{1}{4}, U_{\text{CNOT}_{1,2}} \right\}, \left\{ \frac{1}{4}, U_{\text{CNOT}_{2,1}} \right\} \right\} \quad (4.4)$$

où le premier nombre dans chaque paire dans la liste est la probabilité que le second membre dans la paire soit sélectionné à toutes étapes dans l'algorithme. $\mu(U_2)$ est la mesure de Haar de $U(2)$, et $U_{\text{CNOT}_{i,j}}$ est l'opérateur unitaire associé à l'opération CNOT dans laquelle i est le qubit de contrôle et j est le qubit cible. On peut facilement vérifier que Γ_{UCE} est *2-copy gapped* suivant la terminologie de [HL08]. Cela signifie que l'opérateur

$$G = \int_{U(4)} U \otimes U \otimes U^* \otimes U^* d\mu_{\Gamma}(U) \quad (4.5)$$

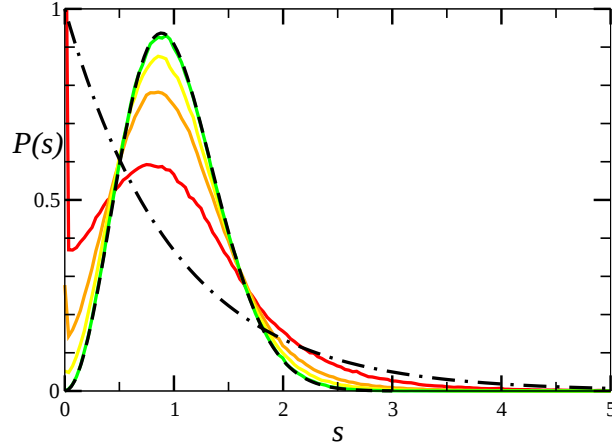


FIGURE 4.1.: $P(s)$ pour $n_g = 10$ (rouge), 15 (orange) 20 (jaune) et 40 (vert) pour UCE avec $n_q = 4$ et $p = 0.5$, $n_r = 10^5$ matrices, comparée à la forme de Wigner $P_W(s)$ (4.6) en pointillés et à la distribution de Poisson $P(s) = \exp(-s)$ en tirés.

définit pour un ensemble général de portes distribuées continuellement sur $U(4)$ avec la mesure $\mu_\Gamma(U)$, a seulement deux valeurs propres de module unité. La différence entre le module d'une de ces deux valeurs propres et le module de la valeurs propre suivante constitue le *gap spectral* Δ . Dans le cas de UCE on peut calculer explicitement l'opérateur G de taille 256 et en le diagonalisant on montre que le gap spectral de UCE vaut $\Delta \simeq 0.232703$ si les portes sont représentées comme des matrices 4×4 . Cependant si les portes sont représentées par des matrice de taille 2^{n_q} , on attend que le gap spectral diminue comme $1/n_q$ tout en restant fini [Har].

II. Convergence de UCE vers CUE

II.1. Distribution des différences entre phases propres voisines

Une des quantités les plus caractéristique des ensembles de matrices aléatoires circulaires est la distribution $P(s)$ des différences s entre phases φ_l propres voisines. Comme cette distribution est vraiment typique dans le domaine des matrices aléatoires il convient tout d'abord de vérifier que l'ensemble UCE se rapproche bien du même genre de statistique quand le nombre de porte n_g devient suffisamment grand. Pour cela on peut comparer la distribution $P(s)$, calculer pour CUE et UCE. Pour N grand et pour une moyenne de s normalisé à l'unité, CUE mène à une distribution $P(s)$ qui est bien approximée par la forme de Wigner [Meh91]

$$P_W(s) = \frac{32s^2}{\pi^2} e^{-4s^2/\pi}. \quad (4.6)$$

Les déviations de $P_W(s)$ par rapport à la forme exacte sont de l'ordre de 10^{-3} [Haa91] et la limite $N \rightarrow \infty$ est déjà atteinte dès $N = 3$. Notons que cette distribution se différencie d'une distribution poissonnienne d'une part par le phénomène de *répulsion des niveaux*, la distribution s'annulant pratiquement dans le cas de phases propres dégénérées $s = 0$ et d'autre part par le comportement gaussien pour s grand (cf. fig 4.1).

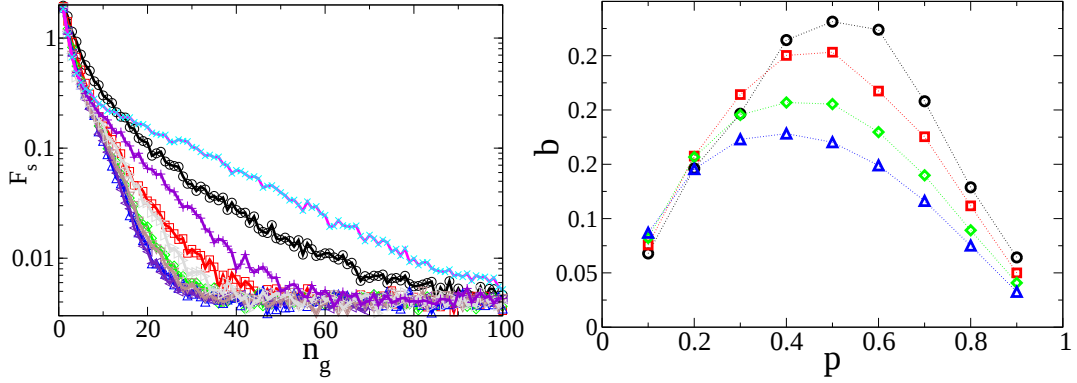


FIGURE 4.2.: Convergence de $P(s)$ pour UCE vers $P_W(s)$ en fonction de n_g pour $n_q = 4$ qubits, $n_r = 10^3$ et différentes valeurs de p : $p = 0.1$ carrés noirs, $p = 0.2$ carrés rouges, $p = 0.3$ losanges verts, $p = 0.4$ triangle hauts bleus, $p = 0.5$ triangles indigo gauche, $p = 0.6$ triangles bas brun, $p = 0.7$ triangle gris droit, $p = 0.8$ violette et $p = 0.9$ magenta (gauche). Taux de convergence b en fonction de la probabilité p pour $n_q = 3$ (cercles), $n_q = 4$ (carrés rouges), $n_q = 5$ (carrés verts), et $n_q = 6$ (triangles).

Pour $n_q = 4$ (c'est à dire $N = 16$), le nombre minimal de portes qui donne approximativement une densité constante de phases propres, de telle sorte qu'il ne soit pas nécessaire de replier le spectre [Meh91], est $n_g \simeq 10$. Pour un petit nombre de portes, la densité d'états connaît certes de gros pics en $\varphi = 0$ et $\varphi = \pi$ qui correspondent à la prédominance de valeurs propres réelle, mais en dehors de ces valeurs précises, la densité est déjà plutôt uniforme. Des indications numériques présentées dans [PZK98] indiquent la même forme pour $P(s)$ dans le cas HOE que dans le cas CUE (et donc que la forme (4.6) est une bonne approximation de $P_{HOE}(s)$), en particulier une répulsion quadratique des niveaux $P(s) \propto s^2$ pour $s \ll 1$. Nous avons ainsi examiné la convergence de OCE vers HOE en nous basant sur celle des distributions d'espacement $P(s) \rightarrow P_W(s)$ et avons trouvé des résultats similaires à ceux présentés sur la figure (4.1). Cependant comme les études sur l'ensemble HOE sont beaucoup moins répandues, la suite des résultats présentés ne portera que sur la convergence vers CUE.

La figure (4.1) montre la distribution $P(s)$ calculée numériquement associée à UCE pour $n_q = 4$ et différentes valeurs de n_g ($n_r = 10^5$ réalisations). Pour n_g petit, $P(s)$ a un gros pics en $s = 0$. Le reste de la distribution ressemble à un mélange de la distribution pour des phases sans corrélations (distribution poissonnienne $P(s) = \exp(-s)$), et la forme de Wigner $P_W(s)$. Le pic en $s = 0$ devient de plus en plus petit quand le nombre de porte augmente, et en même temps le maximum en $s = 1$ devient de plus en plus prononcé, la distribution calculée s'approchant rapidement de $P_W(s)$ (4.6). Pour $n_g = 40$, $P(s)$ est indiscernable de $P_W(s)$. Nous pouvons examiner la convergence de façon plus quantitative à l'aide de la quantité

$$F_s = \int_0^\infty \left(\sqrt{P_{UCE}(s)} - \sqrt{P_W(s)} \right)^2 ds = 2 \left(1 - \int_0^\infty \sqrt{P_{UCE}(s)P_W(s)} ds \right) \quad (4.7)$$

qui mesure le carré d'une distance entre (la racine carré de) la distribution d'espacement $P_{UCE}(s)$ pour UCE avec $P_W(s)$ [BZ06]. La figure (4.2) montre comment F_s évolue en fonction de n_g pour $n_q = 4$ qubits et pour différentes valeurs de p , résultat numérique obtenu à partir de 10^3 algorithmes aléatoires UCE. Pour p différent de 0 et 1, F_s décroît en bonne ap-

proximation de manière exponentielle comme $\sim \exp(-b(n_q, p)n_g)$, avec un taux b qui dépend de n_q et de p , avant de saturer à faible niveau de manière indépendante de p . Tout comme au chapitre 2 pour F_Z , cette saturation n'est due qu'aux fluctuations numériques inhérentes au calcul de $P_{UCE}(s)$ pour une valeur finie de réalisations n_r . On peut s'en assurer en variant le paramètre n_r . De plus, la précision finie de $P_W(s)$ pour $N > 2$ fixe une autre borne inférieure pour les valeurs de F qui peuvent être réalisées. La figure (4.2) montre aussi que le taux de convergence $b(n_q, p)$, obtenu à l'aide d'un fit de $\ln F_s$ par une fonction linéaire de n_g entre $F_2 = 2$ et $F_2 = 0.1$, possède un maximum au alentour de $p = 0.5$ et qu'il décroît quand n_q croît.

III. Efficacité de la convergence

III.1. Distribution relative au éléments de matrice

L'ensemble CUE, mène à une distribution jointe de probabilité $P(U) \equiv P(U_{11}, U_{12}, \dots, U_{NN})$ des éléments de matrice U_{ij} qui caractérise entièrement cet ensemble. La convergence de UCE vers CUE peut être vue comme la convergence de la distribution de probabilité $\tilde{P}(U)$ associée à UCE vers $P(U)$. Cependant, une étude numérique directe de cette distribution jointe est impossible en pratique car sa construction nécessite un nombre d'opérations qui augmente exponentiellement avec le nombre d'arguments indépendants la définissant. Des quantités plus manipulables peuvent être obtenues à partir de la distribution jointe en intégrant sur plusieurs variables, le prix à payer étant bien sûr, le lien univoque existant entre l'ensemble étudié et la quantité le représentant. La quantité de ce type la plus naturelle est la distribution reliée à un seul élément de matrice. Cette distribution dépend d'un seul paramètre complexe, l'élément U_{ij} , et est obtenue en intégrant sur les $N^2 - 1$ autres paramètres complexes,

$$P(U_{ij}) = \int \dots \int \prod_{(k,l) \neq (i,j)} dU_{kl} P(U). \quad (4.8)$$

La première quantité étudiée dans ce chapitre est très proche de celle-ci. C'est la distribution des quantités définies par $l_{ij} = \ln(N|U_{ij}|^2)$. Le choix de la renormalisation par N et d'une échelle logarithmique est toujours relié à l'idée de faciliter les manipulations. La théorie des matrices aléatoires permet de montrer que pour CUE, tous les l_{ij} sont distribués selon la distribution normalisée

$$P(l) = \frac{(N-1)}{N} e^l \left(1 - \frac{e^l}{N}\right)^{N-2}, \quad (4.9)$$

indépendamment des indices i et j de l_{ij} [Haa00]. Cette propriété d'indépendance vis-à-vis des indices, aura son importance dans la suite.

A priori pour UCE, la distribution des éléments de matrice n'est pas indépendante de l'élément choisi tant que le nombre de portes n_g est petit. Mais on s'attend à ce que la distribution s'uniformise pour $n_g \rightarrow \infty$. Ainsi nous pouvons utiliser cette uniformité des distributions sur toute la matrice pour faire deux simplifications permettant d'améliorer l'efficacité numérique :

- Premièrement, nous produisons et propageons seulement la première colonne de chaque matrice. Ceci réduit bien entendu la mémoire requise pour simuler le circuit UCE. Sous cette forme l'action d'une porte CNOT ne nécessite la manipulation que d'un sous-ensemble

de sa matrice représentative. En écrivant l'indice de ligne i d'un élément de matrice U_{i1} en notation binaire comme $i = 1 + \sum_{\alpha=1}^{n_q} \sigma_\alpha 2^\alpha$, l'action d'un CNOT entre le qubit k (contrôle) et l (cible) choisi dans $[1, n_q]$ nécessite seulement l'échange des 2^{n_q-2} éléments en position où $(\sigma_k = 0, \sigma_l = 1)$ avec les 2^{n_q-2} éléments en position où $(\sigma_k = 1, \sigma_l = 1)$. De la même façon pour l'action de la porte U_2 , chaque élément dans la nouvelle colonne devient une combinaison linéaire de deux anciens éléments, avec $c, s, -\bar{s}$, ou $\bar{c}$ comme coefficients.

- Deuxièmement, nous définissons $\tilde{P}(l)$ en moyennant aussi bien sur l'ensemble de réalisations statistiques $\langle \dots \rangle_R$ que sur l'ensemble des distributions (histogrammes) associées à chaque élément le long de la première colonne $\langle \dots \rangle_C$,

$$\tilde{P}(l) = \frac{1}{n_r N} \sum_{r=1}^{n_r} \sum_{i=1}^N \tilde{h} \left(l_{i1}^{(r)} \right) \equiv \langle \langle \tilde{h}(l) \rangle_C \rangle_R \quad (4.10)$$

où $\tilde{h}(l_{i1}^{(r)})$ est l'histogramme pour la $i^{\text{ième}}$ composante de la première colonne de la $r^{\text{ième}}$ matrice. Dans le but d'avoir une bonne statistique, il est important de produire un nombre n_r suffisamment grand de matrices. Cependant, pour un nombre donné de matrices, le rajout d'un qubit au circuit se traduit grossièrement par un doublement du temps de calcul et de l'espace mémoire nécessaire, résultant du doublement de la taille de l'espace de Hilbert. Pour cette raison, et ceci pour un nombre de qubit $n_q \leq 20$, nous considérons un ensemble de $n_r = a 2^{b-n_q}$ matrices (avec a et b entiers). En d'autre terme, pour le rajout d'un qubit, l'augmentation de la taille de l'espace de Hilbert préjudiciable au calcul est compensée par la réduction de la taille de l'ensemble de matrices, ceci sans perte de statistique. Dans nos travaux nous avons fixé $a = 10$ et $b = 20$ menant à un total d'environ 10^7 matrices. Notons que les corrélations entre éléments de matrice de UCE rendent la moyenne sur la première colonne moins efficace que la moyenne sur les réalisations (voir plus bas). De ce fait le bruit des données numériques augmente cependant avec n_q . De plus pour $n_q \geq 20$, les limitations en terme de temps de calcul ($\simeq 300$ heures en parallèle sur 8 processeurs pour $n_q = 28$) obligent de fixer le nombre de réalisations à $n_r = 10$.

La figure (4.3) montre pour $n_q = 4$, comment la convergence de $\tilde{P}(l)$ vers $P(l)$ se manifeste quand n_g augmente. Pour quantifier la loi d'échelle de cette convergence avec le nombre de qubits, nous définissons comme précédemment la quantité

$$\begin{aligned} D_P &= \int_0^\infty \left(\sqrt{\tilde{P}(l)} - \sqrt{P(l)} \right)^2 dl \\ &= 2 \left(1 - \int_0^\infty \sqrt{\tilde{P}(l)P(l)} dl \right) \leq 2 \end{aligned} \quad (4.11)$$

qui représente une distance entre UCE et CUE en terme de leur distribution de la variable l . Cette distance tend vers zéro au fur et à mesure que UCE converge vers CUE pour $n_g \rightarrow \infty$. Dans la définition (4.11), l'utilisation des racines carrées des distributions plutôt que les distributions elles même, est motivé par le fait que de cette manière D_P est bornée (par la valeur 2), ce qui simplifie l'analyse. La figure (4.4) illustre le comportement de D_P en fonction de n_g pour $n_q = 2, 3, \dots, 28$. Comme attendu, cette quantité décroît rapidement avec le nombre de portes n_g , le taux de décroissance se ralentissant avec le nombre de qubit n_q . Comme précédemment, le nombre fini de réalisations statistiques fait que $\tilde{P}(l)$ fluctue toujours autour de $P(l)$. Ainsi la distance D_P ne peut jamais exactement atteindre sa valeur exacte (inférieure) et nous observons de ce fait qu'elle sature pour n_g grand à un niveau fini $d_{\min}$ qui dépend de n_q . On observe

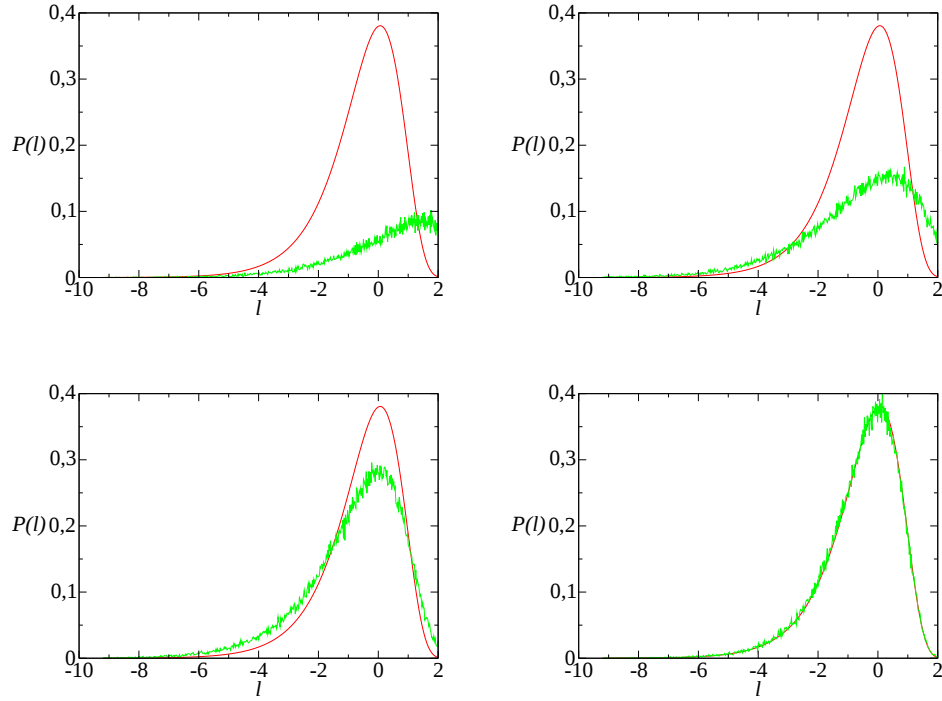


FIGURE 4.3.: Convergence de $\tilde{P}(l)$ vers $P(l)$ (ligne) pour 4 qubits avec $n_g = 5, 10, 20$ et 50 pour un ensemble de 10^4 matrices.

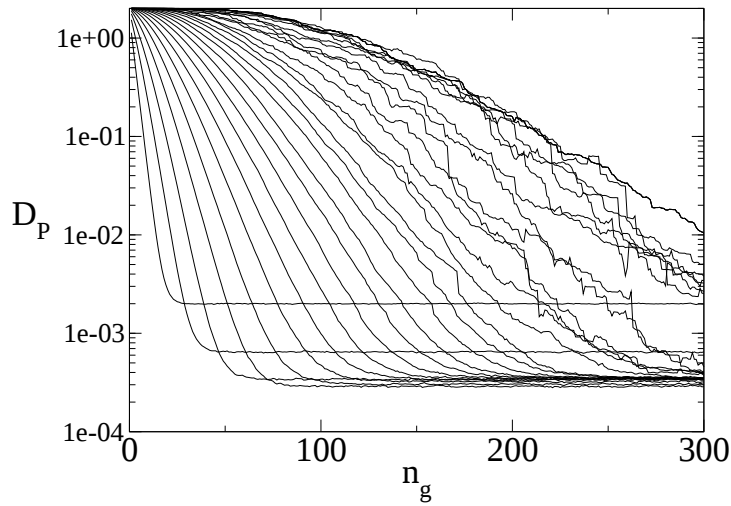


FIGURE 4.4.: Distance $D_P(n_g)$ entre la distribution $P(l)$ et $\tilde{P}(l)$ en fonction du nombre de portes n_g pour $n_q = 2, 3, \dots, 28$ qubits (de gauche à droite).

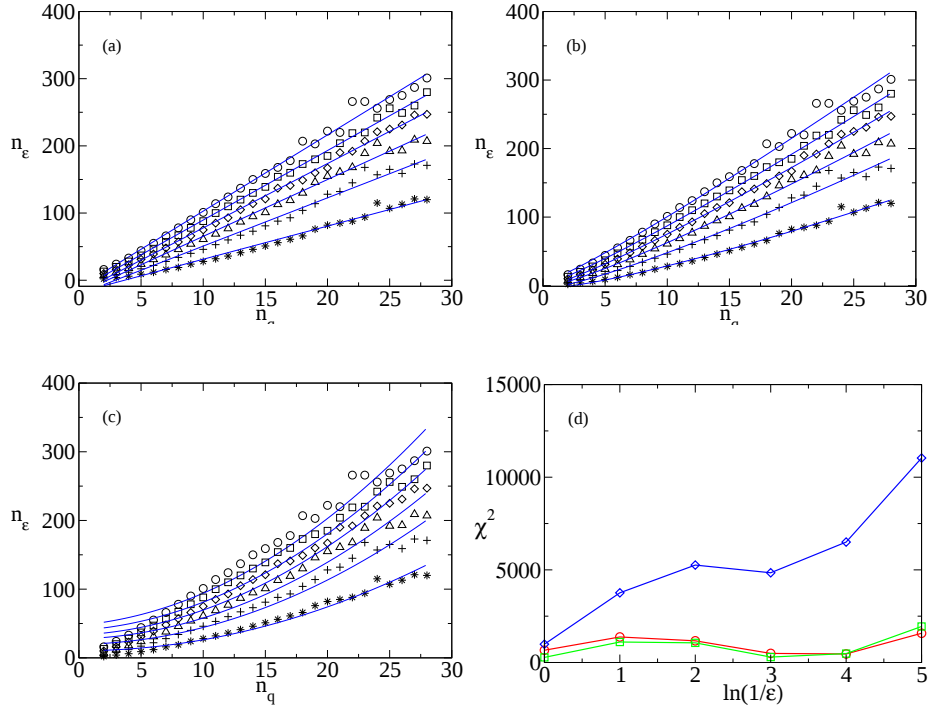


FIGURE 4.5.: Nombre de portes n_ϵ nécessaire pour atteindre $D_p \leq \epsilon$ pour $\ln(\epsilon) = 0, -1, -2, -3, -4$ et -5 (*, +, $\triangle$, $\diamond$, $\square$, $\circ$ respectivement) et $n_q = 2 \dots 28$. Les lignes droites représentent les modèles de fonctions f_1 , f_2 et f_3 (1^{ier}, 2^{ième} et 3^{ième} graphe respectivement). Le dernier graphe représente le χ^2 pour ces modèles (f_1 ($\circ$), f_2 ($\square$)), et f_3 ($\diamond$)).

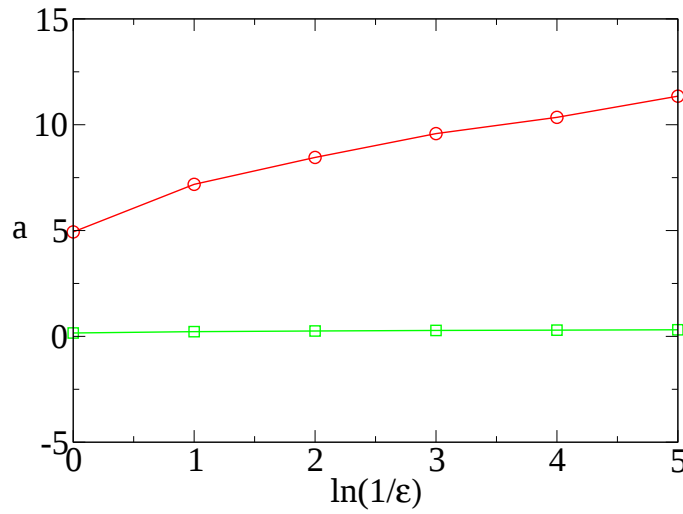


FIGURE 4.6.: coefficients a_1 (cercles rouges) et a_2 (carrés verts) en fonction de $\ln(1/\epsilon)$.

ainsi que le niveau de saturation peut être diminué en augmentant le nombre de réalisations n_r . Ainsi quand D_P sature cela signifie que notre ensemble de circuits devient indiscernable de CUE, à la précision numérique considérée.

Nous pouvons chercher différents modèles pour capturer le comportement de D_P . On observe que pour n_q petit ($n_q \lesssim 12$), D_P est bien fittée par $2e^{-\alpha n_q}$ tandis que pour de plus grandes valeurs de n_q , une composante quadratique dans l'exposant devient prédominante ($D_P \simeq 2e^{-\alpha n_q - \beta n_q^2}$). Ce changement dans la dépendance du modèle en fonction du nombre de qubit rend difficile l'extraction d'une loi d'échelle à partir des exposants α et β . De ce fait il est plus commode (et plus intuitif) de baser l'analyse sur le nombre de porte n_ϵ nécessaire pour atteindre une précision fixée ϵ de D_P ceci pour un nombre donné de qubits. La figure (4.5) montre comment se comporte n_ϵ en fonction de n_q pour six valeurs différentes de ϵ ($\ln(\epsilon)$ entre -5 et 0). Nous avons fitté $n_\epsilon(n_q)$ par trois types de fonctions différentes à 2 paramètres,

$$f_1 = a_1 n_q + b_1, \quad (4.12)$$

$$f_2 = a_2 n_q \ln(n_q/\epsilon) + b_2, \quad (4.13)$$

$$f_3 = a_3 n_q(n_q + \ln(1/\epsilon)) + b_3, \quad (4.14)$$

La fonction linéaire f_1 est un choix évident à la vue des données numériques obtenues. Les fonctions f_2 et f_3 sont motivées par les résultats concernant les 2-designs présentés dans [HL08]. Ses auteurs définissent la convergence des k -designs unitaires par leur action sur une matrice densité test ρ de dimension $k2^{n_q}$. Leur mesure de distance est la norme complètement bornée (the *diamond* norm), de la différence entre l'état $\mathcal{G}_W(\rho) = \sum_i p_i U_i^{\otimes k} \rho (U_i^\dagger)^{\otimes k}$ propagé par le k -design et l'état $\mathcal{G}_H(\rho) = \int_U U^{\otimes k} \rho (U^\dagger)^{\otimes k}$ résultant de la propagation moyennée sur le groupe unitaire. Le résultat majeur de [HL08] est le suivant. Soit un ensemble de portes $\Gamma = \{\{p_i, U_i\}\}$ constitué de matrices unitaires U_i de probabilité d'occurrence p_i , qui forme un ensemble de portes *2-copy gapped*. Alors un circuit aléatoire de taille n_g construit à partir de Γ est un 2-design unitaire ϵ -approximé si $n_g \geq C(n_q(n_q + \log(1/\epsilon)))$ où C est une constante positive qui peut dépendre de l'ensemble de portes. Dans le cas précis où l'ensemble de portes est constitué de portes tirées uniformément sur $U(4)$ (ensemble qui maximise le gap spectral par $\Delta = 1$ puisque l'opérateur G devient un projecteur), il est montré qu'un 2-design unitaire ϵ -approximé est déjà atteint pour $n_g \geq C n_q \log(n_q/\epsilon)$.

Notre ensemble de portes Γ_{UCE} est bien *2-copy gapped* avec un gap spectral $\Delta \simeq 0.232703$, et les résultats de [HL08] sont donc *a priori* applicable en ce qui concerne la convergence de UCE vers CUE. Cependant il convient d'être prudent en comparant ces résultats donnant une borne supérieure basé sur la propagation d'un état test et utilisant la norme *diamond* alors que nos résultats se base sur la distance D_P . Bien sûr, il semble possible que la convergence de la distribution des éléments de matrice du propagateur soit reliée à la convergence d'un état test propagé (voir aussi [TJR07] où est introduit un algorithme quantique efficace pour le twirling). D'après les résultats précédemment cités, le fonction f_3 semble être le meilleur candidat pour modéliser $n_\epsilon(n_q)$. Cependant, il se trouve que la fonction f_2 , relative aux cas dans lesquels le gap spectral vaut $\Delta = 1$, fitte mieux les données numériques. De ce fait, du moins en ce qui concerne la distribution des éléments de matrice (rigoureusement de la quantité l relié aux éléments de matrice), UCE converge vers CUE, plus rapidement que la borne supérieure mentionnée plus haut.

La qualité des fits est mesurée par χ^2 , la somme du carré des déviations (voir la figure (4.5)). On voit que le simple modèle linéaire fonctionne mieux que le modèle quadratique en dépit de la légère déformation vers le haut que connaît les courbes $n_\epsilon(n_q)$. Cette déformation est bien reproduite par le comportement en $n_q \ln n_q$ de f_2 , tandis que le comportement quadratique de

f_3 la reproduit moins bien. De plus la fonction f_2 donne la bonne dépendance en ϵ comme le manifeste la valeur à peu près constante du coefficient $a_2 \simeq 0.2$ (voir figure (4.6)).

Numériquement, une distinction claire entre f_1 et f_2 n'est pas possible dans l'intervalle restreint en nombre de qubits, puisque la qualité des fits (mesurée par χ^2) est similaire. On peut quand même en conclure en ce qui concerne la distribution de l reliée aux éléments de matrice, que UCE peut simuler efficacement CUE, dans le sens où le nombre de portes nécessaires pour atteindre une précision de similitude donnée ϵ , augmente avec le nombre de qubit comme $n_q \ln(n_q/\epsilon)$, et dans tout les cas plus lentement que n_q^2 . Il convient de remarquer que ceci est plutôt surprenant puisque $\tilde{P}(l)$ contient l'information de tous ces moments. De ce fait, cela semble indiquer qu'aucun moment d'un poids appréciable dans la reconstruction de cette distribution, ne nécessite plus que $\mathcal{O}(n_q \ln(n_q/\epsilon))$ portes avant d'approcher sa valeur CUE correspondante d'une distance relative ϵ . Pour tester cette hypothèse de façon plus stricte, il convient d'étudier le comportement des moments d'ordre k directement.

III.2. Moments de la distribution du carré des éléments de matrice

Le moment μ_k d'ordre k de la distribution du carré des éléments de matrice est défini comme $\mu_k = \langle y^k \rangle = N^k \langle |U_{ij}|^{2k} \rangle$. La méthode d'intégration invariant de [AL03] (cf. Annexe 1) permet d'obtenir leurs valeurs analytiques pour CUE

$$\mu_k = N^k F(k) = \frac{k! N^k (N-1)!}{(N+k-1)!}, \quad (4.15)$$

valeurs qui tendent vers $k!$ pour $N \rightarrow \infty$ et k fixé. Pour calculer les moments associés à UCE, de la même façon que pour la distribution de l , on moyenne aussi bien sur l'ensemble des réalisations aléatoire que sur les éléments dans la première colonne. Ainsi nous définissons le k ème moment pour UCE par

$$\tilde{\mu}_k = \frac{1}{n_r} \sum_{r=1}^{n_r} \frac{1}{N} \sum_{i=1}^N \left(y_{i1}^{(r)} \right)^k = \langle \langle y^k \rangle_C \rangle_R. \quad (4.16)$$

où $y_{i1}^{(r)}$ est donné par $N|U(l_{i1}^{(r)})|^2$ pour la i ème composante dans la première colonne de la r ème matrice. Comme mesure de déviation par rapport à CUE nous utilisons la déviation relative

$$D_{\mu_k} = \frac{|\tilde{\mu}_k - \mu_k|}{\mu_k}.$$

Nous avons calculé D_{μ_k} $k=2, 4$ et 8 . Pour des valeurs supérieures de k , le nombre de réalisations statistiques devient trop petit et une saturation similaire à celle observée pour D_P apparaît pour des niveaux trop grands, rendant les données inexploitable. Pour les deux derniers cas ($k=4$ et 8), les calculs sont faits avec $n_r = 10^5$ réalisations pour $n_q = 2, \dots, 14$ et $n_r = 5 \cdot 10^4$ réalisations pour la dernière taille de 15 qubits. La figure (4.7) montre le comportement de $D_{\mu_2}(n_q)$ pour $n_q = 2, \dots, 18$. Les courbes pour D_{μ_4} et D_{μ_8} ne sont pas montrées car elles sont similaires à celles pour D_{μ_2} à l'exception de la saturation apparaissant à un niveau plus élevé.

Sur la figure (4.8) est représenté le nombre n_ϵ nécessaire pour atteindre une petite valeur fixée de $D_{\mu_k} < \epsilon$ ceci pour les trois moments étudiés $k=2, 4, 8$, pour différentes valeurs de $\ln(\epsilon)$ et comparé au modèle linéaire f_1 introduit précédemment. Cette figure illustre bien que $n_\epsilon(n_q)$ est très bien décrit par un modèle linéaire en n_q , avec des valeurs pour les taux de croissance très similaires, pour toutes les valeurs de k considérées. Nous pouvons aussi comparer les

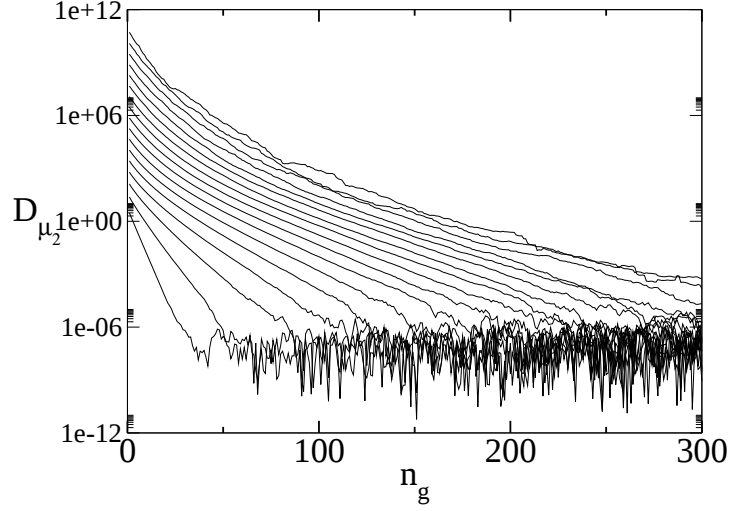


FIGURE 4.7.: $D_{\mu_2}(n_g)$ pour un nombre de qubits n_q variant entre 2 et 18 (de gauche à droite).

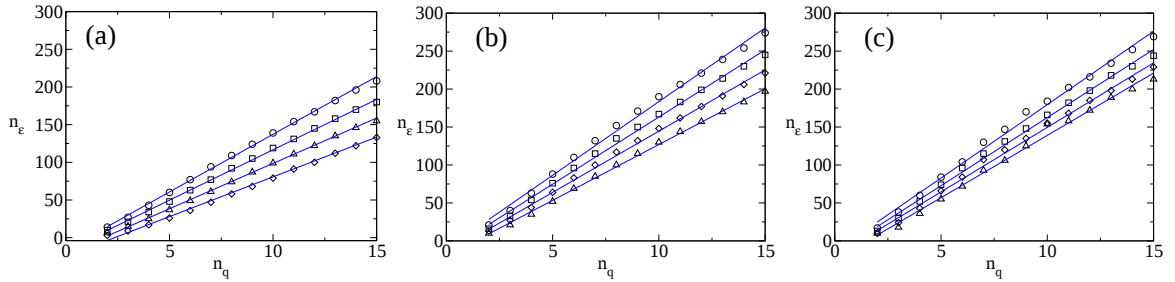


FIGURE 4.8.: Nombre de portes n_ϵ nécessaires pour atteindre $D_{\mu_k} \leq \epsilon$ en fonction du nombre de qubits, comparé au modèle f_1 . Les graphes (a), (b), et (c) correspondent respectivement à $k = 2, 4, 8$. Les différents symboles dans chaque graphe ($\triangle$, $\diamond$, $\square$, $\circ$) représentent différentes valeurs pour ϵ , avec $\ln(\epsilon) = 0.5, -1.5, -2.5$ et -3.5 pour $k = 2$, $\ln(\epsilon) = -1, -2, -3$ et -4 pour $k = 4$, et $\ln(\epsilon) = 1, 0, -1$ et -2 , respectivement pour $k = 8$.

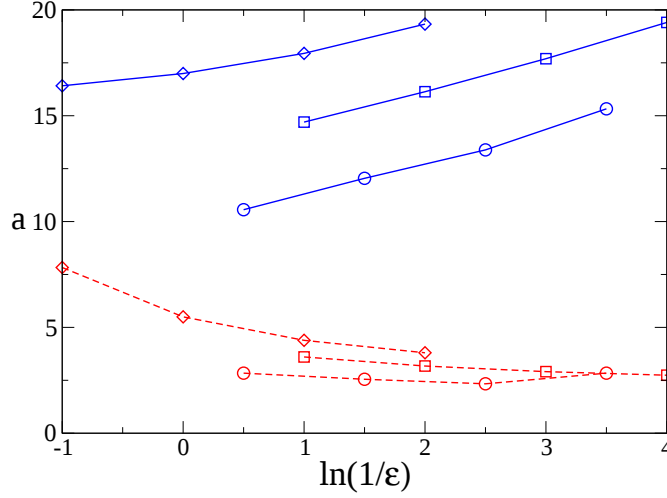


FIGURE 4.9.: Coefficients a_1 (lignes bleues continues) et a_2 (lignes rouges pointillées) pour la convergence de D_{μ_2} ($\circ$), D_{μ_4} ($\square$), et D_{μ_8} ($\diamond$) en fonction des valeurs accessibles pour $\ln(1/\epsilon)$.

données numériques avec le modèle f_2 , mais l'intervalle disponible en terme du nombre de qubit n_q est trop petit pour pouvoir décider quelle fonction entre f_1 et f_2 décrit le mieux la loi d'échelle recherchée. En fait, les données numériques pour $n_\epsilon(n_q)$ relatives à D_{μ_4} et D_{μ_8} montrent une légère courbure négative, qui rend le modèle f_2 moins bon que f_1 . Cependant, le modèle f_2 reproduit mieux la dépendance en terme de ϵ . Ceci est visible sur la figure (4.9) où sont rassemblés les coefficients a_1 et a_2 pour tous les moments considérés. On voit nettement que la dépendance en ϵ est correctement reproduite par la fonction $n_q \ln(n_q/\epsilon)$: a_2 devient pratiquement indépendant de ϵ pour des valeurs suffisamment petites de ce paramètre et ceci quelque soit la valeur de k considérée. Il est intéressant de noter que la vitesse de convergence des trois moments étudiés se trouve être plus ou moins la même. Il convient de préciser que la légère courbure négative qui tend à invalider la modélisation de $n_\epsilon(n_q)$ par f_2 pour D_{μ_4} et D_{μ_8} , n'est qu'un artefact numérique relié au niveau de saturation qui est d'autant plus grand que l'ordre k du moment est grand. De ce fait, les valeurs de ϵ choisies sont plus proches de ce niveau que pour le cas de figure relatif à D_{μ_2} et il en suit que la valeur de $n_\epsilon(n_q)$ est légèrement sous-estimée.

On pourrait aussi se demander si cette légère courbure négative n'est pas un effet de la procédure de moyenne sur la première colonne de la matrice. En effet, alors que pour CUE tous les éléments de matrice sont équivalents dans le sens où $\langle |U_{ij}|^{2k} \rangle$ est indépendant de i et j , et où moyenner sur une colonne est équivalent à moyenner sur un ensemble indépendant de réalisations, ceci n'est pas le cas pour un circuit CUE de taille n_q fini. Par exemple, après seulement une porte, le premier élément U_{11} (où l'indice 1 renvoie à l'état $|0 \dots 0\rangle$ dans la base computationnelle) n'est jamais affecté par une porte CNOT, alors que les autres le sont. Un des effets de la convergence de UCE vers CUE est que ce type d'inhomogénéités disparaît au fur et à mesure que n_q augmente. On pourrait suspecter que moyenner sur la première colonne contribue aussi à réduire ces inhomogénéités et produit de ce fait un mécanisme susceptible d'accélérer de manière effective la convergence observée. En effet comme le nombre de terme entrant dans

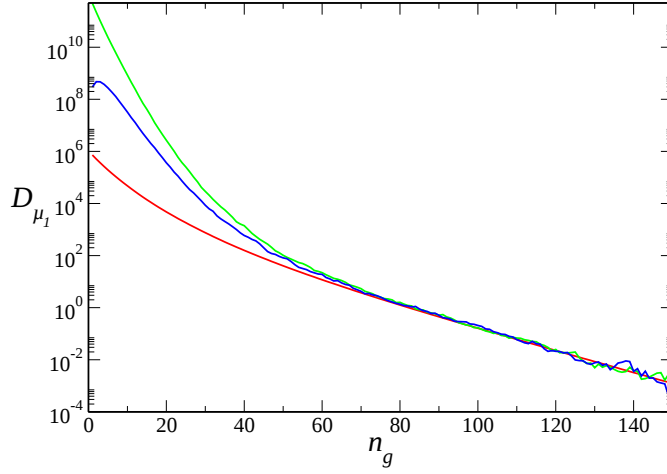


FIGURE 4.10.: Comparaison entre D_{μ_1} (rouge), D'_{μ_1} pour l'élément U_{11} (vert) et D'_{μ_1} pour l'élément U_{31} (bleu). Dans tous les cas, $n_q = 10$ et $n_r = 10^5$.

la moyenne sur la première colonne augmente exponentiellement avec le nombre de qubit, les petites différences dans $\langle |U_{ij}|^{2k} \rangle$ sont rapidement lissées, ce qui suggère une convergence plus rapide que la quantité calculée pour un seul élément de matrice. De plus, cet effet est d'autant plus prononcé pour de grands moments (grand k) pour lesquels les petites différences sont amplifiées.

Pour tester cette hypothèse, les mêmes types de calculs que précédemment peuvent être reproduits pour certaines valeurs des paramètres (en l'occurrence pour $n_q \leq 10$, $n_r = 10^4$) dans le cas du premier, du deuxième, du quatrième et du huitième moment μ'_1 , μ'_2 , μ'_4 et μ'_8 mais ceci pour un élément de matrice donnée sans moyenne sur la première colonne. Nous avons choisi les éléments U_{11} et U_{31} pour calculer les moments définis comme dans (4.16) mais en moyennant seulement sur les réalisations. Pour des valeurs plus grandes de n_q , ce calcul sans la moyenne sur la colonne est malheureusement au delà de nos capacités numériques. Les données correspondantes D'_{μ_k} pour l'élément U_{11} et U_{31} commencent à des valeurs plus grandes que D_{μ_k} (pour $k=1, 2, 4$ et 8). Cependant ils décroissent rapidement pour finalement atteindre leur analogues CUE pour des faibles valeurs de n_g (cf. la figure (4.10) pour $k=1$ et $n_q = 10$). Ainsi quand n_g est suffisamment grand, les deux courbes D'_{μ_k} et D_{μ_k} deviennent indiscernables l'une de l'autre à la précision numérique considérée. Ainsi, moyennner sur une colonne est une procédure parfaitement légitime qui n'entraîne pas de changement significatif dans le comportement de $n_\epsilon(n_q)$, du moins en ce qui concerne les moments d'ordre étudiés. D'un autre côté, on peut vérifier si D_{μ_2} et D_P subissent une légère courbure négative de leur $n_\epsilon(n_q)$ en choisissant une valeur de ϵ très proche du niveau de saturation. En effet, la qualité des fits par f_1 et f_2 est détériorée en pratiquant cette diminution de ϵ et des modèles du type $n_\epsilon(n_q) \simeq \sqrt{n_q}$ semble même très bien fonctionner. Bien sûr, d'un point de vue physique, un comportement sublinéaire paraît totalement impossible car cela signifierait que l'état global d'un grand circuit quantique pourrait s'équilibrer avant même que tout les qubits soient touchés par au moins une porte. Tous ces éléments confirment l'idée selon laquelle la légère courbure négative observée sur les données est bien un artefact numérique.

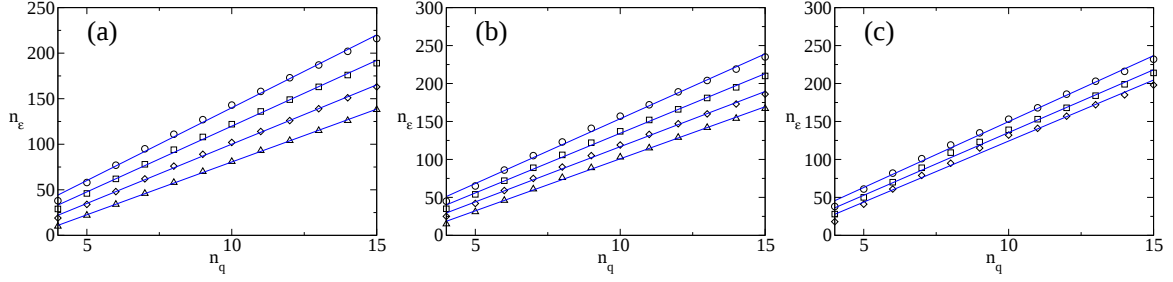


FIGURE 4.11.: Nombre de portes n_ϵ nécessaires pour atteindre $D_{c_k} \leq \epsilon$ en fonction du nombre de qubits, comparé au modèle f_1 . Les graphes 1, 2 et 3 correspondent respectivement à $k = 2, 4, 8$. Pour $k = 2$, $\ln(\epsilon) = -1, -2, -3$ et -4 ($\triangle, \diamond, \square, \circ$). Pour $k = 4$, $\ln(\epsilon) = 0, -1, -2$ et -3 ($\triangle, \diamond, \square, \circ$). Pour $k = 8$, $\ln(\epsilon) = 1, 0$ et -1 ($\diamond, \square, \circ$, respectivement).

Le message le plus important provenant des figures (4.8) et (4.9) est que

1. Tous les moments considérés convergent à la même vitesse.
2. Le nombre de portes nécessaires pour atteindre une précision donnée augmente en bonne approximation linéairement avec le nombre de qubits.
3. La dépendance additionnelle en ϵ est bien modélisée par un comportement de $n_\epsilon(n_q)$ du type $n_q \ln(n_q/\epsilon)$.

III.3. Corrélations entre les éléments de matrice

Même pour CUE, des éléments de matrice différents ne sont pas distribués indépendamment (contrairement au cousin hermitien de CUE qu'est GUE). Une des raisons évidentes pour l'existence de ces corrélations est la contrainte d'orthonormalisation pour les lignes et les colonnes d'une matrice unitaire ($\sum_i |U_{ij}|^2 = \sum_j |U_{ij}|^2 = 1$). Nous définissons les corrélations entre k éléments différents y_{ij} pour une même colonne j par $c_k = \langle \prod_{i=1}^k y_{ij} \rangle = N^k \langle \prod_{i=1}^k |U_{ij}|^2 \rangle$, où la moyenne est faite sur l'ensemble considéré. En ce qui concerne CUE, la méthode d'intégration invariante de [AL03] donne

$$c_k = \frac{N^k (N-1)!}{(N+k-1)!}, \quad (4.17)$$

qui ne diffère de l'expression de μ_k que par un facteur $\frac{1}{k!}$. Ainsi pour k petit, les corrélations sont importantes et comparables aux moments du même ordre.

A nouveau pour UCE, nous définissons les quantités correspondantes en moyennant aussi bien sur la 1^{ère} colonne que sur l'ensemble de réalisations. Cependant pour ne pas créer de corrélations artificielles, chaque élément n'apparaît que dans un seul et unique produit,

$$\tilde{c}_k = \frac{1}{n_r \lfloor \frac{N}{k} \rfloor} \sum_{r=1}^{n_r} \sum_{i=1}^{\lfloor \frac{N}{k} \rfloor} \prod_{j=1}^k y_{(k i - k + j)1}^{(r)}. \quad (4.18)$$

Ici $[x]$ est la partie entière de x . Pour deux qubits ($N=4$) on a par exemples :

$$\begin{aligned} \tilde{c}_2 &= \frac{1}{2n_r} \sum_{r=1}^{n_r} \left(y_{11}^{(r)} y_{21}^{(r)} + y_{31}^{(r)} y_{41}^{(r)} \right), \\ \tilde{c}_4 &= \frac{1}{n_r} \sum_{r=1}^{n_r} \left(y_{11}^{(r)} y_{21}^{(r)} y_{31}^{(r)} y_{41}^{(r)} \right). \end{aligned} \quad (4.19)$$

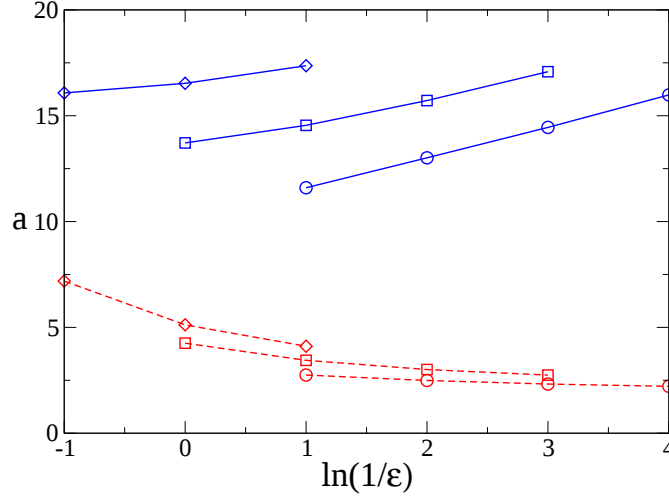


FIGURE 4.12.: Coefficients a_1 (lignes bleues continues) et a_2 (lignes rouges discontinues) pour la convergence de D_{c_2} (o), D_{c_4} (□) et D_{c_8} (◇) en fonction des valeurs accessibles pour $\ln(1/\epsilon)$.

De la même manière que pour les moments, nous définissons la distance par rapport à CUE comme la déviation relative de $\tilde{c}_k$ par rapport à la valeur CUE

$$D_{c_k} = \frac{|\tilde{c}_k - c_k|}{c_k}. \quad (4.20)$$

La figure (4.11) représente les résultats pour l'évolution de $n_\epsilon(n_q)$ (défini de la même manière que pour les moments) dans les cas $k = 2$, $k = 4$ et $k = 8$ comparés au modèle f_1 . Le comportement est essentiellement linéaire et très similaire quelque soit l'ordre k considéré. Les données sont aussi bien fittées par la fonction f_2 mais une fois encore, l'intervalle restreint en nombre de qubits ne permet pas de trancher véritablement entre ces deux modèles possibles. Généralement la fonction f_2 fonctionne moins bien que la fonction f_1 , mais ceci est aussi dû à l'artefact numérique qui engendre une légère courbure négative de $n_\epsilon(n_q)$. De plus, on voit sur la figure (4.12) représentant les coefficients a_1 et a_2 des deux modèles pour les trois cas $k = 2$, $k = 4$ et $k = 8$, que la dépendance en ϵ est correctement décrite par $n_q \ln(n_q/\epsilon)$ dans le sens où le coefficient a_2 est essentiellement indépendant de l'ordre des corrélations pour un ϵ suffisamment petit. De plus, en comparant les figures (4.12) et (4.9), nous voyons que les corrélations c_k convergent typiquement à la même vitesse que les moments μ_k de même ordre, un résultat qui est attendu d'après la théorie sur les k -designs [HL08, DCEL06]. En fait, d'après la définition citée au début du chapitre, un k -design unitaire est caractérisé par le fait que n'importe quel polynôme des éléments de matrice complexes de degrés (m, l) avec $m, l \leq k$ (tels ceux intervenant dans la définition des moments μ_k et des corrélations c_k), ont la même moyenne sur le design unitaire que sur le groupe unitaire dans son ensemble [DCEL06].

IV. Conclusion partielle

Dans ce chapitre, nous avons analysé la convergence de l'ensemble de circuits unitaires (UCE) vers celui de l'ensemble unitaire circulaire. Cette étude est faite en comparant les grandeurs suivantes que sont

- La distribution de la grandeur $l = N|U_{ij}|^2$ reliée à la distribution du module carré des éléments de matrice,
- Les moments de la distribution des modules carrés des éléments de matrice, jusqu'à $\langle |U_{ik}|^{16} \rangle$,
- Les corrélations entre modules carrés des éléments d'une même colonne jusqu'à 16 facteurs $|U_{ik}|$.

Ces quantités ont été calculées analytiquement pour CUE et numériquement pour UCE, en simulant des circuits quantiques jusqu'à 28 qubits pour la distribution de l , 18 qubits pour les moments et 15 qubits pour les corrélations. Nous avons ainsi pu montrer que ces quantités pour CUE sont efficacement reproduites avec une précision ϵ par des circuits quantiques UCE contenant un nombre de portes qui augmente avec le nombre de qubits au plus comme $n_\epsilon \leq C n_q \ln(n_q/\epsilon)$, où C est une constante positive. Une telle convergence reste cependant surprenante puisque l'ensemble de portes de UCE, ayant un gap spectral valant $\Delta \simeq 0.232703$ non-nul, devrait converger avec un nombre de portes se comportant comme $n_\epsilon \leq C n_q (n_q + \ln(1/\epsilon))$ (en limite supérieur) comme il est montré dans [HL08]. Notons dans le fond, il ne faut cependant pas perdre de vue que pour reproduire fidèlement la distribution jointe de probabilités dans son entier, il est nécessaire d'utiliser des circuits quantiques contenant un nombre de portes qui augmente exponentiellement avec le nombre de qubits. Nos résultats suggèrent donc que les quantités statistiques qui ne sont pas reproduites efficacement sont de nature plus sophistiquées que les moments de faible ordre et les corrélations de faible ordre étudiés dans ce chapitre. Précisons que depuis la fin de nos travaux, un résultat analytique important a été publié dans [BV09]. Ce résultat stipule qu'un circuit quantique composé de $\mathcal{O}(n_q)$ porte $U(4)$ distribuées de manière uniforme, reproduit efficacement tout les moments μ_k de CUE. Pour conclure, il est important d'insister sur l'utilité de la méthode de production de matrices aléatoires décrite dans ce chapitre. L'algorithme classique utilisant la paramétrisation de Hurwitz [Hur97, PZK98] nécessite typiquement N^2 multiplications de matrices $N \times N$ soit $\mathcal{O}(N^5)$ opérations arithmétiques élémentaires pour N grand. Bien sûr, étant basé sur la paramétrisation complète de $U(N)$, cet algorithme est parfait dans le sens où il reproduit exactement toutes les quantités attendues pour CUE. Cependant, du moment où on ne s'intéresse qu'à la reproduction de quantités statistiques efficacement reproduites, les circuits aléatoires permettent de créer des matrices de taille N (N étant une puissance de 2) en un nombre d'opération augmentant comme $\mathcal{O}(\log(N) \log(\log(N)))$. La seule difficulté reste l'ordinateur quantique obligatoire pour mettre en œuvre la méthode. Cependant, la simulation de circuits comme UCE sur un ordinateur classique, qui utilise typiquement $\mathcal{O}(n_q \log(N))$ produit de matrices, reste plus efficace que la paramétrisation de Hurwitz et permet, à condition de rester prudent, d'analyser des problèmes numériques difficilement accessible.

Conclusion générale et perspectives

Les résultats exposés dans cette thèse concernent deux thèmes précis dans le domaine de l'information quantique. Le premier thème est le concept d'interférence quantique regardé d'un point de vue de l'information quantique. Le chapitre 1 définit et expose l'idée selon laquelle l'interférence quantique pourrait jouer un rôle important dans le comportement quantique de l'information, en particulier être responsable de l'accélération des algorithmes quantiques. Cette thèse ne répond pas directement à la question de savoir si l'interférence est vraiment nécessaire pour engendrer cette *supériorité* quantique. Elle stipule modestement que l'interférence quantique est une quantité mesurable qui caractérise l'évolution d'un système quantique. Plus précisément dans le chapitre 2, il est démontré que si on choisit un algorithme quantique agissant sur un nombre suffisamment grand de qubits, le choix se faisant au hasard, sans connaissances particulières sur sa structure interne ou sur le rôle qu'il est sensé rendre, alors il est fort probable que cet algorithme contient une quantité d'interférence proche de la valeur maximale admise. Dans le chapitre 3 est étudié l'influence d'un environnement décohérent de taille donnée sur l'interférence présente dans l'évolution d'un système quantique. Comme on s'y attend intuitivement, l'interférence quantique est détruite par la décohérence, d'autant plus que l'environnement est grand. Cependant, la manière dont cette destruction se manifeste n'est pas catastrophique, dans le sens où la diminution d'interférence se fait comme une loi de puissance de la taille de l'environnement. Parmi les perspectives concernant ce thème, les plus intéressantes sont sûrement celles concernant le lien potentiel entre intrication et interférence quantique ainsi que celle reliant directement le concept d'interférence et d'accélération quantique.

Le second thème concerne la possibilité de construire des ensembles de matrices aléatoires en utilisant les ressources de l'information quantique. Construire des matrices aléatoires est un problème essentiellement inefficace sur un ordinateur classique, car le temps de création augmente exponentiellement avec la taille des matrices à créer. Les résultats numériques présentés dans le chapitre 4 montrent que si on considère certaines quantités qui caractérisent bien l'ensemble de matrices aléatoires CUE, alors en créant une séquence contenant un nombre polynomial de portes quantiques, ces portes étant choisies au hasard dans l'ensemble universel $\{U(2), \text{CNOT}\}$, alors on peut imiter ces quantités avec une précision arbitraire, et ce faisant obtenir des ensembles de matrices aléatoires ressemblant fortement à CUE.

Pour conclure, il convient de dire un mot sur le concept mathématique qui est le *ciment* de cette thèse, celui de matrices aléatoires. Le chapitre 2 et 3 l'utilisent comme base pour tenter de répondre à des questions motivées par la physique de l'information quantique. Les analyses numériques poussées aussi bien que les techniques analytiques telles que l'intégration invariante permettent de cerner la puissance de cet outil. Ainsi, il est particulièrement amusant de constater, comme il est évoqué au chapitre 4, que la physique puisse rendre la monnaie de

sa pièce en permettant la production efficace d'ensembles de matrices aléatoires.

Appendices

Méthode d'intégration invariante

I. Introduction

La méthode d'intégration invariante introduite par Aubert et Lam dans [AL03] est une méthode qui permet de calculer des intégrales de monômes d'éléments de matrices unitaires sur le groupe $U(N)$, c'est à dire des intégrales de la forme :

$$\int dU U_{i_1 j_1}^* \dots U_{i_p j_p}^* U_{k_1 l_1} \dots U_{k_q l_q} \equiv \langle U_{i_1 j_1}^* \dots U_{i_p j_p}^* U_{k_1 l_1} \dots U_{k_q l_q} \rangle \quad (\text{A.1})$$

où dU est la mesure de Haar invariante normalisée à $\int dU = 1$, les U_{ij} sont des composantes de matrice d'un élément de $U(N)$ et les U_{ij}^* sont leurs complexes conjugués. Ce type d'expressions est rencontré dans les chapitres 2, 3 et 4 de cette thèse, mais on peut les retrouver dans de nombreux domaines allant de la gravité quantique en 2 dimensions à la QCD en passant par divers problèmes de physique statistique et de matière condensée. Il existe des formules explicites (voir dans [AL03]) permettant de calculer les intégrales de la forme (A.1) mais leurs utilisations se révèlent souvent difficiles.

II. Propriétés et écriture diagrammatique

II.1. Notations et propriétés générales

Il convient de remarquer que (A.1) dépend des indices $I = (i_1 i_2 \dots i_p)$, $J = (j_1 j_2 \dots j_p)$, $K = (k_1 k_2 \dots k_q)$ et $L = (l_1 l_2 \dots l_q)$ et on peut donc introduire la notation condensée

$$\langle I, J | K, L \rangle \equiv \int dU U_{IJ}^* U_{KL} \quad (\text{A.2})$$

avec $U_{IJ}^* = \prod_{a=1}^p U_{i_a j_a}^*$ et $U_{KL} = \prod_{a=1}^q U_{k_a l_a}$. Comme les composantes U_{ij} commutent entres-elles, on a par exemple l'égalité $U_{IJ}^* = U_{I_P J_P}^*$ où I_P et J_P sont obtenues par une permutation $P \in S_p$ de leurs p indices. Il suit

$$\langle I, J | K, L \rangle = \langle I_P, J_P | K_Q, L_Q \rangle \quad (\text{A.3})$$

pour tout $P \in S_p$ et tout $Q \in S_q$.

L'invariance de la mesure de Haar dU entraîne les relations suivantes qui servent de base à la méthode d'intégration invariante :

$$\int dU f(U, U^*) = \int dU f(U^*, U) = \int dU f(U^T, U^{*T}) \quad (\text{A.4})$$

$$= \int dU f(VU, V^*U^*) = \int dU f(UV, U^*V^*) \quad (\text{A.5})$$

pour toute fonction f et pour tout $V \in U(N)$. En choisissant un V particulier on obtient trois propriétés fondamentales pour l'intégrale (A.1) :

- 1 Si V est une matrice de permutation, VU est obtenue en permutant les lignes de U tandis que UV est obtenue en permutant ses colonnes. Ainsi

$$\langle I, J | K, L \rangle = \langle I', J | K', L \rangle = \langle I, J' | K, L' \rangle \quad (\text{A.6})$$

où I' et K' sont obtenus par le même réassignement des valeurs des indices de I et de K (même chose pour J et L). Cette propriété très importante stipule que l'intégrale (A.1) ne dépend pas de la valeur précise des indices mais seulement de leurs relations.

- 2 Si V est de la forme $V_{ij} = e^{i\phi} \delta_{ij}$, on a $f(VU, V^*U^*) = e^{i\phi(q-p)} f(U, U^*)$. Ainsi l'intégrale est non-nulle ssi la phase s'annule c'est à dire si $p = q$. On appelle p l'ordre de l'intégrale.
- 3 Si V est de la forme $V_{ij} = e^{i\phi_i} \delta_{ij}$, on a $f(UV, U^*V^*) = e^{i\xi} f(U, U^*)$ où $\xi = \sum_{a=1}^p (\phi_{l_a} - \phi_{j_a})$ et donc l'intégrale est non-nulle ssi les phases ϕ_l et ϕ_j s'annulent 2 à 2 c'est à dire si $L = J_R$ pour $R \in S_P$. De même on démontre que l'intégrale est non-nulle ssi $K = I_M$ pour $M \in S_P$. En utilisant (A.3) et en posant $Q = RM^{-1}$ il découle que les seules intégrales non-nulles sont de la forme :

$$\langle I, J | I, J_Q \rangle = \int dU U_{IJ}^* U_{IJ_Q} \quad (\text{A.7})$$

pour tout $Q \in S_P$. Désormais nous ne considérerons que des intégrales de cette forme. Notons que dans [AL03] une intégrale telle que $J = J_Q$, est qualifiée de *directe* par opposition aux intégrales dites d'*échange* quand Q diffère de la permutation identité.

Pour mieux cerner ces différentes propriétés voici les exemples suivants :

- $\langle |U_{11}|^2 \rangle = \langle |U_{12}|^2 \rangle = \langle |U_{49}|^2 \rangle = \langle |U_{ij}|^2 \rangle \forall i, j$ (propriété 1 et 2).
- $\langle U_{ij} \rangle = \langle U_{ij}^* \rangle = \langle |U_{ij}|^2 U_{kl} \rangle = \langle U_{kl}^* |U_{ij}|^2 \rangle = 0 \forall i, j, k, l$ (propriété 2).
- $\langle U_{12}^* U_{21}^* U_{11} U_{22} \rangle \neq 0$ mais $\langle U_{11}^* U_{12}^* U_{22} U_{21} \rangle = 0$ (propriété 3).

Ces exemples sont écrits avec la notation en terme d'éléments de matrice. Avec la notation (A.2) ces mêmes exemples s'écrivent :

- $\langle (1), (1) | (1), (1) \rangle = \langle (1), (2) | (1), (2) \rangle = \langle (4), (9) | (4), (9) \rangle = \langle (i), (j) | (i), (j) \rangle \forall i, j$.
- $\langle () () | (i), (j) \rangle = \langle (i), (j) | () () \rangle = \langle (i), (j) | (ik), (jl) \rangle = \langle (ik), (jl) | (i), (j) \rangle = 0 \forall i, j, k, l$.
- $\langle (12), (21) | (12), (12) \rangle \neq 0$ mais $\langle (11), (12) | (22), (21) \rangle = 0$.

II.2. Écriture diagrammatique

Les trois propriétés présentées précédemment permettent une écriture diagrammatique comme sur la figure (A.1) où est représentée l'intégrale $\langle I, J | I, J_Q \rangle$ pour $I = (12334555) = (123^2 45^3)$, $J = (1^5 23)$ et $J_Q = (121^4 32)$. Les points sur la gauche représentent de bas en haut les indices 1, 2, 3, 4 et 5 de I , et ceux de droite représentent les indices 1, 2 et 3 de J : Ainsi les ensembles

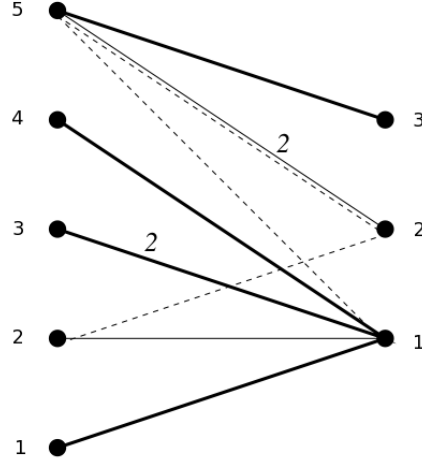


FIGURE A.1.: Représentation diagrammatique de l'intégrale $\langle I, J | I, J_Q \rangle$ pour $I = (123^245^3)$, $J = (1^523)$ et $J_Q = (121^432)$. Ce diagramme peut être obtenu en considérant l'intégrale $\langle |U_{11}|^2 |U_{31}|^4 |U_{41}|^2 |U_{53}|^2 U_{21}^* U_{22} U_{51} U_{52} U_{52}^* \rangle$ (Source [AL03])

I et J sont représentés par des points (ou vertex), chacun représentant les indices i_a et j_b mais leur valeur n'important pas, d'après la première propriété. Les points de I sont reliés aux points de J par l'intermédiaire de lignes fines continues représentant les termes U_{ij}^* et par des lignes pointillées représentant les termes U_{ij} sachant que comme les termes peuvent apparaître avec une certaine puissance, les lignes peuvent être munies d'une certaine multiplicité. Pour ne pas surcharger l'écriture, une seule ligne est dessinée et la multiplicité est écrite à sa proximité. Une autre simplification concerne les termes du type $|U_{ij}|^{2m}$: ils sont représentés par une ligne continue épaisse assignée de la multiplicité m et de manière générale les termes du type $U_{ij}^{*m} U_{ij}^n$ sont représentés par une ligne continue épaisse assignée de la multiplicité (m, n) . Avec cette convention d'écriture, la deuxième propriété entraîne que l'intégrale est non-nulle ssi le nombre de lignes fines continues est égal au nombre de lignes pointillées. La troisième propriété se traduit par le fait que l'intégrale est non-nulle ssi pour chaque point de I et de J le nombre de lignes fines continues et le nombre de lignes pointillées en émergeant sont égaux. En occurrence avec la notation diagrammatique on a par exemple :

$$\langle U_{11} U_{22} U_{12}^* U_{21}^* \rangle = \begin{array}{c} 2 \\ 1 \end{array} \begin{array}{c} \text{---} \\ \diagdown \quad \diagup \\ \text{---} \end{array} \begin{array}{c} 2 \\ 1 \end{array} \neq 0 \quad (\text{A.8})$$

$$\langle U_{11} U_{22}^* U_{12} U_{21}^* \rangle = \begin{array}{c} 2 \\ 1 \end{array} \begin{array}{c} \text{---} \\ \diagup \quad \diagdown \\ \text{---} \end{array} \begin{array}{c} 2 \\ 1 \end{array} = 0 \quad (\text{A.9})$$

III. Calculs des intégrales

Voyons maintenant comment on peut obtenir explicitement des expressions pour (A.1). À proprement parler, la méthode d'intégration invariante ne permet pas d'obtenir une forme directement utilisable pour une valeur donnée de I , J et J_Q . Mais elle permet d'obtenir des relations entre des intégrales de même ordre, et entre des intégrales d'ordre p et d'ordre $p + 1$. Ainsi elle permet d'arriver à une expression finale en procédant de proche en proche à partir de l'intégrale d'ordre 0 normalisée $\int dU = \langle 1 \rangle = 1$. Pour pouvoir trouver ces relations, la méthode

invariante se base d'une part sur l'invariance de la mesure dU illustrée par les formules (A.4), et d'autre part sur la propriété d'unitarité $UU^\dagger = \mathbf{1}$ d'un élément de $U(N)$.

III.1. Rotation

Nous avons déjà vu comment l'invariance de la mesure (A.4) permet d'obtenir diverses propriétés importantes. Voici comment on peut encore l'exploiter quand V est une rotation $SO(2)$ et ainsi obtenir une première relation entre intégrales de même ordre. Choisissons V comme étant une matrice de rotation dans le plan (ab) d'un angle ξ . On peut montrer que le passage de U à VU s'écrit en composantes :

$$\begin{aligned} U_{ia} &\rightarrow cU_{ia} + sU_{ib} \\ U_{ib} &\rightarrow -sU_{ia} + cU_{ib} \\ U_{ij} &\rightarrow U_{ij} \end{aligned} \tag{A.10}$$

en posant $c = \cos(\xi)$ et $s = \sin(\xi)$. Ainsi l'intégrale $\langle I, J | I, J_Q \rangle$ devient une somme de termes du genre $M_e(c^2)^{d-e}(s^2)^e$, les puissances impaires de c et s s'annulant entre elles d'après la propriété 3. Ici d est le nombre total d'indices de colonne dans U_{IJ} et où a, b et e vont de 0 à d . L'invariance (A.4) entraîne :

$$\langle I, J | I, J_Q \rangle = \sum_{e=0}^d M_e(c^2)^{d-e}(s^2)^e$$

Ceci ne peut être vrai pour tout ξ que si

$$M_e = \langle I, J | I, J_Q \rangle \begin{pmatrix} e \\ d \end{pmatrix} = M_0 \begin{pmatrix} e \\ d \end{pmatrix}$$

où $\begin{pmatrix} e \\ d \end{pmatrix} = d!/e!(d-e)!$ sont les coefficients binomiaux. Ensuite l'écriture diagrammatique permet de calculer les termes M_e de la manière suivante en considérant une rotation précise d'un angle $\xi = \pi/2$ pour lequel les relations (A.10) se simplifie en $U_{ia} \rightarrow U_{ib}$, $U_{ib} \rightarrow -U_{ia}$ et $U_{ij} \rightarrow U_{ij}$. D'un point de vue diagrammatique, si on considère 2 points sur la partie droite d'un diagramme quelconque, et correspondant aux valeurs a et b , alors ces relations traduisent le déplacement de lignes entre les points a et b : Un des points doit être connecté à certaines lignes mais l'autre peut ne pas l'être, c'est à dire que ce point peut être externe au diagramme. Le nombre total de lignes fines continues (ou lignes pointillées) attachées au point est d . Maintenant déplaçons e lignes fines continues et e lignes pointillées entre les deux points considérés sans oublier la contrainte qui impose qu'après le déplacement, le nombre de lignes fines continues doit équaler le nombre de lignes pointillées sur chaque point (sinon l'intégrale s'annule). Assignons le poids $+1$ pour chaque ligne déplacée de a vers b et le poids -1 pour chaque ligne déplacée de b vers a . La quantité M_e est simplement la somme de toutes les intégrales obtenues après le mouvement, pondérée par les facteurs ± 1 . Il est important de bien noter que ces relations sont locales et que de ce fait seul importe de considérer les indices a et b de J et J_Q , indépendamment des indices de I . La figure (A.2) illustre graphiquement cette procédure : appliquons une rotation au diagramme (a). On peut de cette façon bouger soit deux lignes ($e = 1$), soit quatre lignes ($e = 2$). En bougeant deux lignes on obtient le diagramme (b-deux lignes fines continues), le diagramme (c-deux lignes pointillées), et deux autres, le diagramme (e-une ligne fine continue

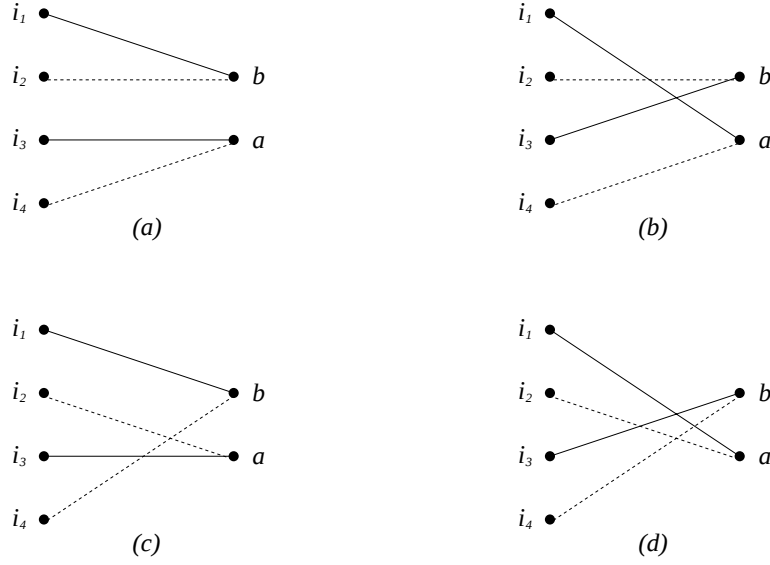


FIGURE A.2.: Illustration diagrammatique de la manière dont l'invariance par rotation de la mesure peut être utilisée pour relier des intégrales entre elles. Voir eq. (A.11) (Source [AL03]).

et une ligne pointillée de a à b) et (f-une ligne fine continue et une ligne pointillée de b à a). Les deux derniers diagrammes (e) et (f) ne sont pas dessinés, mais ils peuvent être obtenus à partir du diagramme (a) en fusionnant les points a et b . De cette manière on obtient $M_1 = -I(b) - I(c) + I(e) + I(f) = 2(-I(b) + I(e))$, où $I(i)$ dénote l'intégrale associée au diagramme i et où (A.3) est utilisée lors de la dernière étape. Grâce à (A.11), on arrive à $M_1 = 2M_0 = 2I(a)$. Finalement on a la relation

$$I(a) = -I(b) + I(e). \quad (\text{A.11})$$

Si on déplace les quatre lignes, on obtient $M_2 = I(d)$. La formule dans (A.11) impose donc $M_2 = M_0$, ou $I(d) = I(a)$. Ceci n'a rien d'étonnant d'après la propriété d'invariance par permutation (A.3).

III.2. Unitarité

Voici un autre moyen, utilisant l'unitarité $UU^\dagger = \mathbf{1}$, permettant de relier des intégrales entres elles, en particulier des intégrales directes de degrés p avec des intégrales directes de degrés $p-1$. En composante l'unitarité s'écrit :

$$\sum_{j=1}^N U_{ij} U_{kj} = \sum_{j=1}^N U_{ji}^* U_{jk} = \delta_{ik} \quad (\text{A.12})$$

A titre d'exemple, voyons comment obtenir l'expression pour l'intégrale $I_{ij} = \langle |U_{ij}|^2 \rangle$. Comme l'intégrale ne dépend pas des indices i et j , en sommant sur l'un d'entre eux, disons l'indice j on obtient :

$$\begin{aligned} \sum_{j=1}^N I_{ij} &= \sum_{j=1}^N \langle |U_{ij}|^2 \rangle = \sum_{j=1}^N \langle U_{ij}^* U_{ij} \rangle = \langle \sum_{j=1}^N U_{ij}^* U_{ij} \rangle = \langle \delta_{ii} \rangle = 1 \\ N \langle |U_{ij}|^2 \rangle &= 1 \implies \langle |U_{ij}|^2 \rangle = \frac{1}{N} \end{aligned}$$

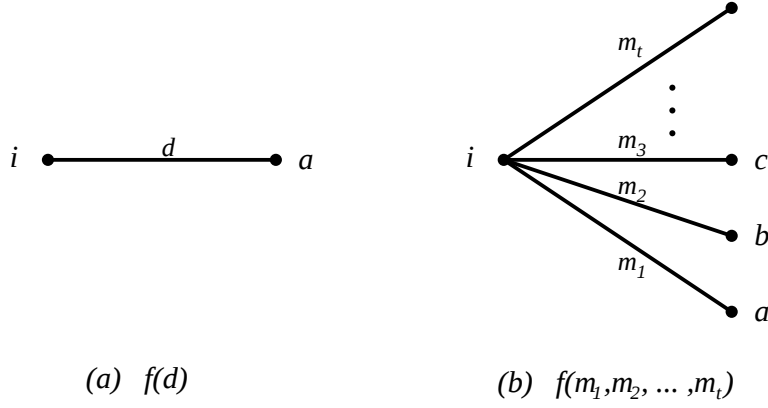


FIGURE A.3.: L'invariance par rotation permet d'étendre l'unique branche (a) vers un ensemble de branches (b) en accord avec la relation (A.14) (Source [AL03]).

En notation diagrammatique, en spécifiant l'indice de la somme par une flèche, les mêmes relations s'écrivent :

$$\sum \text{---} \leftarrow = 1$$

$$N \text{---} = 1 \implies \text{---} = \frac{1}{N}$$

Notons que la diminution de l'ordre p provient de la disparition d'un terme $|U|^2$ lors de l'opération de sommation. Dans [AL04] ceci se généralise à toutes intégrales directes de degré p et se traduit en notation diagrammatique par le fait qu'une somme sur un indice fait disparaître la ligne en gras qui est reliée au point correspondant. Pour les intégrales d'échanges, sommer sur un indice donne zéro. Par exemple on à

$$\sum \rightarrow \text{---} \times \text{---} = (N-1) \text{---} \times \text{---} + \text{---} \text{---} = 0 \quad (\text{A.13})$$

IV. Relations entre intégrales et formules explicites

Voici une liste des relations entre intégrales démontrées dans [AL03, AL04].

IV.1. La fan relation

Introduite dans [AL03], la fan relation permet de relier des intégrales directes de même ordre entre elles. Grâce à la propriété de rotation, on peut relier l'intégrale dont l'intégrand est proportionnel à $|U_{ia}|^{2d}$ à l'intégrale dont l'intégrand est proportionnel à $|U_{ia}|^{2m_1}|U_{ib}|^{2m_2}\dots|U_{iz}|^{2m_t}$. La figure (A.3) permet de voir que l'on passe de la première intégrale à la deuxième en *éclatant* l'unique branche de multiplicité d en t branches distinctes de multiplicités inférieures. Notons que toutes les branches sont reliées au point i , point pouvant être lui-même relié à d'autres lignes non-considérées. En dénotant la première intégrale par $f(d)$, la seconde par $f(m_1, m_2, \dots, m_t)$ et en utilisant de manière répétée la relation (A.11) on obtient la *fan relation*

$$f(m_1, m_2, \dots, m_t) = \frac{(\prod_{i=1}^t m_i!)}{(\sum_{i=1}^t m_i)!} f(d), \quad (\text{A.14})$$

où $d = \sum_{i=1}^t m_i = d$ et $f(d) = f(d, 0, 0, \dots, 0)$ (voir [AL03] pour une démonstration plus détaillée).

IV.2. La double fan relation

Introduite dans [AL04] comme une généralisation de la fan relation, la double fan relation permet de relier des intégrales directes et des intégrales d'échange du même ordre. Elle relie l'intégrale représentée par le digramme (A.4.a) à l'intégrale représentée par le diagramme (A.4.b). Notons que comme dans le cas de la fan relation, les points a et b peuvent être reliés à d'autres lignes non représentées sur le dessin. La première intégrale est notée $[(m_a n_a)(m_b n_b)]$. La deuxième intégrale est notée $[A_a]^{\alpha'_a} [A_b]^{\alpha'_b} [B_a]^{\beta'_a} [B_b]^{\beta'_b}$ où les termes entre crochets correspondent aux motifs élémentaires représentés sur la figure (A.5). Par exemple, avec cette notation, l'intégrale d'échange (A.8) s'écrit :

$$\begin{array}{c} \diagup \quad \diagdown \\ \diagdown \quad \diagup \end{array} = [A_a][A_b] \quad (\text{A.15})$$

avec donc $(\alpha'_a, \alpha'_b, \beta'_a, \beta'_b) = (1, 1, 0, 0)$.

La double fan relation s'écrit

$$[(m_a n_a)(m_b n_b)] = \sum v(\alpha'_a, \alpha'_b, \beta'_a, \beta'_b) [A_a]^{\alpha'_a} [A_b]^{\alpha'_b} [B_a]^{\beta'_a} [B_b]^{\beta'_b} \quad (\text{A.16})$$

où

$$v(\alpha'_a, \alpha'_b, \beta'_a, \beta'_b) = \frac{m_a! n_a! m_b! n_b!}{\alpha'_a! \alpha'_b! \beta'_a! \beta'_b!},$$

la somme étant faite sur tous les indices $(\alpha'_a, \alpha'_b, \beta'_a, \beta'_b)$ solutions de :

$$\begin{aligned} m_a &= \alpha_a + \beta_a, & n_a &= \alpha_b + \beta_a, \\ m_b &= \alpha_b + \beta_b, & n_b &= \alpha_a + \beta_b, \end{aligned}$$

le nombre total de points du diagramme sur la colonne de droite du diagramme (A.4.b) valant $m_a + m_b = n_a + n_b$. A titre d'exemple, utilisons cette relation pour calculer l'intégrale $[(m_a n_a)(m_b n_b)]$ avec $m_a = n_a = m$ et $m_b = n_b = 0$. On a donc d'après (IV.2)

$$[(mm)(00)] = v(00m0)[B_a]^m = m![B_a]^m \quad (\text{A.17})$$

qui équivaut à la fan relation (A.14) pour tout les $m_i = 1$.

IV.3. Formules explicites

Les relations précédentes peuvent être utilisées pour construire des formules explicites pour certains types d'intégrales. La première intégrale explicite est l'intégrale $F(m) = \langle |U_{ij}|^{2m} \rangle$ représentée par le diagramme (A.3.a), le point de gauche n'étant pas relié à d'autres lignes.

$$F(m) = \frac{(N-1)! m!}{(N+m-1)!}$$

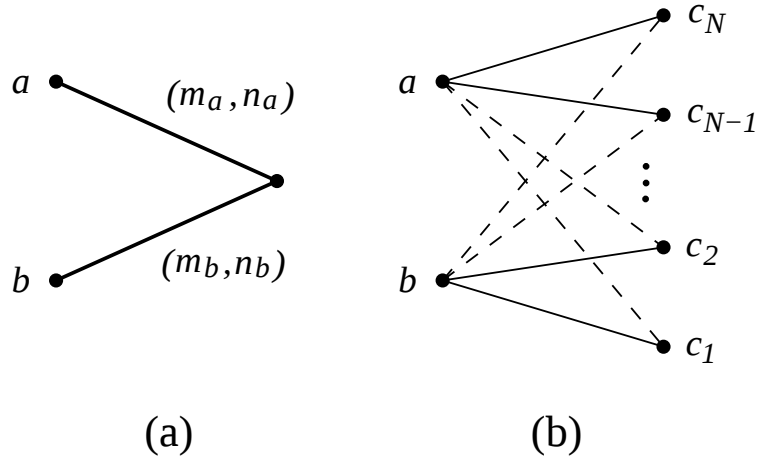


FIGURE A.4.: Intégrales reliées par la double fan relation (Source [AL04]).

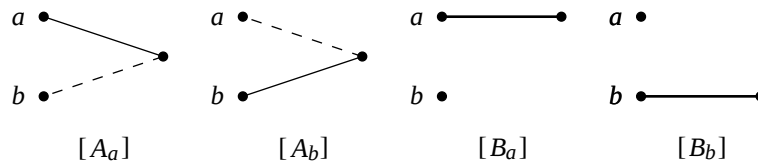


FIGURE A.5.: Motifs élémentaires permettant de construire les intégrales de la figure (A.4.b) (Source [AL04]).

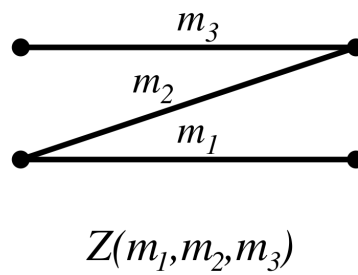


FIGURE A.6.: Diagramme associé à l'intégrale Z (Source [AL03]).

L'intégrale fan , représenté par le diagramme (A.3.b), le point de gauche n'étant pas lui non plus relié à d'autres lignes, vaut

$$F(m_1, m_2, \dots, m_t) = \frac{(\prod_{i=1}^t m_i!)(n-1)!}{(N + \sum_{i=1}^t m_i - 1)!}.$$

L'intégrale Z , représentée sur le diagramme (A.6) vaut

$$Z(m_1, m_2, m_3) = \frac{m_1!m_2!m_3!(N-2)!(N-1)!(N+m_1+m_3-2)!}{(N+m_1-2)!(N+m_3-2)!(N+m_1+m_2+m_3-1)!}.$$

D'autres formules explicites sont démontrées dans [AL03].

V. Conclusion

Ainsi la méthode d'intégration invariante constitue une manière particulièrement élégante pour calculer des intégrales de monômes d'éléments de matrices unitaires sur le groupe $U(N)$. Bien sûr, quand les intégrales deviennent trop grandes, avec une structure d'indices I , J et J_Q compliquée, les avantages de la méthode s'effondrent et l'emploi d'autres techniques devient obligatoire. Parmi ces techniques citons celle combinant l'approche invariante et la théorie des groupes, introduite dans [AL04].

Détails sur le calcul du second moment

Ici sont exposés les détails du calcul des termes A and B apparaissant dans l'expression (3.28) de $\langle \mathcal{I}^2 \rangle$ du chapitre 3.

$$\langle \mathcal{I}^2 \rangle = \frac{1}{Z^4} \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} e^{-x(\mu_5 + \mu_6 + \mu_7 + \mu_8)} \langle U_{11,35} U_{11,45}^* U_{12,36}^* U_{12,46} U_{23,57}^* U_{23,67} U_{24,58} U_{24,68}^* \rangle \quad (\text{B.1})$$

$$= \frac{1}{Z^4} \sum_{\{\alpha_i, \mu_j\}}^{(n,d)} e^{-x(\mu_5 + \mu_6 + \mu_7 + \mu_8)} \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} 6,8 \\ 5,8 \\ 6,7 \\ 5,7 \\ 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{array} . \quad (\text{B.2})$$

I. Le terme A

De l'équation (3.32) on obtient

$$\begin{aligned} A &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \left(\sum_{\{\alpha_1 = \alpha_2\}}^n f(x) \begin{array}{c} 1,4 \\ 1,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} + g(x) \begin{array}{c} 1,4 \\ 1,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ 4 \\ 3 \end{array} \right) \\ &+ \sum_{\{\alpha_1 \neq \alpha_2\}}^n f(x) \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} + g(x) \begin{array}{c} 2,4 \\ 2,3 \\ 1,2 \\ 1,1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ 4 \\ 3 \end{array} \Big) \\ &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \left(n f(x) \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} + n g(x) \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ 4 \\ 3 \end{array} \right) \\ &+ n(n-1) f(x) \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} + n(n-1) g(x) \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \nearrow \\ \nearrow \\ \nearrow \\ \nearrow \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ 4 \\ 3 \end{array} \Big) \\ &= n f(x) A_1 + n g(x) A_2 + n(n-1) f(x) A_3 + n(n-1) g(x) A_4 . \quad (\text{B.3}) \end{aligned}$$

En prenant en compte les contraintes sur les indices α_i il suit

$$\begin{aligned}
 A_1 &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} \\
 &= \sum_{\{\mu_j\}}^m \left(n[3] \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 4n[2] \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 2n[1] \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) \\
 &= n[3]A_{11} + 4n[2]A_{12} + 2n[1]A_{13},
 \end{aligned}$$

avec $n[i] = n(n-1)(n-2)\dots(n-i)$. On vérifie qu'il y a bien les $n[3] + 4n[2] + 2n[1] = n^2(n-1)^2$ configurations correspondant à la somme sur les quatre indices α_j avec les deux contraintes $\alpha_3 \neq \alpha_4$ and $\alpha_5 \neq \alpha_6$. Les termes A_{1k} s'écrivent

$$\begin{aligned}
 A_{11} &= \sum_{\{\mu_j\}}^m \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \\
 &= m[3] \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + m[2] \left(4 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 2 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) \\
 &\quad + m[1] \left(2 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 4 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) + m \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array}, \\
 A_{12} &= \sum_{\{\mu_j\}}^m \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \\
 &= m[3] \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 2m[2] \left(\begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) \\
 &\quad + m[1] \left(\begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) + 4m[1] \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + m \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array}, \\
 A_{13} &= \sum_{\{\mu_j\}}^m \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \\
 &= m[3] \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + m[2] \left(4 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 2 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) \\
 &\quad + m[1] \left(2 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} + 4 \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \right) + m \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array}.
 \end{aligned}$$

Pour A_2 on a directement

$$A_2 = \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \begin{array}{c} 4 \\ 3 \\ 2 \\ 1 \end{array} \begin{array}{c} \diagup \diagdown \\ \diagdown \diagup \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ \bar{4} \\ \bar{3} \end{array} = n^2(n-1)^2 A_{11},$$

tandis que A_3 est donné par

$$\begin{aligned}
 A_3 &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \begin{array}{c} 6 \\ 5 \\ 4 \\ 3 \end{array} \\
 &= \sum_{\{\mu_j\}}^m \left(n[3] \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + 4n[2] \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + 2n[1] \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \right) \\
 &= n[3]A_{31} + 4n[2]A_{32} + 2n[1]A_{33}.
 \end{aligned}$$

Les termes A_{3k} s'écrivent

$$\begin{aligned}
 A_{31} &= \sum_{\{\mu_j\}}^m \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \\
 &= (m[3] + 4m[2] + 2m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (2m[2] + 4m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (m[1] + m) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \\
 A_{32} &= \sum_{\{\mu_j\}}^m \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \\
 &= (m[3] + 4m[2] + 2m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (2m[2] + 4m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (m[1] + m) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \\
 A_{33} &= \sum_{\{\mu_j\}}^m \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \\
 &= (m[3] + 4m[2] + 2m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (2m[2] + 4m[1]) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} + (m[1] + m) \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array}.
 \end{aligned}$$

Le terme A_4 peut s'exprimer en fonction du terme A_{31} ,

$$A_4 = \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \begin{array}{c} \bar{4} \\ \bar{3} \\ 2 \\ 1 \end{array} \begin{array}{c} \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \\ \text{---} \diagup \text{---} \\ \text{---} \diagdown \text{---} \end{array} \begin{array}{c} \bar{6} \\ \bar{5} \\ 4 \\ 3 \end{array} = n^2(n-1)^2 A_{31}.$$

On vérifie qu'il y a bien dans les expressions des termes A_{3i} et du terme A_4 , les $m[3] + 6m[2] + 7m[1] + m = m^4$ configurations correspondant à la somme sur les quatre indices μ_j .

II. Le terme B

De la même manière que pour A , on trouve pour le terme B

$$\begin{aligned}
 B &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\alpha_3, \alpha_4}^n \sum_{\mu_5 \neq \mu_6}^m e^{-2x(\mu_5 + \mu_6)} \begin{matrix} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \begin{matrix} 4,6 \\ 3,6 \\ 4,5 \\ 3,5 \end{matrix} \\
 &= \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \sum_{\alpha_3, \alpha_4}^n \left(g(x) \begin{matrix} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \begin{matrix} \bar{4} \\ \bar{3} \\ 4 \\ 3 \end{matrix} \right) \\
 &= n(n-1)g(x) \sum_{\{\alpha_i, \mu_j\}}^{(n,m)} \begin{matrix} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \\
 &= n(n-1)g(x) \sum_{\{\mu_j\}}^m \sum_{\alpha_1 \alpha_2}^n \left(g(x) \begin{matrix} 2,4 \\ 1,2 \\ 2,3 \\ 1,1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \right) \\
 &= n(n-1)g(x) \sum_{\{\mu_j\}}^m \left(n \begin{matrix} 4 \\ 2 \\ 3 \\ 1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} + n(n-1) \begin{matrix} \bar{4} \\ 2 \\ \bar{3} \\ 1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \right) \\
 &= n^2(n-1)g(x) (B_1 + (n-1)B_2), \tag{B.4}
 \end{aligned}$$

où les termes B_i sont donnés par

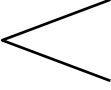
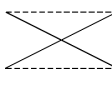
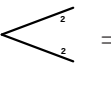
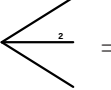
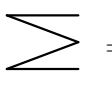
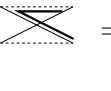
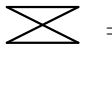
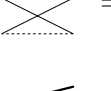
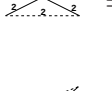
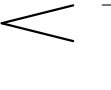
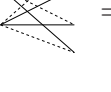
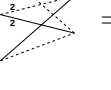
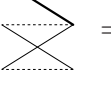
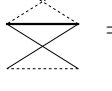
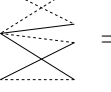
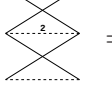
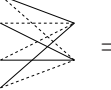
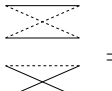
$$\begin{aligned}
 B_1 &= \sum_{\{\mu_j\}}^m \begin{matrix} 4 \\ 3 \\ 2 \\ 1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} = A_{11} \\
 &= m[3] \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + 4m[2] \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + 2m[2] \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} \\
 &\quad + m[1] \left(\begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + 2 \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + 4 \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} \right) + m \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix}, \\
 B_2 &= \sum_{\{\mu_j\}}^m \begin{matrix} 4 \\ 3 \\ 2 \\ 1 \end{matrix} \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \\ \text{diagram 3} \\ \text{diagram 4} \end{matrix} \\
 &= (m[3] + 4m[2] + 2m[1]) \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + (2m[2] + 4m[1]) \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix} + (m[1] + m) \begin{matrix} \text{diagram 1} \\ \text{diagram 2} \end{matrix}.
 \end{aligned}$$

L'ensemble des valeurs exactes de chaque diagramme est répertorié dans l'annexe suivante.

Valeurs des intégrales nécessaires pour le calcul du second moment

Toutes ces intégrales peuvent être calculées par la méthode d'intégration invariante introduite dans [AL03] et explicité dans l'annexe A. Dans la liste suivante, les intégrales (représentées par leur diagramme) sont ordonnées de telle sorte qu'une intégrale donnée ne dépend dans le pire des cas, que des intégrales notées au dessus d'elle. Le commentaire entre parenthèses illustre quelles relations spécifiques sont utilisées pour calculer l'intégrale en question, en occurrence *exp* pour l'utilisation d'une expression directement exploitable issue de [AL03], *unit* pour la relation d'unitarité, *fan* pour la fan relation, $2^{ble} fan$ pour la double fan relation. Cependant, il convient de préciser qu'une intégrale peut être obtenue de plusieurs manières. Notons que l'expression de la dernière intégrale calculée à partir de toutes les autres intégrales, correspond bien à l'expression calculée dans [AL04] par une méthode mixte intégration invariante - théorie des groupes.

ANNEXE C. VALEURS DES INTÉGRALES NÉCESSAIRES POUR LE CALCUL DU SECOND MOMENT

	$= \frac{1}{N(N+1)}$	<i>exp</i>		$= \frac{-1}{N(N^2-1)}$	<i>exp</i>
	$= \frac{4}{(N+3)(N+2)(N+1)N}$	<i>exp</i>		$= \frac{1}{(N+3)(N+2)(N+1)N}$	<i>exp</i>
	$= \frac{2}{(N+3)(N+2)(N+1)N}$	<i>exp</i>		$= \frac{N+1}{(N+3)(N+2)N^2(N-1)}$	<i>exp</i>
	$= \frac{-4}{(N+3)(N+2)(N^2-1)N}$	<i>exp</i>		$= \frac{N^2+N+2}{(N+3)(N+2)(N^2-1)N^2}$	<i>unit</i>
	$= \frac{-2}{(N+3)(N+2)(N^2-1)N}$	<i>unit</i>		$= \frac{8}{(N+3)(N+2)(N^2-1)N^2}$	<i>exp</i>
	$= \frac{1}{(N+3)(N-1)N^2}$	<i>fan</i>		$= \frac{-1}{(N+3)(N+2)(N+1)N^2}$	$2^{ble} fan$
	$= \frac{2}{(N+3)(N+2)(N^2-1)N^2}$	$2^{ble} fan$		$= \frac{-1}{(N+3)(N+2)(N^2-1)N}$	$2^{ble} fan$
	$= \frac{4}{(N+3)(N+2)(N^2-1)N^2}$	<i>unit</i>		$= \frac{-1}{(N+3)(N+2)(N+1)N^2}$	$2^{ble} fan$
	$= \frac{-(N^2+1)}{(N+3)(N^2-4)(N^2-1)N^2}$	<i>unit</i>		$= \frac{3N-1}{(N+3)(N^2-4)(N^2-1)N^2}$	<i>unit</i>
	$= \frac{1}{(N+3)(N+2)(N^2-1)N^2}$	<i>unit</i>		$= \frac{2}{(N+3)(N+2)(N^2-1)N^2}$	<i>unit</i>
	$= \frac{2}{(N+3)(N+2)(N^2-1)N^2}$	$2^{ble} fan$		$= \frac{-(N^2+2N+2)}{(N+3)(N^2-4)(N^2-1)N^2}$	<i>unit</i>
	$= \frac{1}{(N+3)(N+2)(N^2-1)N^2}$	<i>unit</i>		$= \frac{(N^2+6)}{(N^2-9)(N^2-4)(N^2-1)N^2}$	<i>unit</i>

Publications

- [1] L. Arnaud and D. Braun, *Distribution of interference in random quantum algorithms*, quant-ph/0612168, Phys. Rev A **75**, 062314 (2007).
- [2] L. Arnaud and D. Braun, *Efficiency of producing random unitary matrices with quantum circuits*, arXiv :0807.0775, Phys. Rev. A **78**, 062329 (2008).
- [3] L. Arnaud and D. Braun, *Distribution of interference in the presence of decoherence*, arXiv :0910.1826. Phys. Rev. A **80**, 062329 (2009).

Bibliographie

- [ADR82] A. Aspect, J. Dalibard, and G. Roger. Phys. Rev. Lett. **49** (1982).
- [Aha03] D. Aharonov. *quant-ph/0301040* (2003).
- [AL03] S. Aubert and C.S. Lam. J.Math.Phys. **44**, 6112 (2003).
- [AL04] S. Aubert and C.S. Lam. J.Math.Phys.**45**, 3019 (2004).
- [ALM06] D. Aharonov, Z. Landau, and J. Makowsky. *quant-ph/0611156* (2006).
- [AS04] A. Ambainis and A. Smith. In *Proceedings of RANDOM 2004*. Cambridge, MA (2004).
- [Bar95] A. Barenco. Proc. R. Soc. Lond. A. **51**, 1015 (1995).
- [BBC95] A. Barenco, C. H. Bennett, R. Cleve, D. P. DiVincenzo, N. Margolus, P. Shor, T. Sleator, J. A. Smolin, and H. Weinfurter. Phys. Rev. A **52**, 3457 (1995).
- [BCSG08] D. Bigourd, B. Chatel, W.P. Schleich, and B. Girard. Phys. Rev. Lett. **100**, 030202 (2008).
- [BD00] C. H. Bennett and D. P. DiVincenzo. Nature, **404**, 247 (2000).
- [Bel64] J.S. Bell. Physics **1** (1964).
- [BF07] F. Alexander Bais and J. Doyne Farmer. *The physics of information* (2007).
- [BG06] D. Braun and B. Georgeot. Phys. Rev. A **73**, 022314 (2006).
- [BG08] D. Braun and B. Georgeot. Phys. Rev. A **77**, 022318 (2008).
- [Bra06] D. Braun. J. Phys. A : Math. Gen. **39**, 14581 (2006).
- [BV09] W. G. Brown and L. Viola. arXiv :0910.0913v3 (2009).
- [BZ06] I. Bengtsson and K. Życzkowski. *Geometry of quantum states*. Cambridge University Press (2006).
- [DCEL06] C. Dankert, R. Cleve, J. Emerson, and E. Livine. *quant-ph/0606161* (2006).
- [Deu85] D. Deutsch. Proc. Roy. Soc. Lond. A. **400**, 97 (1985).
- [DiV95] D. P. DiVincenzo. Science **270**, 255 (1995).
- [DJ92] D. Deutsch and R. Jozsa. Proc. R. Soc. Lond. A **439**, 553 (1992).
- [DOB07] O. Dahlsten, R. Oliveira, and M. B. Plenio. J. Phys. A : Math. Theor. **40**, 8081 (2007).
- [ELL05] J. Emerson, E. Livine, and S. Lloyd. Phys. Rev. A **72**, 060302 (2005).
- [EPR35] A. Einstein, B. Podolsky, and N. Rosen. Phys. Rev. **47** (1935).

- [EWS03] J. Emerson, Y. Weinstein, M. Saraceno, S. Lloyd, and D. Cory. *Science* **302**, 2098 (2003).
- [Gro97] L. K. Grover. *Phys. Rev. Lett.* **79**, 325 (1997).
- [GSV00] W. G. Brown, Y. S. Weinstein, and L. Viola. *quant-ph/0011067v2* (2000).
- [Haa91] F. Haake. *Quantum Signatures of Chaos*. Springer, Berlin (1991).
- [Haa00] F. Haake. Springer, Berlin, Heidelberg, 2ème edition (2000).
- [Har] A. Harrow. *communication privée*.
- [HHL04] A. Harrow, P. Hayden, and D. Leung. *Phys. Rev. Lett.* **92**, 187901 (2004).
- [HL08] A. Harrow and R. Low. *quant-ph/0802.191v1* (2008).
- [Hur97] A. Hurwitz. *Nachr. Ges. Wiss. Gött. Math.-Phys. Kl.* **71**, 71 (1897).
- [HWW06] P. Hayden, D. W. Leung, and A. Winter. *Commun. Math. Phys.* **265**, 95 (2006).
- [JL03] R. Jozsa and N. Linden. *Proc. R. Soc. Lond. A* **459**, 2011 (2003).
- [Kup03] G. Kuperberg. *quant-ph/0302112v2* (2003).
- [LBB07] A. O. Lyakhov, D. Braun, and C. Bruder. *Phys. Rev. A* **76**, 022321 (2007).
- [MCD02] A. M. Childs, R. Cleve, E. Deotto, E. Farhi, S. Gutmann, and D. A. Spielman. In *Proc. 35th ACM Symposium on Theory of Computing* (2002).
- [MD01] R. Mosseri and R. Dandoloff. *Geometry of entangled states, bloch spheres and hopf fibrations. J. Phys. A* **34**, 10243 (2001).
- [Meh91] M. L. Mehta. *Random Matrices*. Academic Press, New York, 2ème edition (1991).
- [NC00] M. A. Nielsen and I. L. Chuang. *Quantum Computation and Quantum Information*. Cambridge University Press (2000).
- [ODB07] R. Oliveira, O. Dahlsten, and M. B. Plenio. *Phys. Rev. Lett.* **98**, 130502 (2007).
- [PWM02] D. P. DiVincenzo, D. W. Leung, and B. M. Terhal. *IEEE Trans. Inform. Theo.* **48**, 580 (2002).
- [PZK98] M. Pozniak, K. Życzkowski, and M. Kus. *Composed ensembles of random unitary matrices. J. Phys. A* **31**, 1059 (1998).
- [RB08] B. Roubert and D. Braun. *Phys. Rev. A* **78**, 042311 (2008).
- [Sch35] E. Schrödinger. *Die Naturwissenschaften*, **48**, 52 (1935).
- [Sen05] P. Sen. *quant-ph/0512085v1* (2005).
- [Shi02] Y. Shi. *quant-ph/0205115* (2002).
- [Sho94] P. W. Shor. IEEE Computer Society, Los Alamitos, CA (1994).
- [SL04] L. Susskind and J. Lindesay. *An Introduction to black holes, information and the string theory revolution : The Holographic Universe*. World Scientific Publishing Company (2004).
- [Ste98] A. Steane. *Quantum computing. Rep. Prog. Phys.* **61** (1998).
- [SW95] T. Sleator and H. Weinfurter. *Phys. Rev. Lett.* **74**, 4087 (1995).
- [TJR07] G. Toth and J. J. Garcia-Ripoll. *Phys. Rev. A* **75**, 042311 (2007).
- [Ž07] M. Žnidarič. *Phys. Rev. A* **76**, 012318 (2007).
- [WHL09] A. W. Harrow, A. Hassidim, and S. Lloyd. *Phys. Rev. Lett.* **103**, 150502 (2009).
- [WZ82] W. K. Wootters and W. H. Zurek. *Nature* **299**, 802 (1982).
- [Zei97] A. Zeilinger. *Nature* **390**, 575 (1997).
- [ZHSL99] K. Życzkowski, P. Horodecki, A. Sanpera, and M. Lewenstein. *Phys. Rev. A* **58**, 883 (1999).

Index

- k -design, 43
- k -design ϵ -approximé, 43
- algorithme quantique, 5
- base computationnelle, 3
- calcul quantique, 4
- CNOT, 5
 - porte, 5
- CUE, 17, 27, 41
- ensemble de porte 2-copy gapped, 43
- ensemble universel de portes, 5
- ensembles de circuits aléatoires
 - l'ensemble de circuits orthogonaux OCE, 21
 - l'ensemble de circuits unitaires UCE, 21, 43
- ensembles de matrices aléatoire
 - l'ensemble orthogonal de Haar HOE, 17
- ensembles de matrices aléatoires
 - ensembles gaussien GUE, GOE GSE, 18
 - l'ensemble unitaire circulaire CUE, 17, 27, 41
- gap spectral, 44
- Hadamard
 - porte de, 5, 12
- HOE, 17
- interférence, 9
- intrication, 7
- OCE, 21
- ordinateur quantique, 6
- paramétrisation de Hurwitz, 17, 57
- porte quantique, 5
- quantum bit, 2
- registre quantique, 3
- sphère de Bloch, 3
- transformée de Fourier quantique, 12
- UCE, 21, 43

Résumé

Cette thèse présente différents résultats portant sur deux thèmes précis dans le domaine de l'information quantique. Le premier de ces thèmes concerne l'interférence présente dans les algorithmes quantiques. Par l'intermédiaire d'une mesure récemment introduite dans la littérature, il est possible de quantifier l'interférence présente dans tout algorithme quantique, dans le but de vérifier si il existe un lien entre interférence et ressource computationnelle. Pour ce faire, deux types de modèles *statistiques* d'algorithmes quantiques ont été utilisés. Le premier type, issu de la théorie des matrices aléatoires, est l'ensemble circulaire unitaire (CUE) tandis que le second est un ensemble de circuits quantiques, construits par des séquences aléatoires de portes quantiques. Les résultats analytiques et numériques obtenus dans cette thèse montrent qu'*en moyenne* tout algorithme quantique contient une grande quantité d'interférence. L'étape supplémentaire fût d'étudier l'influence de la décohérence engendrée par un bain thermique sur le comportement statistique de l'interférence. Grâce à l'utilisation de méthodes mathématiques d'intégrations sur le groupe unitaire $U(N)$, il est possible de généraliser les résultats analytiques et numériques précédents pour inclure les effets de la décohérence.

Le deuxième thème étudié concerne la possibilité d'utiliser des algorithmes quantiques pour créer *efficacement* des ensembles de matrices aléatoires. Pendant les travaux sur l'interférence, une équivalence entre CUE et le modèle de circuits quantiques aléatoires fût observée. On peut attendre que de tels circuits quantiques aléatoires puissent être utilisés pour construire des matrices aléatoires distribuées comme CUE. Les résultats numériques de cette thèse montrent que certaines quantités statistiques propres à CUE sont bien reproduites par le modèle de séquences aléatoires, ceci de manière efficace. L'efficacité signifie que les séquences sont constituées d'un nombre de portes qui augmente comme le logarithme de la taille des matrices produites. Ces résultats sont en parfait accord avec des travaux analytiques récemment publiés.

Mots-clefs

Mécanique quantique, interférence quantique, information quantique, calcul quantique, théorie des matrices aléatoires, intégration invariante sur les groupes de Lie, simulations numériques.

Abstract

This thesis presents different results about two topics in the field of quantum information. The first of these topics concerns the interference present in quantum algorithms. Thanks to a measure recently introduced in the literature, it is possible to quantify the interference present in any quantum algorithm, in order to check if there is a link between interference and computational resource. To do this, two kinds of *statistical* models of quantum algorithms were used. The first one, which comes from the random matrix theory, is the circular unitary ensemble (CUE) while the second one is a set of quantum circuits, built by random sequences of quantum gates. The analytical and numerical results obtained in this thesis show that *on average* any quantum algorithm contains a big amount of interference. The next step was to study the influence of the decoherence created by a thermal bath on the statistics of the interference. Thanks to mathematical methods of integration over the unitary group $U(N)$, it is possible to generalise the previous analytical and numerical results to include the effect of decoherence.

The second topic is about the possibility of using quantum algorithms to create *efficiently* random matrix ensembles. During the works on interference, an equivalence between CUE and the model of random quantum circuits was noticed. One might expect that such random quantum circuits can be used to build random matrix ensembles drawn as CUE. The numerical results of this thesis show that given CUE quantities are well reproduced by the model of random sequences in an efficient way. The efficiency means that the sequences are built with a number of gates which increases as the logarithm of the size of the produced matrices. These results are in agreement with recently published works.

Keywords

Quantum mechanics, quantum interference, quantum information, quantum computation, random matrix theory, invariant integration over Lie's groups, numerical simulations.
