



Modèles formels du calcul quantique : ressources, machines abstraites et calcul par mesure

Simon Perdrix

► To cite this version:

Simon Perdrix. Modèles formels du calcul quantique : ressources, machines abstraites et calcul par mesure. Informatique [cs]. Institut National Polytechnique de Grenoble - INPG, 2006. Français. NNT : . tel-00275323

HAL Id: tel-00275323

<https://theses.hal.science/tel-00275323>

Submitted on 23 Apr 2008

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.

INSTITUT NATIONAL POLYTECHNIQUE DE GRENOBLE

N° attribué par la bibliothèque

--	--	--	--	--	--	--	--	--	--

THÈSE

pour obtenir le grade de **DOCTEUR DE L'INPG**

Spécialité : « Systèmes et Logiciels »

préparée au laboratoire Leibniz dans le cadre de l'**École Doctorale**
« Mathématiques, Sciences et Technologies de l'Information,
Informatique »

préparée et soutenue publiquement par

Simon PERDRIX

le 11 Décembre 2006

Titre :

**Modèles formels du calcul quantique :
ressources, machines abstraites et calcul
par mesure**

sous la direction de Philippe Jorrand

JURY

Pr. Roger Mohr
Pr. Samson Abramsky FRS
Pr. Hans J. Briegel
Dr. Vincent Danos
Pr. Hubert Comon Lundh
Dr. Philippe Jorrand

Président
Rapporteur
Rapporteur
Examineur
Examineur
Directeur de thèse

Résumé

L'étude des structures fondamentales du traitement de l'information quantique est un défi majeur, dont l'un des objectifs est de mieux cerner les capacités et les limites de l'ordinateur quantique, tout en contribuant à sa réalisation physique notamment en s'intéressant aux ressources du calcul quantique. Les ressources d'un calcul quantique incluent le temps et l'espace mais également la taille des opérations utilisées et la quantité d'intrication.

Cette thèse contribue de plusieurs manières à la recherche de ressources minimales dans le cadre de modèles de calcul quantique ouvrant de prometteuses perspectives de réalisations physiques. Ces modèles sont le calcul par consommation d'intrication et le calcul par mesures projectives. Cette thèse a également permis de réduire les ressources en temps et en espace nécessaires à la préparation de certains états quantiques, les états graphes.

Etudier la réduction des ressources nécessite l'abstraction et la formalisation des modèles de calcul quantique mettant en évidence les structures même du traitement de l'information quantique. Le q-calcul et les machines de Turing contrôlées classiquement, introduits dans cette thèse, ont cet objectif. Des modèles plus spécifiques au calcul par consommation d'intrication, ou au calcul par mesures projectives sont également considérés.

Abstract

The study of foundational structures of quantum information processing is a key issue to gain a deeper insight into what quantum computation is in general, its scope and limits. It also contributes to the physical realisation while minimising the resources of quantum computing. The resources consist of the space and times as well as the size of the operations and the amount of entanglement.

This thesis contributes in several ways to minimise resources for recently developed models of quantum computation which open new promising perspectives of physical realisation. These models are the one-way quantum computation and the measurement-only quantum computation. This thesis has also permitted to reduce the resources in time and space necessary for the preparation of some quantum states called graph states.

The reduction of the resources requires abstraction and formalisation of quantum computing models which point out the structures of the quantum computing processing. The q-calculus and the classically-controlled quantum Turing machines, introduced in this thesis, contribute to this objective. More specific models dedicated to one-way and measurement-only quantum computations are considered as well.

Remerciements

Mes remerciements s'adressent, en premier lieu, à mon directeur de thèse Philippe Jorrand pour son soutien continu, sa disponibilité exceptionnelle et ses conseils avisés. Il a su me témoigner une grande confiance dont je lui suis reconnaissant. Je le remercie également de m'avoir fait découvrir le monde de la recherche et fait part de son expérience pendant ces quatre années où nous avons partagé le même bureau.

Samson Abramsky et Hans J. Briegel m'ont fait l'honneur d'être rapporteurs de cette thèse. Je leur suis reconnaissant d'avoir accepté cette lourde tâche, malgré l'obstacle de la langue.

Je témoigne toute ma gratitude à Vincent Danos et à Hubert Comon Lundh d'avoir accepté de participer à mon jury. Je remercie enfin Roger Mohr d'avoir présidé ce jury, et de m'avoir fait part de précieux conseils pour l'enseignement tout au long du monitorat.

Je souhaite également remercier les autres chercheurs avec qui j'ai eu la chance de travailler. Merci à bon nombre de doctorants, et tous les membres du laboratoire Leibniz, pour les discussions scientifiques, mais aussi pour les moments de détente. Merci également à mes amis, en particulier à Baptiste et à Dimitri pour leur soutien.

Je veux remercier mes parents pour tout ce qu'ils ont fait pour moi, pour leur amour et leurs encouragements. Enfin, je remercie Laurence pour tout le bonheur qu'elle m'apporte.

Préambule

L'informatique quantique

La physique quantique a mis en évidence des phénomènes dans le comportement des particules élémentaires, qui sont désormais considérés sous l'angle de leur exploitation pour représenter, traiter et communiquer l'information.

La rencontre entre la physique quantique et les sciences de l'information débute en 1982 quand Richard Feynman [Fey82, Fey84, Fey86], Prix Nobel de Physique, propose l'utilisation de la physique quantique au lieu de la physique classique comme support matériel de l'information et du calcul. Des problèmes hors de portée de l'informatique actuelle (dite *classique*) pourraient alors être traités efficacement.

Paul Benioff [Ben80, Ben82], puis David Deutsch [Deu85] remarquent que la machine introduite par Alan Turing [Tur36] est fondée sur la physique classique, celle de Newton. La machine de Turing est un modèle abstrait d'ordinateur sur lequel s'appuie l'informatique classique. Une réinterprétation des fonctions calculables amène Deutsch à introduire, en 1985, une généralisation quantique de la classe des machines de Turing, dont il prouve quelques propriétés remarquables car non reproductibles en temps polynomial par les machines classiques.

Dans les années 90, des résultats algorithmiques, théoriques puis expérimentaux sont venus confirmer l'intérêt d'un fondement quantique des sciences de l'information. En 1993, Charles Bennett, Gilles Brassard, Claude Crépeau, Richard Jozsa, Asher Peres et William Wootters [BBC⁺93] publient les principes théoriques d'un protocole de *téléportation* , qui repose sur l'usage d'états quantiques intriqués : l'état d'un système quantique a localisé en A peut, après avoir été détruit, être attribué à un autre système quantique b localisé en B sans que l'état de a ne soit connu ni en A , ni en B , ni qu'il soit transporté dans l'espace sur une trajectoire reliant A et B . Un an plus tard, Peter Shor [Sho94] montre que le calcul quantique permet la factorisation d'un entier en nombres premiers en temps polynomial, alors que le meilleur algorithme classique connu est exponentiel.

En 1996, Lov Grover [Gro96b] publie un algorithme quantique qui réalise la recherche d'un élément dans une base de données non ordonnée de taille n , en $\sqrt{n}$ appels à un oracle, là où le calcul classique requiert $\Theta(n)$ appels à ce même

oracle. En 1997, Anton Zeilinger [BPM⁺97] réalise la première expérience de téléportation de l'état d'un photon. Isaac Chuang [CVZ⁺98, VSB⁺01], de 1998 à 2002, conçoit et réalise un ordinateur quantique possédant 7 bits quantiques. Cet ordinateur quantique confirme que la physique quantique permet de mettre en œuvre expérimentalement les algorithmes pensés par Peter Shor et Lov Grover.

Ces splendides résultats théoriques sur l'information quantique et son traitement, puis leurs confirmations expérimentales, montrent bien que des problèmes hors de portée de l'informatique classique peuvent, en principe, être traités en exploitant ce paradigme de calcul non classique. Cela ouvre des perspectives scientifiques et technologiques lointaines, certes, mais immenses.

À présent, la recherche en informatique quantique ne se limite plus à l'algorithmique quantique et à l'expérimentation. En effet, l'étude formelle des structures du traitement de l'information quantique, l'étude de modèles formels de calcul quantique, et l'utilisation de méthodes de haut niveau pour le calcul quantique sont des sujets émergents dans le domaine, et qui comptent déjà des résultats prometteurs.

Par exemple l'axiomatisation catégorique du traitement de l'information quantique permet une meilleure compréhension des phénomènes quantiques, ainsi que la mise en évidence de nouvelles propriétés [Coe04, Abr04b, AC04a, Sel04c, Sel05, AD06].

Le développement de langages de programmation et de modèles abstraits pour le calcul quantique vise une représentation plus formelle de l'évolution quantique que la représentation traditionnelle par des circuits quantiques, et permet la mise en évidence de propriétés des algorithmes et protocoles quantiques. Différents paradigmes de programmation ont été et sont encore étudiés : λ -calcul ou calcul fonctionnel [Sel04b, vT04, AD04, AG05, SV06], langages séquentiels [Ome03], ou algèbres de processus [JL04, GN05, Lal06, FDJY06]. Le problème central est évidemment l'adéquation de la sémantique de ces langages aux lois de la mécanique quantique. En retour, l'élaboration de modèles sémantiques pour ces langages apporte non seulement une compréhension plus profonde de ce qu'est le calcul quantique, mais jette aussi un éclairage nouveau sur la mécanique quantique elle-même. Une description comparative de ces langages est donnée dans [Gay06].

Thèse

Comme cela vient d'être dit, en plus de l'algorithmique quantique et de l'expérimentation physique, le développement de l'informatique quantique passe par l'élaboration de modèles formels et abstraits. Tout d'abord, des modèles formels sont nécessaires pour les modèles émergents de calcul quantique comme le calcul par consommation d'intrication ou par mesures projectives uniquement. L'émergence de ces nouveaux modèles montre l'intérêt qu'il y a à ne pas limiter le calcul

quantique à son fragment unitaire (c'est-à-dire réversible). Or, la formalisation désormais traditionnelle du calcul quantique se limite souvent à ce fragment. Un des objectifs de cette thèse est de proposer et d'étudier des modèles formels de calcul quantique pour ces modèles émergents.

Il faut aussi remarquer que l'exécution d'un calcul qui utilise des ressources et opérations quantiques s'inscrit dans un cadre classique. En effet l'algorithmique quantique s'attache, le plus souvent, à résoudre des problèmes classiques comme la recherche d'un élément dans une base de données, la factorisation d'un entier, ou encore la résolution de problèmes sur les graphes [DHHM06]. Cette utilisation des phénomènes quantiques pour résoudre des problèmes classiques, avec une entrée et une sortie classiques, inscrit l'ordinateur quantique dans une interaction entre classique et quantique, ou autrement dit à la frontière entre les mondes classique et quantique. L'environnement classique a également pour rôle de contrôler l'évolution quantique. En effet, durant une exécution quantique, des transformations, unitaires ou non, sont *appliquées* à un système quantique. D'un point de vue expérimental, cette application nécessite une intervention classique qui rompt l'éventuel isolement du système quantique. Il s'avère donc nécessaire de formaliser les évolutions quantiques, mais également l'interaction avec l'environnement classique.

Une des applications de la formalisation du calcul quantique contrôlé classiquement est la minimisation des ressources nécessaires à un calcul. La construction d'un ordinateur quantique de taille raisonnable est encore un défi, et chercher à diminuer les ressources nécessaires au calcul quantique contribue à diminuer ainsi les difficultés de réalisations physiques.

Ainsi l'objectif de ma thèse est le développement de modèles formels et abstraits de calcul quantique, permettant une représentation de l'état d'un système quantique, de son évolution, mais aussi de l'évolution de son éventuel environnement classique. Le développement de ces modèles formels et abstraits s'appuie sur des outils d'informatique théorique : réécriture, sémantique, machines de Turing et théorie des graphes.

Plan de la Thèse

Cette thèse se découpe en cinq parties :

- La première partie, introductive, présente les postulats de la mécanique quantique. Cette présentation de postulats vénérables rompt avec la tradition. En effet, les évolutions quantiques y sont présentées via le formalisme des transformations admissibles, formalisme qui permet de représenter l'évolution des systèmes quantiques aussi bien ouverts que fermés. Nous montrons que les transformations admissibles sont alors particulièrement adaptées à

l'informatique quantique, notamment grâce à leur compositionnalité. Une transformation admissible peut être interprétée comme une évolution probabiliste sur les états quantiques, ou comme une transformation linéaire sur les matrices de densités.

Cette partie introductive met en évidence l'existence de relations d'interprétation entre trois domaines sémantiques différents utilisés pour la représentation des évolutions quantiques. Le plus concret est celui des transformations admissibles, le plus abstrait est celui des matrices de densité. Les fonctions probabilistes sur des distributions de probabilités d'états purs forment un domaine intermédiaire.

Les modèles traditionnels du calcul quantique réversible, circuits quantiques et machines de Turing quantiques, sont également présentés dans cette partie.

- La deuxième partie est dédiée à des modèles de calcul plus généraux, dans le cadre du *calcul quantique contrôlé classiquement*, incluant le calcul quantique réversible mais également les modèles alternatifs comme le calcul par mesures projectives ou par consommation d'intrication. Un nouveau modèle formel, le $\mathbf{q}$ -calcul est alors introduit et étudié.

Alors que la machine de Turing quantique introduite par Deutsch [Deu85] puis développée par Bernstein et Vazirani [BV97] est un modèle abstrait de calcul quantique réversible, nous introduisons une machine de Turing quantique possédant un contrôle classique, que nous étudions et comparons avec les machines de Turing quantiques traditionnelles, mais aussi avec d'autres modèles formels du calcul quantique.

- La troisième partie est consacrée au modèle de calcul quantique introduit par Nielsen [Nie03], le calcul quantique par mesures projectives. Le travail effectué sur ce modèle de calcul a un double objectif. Le premier objectif est de diminuer les ressources nécessaires à ce modèle. Ces ressources sont de deux types : d'une part l'ensemble des observables nécessaires au calcul (c'est-à-dire le nombre de mesures différentes effectuées pendant un calcul), d'autre part l'espace nécessaire au calcul.

Le second objectif est l'introduction d'un modèle formel permettant de représenter un calcul quantique par mesures projectives. Le $\mathbf{q}$ -calcul et les machines de Turing contrôlées classiquement s'avèrent être adaptés à ce modèle de calcul, à condition de ne considérer qu'un fragment des modèles formels développés dans le cadre du calcul quantique contrôlé classiquement.

- La quatrième partie traite de la représentation de l'intrication et du calcul par consommation d'intrication. Le formalisme des états graphes permet de représenter certains états quantiques à l'aide d'un graphe. Cette représentation établit des relations entre théorie des graphes et informatique quantique. Ceci permet d'une part de découvrir les propriétés de ces états quantiques

particuliers, grâce à la structure combinatoire des graphes. D'autre part, de nouveaux résultats en théorie des graphes peuvent être obtenus en faisant ce détour quantique par les états graphes.

Le m -calcul, modèle formel pour le calcul par consommation d'intrication, introduit par Danos, Kashefi et Panengaden [DKP04a] est également présenté dans cette partie. Nous nous intéressons tout particulièrement aux conditions permettant de décider si une ressource est suffisante ou non pour effectuer un calcul, les ressources étant ici des état graphes.

- La cinquième et dernière partie résume les résultats obtenus au long de cette thèse, et présente des perspectives. J'évoque aussi d'autres sujets auxquels je me suis intéressé au cours de ces trois années et qui ne sont pas présentés dans ce manuscrit.

Table des matières

I	Introduction	1
1	Les postulats de la mécanique quantique	3
1.1	Vecteur d'état et notations de Dirac	3
1.2	Système composé	4
1.3	Evolution quantique	5
1.3.1	Une évolution probabiliste	6
1.3.2	Composition spatiale	8
1.3.3	Composition temporelle	8
1.3.4	Exemples de transformations admissibles	9
1.4	Transformations admissibles vs transformations unitaires	10
1.5	Conclusion	11
2	Représentations des états quantiques	13
2.1	Etats mixtes et matrices de densité	13
2.1.1	Abstraction	13
2.1.2	Système composé	14
2.1.3	Evolution	15
2.1.4	Indistingabilité	17
2.2	Autres formalismes	18
2.2.1	Distributions sur les sous-espaces de Hilbert	18
2.2.2	Ensemble d'états purs non normés	19
2.3	Conclusion	20
3	Modèles de calcul quantique réversibles	21
3.1	Circuit quantique	22
3.1.1	Qubit	22
3.1.2	Porte unitaire	22
3.1.3	Circuit	24
3.1.4	Universalité	25
3.2	Machine de Turing quantique	25
3.3	Autres modèles réversibles	26

3.4	Conclusion	27
II	Données quantiques - contrôle classique	29
4	Evolution quantique contrôlée classiquement	31
4.1	Motivation	31
4.2	Formalisme	32
4.2.1	Espace d'états	32
4.2.2	Evolutions	32
4.2.3	Composition spatiale	33
4.2.4	Composition temporelle	33
4.3	Interprétation des transformations admissibles contrôlées classiquement	34
4.3.1	Fonction probabiliste	34
4.3.2	Super-opérateur	34
4.4	Conclusion	35
5	q-calcul	37
5.1	Introduction	37
5.2	Termes du q-calcul	39
5.2.1	Définitions	39
5.2.2	Représentation graphique	40
5.3	Sémantiques dénotationnelles	40
5.3.1	Sémantique pure	41
5.3.2	Sémantique observable	43
5.3.3	Sémantique admissible	45
5.3.4	Equivalences	49
5.4	Vers un q-calcul	50
5.4.1	Terminaison et non-confluence	53
5.4.2	Perspectives	54
5.5	Un rôle unificateur	55
5.6	Conclusion	56
6	MTQC	57
6.1	Introduction	57
6.2	Machines de Turing quantiques contrôlées classiquement	58
6.3	MTQC et MT	62
6.4	MTQC multi-rubans	63
6.5	MTQC et les modèles de calcul quantique réversible	66
6.5.1	MTQC et Circuits quantiques	66

6.5.2	MTQC et MTQ	68
6.5.3	Circuits quantiques et MTQ	73
6.6	MTQC à 1 ruban et MTQC à 2 rubans	73
6.7	Conclusion	74
III	Calcul par mesures projectives	77
7	Ressources du calcul par mesures projectives	79
7.1	Introduction	79
7.2	Calcul par mesures projectives à base de téléportation	80
7.2.1	Ressources	82
7.3	Transfert d'état	83
7.3.1	Ressources	87
7.4	Compromis entre observables et qubits auxiliaires ?	89
7.5	Vers les ressources minimales	90
7.6	Conclusion	94
8	Modèles formels du calcul par mesures projectives	95
8.1	Fragment observable du q -calcul	96
8.2	Machine de Turing quantique fondée sur la mesure	98
8.3	Conclusion	102
IV	Représentation de l'intrication et calcul par consommation d'intrication	105
9	Etats Graphes	107
9.1	Introduction	107
9.2	Etats graphes et états graphes signés	108
9.3	Propriétés combinatoires des états graphes	110
9.3.1	Complémentation locale et pivot	111
9.3.2	Transformations	113
9.4	Conclusion	118
10	Calcul quantique par consommation d'intrication	121
10.1	Introduction	121
10.2	m -calcul	122
10.2.1	Syntaxe	122
10.2.2	Sémantique	123
10.2.3	Forme standard	124
10.3	Condition de flots	125

10.3.1	Condition de flot simple	125
10.3.2	Condition de flot généralisé	127
10.4	m -calcul $3P$	129
10.4.1	Définitions	130
10.5	Calcul par mesures dans le plan (X, Z) sur une grille	133
10.5.1	Universalité des mesures dans le plan (X, Z)	134
10.5.2	Pivot mineur	135
10.6	Unification	139
10.6.1	Le secret du calcul par consommation d'intrication est caché dans la préparation de l'état graphe initial	139
10.6.2	Unification et m -calcul	141
10.7	Conclusion	142
11	Préparation des états graphes	143
11.1	Introduction	143
11.2	Préparation d'un état graphe	144
11.3	Préparation fondée sur la mesure	147
11.4	Circuits et complémentation locale	148
11.5	Séparabilité et δ_{loc}	154
11.6	Bornes inférieures sur δ_{loc}	156
11.7	Conclusion	157
V	Conclusion et perspectives	159

Première partie

Introduction

Chapitre 1

Les postulats de la mécanique quantique

La mécanique quantique est un modèle mathématique élaboré par les physiciens de la première moitié du XX^e siècle pour rendre compte de phénomènes mis en évidence par l'expérimentation sur les particules élémentaires. Ce modèle mathématique est construit de manière axiomatique à partir de postulats. A la fin du XX^e siècle, des physiciens, des informaticiens et des mathématiciens ont montré que les phénomènes quantiques, tels qu'ils sont formulés par la mécanique quantique, peuvent être exploités pour représenter, traiter et communiquer l'information. Ceci invite à revisiter les postulats de la mécanique quantique avec un regard d'informaticien : quelles contraintes ces postulats imposent-ils aux objets et opérations que l'on peut effectuer, quelles propriétés leur confèrent-ils ? En bref, à quoi peuvent ressembler les domaines sémantiques du calcul quantique ?

1.1 Vecteur d'état et notations de Dirac

Postulat 1. [Dir47] *A tout système physique isolé est associé un espace de Hilbert séparable appelé espace d'état. Le système est entièrement décrit par son vecteur d'état, vecteur de norme 1 dans l'espace des états associé.*

Un espace de Hilbert $\mathcal{H}$ est un espace vectoriel sur les complexes muni d'un produit scalaire. $\mathcal{H}$ est séparable si et seulement s'il admet une base dénombrable. On notera $\mathcal{H}^1$ la sphère unité de $\mathcal{H}$, *i.e.* $\mathcal{H}^1 = \{|\varphi\rangle \in \mathcal{H} \mid \|\varphi\rangle\| = 1\}$ et $\mathcal{H}^{\leq 1}$ la boule unité de $\mathcal{H}$, *i.e.* $\mathcal{H}^{\leq 1} = \{|\varphi\rangle \in \mathcal{H} \mid \|\varphi\rangle\| \leq 1\}$.

Les vecteurs, les produits scalaire et externe sont exprimés à l'aide de la notation introduite par Dirac en 1920. Les vecteurs sont notés $|\varphi\rangle$ (*ket* φ) ; le produit scalaire de deux vecteurs $|\varphi\rangle$ et $|\psi\rangle$ est noté $\langle\varphi|\psi\rangle$.

Une base orthonormée de cet espace de Hilbert séparable $\mathcal{H}$ est décrite par $\{|\tau\rangle, \tau \in B\}$, où B est un ensemble dénombrable. Ainsi un vecteur d'état quelconque $|\varphi\rangle \in \mathcal{H}^1$ du système physique peut être décrit comme une *superposition* d'états de base :

$$\sum_{\tau \in B} \alpha_{\tau} |\tau\rangle,$$

avec $\sum_{\tau \in B} |\alpha_{\tau}|^2 = 1$ car $|\varphi\rangle$ est un vecteur normé.

Etant donné un ensemble dénombrable B , $\mathcal{H}_B$ désigne l'espace de Hilbert ayant pour base $\{|\tau\rangle, \tau \in B\}$. Si $|\varphi\rangle = \sum_{\tau \in B} \alpha_{\tau} |\tau\rangle$ et $|\psi\rangle = \sum_{\tau \in B} \beta_{\tau} |\tau\rangle$, alors le produit scalaire de $|\varphi\rangle$ et $|\psi\rangle$ est $\langle\varphi|\psi\rangle = \sum_{\tau \in B} \alpha_{\tau}^* \beta_{\tau}$ (où α^* représente le conjugué de α). La partie gauche $\langle\varphi|$ du produit scalaire est un vecteur *bra*, alors que la partie droite $|\varphi\rangle$ est un vecteur *ket*. Un vecteur *bra* est défini comme l'adjoint du vecteur *ket* correspondant : si $|\varphi\rangle = \sum_{\tau \in B} \alpha_{\tau} |\tau\rangle$, alors $\langle\varphi| = |\varphi\rangle^{\dagger} = \sum_{\tau \in B} \alpha_{\tau}^* \langle\tau|$. La notation *bra-ket* permet de représenter le produit scalaire : $\langle\varphi|\psi\rangle = \langle\varphi|\psi\rangle$. Cette notation peut également être utilisée pour décrire le produit externe : $|\varphi\rangle\langle\psi|$ est un opérateur linéaire tel que $(|\varphi\rangle\langle\psi|)|\chi\rangle = \langle\psi|\chi\rangle |\varphi\rangle$.

Etant donné une base B , le phénomène de *superposition* désigne la capacité d'un système quantique à être dans un état qui est une combinaison linéaire d'états de bases.

1.2 Système composé

Postulat 2. *L'espace des états d'un système physique, composé de sous-systèmes, est le produit tensoriel des espaces des états des sous-systèmes.*

Etant donnés deux systèmes S_1 et S_2 dont les espaces d'états sont respectivement $\mathcal{H}_{B_1}^1$ et $\mathcal{H}_{B_2}^1$, l'espace d'état du système S composé de S_1 et S_2 est $\mathcal{H}_{B_1}^1 \otimes \mathcal{H}_{B_2}^1 \cong \mathcal{H}_{B_1 \times B_2}^1$. Si le système S_1 est dans l'état $|\varphi_1\rangle \in \mathcal{H}_{B_1}^1$ et le système S_2 dans l'état $|\varphi_2\rangle \in \mathcal{H}_{B_2}^1$, alors S est dans l'état $|\varphi_1\rangle \otimes |\varphi_2\rangle$, noté parfois $|\varphi_1\rangle |\varphi_2\rangle$ ou encore $|\varphi_1 \varphi_2\rangle$. Un cas particulier est le celui d'un espace d'états de dimension 1, il s'agit alors de l'espace de Hilbert $\mathbb{C}$. L'unique état de base de $\mathbb{C}$ est noté $|\rangle$. $\mathbb{C}$ est un élément neutre pour le produit tensoriel car pour tout espace de Hilbert $\mathcal{H}$,

$$\mathcal{H} \otimes \mathbb{C} \cong \mathcal{H}$$

Il est important de remarquer que l'état $|\varphi\rangle$ d'un système composé de deux sous-systèmes ne peut pas toujours être décomposé en un produit tensoriel de la forme $|\varphi_1\rangle \otimes |\varphi_2\rangle$, où $|\varphi_1\rangle$ serait l'état du premier sous système et $|\varphi_2\rangle$ celui du second. De tels états non séparables sont appelés états *intriqués*.

Propriété 1.1 *Pour tout espace de Hilbert $\mathcal{H}_{B_1}$ et $\mathcal{H}_{B_2}$ de dimension supérieure ou égale à 2, il existe $|\varphi\rangle \in \mathcal{H}_{B_1}^1 \otimes \mathcal{H}_{B_2}^1$ tel que pour tout $(|\varphi_1\rangle, |\varphi_2\rangle) \in \mathcal{H}_{B_1}^1 \times \mathcal{H}_{B_2}^1$, $|\varphi\rangle \neq |\varphi_1\rangle \otimes |\varphi_2\rangle$*

Preuve : L'état $|B_0\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle) \in \mathcal{H}_{\{0,1\}} \otimes \mathcal{H}_{\{0,1\}}$ est intriqué. En effet pour tout $a, b, c, d \in \mathbb{C}$, $|B_0\rangle = (a|0\rangle + b|1\rangle) \otimes (c|0\rangle + d|1\rangle) = ac|00\rangle + ad|01\rangle + bc|10\rangle + bd|11\rangle$. On en déduit que $ad = 0$ or $a \neq 0$ car $ac = 1$ et $d \neq 0$ car $bd = 1$. Donc $|B_0\rangle$ est un état intriqué. Dans le cas général, une analyse sur les dimensions des espaces de Hilbert permet de conclure. $\square$

Cet état $|B_0\rangle$ est appelé état EPR pour Einstein, Podolsky et Rozen [EPR35]. L'état EPR est, historiquement l'état qui a permis la mise en évidence de l'intrication. Les auteurs pensaient alors que la possibilité de décrire de tels états remettait en cause le formalisme de la mécanique quantique car ils violeraient le principe de causalité. Un modèle alternatif, à variables cachées, a même été élaboré [Boh52], concurrent du modèle standard de la mécanique quantique. Alain Aspect [AGR81] a confirmé, par l'expérience, la violation des inégalité de Bell par les états intriqués, démontrant ainsi que la mécanique quantique rend compte de façon adéquate des phénomènes constatés, aussi surprenant soient-ils.

L'intrication est une spécificité des lois quantiques permettant la mise en place de protocole comme celui de la téléportation [BBC⁺93]. D'un point de vue algorithmique, l'intrication joue aussi un rôle important. Elle est même une ressource du calcul quantique dans certains modèles de calcul quantique comme le modèle de calcul par consommation d'intrication, où l'intrication est construite puis consommée par le processus de calcul. Les chapitres 9, 10 et 11 sont dédiés à la représentation et aussi, dans une certaine mesure, à l'utilisation de l'intrication en informatique quantique.

1.3 Evolution quantique

Les évolutions quantiques sont souvent présentées en deux temps : tout d'abord l'évolution des systèmes fermés est traité, puis celle des systèmes ouverts. Les systèmes fermés ont la propriétés d'évoluer de façon réversible. Nous choisissons de présenter l'évolution des systèmes quantiques de façon unifiée, en un seul postulat, s'appuyant sur un seul formalisme, celui des transformations admissibles. Les transformations admissibles présentent l'avantage d'être stable par composition, y compris quand le système passe du statut de système fermé à ouvert et inversement.

Postulat 3. *Tout système quantique évolue selon une transformation admissible¹. Une transformation admissible agissant de $\mathcal{H}$ dans $\mathcal{K}$ est une famille dénombrable $(M_i)_{i \in A}$ d'opérateurs linéaires de $\mathcal{H}$ dans $\mathcal{K}$, vérifiant la condition de complétude :*

$$\sum_{i \in A} M_i^\dagger M_i = Id_{\mathcal{H}}$$

où $Id_{\mathcal{H}}$ est l'identité sur $\mathcal{H}$.

Si une transformation admissible $(M_i)_{i \in A}$ est appliquée à $|\varphi\rangle \in \mathcal{H}^1$, alors avec une probabilité $p(i) = \langle \varphi | M_i^\dagger M_i | \varphi \rangle$, le résultat classique i est observé et l'état du système après la transformation est :

$$\frac{M_i |\varphi\rangle}{\sqrt{\langle \varphi | M_i^\dagger M_i | \varphi \rangle}}$$

1.3.1 Une évolution probabiliste

On remarque que l'évolution quantique est probabiliste. Il est donc naturel d'employer des outils habituellement utilisés en informatique classique pour représenter une évolution probabiliste.

Une distribution discrète, ou évaluation discrète sur un ensemble X est une fonction $\nu : X \rightarrow \overline{\mathbb{R}^+}$. Une distribution discrète définit une unique distribution sur les parties de X : $\forall Y \subseteq X, \nu(Y) = \sum_{y \in Y} \nu(y)$. $V(X)$ est l'ensemble des distributions discrètes sur X . Une distribution ν sur X vérifiant $\nu(X) = 1$ (resp. $\nu(X) \leq 1$) est appelée distribution de probabilité (resp. distribution de sous-probabilité). L'ensemble des distributions de probabilité (resp. sous-probabilité) sur X est noté $V^1(X)$ (resp. $V^{\leq 1}(X)$).

Une évolution quantique est donc naturellement une fonction de $\mathcal{H}^1 \rightarrow V^1(\mathcal{K}^1)$. Pourtant, un domaine très différent de $\mathcal{H}^1 \rightarrow V^1(\mathcal{K}^1)$ est utilisé pour représenter les évolutions quantiques que sont les transformations admissibles.

En effet, d'après le troisième postulat, une transformation admissible est entièrement décrite par famille dénombrable $(M_i)_{i \in A}$ d'opérateurs linéaires de $\mathcal{H}$ dans $\mathcal{K}$ vérifiant la condition de complétude.

Définition 1.1 (Transformation admissible)

- $T(\mathcal{H}, \mathcal{K}) = \{(M_i)_{i \in A} \mid A \text{ dénombrable} \wedge \forall i \in A, M_i \in \mathbf{L}(\mathcal{H}, \mathcal{K})\}$ est l'ensemble des familles dénombrables d'opérateurs linéaires de $\mathcal{H}$ dans $\mathcal{K}$.
- $T^1(\mathcal{H}, \mathcal{K})$ est l'ensemble des transformations admissibles, i.e. des familles $(M_i)_{i \in A} \in T(\mathcal{H}, \mathcal{K})$ vérifiant $\sum_{i \in A} M_i^\dagger M_i = Id_{\mathcal{H}}$.

¹Les transformations admissibles sont parfois appelées mesures généralisées

- $T^{\leq 1}(\mathcal{H}, \mathcal{K})$ est l'ensemble des transformations sous-admissibles, i.e. des familles $(M_i)_{i \in A} \in T(\mathcal{H}, \mathcal{K})$ vérifiant $\sum_{i \in A} M_i^\dagger M_i \leq Id_{\mathcal{H}}^2$.

Nous pouvons introduire une fonction d'interprétation $\mathcal{X}$ permettant de faire le lien entre le domaine utilisé en mécanique quantique pour représenter les transformations admissibles, à savoir $T(\mathcal{H}, \mathcal{K})$, et un domaine plus naturel d'un point de vue computationnel qui est $\mathcal{H}^1 \rightarrow V(\mathcal{K}^1)$:

$$\begin{aligned} \mathcal{X} : T(\mathcal{H}, \mathcal{K}) &\rightarrow (\mathcal{H}^1 \rightarrow V(\mathcal{K}^1)) \\ (M_i)_{i \in A} &\mapsto \lambda |\varphi\rangle \cdot \sum_{i \in A} \langle \varphi | M_i^\dagger M_i | \varphi \rangle \cdot \eta_{\frac{M_i |\varphi\rangle}{\sqrt{\langle \varphi | M_i^\dagger M_i | \varphi \rangle}}} \end{aligned}$$

Ici η_x pour $x \in X$ représente une distribution de probabilité concentrée en x . Plus précisément,

$$\begin{aligned} \eta_x : X &\rightarrow \overline{\mathbb{R}^+} \\ y &\mapsto \begin{cases} 1 & \text{si } x = y \\ 0 & \text{sinon} \end{cases} \end{aligned}$$

Les transformations admissibles sont exactement les familles d'opérateurs linéaires dont l'image par $\mathcal{X}$ est une fonction probabiliste :

Propriété 1.2

$$\begin{aligned} T^1(\mathcal{H}, \mathcal{K}) &= \mathcal{X}^{-1}(\mathcal{H}^1 \rightarrow V^1(\mathcal{K}^1)) \\ T^{\leq 1}(\mathcal{H}, \mathcal{K}) &= \mathcal{X}^{-1}(\mathcal{H}^1 \rightarrow V^{\leq 1}(\mathcal{K}^1)) \end{aligned}$$

Preuve : Soit $(M_i)_{i \in A} \in \mathcal{X}^{-1}(\mathcal{H}^1 \rightarrow V^1(\mathcal{K}^1))$, pour tout $|\varphi\rangle \in \mathcal{H}^1$, $\mathcal{X}((M_i)_{i \in A})(|\varphi\rangle)(\mathcal{K}^1) = 1$, donc $\forall |\varphi\rangle \in \mathcal{H}^1$, $\sum_{i \in A} \langle \varphi | M_i^\dagger M_i | \varphi \rangle = 1$. Ainsi $\sum_{i \in A} M_i^\dagger M_i = Id_{\mathcal{H}}$.

Inversement, si $(M_i)_{i \in A} \in T^1(\mathcal{H}, \mathcal{K})$ alors $\sum_{i \in A} M_i^\dagger M_i = Id_{\mathcal{H}}$. Donc pour tout vecteur normé $|\varphi\rangle \in \mathcal{H}^1$, $\mathcal{X}((M_i)_{i \in A})(|\varphi\rangle)(\mathcal{K}) = \sum_{i \in A} \langle \varphi | M_i^\dagger M_i | \varphi \rangle = 1$.

Une preuve similaire peut être obtenue dans le cas des transformations sous-admissibles. $\square$

La fonction d'interprétation $\mathcal{X}$ n'est pas injective, ainsi des transformations admissibles différentes peuvent avoir la même image par $\mathcal{X}$, i.e. la même action. De telles transformations admissibles sont dites équivalentes :

Définition 1.2 Pour toute famille $F, G \in T(\mathcal{H}, \mathcal{K})$, $F \equiv G$ si et seulement si $\mathcal{X}(F) = \mathcal{X}(G)$.

Il existe notamment des familles infinies équivalentes à des familles finies. En revanche, certaines familles infinies n'ont pas d'équivalent fini.

²si A et B sont des matrices positives, $A \leq B$ si et seulement si $B - A$ est une matrice positive.

1.3.2 Composition spatiale

Avant de considérer la composition spatiale de transformations admissibles, revenons sur le deuxième postulat régissant l'état des systèmes composés, dans le cas où l'état du système est une distribution de probabilité. Etant donné un système S dont l'espace de Hilbert associé est $\mathcal{H}$, on appelle états ou états purs de S les éléments de $\mathcal{H}^1$ et distributions sur les états purs, ou simplement états de S les éléments de $V(\mathcal{H}^1)$.

Si un système S_1 est dans l'état $\nu_1 \in V(\mathcal{H}_{B_1}^1)$ et un système S_2 dans l'état $\nu_2 \in V(\mathcal{H}_{B_2}^1)$, alors le système composé de S_1 et S_2 est dans la distribution $\nu_1 \otimes_p \nu_2 \in V(\mathcal{H}_{B_1 \times B_2}^1)$:

Définition 1.3 Soit $\nu_1 \in V^1(\mathcal{H}_{B_1}^1)$ et $\nu_2 \in V^1(\mathcal{H}_{B_2}^1)$, $\forall |\varphi\rangle \in \mathcal{H}_{B_1 \times B_2}^1$,

$$\nu_1 \otimes_p \nu_2(|\varphi\rangle) = \sum_{\substack{(|\varphi_1\rangle, |\varphi_2\rangle) \in \mathcal{H}_{B_1}^1 \times \mathcal{H}_{B_2}^1 \\ t.q. |\varphi\rangle = |\varphi_1\rangle \otimes |\varphi_2\rangle}} \nu_1(|\varphi_1\rangle) \cdot \nu_2(|\varphi_2\rangle)$$

Si un système S_1 évolue selon $F_1 \in T(\mathcal{H}_1, \mathcal{K}_1)$ et un système S_2 selon $F_2 \in T(\mathcal{H}_2, \mathcal{K}_2)$, alors le système composé de S_1 et S_2 évolue selon $F_1 \otimes F_2$:

Définition 1.4 Soit $\otimes : T(\mathcal{H}_1, \mathcal{K}_1) \times T(\mathcal{H}_2, \mathcal{K}_2) \rightarrow T(\mathcal{H}_1 \otimes \mathcal{H}_2, \mathcal{K}_1 \otimes \mathcal{K}_2)$, tel que

$$(M_i)_{i \in A} \otimes (N_j)_{j \in A'} = (M_i \otimes N_j)_{(i,j) \in A \times A'}.$$

1.3.3 Composition temporelle

La composition (temporelle) de fonctions probabilistes peut être exprimée à l'aide de l'extension de Kleisli :

Définition 1.5 (Extension de Kleisli) Soit $f : X \rightarrow V(Y)$, la fonction $f^\diamond : V(X) \rightarrow V(Y)$ est définie par :

$$f^\diamond = \lambda \nu. \lambda y. \sum_{x \in X} \nu(x) \cdot f(x)(y)$$

Ainsi la composition de $f : \mathcal{H}^1 \rightarrow V(\mathcal{K}^1)$ avec $g : \mathcal{K}^1 \rightarrow V(\mathcal{L}^1)$ est la fonction $g^\diamond \circ f : \mathcal{H}^1 \rightarrow V(\mathcal{L}^1)$.

Si un système S évolue selon une transformation admissible $F \in T(\mathcal{H}, \mathcal{K})$, puis selon $G \in T(\mathcal{K}, \mathcal{L})$, alors l'évolution totale du système est selon $G \circ F : T(\mathcal{H}, \mathcal{L})$:

Définition 1.6 Soit $\circ : T(\mathcal{K}, \mathcal{L}) \times T(\mathcal{H}, \mathcal{K}) \rightarrow T(\mathcal{H}, \mathcal{L})$, tel que

$$(M_i)_{i \in A} \circ (N_j)_{j \in A'} = (M_i N_j)_{(i,j) \in A \times A'}.$$

L'opérateur de composition sur les fonctions probabilistes est une interprétation exacte [Cou97] de l'opérateur de composition sur les transformations admissibles :

Propriété 1.3 $\forall F \in T(\mathcal{H}, \mathcal{K}), \forall G \in T(\mathcal{K}, \mathcal{L}),$

$$\mathcal{X}(G)^\diamond \circ \mathcal{X}(F) = \mathcal{X}(G \circ F)$$

Preuve : Soient $F = (M_i)_{i \in A}$ et $G = (N_j)_{j \in A'}$, pour tout $|\varphi_1\rangle \in \mathcal{H}^1$,

$$\begin{aligned} \mathcal{X}(G)^\diamond(\mathcal{X}(F)(|\varphi_1\rangle)) &= \sum_{|\varphi_2\rangle \in \mathcal{K}^1} \mathcal{X}(F)(|\varphi_1\rangle)(|\varphi_2\rangle) \cdot \mathcal{X}(G)(|\varphi_2\rangle) \\ &= \sum_{i \in A} \langle \varphi_1 | M_i^\dagger M_i | \varphi_1 \rangle \cdot \mathcal{X}(G)\left(\frac{M_i |\varphi_1\rangle}{\sqrt{\langle \varphi_1 | M_i^\dagger M_i | \varphi_1 \rangle}}\right) \\ &= \sum_{(i,j) \in A \times A'} \langle \varphi_1 | M_i^\dagger N_j^\dagger N_j M_i | \varphi_1 \rangle \cdot \eta_{\frac{N_j M_i |\varphi_1\rangle}{\sqrt{\langle \varphi_1 | M_i^\dagger N_j^\dagger N_j M_i | \varphi_1 \rangle}}} \\ &= \mathcal{X}(G \circ F)(|\varphi_1\rangle) \end{aligned}$$

□

La relation d'interprétation exacte entre la composition de transformations admissibles et celle de fonctions probabilistes est représentée par le diagramme suivant :

$$\begin{array}{ccc} (V(\mathcal{K}^1) \rightarrow V(\mathcal{L}^1)) \times (\mathcal{H}^1 \rightarrow V(\mathcal{K}^1)) & \xrightarrow{\circ} & \mathcal{H}^1 \rightarrow V(\mathcal{L}^1) \\ \uparrow (\lambda F. \mathcal{X}(F)^\diamond) \times \mathcal{X} & & \uparrow \mathcal{X} \\ T(\mathcal{K}, \mathcal{L}) \times T(\mathcal{H}, \mathcal{K}) & \xrightarrow{\circ} & T(\mathcal{H}, \mathcal{L}) \end{array}$$

1.3.4 Exemples de transformations admissibles

Parmi les transformations admissibles, certaines jouent un rôle particulier :

- Les transformations admissibles représentées par une famille à un seul élément non nul U . La condition de complétude devient $U^\dagger U = Id_{\mathcal{H}}$, donc U est une *isométrie*. De plus, si U est surjectif alors U est *unitaire*. Des exemples de transformations unitaires sont données dans le chapitre 3, dans le cadre des circuits quantiques, modèle de calcul fondé exclusivement sur ce type de transformation.
- Un autres exemple d'isométrie est l'initialisation. En effet étant donné un état $|\varphi_0\rangle \in \mathcal{H}^1$, $(|\varphi_0\rangle \langle |)$ est une transformation admissible formée d'un unique opérateur linéaire de $\mathbb{C}$ dans $\mathcal{H}^1$ réalisant une initialisation dans l'état $|\varphi_0\rangle$.

- Les transformations admissibles représentées par une famille de projecteurs orthogonaux de $\mathcal{H}$ dans lui-même. La condition de complétude devient $\sum_{i \in A} P_i = Id_{\mathcal{H}}$. Une telle transformation est appelée *mesure projective* et est entièrement déterminée par un observable $\mathcal{O} = \sum_{i \in A} \lambda_i \cdot P_i$, où $\lambda_i \in \mathbb{R}^+$ avec $\lambda_i = \lambda_j \Rightarrow i = j$.
- Les transformations admissibles représentées par une famille d'opérateurs linéaires de $\mathcal{H}$ dans $\mathbb{C}$. Une telle transformation est appelée *mesure destructrice*.

1.4 Transformations admissibles vs transformations unitaires

Nous avons choisi de présenter le postulat 3 sur les évolutions quantiques sous sa forme unifiée, i.e. en ne présentant qu'un seul type de transformation : les transformations admissibles. Or, une autre présentation, non unifiée, de l'évolution d'un système quantique est souvent donnée : un premier postulat 3' décrit l'évolution des systèmes isolés, de tels systèmes ont une évolution unitaire ; puis un deuxième postulat 4' introduit les mesures projectives.

Dans cette section nous comparons rapidement les deux présentations possibles des postulats de la mécanique quantique autour de deux critères : l'équivalence des deux formalismes et la compositionnalité des opérations quantiques.

Les exemples de transformations admissibles présentés en section 1.3.4 incluent les transformations unitaires et les mesures projectives. Inversement, toute transformation admissible peut être décomposée en une initialisation, une transformation unitaire et une mesure destructrice.

Lemme 1.1 *Pour toute transformation admissible $F = (M_i)_{i \in A} \in T(\mathcal{H}_{B_1}, \mathcal{H}_{B_2})$, il existe une initialisation $I_F \in T(\mathbb{C}, \mathcal{H}_{A \times B_2})$, une transformation unitaire $U_F \in T(\mathcal{H}_{B_1 \times A \times B_2}, \mathcal{H}_{B_1 \times A \times B_2})$ et une mesure destructrice $D_F \in T(\mathcal{H}_{B_1 \times A}, \mathbb{C})$ telles que*

$$F = (D_F \otimes Id_{\mathcal{H}_{B_2}}) \circ U_F \circ (Id_{\mathcal{H}_{B_1}} \otimes I_F)$$

Preuve : Soient $|\varphi_0\rangle \in \mathcal{H}_{B_2}$ et $|\varphi_2\rangle \in \mathcal{H}_{B_2}$ deux états fixés. Soit U_F tel que $\forall |\varphi\rangle \in \mathcal{H}_{B_1}, U_F(|\varphi\rangle |\varphi_0\rangle |\varphi_2\rangle) = \sum_{i \in A} (|\varphi_1\rangle |i\rangle (M_i |\varphi\rangle))$. U_F peut être complétée en une transformation unitaire de $\mathcal{H}_{B_1 \times A \times B_2}$ dans $\mathcal{H}_{B_1 \times A \times B_2}$. Soit $I_F = (|\varphi_0 \varphi_2\rangle \langle |) \in T(\mathbb{C}, \mathcal{H}_{A \times B_2})$, enfin soit $D = (d_i)_{i \in A} \in T(\mathcal{H}_{B_1 \times A}, \mathbb{C})$ avec $d_i = | \rangle \langle \varphi_1 i |$.

On a alors $U_F \circ (Id_{\mathcal{H}_{B_1}} \otimes I_F) = \left(\sum_{i \in A, \tau \in B_1} |\varphi_1\rangle |i\rangle (M_i |\tau\rangle) \langle \tau| \right)$, donc $(D_F \otimes Id_{\mathcal{H}_{B_2}}) \circ U_F \circ (Id_{\mathcal{H}_{B_1}} \otimes I_F) = \left(\sum_{\tau \in B_1} M_i |\tau\rangle \langle \tau| \right)_{i \in A} = F$. $\square$

Les deux présentations sont donc équivalentes, même si les transformations admissibles permettent une présentation formelle des initialisations et des mesures

destructrices qui sont souvent traitées de façon ad hoc en l'absence du formalisme des transformations admissibles.

Dans notre cas, d'un point de vue informatique, un avantage décisif des transformations admissibles est leur compositionnalité. En effet la composition temporelle ($\circ$) ou spatiale ($\otimes$) de deux transformations admissibles est une transformation admissible. Cette propriété est également vraie pour les transformations unitaires, mais n'est plus vérifiée dans le cas des mesures projectives, ni dans le cas où mesures projectives et transformations unitaires sont composées.

Ainsi une présentation non unifiée est souvent suffisante dans le cadre des modèles de calcul quantique réversibles, fondés sur les transformations unitaires. En revanche, l'étude des modèles de calcul quantique contrôlé classiquement est facilitée par l'utilisation du formalisme des transformations admissibles.

1.5 Conclusion

Les postulats sont une véritable pierre angulaire de la mécanique quantique. Ils sont également le fondement de l'informatique quantique. Leur présentation de ce chapitre a permis de mettre en évidence des propriétés spécifiques tels que les phénomènes de superposition et d'intrication.

Les évolutions quantiques ont également des propriétés fortes, mêlant linéarité, non déterminisme et probabilisme. Une dualité dans la représentation des évolutions quantiques a été mise en avant : d'un côté les fonctions probabilistes, de l'autre les transformations admissibles mettant en évidence l'aspect non déterministe et linéaire de la mécanique quantique. Une relation originale d'interprétation exacte a été mise en évidence entre les deux représentations.

La représentation des évolutions quantiques présentée dans ce chapitre constitue le fondement de la définition d'une sémantique quantique. En effet, dans le chapitre 5, un modèle formel de calcul quantique, le $\mathbf{q}$ -calcul est introduit. A chaque terme est associé une transformation admissible et une fonction probabiliste représentant la sémantique, i.e. l'action de ce terme.

Il existe une représentation alternative pour les états quantiques : les matrices de densité. Elle induit une troisième représentation pour les évolutions quantiques. L'objectif du chapitre suivant est d'introduire le formalisme des matrices de densité, tout en le situant par rapport aux formalismes vus dans ce chapitre.

Chapitre 2

Représentations des états quantiques

2.1 Etats mixtes et matrices de densité

Dans le cas général, l'état d'un système quantique est une distribution de probabilité sur les états purs, i.e un élément de $V^1(\mathcal{H}^1)$. Une notation alternative, introduite par von Neumann, permet de décrire ces distributions de probabilité sous forme d'endomorphismes de $\mathcal{H}$, appelés matrices de densité. Dans cette section, le formalisme des matrices de densité est présenté comme une abstraction des distributions de probabilité. En effet, le formalisme des matrices de densité n'est pas une représentation équivalente aux distributions de probabilité, mais correspond à une abstraction. Cela implique que deux distributions différentes d'états purs peuvent être abstraites en une même matrice de densité. Nous verrons que cette abstraction est justifiée par l'*indistingabilité* 2.1.4 des distributions abstraites en une même matrice de densité. Les postulats de la mécanique quantique sont présentés dans cette section sous leur forme abstraite.

2.1.1 Abstraction

Etant donné un espace de Hilbert $\mathcal{H}$, une matrice de densité ρ de $\mathcal{H}$ est une matrice Hermitienne (i.e. $\rho^\dagger = \rho$) satisfaisant $Tr(\rho) \leq 1$ ¹. L'ensemble des matrices de densité sur $\mathcal{H}$ est noté $D(\mathcal{H})$. Le sous-ensemble des matrices de densité $\rho \in D(\mathcal{H})$ telles que $Tr(\rho) = 1$ est noté $D^1(\mathcal{H})$.

A toute distribution de sous-probabilité sur les états purs de $\mathcal{H}$, est associée une matrice de densité à l'aide de la fonction d'abstraction α suivante :

¹ $Tr(\rho)$ est la trace de la matrice ρ .

$$\begin{aligned} \alpha &: V^{\leq 1}(\mathcal{H}^1) \rightarrow D(\mathcal{H}) \\ \nu &\mapsto \sum_{|\varphi\rangle \in \mathcal{H}} \nu(|\varphi\rangle) |\varphi\rangle\langle\varphi| \end{aligned}$$

Cette fonction d'abstraction est surjective. En revanche, elle n'est pas injective, deux distributions différentes peuvent avoir la même abstraction. Par exemple, dans l'espace de Hilbert $\mathcal{H}_{\{a,b\}}$ de dimension 2, considérons les distributions de probabilité ν_1 et ν_2 telles que $\nu_1(|a\rangle) = 1/2$, $\nu_1(|b\rangle) = 1/2$, $\nu_2(\frac{|a\rangle+|b\rangle}{\sqrt{2}}) = 1/2$ et $\nu_2(\frac{|a\rangle-|b\rangle}{\sqrt{2}}) = 1/2$. Ces deux distributions différentes ont la même abstraction : $\alpha(\nu_1) = \alpha(\nu_2)$.

2.1.2 Système composé

La version abstraite du deuxième postulat de la mécanique quantique décrit la matrice de densité ρ d'un système composé de deux sous-systèmes S_1 et S_2 dont on connaît les matrices de densité ρ_1 et ρ_2 .

On remarque que le produit tensoriel $\otimes$ sur les matrices de densité est l'exacte α -abstraction de l'opérateur $\otimes_p$ sur les distributions de probabilités :

Propriété 2.1 $\forall \nu_1 \in V^{\leq 1}(\mathcal{H}_1^1), \forall \nu_2 \in V^{\leq 1}(\mathcal{H}_2^1),$

$$\alpha(\nu_1) \otimes \alpha(\nu_2) = \alpha(\nu_1 \otimes_p \nu_2)$$

Preuve : Soient $\nu_1 \in V^{\leq 1}(\mathcal{H}_1^1)$, et $\nu_2 \in V^{\leq 1}(\mathcal{H}_2^1)$,

$$\begin{aligned} \alpha(\nu_1 \otimes_p \nu_2) &= \sum_{|\varphi\rangle \in (\mathcal{H}_1 \otimes \mathcal{H}_2)^1} \sum_{\substack{(|\varphi_1\rangle, |\varphi_2\rangle) \in \mathcal{H}_{B_1}^1 \times \mathcal{H}_{B_2}^1 \\ t.q. |\varphi\rangle = |\varphi_1\rangle \otimes |\varphi_2\rangle}} \nu_1(|\varphi_1\rangle) \cdot \nu_2(|\varphi_2\rangle) |\varphi\rangle\langle\varphi| \\ &= \sum_{(|\varphi_1\rangle, |\varphi_2\rangle) \in \mathcal{H}_{B_1}^1 \times \mathcal{H}_{B_2}^1} \nu_1(|\varphi_1\rangle) \cdot \nu_2(|\varphi_2\rangle) |\varphi_1\varphi_2\rangle\langle\varphi_1\varphi_2| \\ &= \alpha(\nu_1) \otimes \alpha(\nu_2) \end{aligned}$$

□

La relation d'abstraction entre $\otimes_p$ et $\otimes$ est représentée par le diagramme suivant :

$$\begin{array}{ccc} D(\mathcal{H}_1) \times D(\mathcal{H}_2) & \xrightarrow{\quad \otimes \quad} & D(\mathcal{H}_1 \otimes \mathcal{H}_2) \\ \alpha \times \alpha \uparrow & & \uparrow \alpha \\ V^{\leq 1}(\mathcal{H}_1^1) \times V^{\leq 1}(\mathcal{H}_2^1) & \xrightarrow{\quad \otimes_p \quad} & V^{\leq 1}((\mathcal{H}_1 \otimes \mathcal{H}_2)^1) \end{array}$$

On peut donc énoncer une version abstraite du deuxième postulat de la mécanique quantique :

Postulat 2'. *L'espace des états d'un système physique composé de sous-systèmes est le produit tensoriel des espaces des états de ses sous-systèmes.*

2.1.3 Evolution

Le troisième postulat indique que toute transformation quantique peut être représentée par une transformation admissible, i.e. une famille d'opérateurs linéaires.

Postulat 3'. *Tout système quantique évolue selon une transformation admissible. Une transformation admissible agissant sur $\mathcal{H}$ est décrite par une famille dénombrable $(M_i)_{i \in A}$ d'opérateurs linéaires vérifiant la condition de complétude :*

$$\sum_{i \in A} M_i^\dagger M_i = Id_{\mathcal{H}}$$

Ainsi, si une transformation admissible $(M_i)_{i \in A}$ est appliquée à $\rho \in D(\mathcal{H})$, alors avec probabilité $p(i) = \text{Tr}(M_i^\dagger M_i \rho)$ le résultat classique i est observé et l'état du système après la transformation est :

$$\frac{M_i \rho M_i^\dagger}{p(i)}$$

La fonction d'interprétation $\mathcal{X}$ permet d'associer à chaque transformation admissible une fonction probabiliste. Une seconde fonction d'interprétation $\mathcal{X}^\natural$ associe à chaque transformation sous-admissible une application linéaire sur les matrices de densité, appelée *super-opérateur* :

$$\begin{aligned} \mathcal{X}^\natural : T^{\leq 1}(\mathcal{H}, \mathcal{K}) &\rightarrow \mathbf{L}(D(\mathcal{H}), D(\mathcal{K})) \\ (M_i)_{i \in A} &\mapsto \lambda \rho. \sum_{i \in A} M_i \rho M_i^\dagger \end{aligned}$$

Les opérations de compositions spatiale et temporelle de super-opérateurs peuvent être vues comme des interprétations exactes des opérations de compositions sur les transformations admissibles :

Le produit tensoriel sur les super-opérateurs est une interprétation exacte du produit tensoriel sur les transformations sous-admissibles :

Propriété 2.2 *Pour tout $F : T(\mathcal{H}_1, \mathcal{K}_1)$ et tout $G : T(\mathcal{H}_2, \mathcal{K}_2)$,*

$$\mathcal{X}^\natural(F) \otimes \mathcal{X}^\natural(G) = \mathcal{X}^\natural(F \otimes G)$$

$$\begin{array}{ccc}
L(D(\mathcal{H}_1), D(\mathcal{K}_1)) \times L(D(\mathcal{H}_2), D(\mathcal{K}_2)) & \xrightarrow{\otimes} & L(D(\mathcal{H}_1 \otimes \mathcal{H}_2), D(\mathcal{K}_1 \otimes \mathcal{K}_2)) \\
\uparrow \mathcal{X}^\natural \times \mathcal{X}^\natural & & \uparrow \mathcal{X}^\natural \\
T(\mathcal{H}_1, \mathcal{K}_1) \times T(\mathcal{H}_2, \mathcal{K}_2) & \xrightarrow{\otimes} & T(\mathcal{H}_1 \otimes \mathcal{H}_2, \mathcal{K}_1 \otimes \mathcal{K}_2)
\end{array}$$

Preuve : Soient $F = (M_i)_{i \in A}$ et $G = (N_j)_{j \in A'}$,

$$\begin{aligned}
\mathcal{X}^\natural(F) \otimes \mathcal{X}^\natural(G)(\rho_1 \otimes \rho_2) &= \sum_{i \in A} M_i \rho_1 M_i^\dagger \otimes \sum_{j \in A'} N_j \rho_2 N_j^\dagger \\
&= \sum_{(i,j) \in A \times A'} (M_i \otimes N_j)(\rho_1 \otimes \rho_2)(M_i^\dagger \otimes N_j^\dagger) \\
&= \mathcal{X}(F \otimes G)(\rho_1 \otimes \rho_2)
\end{aligned}$$

Par linéarité, on conclut que pour tout $\rho \in D(\mathcal{H}_1 \otimes \mathcal{H}_2)$, $\mathcal{X}^\natural(F) \otimes \mathcal{X}^\natural(G)(\rho) = \mathcal{X}^\natural(F \otimes G)(\rho)$. $\square$

L'opérateur de composition sur les super-opérateurs est une interprétation exacte de l'opérateur de composition sur les transformations admissibles :

Propriété 2.3 Pour tout $F \in T(\mathcal{H}, \mathcal{K})$, et tout $G \in T(\mathcal{K}, \mathcal{L})$,

$$\mathcal{X}^\natural(G) \circ \mathcal{X}^\natural(F) = \mathcal{X}^\natural(G \circ F)$$

$$\begin{array}{ccc}
L(D(\mathcal{K}), D(\mathcal{L})) \times L(D(\mathcal{H}), D(\mathcal{K})) & \xrightarrow{\circ} & L(D(\mathcal{H}), D(\mathcal{L})) \\
\uparrow \mathcal{X}^\natural \times \mathcal{X}^\natural & & \uparrow \mathcal{X}^\natural \\
T(\mathcal{K}, \mathcal{L}) \times T(\mathcal{H}, \mathcal{K}) & \xrightarrow{\circ} & T(\mathcal{H}, \mathcal{L})
\end{array}$$

Preuve : Soient $F = (M_i)_{i \in A}$ et $G = (N_j)_{j \in A'}$, pour tout $\rho \in D(\mathcal{H})$,

$$\begin{aligned}
\mathcal{X}^\natural(G)(\mathcal{X}^\natural(F)(\rho)) &= \sum_{j \in A'} N_j \left(\sum_{i \in A} M_i \rho M_i^\dagger \right) N_j^\dagger \\
&= \sum_{(i,j) \in A \times A'} N_j M_i \rho M_i^\dagger N_j^\dagger \\
&= \mathcal{X}^\natural(G \circ F)(\rho)
\end{aligned}$$

$\square$

Enfin la fonction d'interprétation $\mathcal{X}^\natural$ est elle même une α -abstraction de $\mathcal{X}$:

Propriété 2.4 *Pour toute transformation sous-admissible $F \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$, $\mathcal{X}^{\natural}(F)$ est une exacte α -abstraction de $\mathcal{X}(F)$:*

$$\begin{array}{ccc} D(\mathcal{H}) & \xrightarrow{\mathcal{X}^{\natural}(F)} & D(\mathcal{K}) \\ \alpha \uparrow & & \uparrow \alpha \\ V^{\leq 1}(\mathcal{H}^1) & \xrightarrow{\mathcal{X}(F)^{\diamond}} & V^{\leq 1}(\mathcal{K}^1) \end{array}$$

Preuve : Soit $F = (M_i)_{i \in A}$, pour tout $\nu \in V^{\leq 1}(\mathcal{H}^1)$,

$$\begin{aligned} \alpha(\mathcal{X}(F)^{\diamond}(\nu)) &= \alpha(\lambda |\varphi_2\rangle \cdot \sum_{|\varphi_1\rangle \in \mathcal{H}^1} \nu(|\varphi_1\rangle) \cdot \mathcal{X}(F)(|\varphi_1\rangle)(|\varphi_2\rangle)) \\ &= \sum_{|\varphi_2\rangle \in \mathcal{K}^1} \sum_{|\varphi_1\rangle \in \mathcal{H}^1} \nu(|\varphi_1\rangle) \cdot (\sum_{i \in A} \langle \varphi_1 | M_i^{\dagger} M_i | \varphi_1 \rangle \eta_{\frac{M_i |\varphi_1\rangle}{\sqrt{\langle \varphi_1 | M_i^{\dagger} M_i | \varphi_1 \rangle}}}(|\varphi_2\rangle)) \\ &= \sum_{|\varphi_1\rangle \in \mathcal{H}^1} \sum_{i \in A} \nu(|\varphi_1\rangle) \cdot M_i |\varphi_1\rangle \langle \varphi_1 | M_i^{\dagger} \\ &= \mathcal{X}^{\natural}(F)(\alpha(\nu)) \end{aligned}$$

□

2.1.4 Indistingabilité

Le formalisme des matrices de densité est une abstraction des distributions de probabilité sur les états purs ; ainsi deux distributions différentes ν_1, ν_2 peuvent être abstraites en une même matrice de densité. Cette perte d'information est justifiée par l'indistingabilité de ν_1 et ν_2 . En effet, si les deux distributions sont abstraites en une même matrice de densité, alors aucune opération quantique ne permettra de distinguer ces deux distributions :

Propriété 2.5 *Soient ν_1 et $\nu_2 \in V^{\leq 1}(\mathcal{H})$ telles que $\alpha(\nu_1) = \alpha(\nu_2)$. Pour toute transformation sous-admissible $(M_i)_{i \in A} \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$, et pour tout $i \in A$, les probabilités d'obtenir le résultat i si le système est dans l'état ν_1 ou ν_2 sont identiques.*

Preuve : Etant donnée une transformation sous-admissible $(M_i)_{i \in A} \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$, la probabilité d'obtenir un résultat $i \in A$, si le système est dans un état ν , est :

$$\begin{aligned}
p(i) &= \sum_{|\varphi\rangle \in \mathcal{H}^1} \nu(|\varphi\rangle) \langle \varphi | M_i^\dagger M_i | \varphi \rangle \\
&= \text{Tr}(\sum_{|\varphi\rangle \in \mathcal{H}^1} \nu(|\varphi\rangle) \langle \varphi | M_i^\dagger M_i | \varphi \rangle) \\
&= \text{Tr}(M_i^\dagger M_i (\sum_{|\varphi\rangle \in \mathcal{H}^1} \nu(|\varphi\rangle) |\varphi\rangle \langle \varphi|)) \\
&= \text{Tr}(M_i^\dagger M_i \alpha(\nu))
\end{aligned}$$

Ainsi la probabilité d'obtenir un résultat classique dépend uniquement de $\alpha(\nu)$. $\square$

Cette indistingabilité induit une équivalence *observationnelle* sur les distributions de probabilité :

Définition 2.1 Pour tout $\nu, \mu \in V^{\leq 1}(\mathcal{H})$, $\nu \equiv_{obs} \mu$ si et seulement si $\alpha(\nu) = \alpha(\mu)$.

En suivant le même modèle, on peut également définir une équivalence observationnelle sur les transformations admissibles :

Définition 2.2 Pour tout $F, G \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$, $F \equiv_{obs} G$ si et seulement si $\mathcal{X}^{\natural}(F) = \mathcal{X}^{\natural}(G)$.

2.2 Autres formalismes

Outre les distributions de probabilité et les matrices de densité, d'autres formalismes sont utilisés en informatique quantique pour la représentation des états quantiques. Nous dressons un panorama rapide de ces formalismes. Certains sont équivalents aux formalismes déjà évoqués, d'autres seront approfondis dans des chapitres spécifiques dans cette thèse.

2.2.1 Distributions sur les sous-espaces de Hilbert

Plutôt que de travailler avec des distributions de probabilités sur la sphère $\mathcal{H}^1$, les distributions de probabilité sur les sous-espaces vectoriels de $\mathcal{H}$ ont l'avantage de tenir compte de la structure d'espace vectoriel de l'espace d'états.

Ainsi, alors qu'une distribution de probabilité ν sur l'ensemble $\mathcal{H}^1$ doit vérifier la condition suivante : $\forall X, Y \subseteq \mathcal{H}^1$, tels que $X \cap Y = \emptyset$,

$$\nu(X \cup Y) = \nu(X) + \nu(Y)$$

la condition vérifiée par une distribution de probabilité μ sur les sous-espace vectoriels de $\mathcal{H}$ est : pour tout sous-espace vectoriel X, Y de $\mathcal{H}$, tels que X et Y sont orthogonaux,

$$\mu(X \oplus Y) = \mu(X) + \mu(Y)$$

où $\oplus$ est la somme directe d'espaces vectoriels.

De telles distributions sont utilisées en logique quantique [BvN36] mais aussi dans le cadre de la définition de domaine sémantique spécifique à l'informatique quantique [Kas03].

Le théorème de Gleason [Gle57] permet de montrer que les formalismes des matrices de densité et celui des distributions sur les sous-espaces vectoriels sont équivalents.

2.2.2 Ensemble d'états purs non normés

Une représentation prometteuse est d'encoder la probabilité associée à un état pur dans sa norme. Cet encodage est par exemple utilisé dans [DKP04a]. En effet, un état pur $|\varphi\rangle$ est un vecteur normé d'un certain espace de Hilbert $\mathcal{H}$; si la probabilité associée à $|\varphi\rangle$ est p , alors le vecteur $\sqrt{p}|\varphi\rangle$ permet d'encoder efficacement la probabilité dans la norme de l'état. L'espace des états purs devient donc dans ce formalisme la boule unité $\mathcal{H}^{\leq 1}$ de l'espace de Hilbert associé au système.

Une distribution discrète est encodée par un sous ensemble de $\mathcal{H}$, à l'aide de la fonction suivante :

$$\begin{aligned} e : V(\mathcal{H}^1) &\rightarrow \mathcal{P}(\mathcal{H}) \\ \nu &\mapsto \bigcup_{|\varphi\rangle \in \mathcal{H}^1} \{\sqrt{\nu(|\varphi\rangle)}|\varphi\rangle\} \end{aligned}$$

Un autre avantage réside dans la formulation du troisième postulat dans ce formalisme. En effet, puisque l'aspect probabiliste est encodé dans la norme des états, une évolution quantique se réduit alors à une transformation non déterministe, formalisée naturellement par une fonction sur les parties de $\mathcal{H}^{\leq 1}$:

$$\begin{aligned} \mathcal{X}' : T(\mathcal{H}, \mathcal{K}) &\rightarrow (\mathcal{P}(\mathcal{H}) \rightarrow \mathcal{P}(\mathcal{K})) \\ (M_i)_{i \in A} &\mapsto \lambda E. \uplus_{i \in A} \uplus_{|\varphi\rangle \in E} \{M_i |\varphi\rangle\} \end{aligned}$$

où $\uplus$ ne représente pas l'union habituelle sur les ensembles (voir définition 2.3). Avant de donner la définition de $\uplus$, remarquons que l'union usuelle $\cup$ ne respecte pas l'encodage des probabilités. Illustration dans l'exemple suivant :

Exemple 2.1 Dans l'espace $\mathcal{H}_{\{a,b\}}$, la famille $F = (|a\rangle\langle a|, |a\rangle\langle b|)$ est une transformation admissible. L'image de $\frac{|a\rangle+|b\rangle}{\sqrt{2}}$ par F est :

$$\mathcal{X}'(F) \left(\left\{ \frac{|a\rangle+|b\rangle}{\sqrt{2}} \right\} \right) = \left\{ \frac{|a\rangle}{\sqrt{2}} \right\} \uplus \left\{ \frac{|a\rangle}{\sqrt{2}} \right\}$$

Le résultat attendu est $\{|a\rangle\}$ car, avec probabilité 1, le résultat de l'application de F sur l'état $\frac{|a\rangle+|b\rangle}{\sqrt{2}}$ produit l'état $|a\rangle$. Or, d'après la définition de l'union usuelle $\cup$, $\left\{ \frac{|a\rangle}{\sqrt{2}} \right\} \cup \left\{ \frac{|a\rangle}{\sqrt{2}} \right\} = \left\{ \frac{|a\rangle}{\sqrt{2}} \right\}$, c'est pourquoi nous définissons l'opérateur $\uplus$:

Définition 2.3 Pour tout $E_1, E_2 \in \mathcal{P}(\mathcal{H})$, $|\varphi\rangle \in E_1 \uplus E_2$ si et seulement si l'une des conditions suivantes est vérifiée :

- (a) $|\varphi\rangle \in E_1$ et $|\varphi\rangle \notin E_2$
- (b) $|\varphi\rangle \in E_2$ et $|\varphi\rangle \notin E_1$
- (c) $\frac{|\varphi\rangle}{\sqrt{2}} \in E_1$ et $\frac{|\varphi\rangle}{\sqrt{2}} \in E_2$

Ce formalisme se révèle être équivalent aux distributions sur les vecteurs de $\mathcal{H}^1$. En effet, la fonction d suivante est l'inverse de e :

$$\begin{aligned} d : \mathcal{P}(\mathcal{H}) &\rightarrow V(\mathcal{H}^1) \\ E &\mapsto \lambda |\varphi\rangle \cdot \sum_{|\psi\rangle \in E, \|\psi\rangle = \sqrt{\| |\varphi\rangle \|}} \|\psi\rangle \end{aligned}$$

Le formalisme des distributions discrètes sur la sphère $\mathcal{H}^1$ et le formalisme des sous-ensembles de la boule unité $\mathcal{H}^{\leq 1}$ sont donc équivalents. La stabilité de l'encodage des probabilités dans le second cas nécessite la définition de fonctions adaptées, c'est pourquoi la représentation à l'aide de distributions discrètes est favorisée dans cette thèse.

2.3 Conclusion

Les matrices de densité sont une représentation alternative aux distributions de probabilité des états quantiques. Elles en sont une abstraction, justifiée par l'indistingabilité des distributions d'états purs décrites par une même matrice de densité. Les opérateurs de compositions sur les matrices de densité sont donc des abstractions des opérations correspondantes sur les distributions d'états purs.

Cette abstraction apporte également une troisième alternative à la représentation des évolutions quantiques : les *super-opérateurs*. Les transformations admissibles peuvent être interprétées, via $\mathcal{X}$, comme des fonctions probabilistes, elles-mêmes abstraites, via α , en des super-opérateurs. Les transformations admissibles peuvent également être directement interprétées, via $\mathcal{X}^\natural$, comme des super-opérateurs.

Enfin, la situation est complétée par deux formalismes, les distributions sur les sous-espaces vectoriels et les sous-ensembles de la boule unité, qui se révèlent être équivalents respectivement aux matrices de densité et aux distributions sur les états purs.

Chapitre 3

Modèles de calcul quantique réversibles

L'objectif de l'informatique quantique est de mettre à profit les lois de la mécanique quantique afin de mener un calcul. Ce chapitre présente deux modèles standards du calcul quantique : les *circuits quantiques* et la *machine de Turing quantique*. Un des objectifs de cette thèse est de montrer que ces modèles déjà traditionnels du calcul quantique n'exploitent qu'une partie du potentiel de la mécanique quantique : la partie *réversible*. Des modèles ayant une plus grande expressivité (*q-calcul*, *machine de Turing quantique à contrôle classique*) sont présentés dans les chapitres 5 et 6.

Un algorithme se décompose traditionnellement en trois étapes :

Initialisation $\rightarrow$ **Transformation** $\rightarrow$ **Observation**

Le calcul quantique *réversible* correspond au cas où l'étape de transformation est assurée uniquement par des transformations unitaires qui sont des transformations admissibles particulières ayant la propriété d'être réversibles (voir chapitre 1). L'étape d'initialisation consiste à mettre le système dans un des états d'une base donnée, appelée base standard. Enfin, l'étape d'observation consiste en une mesure de tout ou partie du système selon cette même base standard.

Bien que représentant une restriction forte, la limitation à un calcul quantique réversible se justifie par la capacité de simuler les modèles de calcul quantique non réversibles, comme l'illustre le théorème 6.5.

Le modèle des circuits quantiques est un modèle de bas niveau, dans lequel l'information est encodée dans un système quantique élémentaire à deux états de base appelé *qubit* (quantum bit), et où les transformations unitaires sont décrites par des graphes acycliques où chaque nœud est une transformation unitaire élémentaire appelée porte.

Plus abstrait, le modèle des machines de Turing quantiques est une généralisation de la machine de Turing probabiliste, où la fonction de transition devient quantique et transforme une configuration de la machine en une superposition de configurations.

3.1 Circuit quantique

Les circuits quantiques [KSV02, NC00, Gru99], encore largement utilisés pour décrire les algorithmes quantiques malgré leur limitation en terme d'expressivité, sont une généralisation des circuits booléens. La brique de base de l'information est donc une généralisation quantique du bit, le qubit.

3.1.1 Qubit

Un qubit est un système physique (ou du moins une représentation abstraite d'un système physique) dont l'espace d'état est de dimension deux et possède une base particulière $\{|0\rangle, |1\rangle\}$ appelée base standard. Ainsi l'état $|\varphi\rangle$ d'un qubit s'écrit :

$$|\varphi\rangle = \alpha |0\rangle + \beta |1\rangle,$$

avec $\alpha, \beta \in \mathbb{C}$ et $|\alpha|^2 + |\beta|^2 = 1$. L'espace de Hilbert d'un qubit est noté $\mathcal{H}_{\{0,1\}}$, ou simplement $\mathbb{C}^2$.

L'état $|\psi\rangle$ d'un registre de n qubits est un vecteur normé de $\mathcal{H}_{\{0,1\}^n}$ ou encore $\mathbb{C}^{2^n}$:

$$|\psi\rangle = \sum_{x \in \{0,1\}^n} \alpha_x |x\rangle$$

avec $\sum_{x \in \{0,1\}^n} |\alpha_x|^2 = 1$.

3.1.2 Porte unitaire

Une porte unitaire d'ordre k est une application linéaire $U \in \mathbf{L}(\mathcal{H}_{\{0,1\}^k}, \mathcal{H}_{\{0,1\}^k})$ telle que $U^\dagger U = Id_{\{0,1\}^k}$. Une porte unitaire d'ordre k peut être représentée par une matrice $2^k \times 2^k$, ou encore à l'aide de la représentation de Dirac.

Exemple 3.1

– **Hadamard.** La porte Hadamard est une porte d'ordre 1 :

$$H = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} = \frac{1}{\sqrt{2}} (|0\rangle\langle 0| + |0\rangle\langle 1| + |1\rangle\langle 0| - |1\rangle\langle 1|)$$

La porte Hadamard permet de générer une superposition uniforme $\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ à partir de l'état $|0\rangle$.

- **Pauli.** Les portes de Pauli sont des portes d'ordre 1, transformant un état de base en un autre état de base :

$$X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} = |0\rangle\langle 1| + |1\rangle\langle 0|$$

$$Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} = i|1\rangle\langle 0| - i|0\rangle\langle 1|$$

$$Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = |0\rangle\langle 0| - |1\rangle\langle 1|$$

X , Y et Z sont parfois notés σ_x , σ_y et σ_z respectivement. Un opérateur de Pauli est un élément du groupe multiplicatif engendré par X , Y et Z .

- **Rotation.** Par définition, on appellera rotation autour de l'axe z^1 d'un angle θ , la porte décrite par :

$$R_z(\theta) = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\theta} \end{pmatrix} = |0\rangle\langle 0| + e^{i\theta} |1\rangle\langle 1|$$

En particulier si $\theta = \pi/4$:

$$T = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\pi/4} \end{pmatrix} = |0\rangle\langle 0| + e^{i\pi/4} |1\rangle\langle 1|$$

- **Porte Contrôlée.** Si U est une porte d'ordre k , alors ΛU est une porte d'ordre $k + 1$:

$$\Lambda U = \begin{pmatrix} Id_{\mathbb{C}^{2^k}} & 0 \\ 0 & U \end{pmatrix} = |0\rangle\langle 0| \otimes Id_{\mathbb{C}^{2^k}} + |1\rangle\langle 1| \otimes U$$

En particulier, ΛX est appelée "controlled-Not", ΛZ "controlled-Z" et $\Lambda(\Lambda X)$ (noté $\Lambda^2 X$) "porte de Toffoli" :

$$\Lambda X = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} = |00\rangle\langle 00| + |01\rangle\langle 01| + |10\rangle\langle 11| + |11\rangle\langle 10|$$

¹Cette vision géométrique des opérations unitaires s'appuie sur la représentation des états quantiques dans la sphère de Bloch [NC00].

$$\Lambda Z = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix} = |00\rangle\langle 00| + |01\rangle\langle 01| + |10\rangle\langle 10| - |11\rangle\langle 11|$$

3.1.3 Circuit

Etant donné un ensemble fini (ou registre) R de qubits, si G est une porte unitaire d'ordre k et $A \subseteq R$ est un sous-ensemble ordonné de taille k , alors $G[A]$ désigne la transformation unitaire s'appliquant sur les qubits de A .

Définition 3.1 (Circuit quantique) [NO02a, KSV02] Soit $\mathcal{G}$ un ensemble fixé de portes unitaires ($\mathcal{G}$ est appelé base) et R un registre de n qubits. Un circuit quantique C_n sur n qubits fondé sur $\mathcal{G}$ est une suite finie $(G_1, A_1), \dots, (G_s, A_s)$, où $G_j \in \mathcal{G}$, $A_j \subseteq R$ est un ensemble ordonné de qubits, et s est la taille de C_n pour $\mathcal{G}$. La transformation unitaire réalisée par le circuit est $U = G_L[A_L] \dots G_1[A_1]$.

Exemple 3.2 Etant donné l'ensemble $\mathcal{G} = \{\Lambda X, H\}$ de portes unitaires et un registre $R = \{1, 2\}$ de deux qubits, $(H, \{1\}), (\Lambda X, \{1, 2\})$ est un circuit réalisant la transformation unitaire $\Lambda X[1, 2]H[1]$ (voir figure 3.1).

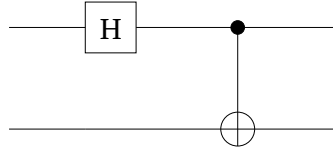


FIG. 3.1 – Circuit réalisant H et ΛX

Un circuit quantique permet de décrire un algorithme pour une taille de registre (nombre de qubits) fixée. Dans le cas général, un algorithme peut être décrit par une ensemble de circuits, un pour chaque taille de registre possible :

Une famille uniforme de circuits quantiques est un ensemble de circuits $F = \{C_n\}$ tel qu'il existe une machine de Turing (classique) produisant une description de C_n à partir de l'entrée n en temps polynomial en la taille $s(n)$ du circuit C_n . Une famille semi-uniforme de circuits quantiques $\{C_n\}$ est une famille uniforme définie sur une base finie $\mathcal{G}$ de portes unitaires.

3.1.4 Universalité

Dans le cadre des circuits booléens (classiques), il existe des jeux de portes universels, comme la porte **OR** associé à la porte **NOT** ou comme la porte **NAND** seule, permettant de décrire n'importe quelle fonction booléenne. L'équivalent pour les circuits quantiques se décline en deux versions : l'*universalité exacte* et l'*universalité approchée*.

Définition 3.2 (Universalité exacte) *Un ensemble de portes unitaires E est universel si et seulement si pour tout n , pour toute transformation unitaire U d'ordre n , il existe un circuit C_n fondé sur E tel que l'action de C_n est U .*

Il existe plusieurs ensembles universels de portes unitaires :

- l'ensemble de toutes les portes unitaires d'ordre 1 et d'ordre 2 [KSV02] ;
- l'ensemble composé de ΛX et de toutes les portes unitaire d'ordre 1 ;
- l'ensemble composé de ΛX , H et de toutes les rotations autour de z .

Un ensemble universel est nécessairement non-dénombrable, il existe donc une notion plus faible d'ensemble *approximativement universel*.

Définition 3.3 (Universalité approchée) *Un ensemble de portes unitaires E est approximativement universel si et seulement si pour tout n , pour toute transformation unitaire U d'ordre n , et pour tout $\epsilon > 0$, il existe un circuit C_n fondé sur E en n tel que l'action $\tilde{U}$ de C_n vérifie*

$$\|U - \tilde{U}\| \leq \epsilon$$

où $\|U - V\| = \sup_{|\psi\rangle \in \mathcal{H}^1} \|(U - V)|\psi\rangle\|$.

Il existe plusieurs ensembles approximativement universels de portes unitaires :

- $\{H, T, \Lambda X\}$;
- $\{\Lambda^2 X, H\}$ [Shi03].

3.2 Machine de Turing quantique

La machine de Turing quantique a été introduite par Benioff [Ben80], puis étudiée par Deutsch [Deu85] puis Bernstein et Vazirani [BV93, BV97].

Définition 3.4 (MTQ) *Une machine de Turing quantique M est définie par un triplet (K, Σ, δ) où Σ est un alphabet fini avec un symbole neutre $\#$, K est un ensemble fini d'états dont un état initial q_0 et un état final $q_f \neq q_0$ et δ la fonction de transition quantique.*

$$\delta : K \times \Sigma \rightarrow \mathcal{H}_{K \times \Sigma \times \{\leftarrow, -, \rightarrow\}}$$

Le scalaire $\langle q, \tau, d | \delta(p, \sigma) \rangle$, appelé amplitude de $|q, \tau, d\rangle$, est noté $\delta(p, \sigma, q, \tau, d)$. De plus, les mouvements $\leftarrow, -$ et $\rightarrow$ seront parfois notés respectivement $-1, 0$ et 1 .

Une *configuration* est un vecteur normé de $\mathcal{H}_{K \times \Sigma^* \times \mathbb{Z}}$. $T \in \Sigma^*$ est un état de base du ruban de M , où seul un nombre fini de symboles de T ne sont pas le symbole neutre, $x \in \mathbb{Z}$ désigne la position de la tête, et $T_x \in \Sigma$ désigne le symbole de base pointé par la tête.

L'opérateur d'évolution d'une MTQ M est un opérateur linéaire U_M sur $\mathcal{H}_{K \times \Sigma^* \times \mathbb{Z}}$ tel que

$$U_M |p, T, x\rangle = \sum_{q \in K, \tau \in \Sigma, d \in \{-1, 0, 1\}} \delta(p, T_x, q, \tau, d) |q, T_x^\tau, x + d\rangle$$

où $T_x^\tau \in \Sigma^*$ désigne T où l'élément indexé par x est remplacé par τ .

Une machine de Turing quantique est *bien formée* si et seulement si U_M est unitaire.

Une présentation plus détaillée des machines de Turing quantiques est donnée dans le chapitre 6. Y sont présentées notamment les machines de Turing quantiques multi-rubans, ainsi que le problème de l'arrêt.

Le pouvoir des machines de Turing quantiques bien fondées est équivalent à celui des circuits quantiques. En effet, toute machine de Turing quantique bien fondée peut être simulée par une famille semi-uniforme de circuits quantiques. Inversement, toute famille semi-uniforme de circuits quantiques peut être simulée par une machine de Turing quantique bien fondée. Ces propriétés sont détaillées dans le chapitre 6 avec notamment une évaluation de la complexité de ces simulations.

Les machines de Turing quantiques bien formées ont une évolution unitaire, et constituent donc un modèle de calcul quantique réversible. Une version quantique de la machine de Turing, ne se limitant pas au calcul quantique réversible, est introduite dans le chapitre 6, il s'agit des machines de Turing quantiques à contrôle classique.

3.3 Autres modèles réversibles

D'autres modèles de calcul quantique s'attachent à décrire le fragment réversible de la mécanique quantique. Ainsi van Tonder [vT04] a introduit un λ -calcul quantique. Les termes du λ -calcul usuel représentent les termes de base, un λ -terme quantique étant une superposition de λ -termes classiques. L'objectif de van

Tonder est de proposer une β -réduction unitaire en utilisant une mémoire afin de garantir la réversibilité de la β -réduction.

Enfin, un λ -calcul algébriquement linéaire a été proposé par Arrighi [AD04, AD05].

3.4 Conclusion

Les modèles de calcul quantique réversibles n'utilisent qu'un fragment des possibilités offertes par la mécanique quantique pour effectuer un calcul : les transformations unitaires.

Les modèles standards du calcul quantiques sont des modèles réversibles : circuits quantiques, machines de Turing quantiques. Pourtant ces modèles standards souffrent de plusieurs lacunes :

- manque d'expressivité en n'autorisant que des transformations unitaires ;
- impossibilité de représenter un calcul "mixte" mélangeant calculs quantique et classique ;
- expressivité limitée du modèle des circuits quantiques qui n'autorise que la manipulation des qubits ;
- contrainte de bonne formation difficile à vérifier pour les MTQ.

Des modèles alternatifs, comme le calcul par consommation d'intrication proposé par les physiciens Briegel et Raussendorf [RB00], ou encore le calcul par mesures projectives introduit par Nielsen [Nie03], marquent une rupture avec ces modèles standards du calcul quantique réversible.

Il est donc important de développer des outils spécifiques à ces nouveaux modèles, comme le *m-calcul* [DKP04a]. Mais il faut également développer des modèles permettant de décrire et comparer les différents modes de calcul quantique dans un même formalisme autour de ce qui est communément admis comme étant la spécificité du calcul quantique : des données quantiques, un contrôle classique. Ceci est l'objet des trois prochains chapitres.

Deuxième partie

Données quantiques - contrôle classique

Chapitre 4

Evolution quantique contrôlée classiquement

4.1 Motivation

Exploiter les propriétés de la mécanique quantique afin de mener un calcul, voilà l'un des objectifs de l'informatique quantique. Un système quantique évolue, d'après le troisième postulat de la mécanique quantique, selon une transformation admissible. Des modèles de calcul quantique, dit réversibles, exploitent uniquement le fragment des évolutions quantiques formé par les transformations unitaires. Nous allons voir dans les prochains chapitres comment l'ensemble des transformations admissibles peut être utilisé pour mener un calcul.

Lorsqu'un système évolue selon une transformation admissible quelconque, l'état du système quantique change, mais de surcroît, un résultat classique est produit. La possibilité d'utiliser ce résultat classique afin d'influer sur l'évolution du système est donc offerte à un utilisateur extérieur au système quantique.

Cet utilisateur extérieur peut décider d'appliquer telle ou telle transformation admissible au système quantique en fonction d'un résultat classique précédemment obtenu. Il s'agit alors d'une composition conditionnelle de transformations admissibles contrairement à la composition inconditionnelle de transformations admissibles vue dans les chapitres précédents.

L'évolution du système quantique devient une évolution quantique contrôlée classiquement.

4.2 Formalisme

4.2.1 Espace d'états

Un système quantique contrôlé classiquement est formé d'une part d'un système quantique, mais également d'un environnement classique : à chaque état e de l'environnement classique est associé un espace de Hilbert $\mathcal{H}_e$. Cela signifie que si l'environnement classique est dans un état e alors le système quantique est dans un état $|\varphi\rangle \in \mathcal{H}_e^1$. Ainsi, étant donné un ensemble E , un système contrôlé classiquement dont l'espace des états classiques est E , a pour espace d'état $\mathcal{S}_E$:

$$\mathcal{S}_E = \{(e, |\varphi\rangle) \mid e \in E \wedge |\varphi\rangle \in \mathcal{H}_e^1\}$$

Etant donnés deux systèmes quantiques contrôlés classiquement S_1 et S_2 dont les espaces d'états sont respectivement $\mathcal{S}_D$ et $\mathcal{S}_E$, l'espace d'états du système S composé de S_1 et S_2 est $\mathcal{S}_{D \times E} = \{((d, e), |\varphi\rangle) \mid (d, e) \in D \times E \wedge |\varphi\rangle \in (\mathcal{H}_d \otimes \mathcal{H}_e)^1\}$. Si le système S_1 est dans l'état $v_1 \in \mathcal{S}_D$ et S_2 dans l'état $v_2 \in \mathcal{S}_E$, alors S est dans l'état $v_1 \otimes v_2 \in \mathcal{S}_{D \times E}$, où $\otimes$ est naturellement étendu à l'environnement classique :

$$(d, |\varphi\rangle) \otimes (e, |\psi\rangle) = ((d, e), |\varphi\rangle \otimes |\psi\rangle)$$

4.2.2 Evolutions

Afin de représenter une évolution quantique contrôlée classiquement, nous introduisons un formalisme qui étend le formalisme des transformations admissibles.

Une transformation admissible est une famille de transformations linéaires représentant chacune une évolution possible du système quantique. En présence de contrôle classique, l'évolution de l'état de l'environnement classique doit également être représentée. Ainsi une transformation admissible contrôlée classiquement est une collection de couples. Chaque couple est formé :

- d'une paire (d, e) d'états, signifiant que si l'environnement classique est dans l'état d avant la transformation, alors il sera dans l'état e après la transformation ;
- et d'un opérateur linéaire représentant l'évolution du système quantique.

Ainsi, chaque couple représente une évolution globale possible du système quantique et de l'environnement classique.

Définition 4.1 (Transformation admissible contrôlée classiquement)

- $T_{D,E} = \{(t_i, M_i)_{i \in A} \mid A \text{ dénombrable} \wedge \forall i \in A, t_i = \langle t_i^{(1)}, t_i^{(2)} \rangle \in D \times E \wedge M_i \in \mathcal{L}(\mathcal{H}_{t_i^{(1)}}, \mathcal{H}_{t_i^{(2)}})\}$ est l'ensemble des familles dénombrables de couples formés d'une paire associant à un élément de D un élément de E , et d'un opérateur linéaire sur les espaces de Hilbert correspondants.

- $T_{D,E}^1$ est l'ensemble des transformations admissibles contrôlées classiquement, i.e. des familles $(t_i, M_i)_{i \in A} \in T_{D,E}$ vérifiant pour tout $d \in D$, $\sum_{i \in A | t_i^{(1)} = d} M_i^\dagger M_i = Id_{\mathcal{H}_d}$.
- $T_{D,E}^{\leq 1}$ est l'ensemble des transformations sous-admissibles contrôlées classiquement, i.e. des familles $(t_i, M_i)_{i \in A} \in T_{D,E}$ vérifiant pour tout $d \in D$, $\sum_{i \in A | t_i^{(1)} = d} M_i^\dagger M_i \leq Id_{\mathcal{H}_d}$.

4.2.3 Composition spatiale

Si un système S_1 évolue selon $F_1 \in T_{D_1, E_1}$ et un système S_2 selon $F_2 \in T_{D_2, E_2}$, alors le système composé de S_1 et S_2 évolue selon $F_1 \otimes F_2$:

Définition 4.2 Soit $\otimes : T_{D_1, E_1} \times T_{D_2, E_2} \rightarrow T_{D_1 \times D_2, E_1 \times E_2}$, tel que :

$$(f_i, M_i)_{i \in A} \otimes (g_j, N_j)_{j \in B} = \left(((f_i^{(1)}, g_j^{(1)}), (f_i^{(2)}, g_j^{(2)})), M_i \otimes N_j \right)_{(i,j) \in A \times B}$$

4.2.4 Composition temporelle

Si un système contrôlé classiquement évolue selon $F \in T_{C,D}$, puis selon $G \in T_{D,E}$, alors l'évolution totale du système est selon $G \circ F \in T_{C,E}$:

Définition 4.3 Soit $\circ : T_{D,E} \times T_{C,D} \rightarrow T_{C,E}$, tel que :

$$(f_i, M_i)_{i \in A} \circ (g_j, N_j)_{j \in B} = \left((g_j^{(1)}, f_i^{(2)}), M_i N_j \right)_{(i,j) \in A \times B | g_j^{(2)} = f_i^{(1)}}$$

Propriété 4.1 (Linéarité) Pour tout $T \in T_{D,E}$, et $T_1, T_2 \in T_{C,D}$,

$$T \circ (T_1 \cup T_2) = (T \circ T_1) \cup (T \circ T_2)$$

Remarque 4.1 Dans la propriété 4.1, $\cup$ représente l'union sur les familles et non sur les ensembles. $\cup$ peut être définie ainsi : $(t_i)_{i \in A} \cup (u_j)_{j \in B} = (v_k)_{k \in A \cup h(B)}$ avec

$$h \text{ une fonction injective telle que } h(B) \cap A = \emptyset \text{ et } v_k = \begin{cases} t_k & \text{si } k \in A \\ u_{h^{-1}(k)} & \text{si } k \in h(B) \end{cases}$$

Preuve de la propriété 4.1 : Si $T = (f_i, M_i)_{i \in A}$, $T_1 = (g'_j, N'_j)_{j \in B_1}$ et $T_2 = (g''_j, N''_j)_{j \in B_2}$, alors $T_1 \cup T_2 = (g_k, N_k)_{k \in B_1 \cup h(B_2)}$ avec h une fonction injective.

On remarque que $(A \times B_1) \cap (A \times h(B_2)) = \emptyset$ et que $(A \times B_1) \cup (A \times h(B_2)) = A \times (B_1 \cup h(B_2))$. On en déduit que $(T \circ T_1) \cup (T \circ T_2) = T \circ (T_1 \cup T_2)$. $\square$

4.3 Interprétation des transformations admissibles contrôlées classiquement

Une transformation admissible est une famille d'opérateurs linéaires. Cette famille peut être interprétée en une fonction probabiliste, via la fonction $\mathcal{X}$, ou en un super-opérateur agissant sur des matrices de densité via $\mathcal{X}^\natural$ (cf chapitre 2). Dans cette section, les fonctions d'interprétations $\mathcal{X}$ et $\mathcal{X}^\natural$ sont étendues aux transformations admissibles contrôlées classiquement.

4.3.1 Fonction probabiliste

La fonction d'interprétation $\mathcal{X}$ permettant de faire le lien entre le formalisme des transformations admissibles et celui des fonctions probabilistes peut être étendue au cas des évolutions contrôlées classiquement :

Définition 4.4

$$\begin{aligned} \mathcal{X}_c : T_{D,E} &\rightarrow (\mathcal{S}_D \rightarrow V(\mathcal{S}_E)) \\ (t_i, M_i)_{i \in A} &\mapsto \lambda(d, |\varphi\rangle) \cdot \sum_{i \in A | t_i^{(1)} = d} \langle \varphi | M_i^\dagger M_i | \varphi \rangle \cdot \eta \left(t_i^{(2)}, \frac{M_i |\varphi\rangle}{\sqrt{\langle \varphi | M_i^\dagger M_i | \varphi \rangle}} \right) \end{aligned}$$

L'interprétation des transformations admissibles contrôlées classiquement est donc la suivante : si l'environnement classique est dans l'état $d \in D$ et que le système quantique est dans l'état $|\varphi\rangle \in \mathcal{H}_d^1$, alors l'application d'une transformation admissible contrôlée classiquement $(t_i, M_i)_{i \in A} \in T_{D,E}^1$ correspond à l'application de la transformation admissible $(M_i)_{i \in A | t_i^{(1)} = d}$ sur la partie quantique du système. De plus, si le résultat classique i est observé, alors l'état de l'environnement classique devient $t_i^{(2)}$.

En d'autres termes, si l'état global du système avant la transformation est $(d, |\varphi\rangle)$, alors après la transformation, l'état global sera $(t_i^{(2)}, \frac{M_i |\varphi\rangle}{\sqrt{\langle \varphi | M_i^\dagger M_i | \varphi \rangle}})$ avec probabilité $\langle \varphi | M_i^\dagger M_i | \varphi \rangle \delta_{d, t_i^{(1)}}$ ¹.

4.3.2 Super-opérateur

Alors que toute transformation admissible peut être interprétée, via la fonction $\mathcal{X}^\natural$ en un super-opérateur, la généralisation de cette fonction d'interprétation $\mathcal{X}^\natural$ au cas des transformations admissibles contrôlées classiquement s'avère délicate. La principale limite imputable au formalisme des super-opérateurs est que l'encodage des probabilités en des matrices de densité ne peut pas s'étendre naturellement à

¹ $\delta_{x,y} = 1$ si $x = y$ et 0 sinon.

la partie classique. Une solution efficace a cependant été introduite par Selinger [Sel04b] dans le cas où l'environnement classique est fini. Nous nous inspirons de cette méthode pour étendre la fonction $\mathcal{X}^\natural$ dans le cas où l'environnement est fini.

L'espace des états est un n -uplet de matrices de densité, où n est la taille de l'environnement classique :

$$\mathcal{S}_{\{e_1, \dots, e_n\}}^\natural = D(\mathcal{H}_{e_1}) \times \dots \times D(\mathcal{H}_{e_n})$$

Les éléments de $\mathcal{S}_E^\natural$ sont représentés en utilisant une notation fonctionnelle. En effet, un élément de $\mathcal{S}_E^\natural$ peut être vu comme une fonction f associant à chaque indice $e \in E$, une matrice de densité $f(e)$ appartenant à $D(\mathcal{H}_e)$.

La fonction d'interprétation $\mathcal{X}^\natural$ associe à chaque transformation admissible contrôlée classiquement de $T_{D,E}$ un opérateur linéaire de $\mathcal{S}_D^\natural$ dans $\mathcal{S}_E^\natural$:

Définition 4.5

$$\begin{aligned} \mathcal{X}_c^\natural : T_{D,E} &\rightarrow L(\mathcal{S}_D^\natural, \mathcal{S}_E^\natural) \\ (t_i, M_i)_{i \in A} &\mapsto \lambda f. \lambda e. \sum_{i \in A | t_i^{(2)} = e} M_i f(t_i^{(1)}) M_i^\dagger \end{aligned}$$

On peut vérifier que pour toute transformation admissible F , $\mathcal{X}_c(F)$ est une fonction linéaire.

4.4 Conclusion

Nous avons vu dans ce chapitre que la représentation des évolutions quantiques contrôlées classiquement nécessite la mise en place d'un formalisme adapté, généralisant ceux vus dans les chapitres précédents.

La mise en place de ce formalisme est nécessaire à l'étude de modèles de calcul quantique non réversible dans lesquels le contrôle classique a un rôle majeur. Dans le chapitre suivant, un modèle de calcul quantique contrôlé classiquement est introduit, dans lequel toute transformation admissible contrôlée classiquement peut être représentée, dans la limite où l'environnement classique associé est de taille finie. Il s'agit du $\mathbf{q}$ -calcul.

Chapitre 5

q-calcul

5.1 Introduction

Il existe différents modèles pour le calcul quantique réversible, comme le modèle des circuits quantiques ou celui des machines de Turing quantiques. Qu'en est-il pour le calcul quantique "irréversible" ?

Il convient tout d'abord de distinguer deux caractéristiques du calcul quantique non réversible :

- L'utilisation de transformations non réversibles, comme les mesures projectives, ou plus généralement certaines transformations admissibles ;
- La présence d'un contrôle classique, permettant l'utilisation des résultats classiques pour conditionner certaines étapes du calcul.

Par exemple, le modèle des circuits quantiques généralisés [Aha99] est un modèle non réversible car dans ce modèle, les portes sont des super-opérateurs qui ne sont pas nécessairement réversibles, bien que la succession des portes reste inconditionnelle, sans contrôle classique. Un autre modèle formel de calcul quantique non réversible sans contrôle classique est celui des machines de Turing quantiques linéaires [IOV04].

Le contrôle classique est utilisé dans au moins deux modèles de calcul quantique : le calcul quantique par consommation d'intrication [RBB02] et le calcul par mesures projectives [Nie03]. Il est également présent dans le protocole de téléportation [BBC⁺93], où une opération unitaire, dite de correction, appliquée en fin de protocole est contrôlée classiquement. En effet, cette correction dépend du résultat classique des mesures effectuées plus tôt dans le protocole. Une telle dépendance nécessite un contrôle classique.

Le *m*-calcul (*measurement calculus*) [DKP04a] est un modèle formel pour le calcul par consommation d'intrication (*One-way quantum computer* [RB00]), dans lequel des mesures agissant sur un qubit sont appliquées sur un état intriqué appelé

état graphe. Les états graphes et leur utilisation seront étudiés plus loin, dans les chapitres 9, 10 et 11. Mais le m -calcul ne permet de représenter de façon formelle qu'un fragment du calcul quantique contrôlé classiquement. En effet, les seules évolutions quantiques non réversibles représentables dans le m -calcul sont les mesures sur un qubit. De plus, en terme de contrôle classique, des structures conditionnelles sont possibles dans le m -calcul, mais pas de schéma récursif (i.e. de type boucle *tant que* par exemple).

Un autre modèle de calcul irréversible est le calcul par mesures projectives introduit par Nielsen [Nie03]. Les ressources nécessaires à ce type de calcul quantique, composées de mesures multiqubits et de qubits auxiliaires, sont étudiées dans le chapitre 7. Ce calcul à base de mesures projectives ne peut pas être décrit à l'aide des outils formels développés pour le calcul quantique réversible et pour le calcul par consommation d'intrication. D'une part, les mesures sur plusieurs qubits ne peuvent pas être représentées par des circuits quantiques, ni par des termes du m -calcul, et d'autre part, le contrôle classique nécessaire, ayant un caractère récursif, ne peut pas être formalisé à l'aide des modèles formels décrits ci-dessus.

Un modèle formel plus général est donc nécessaire. Plutôt que de développer un modèle formel spécifique au calcul par mesures projectives, nous avons choisi de développer un modèle *unificateur* permettant de représenter les trois types de calcul quantique : réversible, par consommation d'intrication et par mesures projectives. Ce modèle de calcul est le q -calcul.

Dans le q -calcul, sont autorisées : les transformations unitaires, les mesures sur un qubit (ingrédients de base du m -calcul) et les mesures projectives multi-qubits, mais aussi les transformations admissibles en général. De plus, des structures de contrôle classiques peuvent être représentées dans ce modèle, incluant des structures conditionnelles, et également des *boucles* de contrôle (ou contrôle *récursif*).

La possibilité de représenter les transformations admissibles donne au q -calcul un grand pouvoir expressif, mais l'utilisation des transformations admissibles est également justifiée par l'absence de stabilité par composition des mesures projectives. En effet, la composition de deux mesures projectives n'est pas nécessairement une mesure projective. Or, un ensemble stable par composition et contenant les mesures projectives, est précisément celui des transformations admissibles.

Dans ce chapitre, nous introduisons le q -calcul. Une sémantique dénotationnelle *pure* fondée sur les domaines probabilistes [JP89] est également introduite, associant à chaque terme du q -calcul une fonction probabiliste. Cette sémantique dénotationnelle ne prend pas en compte des propriétés comme l'indistingabilité de certains états quantiques, c'est pourquoi une sémantique dénotationnelle *observable* est également introduite, s'appuyant sur les solutions développées par Kashefi [Kas03] et Selinger [Sel04b]. De plus, une relation d'abstraction exacte est établie entre ces deux sémantiques dénotationnelles. Enfin une sémantique admissible as-

socie directement à chaque terme une transformation admissible contrôlée classiquement. Les sémantiques pure et observable s'avèrent être des interprétations de la sémantique admissible, complétant la comparaison de ces trois sémantiques. Une hiérarchie de sémantiques peut alors être établie.

Le $\mathbf{q}$ -calcul est un ensemble de règles agissant sur les $\mathbf{q}$ -termes. Ces règles préservent tout ou partie de la hiérarchie sémantique. De plus, les règles du $\mathbf{q}$ -calcul forment un système de réécriture terminant et confluent.

5.2 Termes du $\mathbf{q}$ -calcul

5.2.1 Définitions

Le formalisme utilisé pour la syntaxe des termes du $\mathbf{q}$ -calcul s'inspire des algèbres de processus, ainsi la brique de base est l'action. Une action a de $\mathcal{H}$ dans $\mathcal{H}'$ est de la forme :

$$a := M \mid M, a$$

où $\mathcal{H}$ et $\mathcal{H}'$ sont des espaces de Hilbert et $M \in \mathbf{L}(\mathcal{H}, \mathcal{H}')$.

Définition 5.1 (Terme du $\mathbf{q}$ -calcul) *Un terme $\mathcal{P}$ est un quadruplet (K, I, F, R) , où K est un ensemble fini de processus, $I, F \subseteq K$ sont respectivement les états initiaux et finaux, et R est un ensemble fini de définitions de processus de la forme :*

$$\mathbf{q} = [a].\mathbf{q} \ (+ \ [a].\mathbf{q})^*$$

où chaque $\mathbf{q} \in K \setminus F$ apparaît une fois et une seule dans la partie gauche des définitions, de plus, tous les processus apparaissant dans R sont des éléments de K . Enfin, il existe un ensemble d'espaces de Hilbert $\{\mathcal{H}_{\mathbf{q}}, \mathbf{q} \in K\}$ tel que pour toute définition de processus $\mathbf{q} = \sum_i [a_i].\mathbf{q}_i$ de R , chaque a_i est une action de $\mathcal{H}_{\mathbf{q}}$ dans $\mathcal{H}_{\mathbf{q}_i}$. De plus, une condition de complétude $\sum_i a_i^\dagger = \text{Id}_{\mathcal{H}_{\mathbf{q}}}$ doit également être vérifiée, où $a^\dagger$ est défini par :

$$\begin{aligned} M^\dagger &= M^\dagger M \\ (M, a)^\dagger &= M^\dagger + a^\dagger \end{aligned}$$

Exemple 5.1 *Pour toute transformation unitaire $U \in \mathbf{L}(\mathcal{K}, \mathcal{L})$, soit $\mathcal{P}_U = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{K}$, $\mathcal{H}_{\mathbf{f}} = \mathcal{L}$ et R :*

$$\mathbf{i} = [U].\mathbf{f}$$

La condition de complétude est vérifiée, $\mathcal{P}_U$ est donc un terme du $\mathbf{q}$ -calcul. La sémantique de ce $\mathbf{q}$ -terme est donnée dans l'exemple 5.3.

Exemple 5.2 Soit $\mathcal{P} = (\{i, q, f\}, \{i\}, \{f\}, R)$, avec $\mathcal{H}_i = \mathcal{H}_q = \mathcal{H}_f = \mathcal{H}_{\{0,1\}}$ et R :

$$\begin{aligned} i &= [|0\rangle\langle 0|].f + [|1\rangle\langle 1|].q \\ q &= [|+\rangle\langle +|, |-\rangle\langle -|].i \end{aligned}$$

Les conditions de complétude sont vérifiées : $|0\rangle\langle 0| + |1\rangle\langle 1| = Id_{\mathcal{H}_{\{0,1\}}}$ et $|+\rangle\langle +| + |-\rangle\langle -| = Id_{\mathcal{H}_{\{0,1\}}}$ ¹.

La sémantique de ce q -terme est donnée dans l'exemple 5.4.

5.2.2 Représentation graphique

Un q -terme peut être représenté graphiquement par un automate d'états fini. En effet à tout q -terme, on peut associer un automate d'états fini dont les états sont les processus du q -terme, les états initiaux (finaux) sont les processus initiaux (finaux), les transitions de q vers p sont étiquetées par les actions a_i telles que $q = \dots + [a_i].p + \dots$ est une définition apparaissant dans le q -terme. La représentation graphique du q -terme de l'exemple 5.2 est donné en figure 5.1.

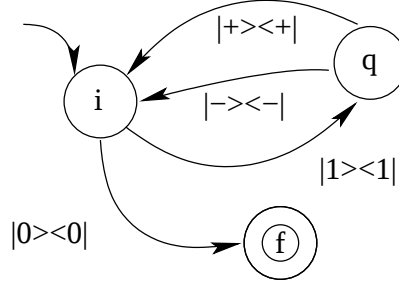


FIG. 5.1 – Représentation graphique du q -terme de l'exemple 5.2.

Cette représentation à base d'automate peut se substituer aux règles de définitions des processus, en revanche elle ne permet pas de décrire quels sont les espaces de Hilbert associés à chaque processus.

5.3 Sémantiques dénotationnelles

Une sémantique dénotationnelle permet d'associer à chaque terme un objet mathématique représentant l'action de ce terme. L'action d'un q -terme est une évolution quantique contrôlée classiquement. Cependant, plusieurs formalismes existent pour représenter de telles évolutions (cf chapitre 4) :

¹ $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ et $|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$

- les fonctions probabilistes : l'état d'un système quantique contrôlé classiquement peut être caractérisé par une distribution de probabilité sur des couples formés de l'état de l'environnement classique d'une part, et de l'état pur du système quantique d'autre part. Dans ce cadre, une évolution est une fonction agissant sur des distributions de probabilité.
- les super-opérateurs : l'état d'un système quantique contrôlé classiquement peut également être caractérisé par un n -uplet de matrices de densité, une évolution est alors une transformation linéaire sur ce domaine.
- Enfin les transformations admissibles contrôlées classiquement sont aussi une bonne caractérisation d'une évolution quantique contrôlée classiquement.

Plutôt que de choisir l'un de ces domaines sémantiques, trois sémantiques différentes sont introduites pour le $\mathbf{q}$ -calcul. Ces sémantiques sont ensuite comparées. Assez peu de sémantiques dénotationnelles ont été proposées pour le calcul quantique, à noter cependant le travail d'Abramsky [Abr04a] et ceux de Selinger [Sel04b] et Kashefi [Kas03].

Etant donné que des définitions récursives de processus sont autorisées, les domaines sémantiques utilisés devront nécessairement être complets (c.f. par exemple [AJ94]) :

Définition 5.2 *Etant donné un ordre partiel $(D, \sqsubseteq)$, un ensemble dirigé X de D est un sous-ensemble non vide de D tel que $\forall x, y \in X, \exists z \in X, x, y \sqsubseteq z$.*

Définition 5.3 *Un ordre partiel complet (CPO) est un ordre partiel avec un plus petit élément (noté $\perp$), et dont tout sous-ensemble dirigé X a une plus petite borne supérieure (noté $\sqcup X$).*

Définition 5.4 *Si D, E sont deux CPO, une fonction $f : D \rightarrow E$ est continue si et seulement si :*

- f est monotone : $x \sqsubseteq y \implies f(x) \sqsubseteq f(y)$,
- pour tout ensemble dirigé $X \subseteq D$, $f(\sqcup X) = \sqcup f(X)$.

Théorème 5.1 (Point fixe) *Soient D un CPO et $f : D \rightarrow D$ une fonction continue. Alors f a un plus petit point fixe, i.e. l'ensemble des $d \in D$ tels que $f(d) = d$ est non vide et admet un minimum.*

5.3.1 Sémantique pure

La sémantique pure du $\mathbf{q}$ -calcul associe à chaque terme une fonction agissant sur des distributions de probabilités sur des états purs. La définition de cette sémantique dénotationnelle des termes du $\mathbf{q}$ -calcul suit l'approche traditionnelle pour ce type de sémantique. Puisque l'évolution d'un système quantique est en

général probabiliste, il est naturel d'introduire une sémantique dénotationnelle fondée sur les *domaines probabilistes des évaluations* [JP89].

$V(X)$ est l'ensemble des évaluations discrètes ν sur l'ensemble X . De plus, $V^{\leq 1}(X)$ est l'ensemble des évaluations vérifiant $\nu(X) \leq 1$, de tels ν sont aussi appelés distributions de probabilité (cf chapitre 2).

Propriété 5.1 [JP89] $(V^{\leq 1}(X), \sqsubseteq)$ est un CPO ayant comme plus petit élément la fonction nulle.

La sémantique dénotationnelle pure $\llbracket \cdot \rrbracket$ du $\mathbf{q}$ -calcul peut maintenant être définie, associant à chaque $\mathbf{q}$ -terme $\mathcal{P}$ une fonction $\llbracket \mathcal{P} \rrbracket : \mathcal{S}_I \rightarrow V^{\leq 1}(\mathcal{S}_F)$, avec $\mathcal{S}_E = \{(e, |\varphi\rangle) \mid e \in E \wedge |\varphi\rangle \in \mathcal{H}_e^1\}$ (cf chapitre 4).

Définition 5.5 (Sémantique dénotationnelle pure) Etant donné un $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$:

- Pour toute action a de $\mathcal{H}$ dans $\mathcal{H}'$, $\llbracket a \rrbracket : \mathcal{H} \rightarrow V^{\leq 1}(\mathcal{H}')$ est :

$$\begin{aligned} \llbracket M \rrbracket &= \lambda |\varphi\rangle . \langle \varphi | M^\dagger M | \varphi \rangle \eta_{\frac{M|\varphi\rangle}{\sqrt{\langle \varphi | M^\dagger M | \varphi \rangle}}} \\ \llbracket M, a \rrbracket &= \lambda |\varphi\rangle . (\llbracket M \rrbracket(|\varphi\rangle) + \llbracket a \rrbracket(|\varphi\rangle)) \end{aligned}$$

On remarque que $\llbracket M \rrbracket$ est définie même pour $|\varphi\rangle$ tel que $\langle \varphi | M^\dagger M | \varphi \rangle = 0$ par $\llbracket M \rrbracket(|\varphi\rangle) = 0$.

- $\forall \mathbf{q} \in F, \llbracket \mathbf{q} \rrbracket : \mathcal{H}_{\mathbf{q}}^1 \rightarrow V^{\leq 1}(\mathcal{S}_F)$ est :

$$\llbracket \mathbf{q} \rrbracket = \lambda |\varphi\rangle . \eta_{(\mathbf{q}, |\varphi\rangle)}$$

On remarque que pour tout $\mathbf{q} \in F$, $\llbracket \mathbf{q} \rrbracket$ est une fonction continue.

- $\forall \mathbf{q} \in K \setminus F$, soit $\mathcal{E}_{\mathbf{q}} = [\mathcal{H}_{\mathbf{q}}^1 \rightarrow V^{\leq 1}(\mathcal{S}_F)]$ l'ensemble des fonctions continues de $\mathcal{H}_{\mathbf{q}}^1$ dans $V^{\leq 1}(\mathcal{S}_F)$. Soit $\mathcal{E}$ le produit cartésien de tous les $\mathcal{E}_{\mathbf{q}}$ pour $\mathbf{q} \in K \setminus F$. Les éléments de $\mathcal{E}$ sont des $|K \setminus F|$ -uplets $\langle g_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F}$ de fonctions continues telles que $g_{\mathbf{q}} \in \mathcal{E}_{\mathbf{q}}$. Pour tout $\mathbf{q} \in K \setminus F$, si $\mathbf{q} = \sum_i [a_i].\mathbf{q}_i$ est une définition de R , alors $\chi_{\mathbf{q}} : \mathcal{E} \rightarrow \mathcal{E}_{\mathbf{q}}$ est définie par :

$$\chi_{\mathbf{q}} = \lambda \langle g_{\mathbf{p}} \rangle_{\mathbf{p} \in K \setminus F} . \left(\sum_{i | \mathbf{q}_i \in K \setminus F} g_{\mathbf{q}_i}^\diamond \circ \llbracket a_i \rrbracket + \sum_{i | \mathbf{q}_i \in F} \llbracket \mathbf{q}_i \rrbracket^\diamond \circ \llbracket a_i \rrbracket \right)$$

où $f^\diamond$ est l'extension de Kleisli de f (voir définition 1.5 du chapitre 1).

Soit $\Psi : \mathcal{E} \rightarrow \mathcal{E}$ la fonction :

$$\Psi = \lambda X . \langle \chi_{\mathbf{q}}(X) \rangle_{\mathbf{q} \in K \setminus F}$$

Puisque la structure de CPO est transmise pour tout $\mathbf{q} \in K \setminus F$ à l'ensemble des fonctions continues $\mathcal{E}_{\mathbf{q}}$, et qu'elle est également conservée par le produit cartésien, $(\mathcal{E}, \sqsubseteq)$ est un CPO où $\langle f_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F} \sqsubseteq \langle g_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F}$ si pour tout $\mathbf{q} \in K \setminus F$, et pour tout $|\varphi\rangle \in \mathcal{H}_{\mathbf{q}}^1$, $f_{\mathbf{q}}(|\varphi\rangle) \sqsubseteq g_{\mathbf{q}}(|\varphi\rangle)$.

De plus, Ψ est continue, donc selon le théorème du point fixe, pour tout $\mathbf{q} \in K \setminus F$, soit $\llbracket \mathbf{q} \rrbracket : \mathcal{E}_{\mathbf{q}}$ telle que :

$$\langle \llbracket \mathbf{q} \rrbracket \rangle_{\mathbf{q} \in K \setminus F} = \mathbf{Fix}(\Psi)$$

Soit $(X_n)_{n \in \mathbb{N}}$ une suite croissante telle que $X_0 = \perp$ et $X_{n+1} = \Psi(X_n)$, alors

$$\langle \llbracket \mathbf{q} \rrbracket \rangle_{\mathbf{q} \in K \setminus F} = \lim_{n \rightarrow \infty} X_n$$

– $\llbracket \mathcal{P} \rrbracket : \mathcal{S}_I \rightarrow V^{\leq 1}(\mathcal{S}_F)$ est

$$\llbracket \mathcal{P} \rrbracket = \lambda(\mathbf{q}, |\varphi\rangle) . \llbracket \mathbf{q} \rrbracket(|\varphi\rangle)$$

5.3.2 Sémantique observable

La sémantique pure introduite dans la section précédente ne prend pas en compte des propriétés quantiques comme l'indistingabilité de certaines distributions de probabilité. Par exemple, la distribution 1/2 sur l'état $|0\rangle$ et 1/2 sur l'état $|1\rangle$ est indistingable de la distribution 1/2 sur l'état $|+\rangle = (|0\rangle + |1\rangle)/\sqrt{2}$ et 1/2 sur l'état $|-\rangle = (|0\rangle - |1\rangle)/\sqrt{2}$. On remarque que les sémantiques proposées par Kashefi [Kas03] et Selinger [Sel04b] tiennent compte de ce phénomène. A la lumière du travail de Selinger, nous introduisons une sémantique dénotationnelle *observable* fondée sur le domaine des matrices de densité.

Nous rappelons que l'ensemble des matrices de densité sur $\mathcal{H}$ est noté $D(\mathcal{H})$, et que $\mathcal{S}_{e_1, \dots, e_n} = D(\mathcal{H}_{e_1}) \times \dots \times D(\mathcal{H}_{e_n})$ est un ensemble de n -uplets de matrices de densité. Pour tout $\rho_1, \rho_2 \in D(\mathcal{H})$, $\rho_1 \sqsubseteq \rho_2$ si et seulement si $\rho_2 - \rho_1$ est une matrice positive.

Théorème 5.2 [Sel04b] $(D(\mathcal{H}), \sqsubseteq)$ est un CPO.

On en déduit que, tant que E est fini, $(\mathcal{S}_E^1, \sqsubseteq)$ est un CPO, où $\sqsubseteq$ est défini point à point.

Dans la définition 5.6, une sémantique observable $\llbracket \cdot \rrbracket^{\sharp}$ des $\mathbf{q}$ -termes est donnée :

Définition 5.6 (Sémantique observable) Pour tout $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$:
– Pour toute action a de $\mathcal{H}$ dans $\mathcal{H}'$, $\llbracket a \rrbracket^{\sharp} : D(\mathcal{H}) \rightarrow D(\mathcal{H}')$ est :

$$\llbracket M \rrbracket^{\sharp} = \lambda \rho . M \rho M^{\dagger}$$

$$\llbracket M, a \rrbracket^{\sharp} = \lambda \rho . (\llbracket M \rrbracket^{\sharp}(\rho) + \llbracket a \rrbracket^{\sharp}(\rho))$$

$$- \forall \mathbf{q} \in F, \llbracket \mathbf{q} \rrbracket^{\natural} : D(\mathcal{H}_{\mathbf{q}}) \rightarrow \mathcal{S}_F^{\natural},$$

$$\llbracket \mathbf{q} \rrbracket^{\natural} = \lambda \rho. \lambda \mathbf{p}. \delta_{\mathbf{q}, \mathbf{p}} \rho$$

où δ est l'opérateur de Kronecker.

- $\forall \mathbf{q} \in K \setminus F$, soit $\mathcal{E}_{\mathbf{q}}^{\natural} = [D(\mathcal{H}_{\mathbf{q}}) \rightarrow \mathcal{S}_F^{\natural}]$ l'ensemble des fonctions continues de $D(\mathcal{H}_{\mathbf{q}})$ dans $\mathcal{S}_F^{\natural}$. Soit $\mathcal{E}^{\natural}$ le produit cartésien de tous les $\mathcal{E}_{\mathbf{q}}^{\natural}$ pour $\mathbf{q} \in K \setminus F$. Les éléments de $\mathcal{E}^{\natural}$ sont des $|K \setminus F|$ -uplets $\langle g_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F}$ de fonctions continues telles que $g_{\mathbf{q}} \in \mathcal{E}_{\mathbf{q}}^{\natural}$. Pour tout $\mathbf{q} \in K \setminus F$, si $\mathbf{q} = \sum_i [a_i].\mathbf{q}_i$ est dans R , alors $\chi_{\mathbf{q}}^{\natural} : \mathcal{E}^{\natural} \rightarrow \mathcal{E}_{\mathbf{q}}^{\natural}$ est définie par :

$$\chi_{\mathbf{q}}^{\natural} = \lambda \langle g_{\mathbf{p}} \rangle_{\mathbf{p} \in K \setminus F}. \left(\sum_{i | \mathbf{q}_i \in K \setminus F} g_{\mathbf{q}_i} \circ \llbracket a_i \rrbracket^{\natural} + \sum_{i | \mathbf{q}_i \in F} \llbracket \mathbf{q}_i \rrbracket^{\natural} \circ \llbracket a_i \rrbracket^{\natural} \right)$$

Soit $\Psi^{\natural} : \mathcal{E}^{\natural} \rightarrow \mathcal{E}^{\natural}$ la fonction :

$$\Psi^{\natural} = \lambda X. \langle \chi_{\mathbf{q}}^{\natural}(X) \rangle_{\mathbf{q} \in K \setminus F}$$

Puisque la structure de CPO est transmise pour tout $\mathbf{q} \in K \setminus F$ à l'ensemble des fonctions continues $\mathcal{E}_{\mathbf{q}}^{\natural}$, et qu'elle est également conservée par le produit cartésien, $(\mathcal{E}^{\natural}, \sqsubseteq)$ est un CPO où $\langle f_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F} \sqsubseteq \langle g_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F}$ si pour tout $\mathbf{q} \in K \setminus F$, et pour tout $\rho \in D(\mathcal{H}_{\mathbf{q}})$, $f_{\mathbf{q}}(\rho) \sqsubseteq g_{\mathbf{q}}(\rho)$.

De plus, $\Psi^{\natural}$ est continue, donc selon le théorème du point fixe, pour tout $\mathbf{q} \in K \setminus F$, soit $\llbracket \mathbf{q} \rrbracket^{\natural} : \mathcal{E}_{\mathbf{q}}^{\natural}$ telle que :

$$\langle \llbracket \mathbf{q} \rrbracket^{\natural} \rangle_{\mathbf{q} \in K \setminus F} = \mathbf{Fix}(\Psi^{\natural})$$

Soit $(X_n)_{n \in \mathbb{N}}$ une suite croissante telle que $X_0 = \perp$ et $X_{n+1} = \Psi^{\natural}(X_n)$, alors

$$\langle \llbracket \mathbf{q} \rrbracket^{\natural} \rangle_{\mathbf{q} \in K \setminus F} = \lim_{n \rightarrow \infty} X_n$$

$$- \llbracket \mathcal{P} \rrbracket^{\natural} : \mathcal{S}_I^{\natural} \rightarrow \mathcal{S}_F^{\natural} \text{ est}$$

$$\llbracket \mathcal{P} \rrbracket^{\natural} = \lambda s. \sum_{\mathbf{q} \in I} \llbracket \mathbf{q} \rrbracket^{\natural}(s(\mathbf{q}))$$

Tout comme dans la définition 5.5, on peut vérifier que les combinateurs utilisés sont continus, et que, par conséquent, le plus petit point fixe utilisé pour définir $\llbracket . \rrbracket^{\natural}$ existe.

La relation entre les sémantiques pure $\llbracket . \rrbracket$ et observable $\llbracket . \rrbracket^{\natural}$ est établie au moyen d'une fonction d'abstraction :

Définition 5.7 *Etant donné un $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$, et pour tout $E \subseteq K$, soit $\alpha_E : V^{\leq 1}(\mathcal{S}_E) \rightarrow \mathcal{S}_E^\natural$ une fonction d'abstraction telle que :*

$$\alpha_E = \lambda \nu. \lambda \mathbf{q}. \sum_{|\varphi\rangle \in \mathcal{H}_{\mathbf{q}}} \nu((\mathbf{q}, |\varphi\rangle)) |\varphi\rangle \langle \varphi|$$

Lemme 5.1 $\llbracket \cdot \rrbracket^\natural$ est une α -abstraction exacte de $\llbracket \cdot \rrbracket^\diamond$, i.e. pour tout $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$,

$$\llbracket \mathcal{P} \rrbracket^\natural \circ \alpha_I = \alpha_F \circ \llbracket \mathcal{P} \rrbracket^\diamond$$

$$\begin{array}{ccc} \mathcal{S}_I^\natural & \xrightarrow{\llbracket \cdot \rrbracket^\natural} & \mathcal{S}_F^\natural \\ \alpha_I \uparrow & & \uparrow \alpha_F \\ V^{\leq 1}(\mathcal{S}_I) & \xrightarrow{\llbracket \cdot \rrbracket^\diamond} & V^{\leq 1}(\mathcal{S}_F) \end{array}$$

La preuve est fondée sur la continuité de α_E pour tout $E \subseteq K$.

5.3.3 Sémantique admissible

Dans le chapitre 4, un nouveau formalisme a été introduit, il s'agit des transformations admissibles contrôlées classiquement. Elles permettent de représenter les évolutions quantiques contrôlées classiquement. Une sémantique admissible $\llbracket \cdot \rrbracket^\flat$, associant directement à chaque $\mathbf{q}$ -terme une transformation admissible contrôlée classiquement est introduite.

Le domaine sémantique considéré ici est donc celui des transformations admissibles contrôlées classiquement. Nous démontrons tout d'abord que les transformations admissibles forment un CPO :

Théorème 5.3 $(T^{\leq 1}(\mathcal{H}, \mathcal{K}), \sqsubseteq)$ est un CPO, où $(M_i)_{i \in A} \sqsubseteq (N_j)_{j \in B}$ si et seulement s'il existe une fonction injective $h : A \rightarrow B$ telle que $h(A) \subseteq B$ et $\forall i \in A, M_i = N_{h(i)}$.

Preuve : Soit $F_0 \sqsubseteq F_1 \dots$ une suite croissante telle que chaque $F_i \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$. On suppose que chaque $F_i = (M_j)_{j \in I_i}$ avec I_i un intervalle de $\mathbb{N}$ et que $\forall i, I_i \subseteq I_{i+1} \wedge 0 \in I_i$. La suite croissante $I_0 \subseteq I_1 \dots$ admet un plus petit majorant $J \subseteq \mathbb{N}$, ainsi le plus petit majorant des F_i est $F = (M_j)_{j \in J}$. De plus $\sum_{j \in J} M_i^\dagger M_j = \lim_{n \rightarrow \infty} \sum_{j \in I_n} M_j^\dagger M_j \leq Id$, donc $F \in T^{\leq 1}(\mathcal{H}, \mathcal{K})$. $\square$

Théorème 5.4 $(T_{D,E}^{\leq 1}, \sqsubseteq)$ est un CPO, où $(f_i, M_i)_{i \in A} \sqsubseteq (g_j, N_j)_{j \in B}$ si et seulement s'il existe une fonction injective $h : A \rightarrow B$ telle que $h(A) \subseteq B$ et $\forall i \in A, M_i = N_{h(i)} \wedge f_i = g_{h(i)}$.

Preuve : La preuve est similaire à celle du théorème 5.3. Soit $F_0 \sqsubseteq F_1 \dots$ une suite croissante telle que chaque $F_i \in T_{D,E}^{\leq 1}$. On suppose que chaque $F_i = (f_j, M_j)_{j \in I_i}$ avec I_i un intervalle de $\mathbb{N}$ et que $\forall i, I_i \subseteq I_{i+1} \wedge 0 \in I_i$. La suite croissante $I_0 \subseteq I_1 \dots$ admet un plus petit majorant $J \subseteq \mathbb{N}$, ainsi le plus petit majorant des F_i est $F = (f_j, M_j)_{j \in J}$. De plus $\sum_{j \in J} M_i^\dagger M_j = \lim_{n \rightarrow \infty} \sum_{j \in I_n} M_j^\dagger M_j \leq Id$, donc $F \in T_{D,E}^{\leq 1}$. $\square$

Définition 5.8 (Sémantique admissible) Pour tout $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R) :$
 – Pour toute action a , $\llbracket a \rrbracket^b \in (K \setminus F) \times K \rightarrow T_{K \setminus F, K}^{\leq 1}$ est² :

$$\llbracket M \rrbracket^b = \lambda (\mathbf{p}, \mathbf{q}) . ((\mathbf{p}, \mathbf{q}), M)$$

$$\llbracket M, a \rrbracket^b = \lambda (\mathbf{p}, \mathbf{q}) . (\llbracket M \rrbracket_{\mathbf{p}, \mathbf{q}}^b \cup \llbracket a \rrbracket_{\mathbf{p}, \mathbf{q}}^b)$$

$$- \forall \mathbf{q} \in F, \llbracket \mathbf{q} \rrbracket^b \in \mathcal{E}_{\mathbf{q}}^b, \text{ où } \mathcal{E}_{\mathbf{q}}^b = T_{\{\mathbf{q}\}, F}^{\leq 1}$$

$$\llbracket \mathbf{q} \rrbracket^b = ((\mathbf{q}, \mathbf{q}), Id_{\mathcal{H}_{\mathbf{q}}})$$

– Soit $\mathcal{E}^b$ le produit cartésien de tous les $\mathcal{E}_{\mathbf{q}}^b$ pour $\mathbf{q} \in K \setminus F$.

Pour tout $\mathbf{q} \in K \setminus F$, si $\mathbf{q} = \sum_i [a_i] . \mathbf{q}_i$ est une définition de R , soit $\chi_{\mathbf{q}}^b : \mathcal{E}^b \rightarrow \mathcal{E}_{\mathbf{q}}^b :$

$$\chi_{\mathbf{q}}^b = \lambda \langle g_{\mathbf{p}} \rangle_{\mathbf{p} \in K \setminus F} . \left(\bigcup_{i | \mathbf{q}_i \in K \setminus F} g_{\mathbf{q}_i} \circ \llbracket a_i \rrbracket_{\mathbf{q}, \mathbf{q}_i}^b \cup \bigcup_{i | \mathbf{q}_i \in F} \llbracket a_i \rrbracket_{\mathbf{q}, \mathbf{q}_i}^b \right)$$

Soit $\Psi^b : \mathcal{E}^b \rightarrow \mathcal{E}^b$ la fonction :

$$\Psi^b = \lambda X . \langle \chi_{\mathbf{q}}^b(X) \rangle_{\mathbf{q} \in K \setminus F}$$

puisque la structure de CPO est conservée par le produit cartésien, $(\mathcal{E}^b, \sqsubseteq)$ est un CPO où $\langle T_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F} \sqsubseteq \langle T'_{\mathbf{q}} \rangle_{\mathbf{q} \in K \setminus F}$ si et seulement si pour tout $\mathbf{q} \in K \setminus F$, $T_{\mathbf{q}} \sqsubseteq T'_{\mathbf{q}}$.

De plus, Ψ^b est continue, donc, d'après le théorème du point fixe, pour tout $\mathbf{q} \in K \setminus F$, soit $\llbracket \mathbf{q} \rrbracket^b : \mathcal{E}_{\mathbf{q}}^b$ telle que :

$$\langle \llbracket \mathbf{q} \rrbracket^b \rangle_{\mathbf{q} \in K \setminus F} = \mathbf{Fix}(\Psi^b)$$

Soit $(X_n)_{n \in \mathbb{N}}$ une suite croissante telle que $X_0 = \perp$ et $X_{n+1} = \Psi^b(X_n)$, alors

$$\langle \llbracket \mathbf{q} \rrbracket^b \rangle_{\mathbf{q} \in K \setminus F} = \lim_{n \rightarrow \infty} X_n$$

² $\llbracket a \rrbracket^b(\mathbf{q}, \mathbf{p})$ est noté $\llbracket a \rrbracket_{\mathbf{q}, \mathbf{p}}^b$.

– $\llbracket \mathcal{P} \rrbracket^b : T_{I,F}^{\leq 1} :$

$$\llbracket \mathcal{P} \rrbracket^b = \bigcup_{q \in I} \llbracket q \rrbracket^b$$

Exemple 5.3 Le terme $\mathcal{P}_U$ est défini dans l'exemple 5.1 par $\mathcal{P}_U = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{K}$, $\mathcal{H}_{\mathbf{f}} = \mathcal{L}$ et $R :$

$$\mathbf{i} = U.\mathbf{f}$$

où U une transformation unitaire de $\mathcal{K}$ dans $\mathcal{L}$.

On a $\chi_{\mathbf{i}}^b = \lambda g_{\mathbf{i}}. \llbracket U \rrbracket_{\mathbf{i}, \mathbf{f}}^b$, donc $\llbracket \mathbf{i} \rrbracket^b = \llbracket U \rrbracket_{\mathbf{i}, \mathbf{f}}^b = ((\mathbf{i}, \mathbf{f}), U)$. Ainsi,

$$\llbracket P_U \rrbracket^b = ((\mathbf{i}, \mathbf{f}), U)$$

Exemple 5.4 Le terme $\mathcal{P} = (\{\mathbf{i}, \mathbf{q}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{q}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et $R :$

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle 0|].\mathbf{f} + [|1\rangle\langle 1|].\mathbf{q} \\ \mathbf{q} &= [|+\rangle\langle +|, |-\rangle\langle -|].\mathbf{i} \end{aligned}$$

On a

$$\begin{aligned} \llbracket \mathcal{P} \rrbracket^b &= \llbracket |0\rangle\langle 0| \rrbracket_{\mathbf{i}, \mathbf{f}}^b \circ \bigcup_{n \in \mathbb{N}} ([|+\rangle\langle +|, |-\rangle\langle -| \rrbracket_{\mathbf{q}, \mathbf{i}}^b \circ \llbracket |1\rangle\langle 1| \rrbracket_{\mathbf{i}, \mathbf{q}}^b)^n \\ &= ((\mathbf{i}, \mathbf{f}), |0\rangle\langle 0|) \cup \bigcup_{n \in \mathbb{N}^*} \bigcup_{k=1}^{2^{n-1}} ((\mathbf{i}, \mathbf{f}), 2^{-n} \cdot |0\rangle\langle 1|), ((\mathbf{i}, \mathbf{f}), 2^{-n} \cdot |0\rangle\langle 1|) \end{aligned}$$

La relation entre la sémantique pure et la sémantique admissible est une relation d'interprétation, la fonction d'interprétation étant $\mathcal{X}_c :$

Propriété 5.2 (Interprétation exacte) Etant donné un $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$,

$$\mathcal{X}_c(\llbracket \mathcal{P} \rrbracket^b) = \llbracket \mathcal{P} \rrbracket$$

où $\mathcal{X}_c$ est la fonction d'interprétation définie dans le chapitre 4 :

$$\begin{aligned} \mathcal{X}_c : T_{I,F} &\rightarrow (\mathcal{S}_I \rightarrow V(\mathcal{S}_F)) \\ (t_i, M_i)_{i \in A} &\mapsto \lambda (d, |\varphi\rangle) \cdot \sum_{i \in A | t_i^{(1)} = d} \langle \varphi | M_i^\dagger M_i | \varphi \rangle \cdot \eta \left(t_i^{(2)}, \frac{M_i |\varphi\rangle}{\sqrt{\langle \varphi | M_i^\dagger M_i | \varphi \rangle}} \right) \end{aligned}$$

La preuve est fondée sur la continuité de $\mathcal{X}_c$.

Exemple 5.5 Soit $\mathcal{P}_0 = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et $R :$

$$\mathbf{i} = [|0\rangle\langle 0|].\mathbf{f} + [|0\rangle\langle 1|].\mathbf{f}$$

On peut vérifier que $\llbracket \mathcal{P}_0 \rrbracket^b = (((i, f), |0\rangle\langle 0|), ((i, f), |0\rangle\langle 1|))$. De plus,

$$\begin{aligned} \llbracket \mathcal{P}_0 \rrbracket &= \mathcal{X}_c(\llbracket \mathcal{P}' \rrbracket^b) \\ &= \lambda(d, |\varphi\rangle) \cdot (\langle \varphi | |0\rangle\langle 0| |0\rangle\langle 0| |\varphi\rangle \cdot \eta_{(f, |0\rangle)} + \langle \varphi | |1\rangle\langle 0| |0\rangle\langle 1| |\varphi\rangle \cdot \eta_{(f, |0\rangle)}) \\ &= \lambda(d, |\varphi\rangle) \cdot (\langle \varphi | (|0\rangle\langle 0| + |1\rangle\langle 1|) |\varphi\rangle \cdot \eta_{(f, |0\rangle)}) \\ &= \lambda(d, |\varphi\rangle) \cdot \eta_{(f, |0\rangle)} \end{aligned}$$

Ainsi $\mathcal{P}_0$ est un terme permettant d'initialiser un qubit dans l'état $|0\rangle$.

Il existe également une relation d'interprétation entre la sémantique admissible et la sémantique observable :

Propriété 5.3 (Interprétation exacte) *Etant donné un $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$,*

$$\mathcal{X}_c^{\natural}(\llbracket \mathcal{P} \rrbracket^b) = \llbracket \mathcal{P} \rrbracket^{\natural}$$

où $\mathcal{X}_c^{\natural}$ est la fonction d'interprétation définie dans le chapitre 4 :

$$\begin{aligned} \mathcal{X}_c^{\natural} : T_{I, K} &\rightarrow L(\mathcal{S}_I^{\natural}, \mathcal{S}_F^{\natural}) \\ (t_i, M_i)_{i \in A} &\mapsto \lambda f. \lambda e. \sum_{i \in A | t_i^{(2)} = e} M_i f(t_i^{(1)}) M_i^{\dagger} \end{aligned}$$

Nous avons mis en évidence dans cette section qu'il n'existe non pas un, mais au moins trois domaines quantiques permettant de décrire l'évolution d'un programme quantique. Ces domaines quantiques ne sont pas équivalents, il est donc intéressant de définir trois sémantiques complémentaires pour les $\mathbf{q}$ -termes. La complémentarité de ces domaines est illustrée par les relations d'abstraction et d'interprétation qui existent entre les différentes sémantiques introduites. Le domaine le plus naturel pour définir une sémantique est sans doute celui des distributions de probabilité. L'indistingabilité de certaines distributions de probabilité justifie l'utilisation d'un domaine plus *abstrait* : le domaine des matrices de densité. L'abstraction qui existe entre les sémantiques pure et observable est une abstraction exacte, ce qui signifie que bien que plus abstraite, la sémantique observable n'est pas une approximation de la sémantique pure.

A l'opposé de la sémantique observable, une sémantique moins abstraite que la sémantique pure peut également être utilisée, il s'agit de la sémantique admissible. Elle s'inspire de la représentation originale utilisée dans le troisième postulat de la mécanique quantique à l'aide de transformations admissibles, i.e. des familles d'opérateurs linéaires. Alors que la sémantique pure est fondée sur l'aspect probabiliste des évolutions quantiques, la sémantique admissible met en évidence son aspect non déterministe. L'avantage de cette sémantique est la possibilité de représenter facilement l'évolution quantique à l'aide d'une famille d'opérateurs. En revanche, plusieurs familles différentes peuvent représenter la même évolution quantique (en effet la fonction d'interprétation $\mathcal{X}_c$ n'est pas injective).

5.3.4 Equivalences

L'existence d'une sémantique permet de définir une relation d'équivalence sur les termes du langage. Nous avons établi dans la section précédente non pas une sémantique mais trois : les sémantiques pure, observable et admissible. Ainsi trois relations d'équivalence sont introduites. Les relations d'interprétation et d'abstraction mise en évidence entre les sémantiques impliquent une hiérarchie dans les relations d'équivalences.

Définition 5.9 *Etant donnés deux $\mathbf{q}$ -termes $\mathcal{P}$ et $\mathcal{P}'$,*

– *Equivalence admissible :*

$$\mathcal{P} \equiv_{adm} \mathcal{P}' \iff \llbracket \mathcal{P} \rrbracket^b = \llbracket \mathcal{P}' \rrbracket^b$$

– *Equivalence (pure) :*

$$\mathcal{P} \equiv \mathcal{P}' \iff \llbracket \mathcal{P} \rrbracket = \llbracket \mathcal{P}' \rrbracket$$

– *Equivalence observable :*

$$\mathcal{P} \equiv_{obs} \mathcal{P}' \iff \llbracket \mathcal{P} \rrbracket^{\natural} = \llbracket \mathcal{P}' \rrbracket^{\natural}$$

Ces trois équivalences forment une hiérarchie sémantique :

Propriété 5.4 (Hiérarchie sémantique) *Etant donnés deux $\mathbf{q}$ -termes $\mathcal{P}$ et $\mathcal{P}'$,*

$$\begin{aligned} \mathcal{P} \equiv_{adm} \mathcal{P}' &\implies \mathcal{P} \equiv \mathcal{P}' \\ \mathcal{P} \equiv \mathcal{P}' &\implies \mathcal{P} \equiv_{obs} \mathcal{P}' \end{aligned}$$

Preuve :

- Si $\mathcal{P} \equiv_{adm} \mathcal{P}'$ alors $\llbracket \mathcal{P} \rrbracket^b = \llbracket \mathcal{P}' \rrbracket^b$, donc $\chi_c(\llbracket \mathcal{P} \rrbracket^b) = \chi_c(\llbracket \mathcal{P}' \rrbracket^b)$ d'où $\llbracket \mathcal{P} \rrbracket = \llbracket \mathcal{P}' \rrbracket$.
- si $\mathcal{P} \equiv \mathcal{P}'$ alors $\llbracket \mathcal{P} \rrbracket^{\diamond} = \llbracket \mathcal{P}' \rrbracket^{\diamond}$, donc $\alpha_F \circ \llbracket \mathcal{P} \rrbracket^{\diamond} = \alpha_F \circ \llbracket \mathcal{P}' \rrbracket^{\diamond}$, d'où $\llbracket \mathcal{P} \rrbracket^{\natural} \circ \alpha_I = \llbracket \mathcal{P}' \rrbracket^{\natural} \circ \alpha_I$. Or α_I est une fonction surjective de $V^{\leq 1}(\mathcal{S}_I) \rightarrow \mathcal{S}_I^{\natural}$ donc $\llbracket \mathcal{P} \rrbracket^{\natural} = \llbracket \mathcal{P}' \rrbracket^{\natural}$.

□

Exemple 5.6 *Soit $\mathcal{P}_1 = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et R :*

$$\mathbf{i} = [|0\rangle\langle 0|, |0\rangle\langle 1|].\mathbf{f}$$

On vérifie facilement que $\llbracket \mathcal{P}_1 \rrbracket^b = \llbracket \mathcal{P}_0 \rrbracket^b$ (cf exemple 5.5), on en déduit que $\mathcal{P}_0 \equiv_{adm} \mathcal{P}_1$, donc $\mathcal{P}_0 \equiv \mathcal{P}_1$, et enfin $\mathcal{P}_0 \equiv_{obs} \mathcal{P}_1$.

Exemple 5.7 *Les termes $\mathcal{P}$ et $\mathcal{P}_0$ définis dans les exemples 5.4 et 5.5, ont des sémantiques admissibles différentes, donc $\mathcal{P} \not\equiv_{adm} \mathcal{P}_0$.*

En revanche, $\llbracket \mathcal{P} \rrbracket = \chi_c(\llbracket \mathcal{P} \rrbracket^b) = \lambda(d, |\varphi\rangle).\eta_{(\mathbf{f}, |0\rangle)}$, donc $\mathcal{P} \equiv \mathcal{P}_0$ et $\mathcal{P} \equiv_{obs} \mathcal{P}_0$. Ainsi $\mathcal{P}$ permet également d'initialiser un qubit dans l'état $|0\rangle$. On remarque que dans $\mathcal{P}$, contrairement à $\mathcal{P}_0$ ou $\mathcal{P}_1$, seules des mesures projectives sont utilisées.

Exemple 5.8 Soient $\mathcal{P}_2$ et $\mathcal{P}_3$ deux $\mathbf{q}$ -termes tels que :

- $\mathcal{P}_2 = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\mathbf{i} = [|0\rangle\langle 0|, |1\rangle\langle 1|].\mathbf{f}$$

- $\mathcal{P}_3 = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\mathbf{i} = [|+\rangle\langle +|, |-\rangle\langle -|].\mathbf{f}$$

On peut vérifier que $\mathcal{P}_2 \not\equiv \mathcal{P}_3$, donc $\mathcal{P}_2 \not\equiv_{adm} \mathcal{P}_3$. En revanche $\mathcal{P}_2 \equiv_{obs} \mathcal{P}_3$.

5.4 Vers un $\mathbf{q}$ -calcul

Dans cette section, nous introduisons des règles de transformations sur les $\mathbf{q}$ -termes. Ces règles permettent de simplifier les termes sans changer leur sémantique.

En plus de l'associativité des opérateurs “+” et “,”, trois règles naturelles de transformations des $\mathbf{q}$ -termes sont introduites. En ce qui concerne les définitions des processus, une règle de *substitution* permet de remplacer un processus par sa définition. En ce qui concerne les actions, une règle de *factorisation* et une seconde de *regroupement* sont introduites.

Les règles de substitution et de factorisation préservent les sémantiques admissible, pure et observable. En revanche, le regroupement préserve toujours la sémantique observable mais jamais la sémantique admissible. La préservation de la sémantique pure par cette règle de regroupement est caractérisée.

Ces trois règles forment une première version d'un $\mathbf{q}$ -calcul permettant de transformer un $\mathbf{q}$ -terme pour obtenir une hypothétique forme normale. Malheureusement, les trois règles forment un système de réécriture terminant mais non confluent. Différentes approches, qui pourraient permettre d'obtenir un $\mathbf{q}$ -calcul confluent, sont évoquées.

Définition 5.10 (Règles de réécriture)

- **Factorisation** : Tout d'abord au niveau de la définition d'un processus, l'expression $[a_1].\mathbf{q} + [a_2].\mathbf{q}$ peut être factorisée en $[a_1, a_2].\mathbf{q}$:

$$[a_1].\mathbf{q} + [a_2].\mathbf{q} \rightarrow_f [a_1, a_2].\mathbf{q}$$

- **Substitution** : Au niveau de l'ensemble des définitions de processus, si un processus $\mathbf{q}_0$ n'est pas un processus initial et est défini de façon non récursive, i.e. $\mathbf{q}_0 = \sum_{i>0} [a_i].\mathbf{q}_i$ avec $\forall i > 0, \mathbf{q}_0 \neq \mathbf{q}_i$, alors $\mathbf{q}_0$ peut être remplacé par sa définition dans toutes les autres définitions de processus :

$$R \cup \{\mathbf{q}_0 = \sum_{i>0} [a_i].\mathbf{q}_i\} \rightarrow_s R[\mathbf{q}_0 \leftarrow \sum_{i>0} [a_i].\mathbf{q}_i]$$

où $R[\mathbf{q}_0 \leftarrow \sum_{i>0} [a_i].\mathbf{q}_i]$ signifie que toutes les occurrences de $\mathbf{q}_0$ dans R sont remplacées par l'expression $\sum_{i>0} [a_i].\mathbf{q}_i$ avec en plus les conventions suivantes :

- l'opérateur $.$ est distributif par rapport à l'opérateur $+$, ainsi $[a].(\sum_i [a_i].\mathbf{q}_i)$ est réécrit en $\sum_i [a].[a_i].\mathbf{q}_i$;
- les expressions de la forme $[M_0, \dots, M_m].[N_0, \dots, N_n]$ sont réécrites en $[N_0M_0, \dots, N_nM_0, \dots, N_0M_m, \dots, N_nM_m]$, où les M_i et N_j sont des opérateurs linéaires.
- **Regroupement** : Au niveau des opérateurs linéaires, deux opérateurs colinéaires M et αM peuvent être remplacés par l'opérateur $\sqrt{1 + |\alpha|^2}M$:

$$M, \alpha M \xrightarrow{\alpha}_r \sqrt{1 + |\alpha|^2}M$$

Si l'application d'une règle $\rightarrow_x$ avec $x \in \{s, r, f\}$ transforme un $\mathbf{q}$ -terme $\mathcal{P}$ en $\mathcal{P}'$, on écrit alors $\mathcal{P} \rightarrow_x \mathcal{P}'$.

Exemple 5.9

- $\mathcal{P}_0 \rightarrow_f \mathcal{P}_1$, où $\mathcal{P}_0$ et $\mathcal{P}_1$ sont les $\mathbf{q}$ -termes définis dans les exemples 5.5 et 5.6 ;
- Soit $\mathcal{P}_4 = (\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\mathbf{i} = [|0\rangle\langle 0|, \frac{1}{\sqrt{2}}|0\rangle\langle 1|, \frac{1}{\sqrt{2}}|0\rangle\langle 1|].\mathbf{f}$$

La règle de regroupement $\xrightarrow{\alpha}_r$ peut être appliquée :

$$\mathcal{P}_4 \xrightarrow{1}_r \mathcal{P}_1$$

- Soit $\mathcal{P}_5 = (\{\mathbf{i}, \mathbf{q}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathcal{H}_{\{0,1\}}$, $\mathcal{H}_{\mathbf{q}} = \mathbb{C}$ et R :

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle 0|].\mathbf{f} + [|1\rangle\langle 1|].\mathbf{q} \\ \mathbf{q} &= [|0\rangle\langle 1|].\mathbf{f} \end{aligned}$$

La règle de substitution puis de factorisation peuvent être appliquées, transformant $\mathcal{P}_5$ en $\mathcal{P}_1$:

$$\mathcal{P}_5 \rightarrow_s \mathcal{P}' \rightarrow_f \mathcal{P}_1$$

Le théorème suivant montre que l'application de la règle de factorisation préserve la sémantique admissible, et par conséquent les sémantiques pure et observable également.

Théorème 5.5 Si $\mathcal{P}_1 \rightarrow_f \mathcal{P}_2$ alors $\mathcal{P}_1 \equiv_{adm} \mathcal{P}_2$.

Preuve : Soient $\mathcal{P}$ et $\mathcal{P}'$ tels que $\mathcal{P} \rightarrow_f \mathcal{P}'$, on suppose que la définition $\mathbf{q}_0 = A + [a].\mathbf{p} + B + [a'].\mathbf{p} + C$ de $\mathcal{P}$ est réécrite en $\mathbf{q}_0 = A + [a, a'].\mathbf{p} + B + C$ dans $\mathcal{P}'$.

En utilisant les notations de la définition 5.8, on remarque que $\forall \mathbf{q}, \mathbf{q}', \llbracket a, a' \rrbracket_{\mathbf{q}, \mathbf{q}'}^b = \llbracket a \rrbracket_{\mathbf{q}, \mathbf{q}'}^b \cup \llbracket a' \rrbracket_{\mathbf{q}, \mathbf{q}'}^b$. La propriété 4.1 sur la linéarité des transformations admissibles contrôlées classiquement, permet de conclure que $\chi_{\mathbf{q}_0}^{b(1)} = \chi_{\mathbf{q}_0}^{b(2)}$, donc $\llbracket \mathcal{P}_1 \rrbracket^b = \llbracket \mathcal{P}_2 \rrbracket^b$. $\square$

La sémantique admissible est également préservée par la règle de substitution :

Théorème 5.6 *Si $\mathcal{P}_1 \rightarrow_s \mathcal{P}_2$ alors $\mathcal{P}_1 \equiv_{adm} \mathcal{P}_2$.*

Preuve : Soit E un ensemble fini, si X est un $|E|$ -uplet indexé par E , alors pour tout $A \subseteq E$, $X^{[A]}$ représente le sous $|A|$ -uplet de X .

En utilisant les notations de la définition 5.8, on remarque que pour tout $X \in \mathcal{E}^{b(1)}$ et tout $\mathbf{q} \in K$,

$$\chi_{\mathbf{q}}^{b(2)}(X^{[K]}) = \chi_{\mathbf{q}}^{b(1)}(X^{[K]}, \chi_{\mathbf{q}_0}^{b(1)}(X))$$

Soient $(Y_n)_{n \in \mathbb{N}}$ et $(Z_n)_{n \in \mathbb{N}}$ deux suites croissantes telles que $Y_0 = Z_0 = \perp$, $Y_{n+1} = \Psi^{b(1)}(Y_n)$ et $Z_{n+1} = \Psi^{b(2)}(Z_n)$. Pour tout $n \in \mathbb{N}$, $Y_n^{[K]} \subseteq Z_n \subseteq Y_{2n+1}^{[K]}$. En effet, $Y_0^{[K]} \subseteq Z_0 \subseteq Y_1^{[K]}$. De plus, la continuité de $\Psi^{b(1)}$ implique que pour tout $X \in \mathcal{E}^{b(1)}$, $X^{[\mathbf{q}_0]} \subseteq \chi_{\mathbf{q}_0}^{b(1)}(X)$. Donc, pour n fixé quelconque, supposons que $Y_n^{[K]} \subseteq Z_n$, on a alors $\forall \mathbf{q} \in K$, $\chi_{\mathbf{q}}^{b(1)}(Y_n) \subseteq \chi_{\mathbf{q}}^{b(1)}(Y_n^{[K]}, \chi_{\mathbf{q}}^{b(1)}(Y_n)) = \chi_{\mathbf{q}}^{b(2)}(Y_n^{[K]}) \subseteq \chi_{\mathbf{q}}^{b(2)}(Z_n)$, donc $Y_{n+1}^{[K]} \subseteq Z_{n+1}$.

On remarque également que pour tout X , $(X^{[K]}, \chi_{\mathbf{q}_0}^{b(1)}(X)) \subseteq \Psi^{b(1)}(X)$, donc pour tout $\mathbf{q} \in K$, $\chi_{\mathbf{q}}^{b(2)}(X^{[K]}) \subseteq \chi_{\mathbf{q}}^{b(1)}(\Psi^{b(1)}(X))$, d'où $\Psi^{b(2)}(X^{[K]}) \subseteq (\Psi^{b(1)} \circ \Psi^{b(1)}(X))^{[K]}$. Ainsi, pour tout n fixé quelconque, supposons $Z_n \subseteq Y_{2n+1}^{[K]}$, on a alors $Z_{n+1} = \Psi^{b(2)}(Z_n) \subseteq \Psi^{b(2)}(Y_{2n+1}^{[K]}) \subseteq (\Psi^{b(1)} \circ \Psi^{b(1)}(Y_{2n+1}^{[K]}))^{[K]} = Y_{2(n+1)+2}^{[K]}$.

Les suites $(Y_n^{[K]})_{n \in \mathbb{N}}$ et $(Z_n)_{n \in \mathbb{N}}$ ont donc les mêmes limites. De plus $\mathbf{q}_0 \notin I$, on en conclut que $\llbracket \mathcal{P}_1 \rrbracket^b = \llbracket \mathcal{P}_2 \rrbracket^b$. $\square$

Théorème 5.7

-(a) *Si $\mathcal{P}_1 \xrightarrow{\alpha}_r \mathcal{P}_2$ et $\alpha \in \mathbb{C}$ alors $\mathcal{P}_1 \equiv_{obs} \mathcal{P}_2$.*

-(b) *Si $\mathcal{P}_1 \xrightarrow{\alpha}_r \mathcal{P}_2$ et $\alpha \in \mathbb{R}^+$ alors $\mathcal{P}_1 \equiv \mathcal{P}_2$.*

Preuve :

– (a) :

$$\begin{aligned} \llbracket M, a, \alpha M \rrbracket^{\natural} &= \lambda \rho. (\llbracket M \rrbracket^{\natural}(\rho) + \llbracket a \rrbracket^{\natural}(\rho) + \llbracket \alpha M \rrbracket^{\natural}(\rho)) \\ &= \lambda \rho. (M \rho M^{\dagger} + \llbracket a \rrbracket^{\natural}(\rho) + |\alpha|^2 M \rho M^{\dagger}) \\ &= \llbracket \sqrt{1 + |\alpha|^2} M, a \rrbracket^{\natural} \end{aligned}$$

Donc $\llbracket \mathcal{P}_1 \rrbracket^{\natural} = \llbracket \mathcal{P}_2 \rrbracket^{\natural}$.

– (b) :

$$\begin{aligned}
\llbracket M, a, \alpha M \rrbracket &= \lambda |\varphi\rangle . (\llbracket M \rrbracket(|\varphi\rangle) + \llbracket a \rrbracket(|\varphi\rangle) + \llbracket \alpha M \rrbracket(|\varphi\rangle)) \\
&= \lambda |\varphi\rangle . (\langle \varphi | M^\dagger M | \varphi \rangle \eta_{\frac{M|\varphi\rangle}{\sqrt{\langle \varphi | M^\dagger M | \varphi \rangle}}} + \llbracket a \rrbracket(|\varphi\rangle)) \\
&\quad + \alpha^2 \langle \varphi | M^\dagger M | \varphi \rangle \eta_{\frac{\alpha M|\varphi\rangle}{\alpha \sqrt{\langle \varphi | M^\dagger M | \varphi \rangle}}} \\
&= \lambda |\varphi\rangle . ((1 + \alpha^2) \langle \varphi | M^\dagger M | \varphi \rangle \eta_{\frac{M|\varphi\rangle}{\sqrt{\langle \varphi | M^\dagger M | \varphi \rangle}}} + \llbracket a \rrbracket(|\varphi\rangle)) \\
&= \llbracket \sqrt{1 + \alpha^2} M, a \rrbracket
\end{aligned}$$

Donc $\llbracket \mathcal{P}_1 \rrbracket = \llbracket \mathcal{P}_2 \rrbracket$.

□

Le théorème 5.7 montre que la règle de regroupement $\xrightarrow{\alpha}_r$ préserve la sémantique observable. De même, si α est un réel positif, la sémantique pure est préservée. En revanche, on remarque que pour tout α la sémantique admissible n'est pas préservée, de plus si α n'est pas un réel positif, alors la sémantique pure n'est pas préservée non plus.

5.4.1 Terminaison et non-confluence

Trois systèmes de réécriture sont introduits $\mathfrak{R} = \{\rightarrow_f, \rightarrow_s\}$, $\mathfrak{R}_{\mathbb{R}^+} = \{\rightarrow_f, \rightarrow_s, \xrightarrow{\alpha}_r \mid \alpha \in \mathbb{R}^+\}$ et $\mathfrak{R}_{\mathbb{C}} = \{\rightarrow_f, \rightarrow_s, \xrightarrow{\alpha}_r \mid \alpha \in \mathbb{C}\}$. Ainsi $\mathfrak{R}$ préserve la sémantique admissible, $\mathfrak{R}_{\mathbb{R}^+}$ la sémantique pure et $\mathfrak{R}_{\mathbb{C}}$ la sémantique observable. Dans cette section, la terminaison et la non confluence de ces systèmes de réécriture sont établies.

Lemme 5.2 (Terminaison) $\mathfrak{R}_{\mathbb{C}}$ est noetherien.

Preuve : Soit w une fonction de pondération associant à chaque $\mathfrak{q}$ -terme $\mathcal{P} = (K, I, F, R)$ un poids $w(\mathcal{P}) \in \mathbb{N}^3$: $w(\mathcal{P}) = (w_1, w_2, w_3)$ où w_1 est la taille de K , w_2 est le nombre de " + " apparaissant dans R et w_3 est le nombre d'opérateurs linéaires apparaissant dans R .

Pour tout $\mathfrak{q}$ -termes $\mathcal{P}, \mathcal{P}'$, si $\mathcal{P} \rightarrow_x \mathcal{P}'$ avec $x \in \{s, f, r\}$ alors $w(\mathcal{P}') \prec w(\mathcal{P})$, où $\prec$ est l'ordre lexicographique. En effet,

- si $\mathcal{P} \xrightarrow{\alpha}_r \mathcal{P}'$ alors le nombre de processus et le nombre de + de $\mathcal{P}$ et $\mathcal{P}'$ sont identiques. En revanche, le nombre d'opérateurs linéaires est inférieur dans $\mathcal{P}'$, donc $w(\mathcal{P}') \prec w(\mathcal{P})$;
- si $\mathcal{P} \rightarrow_f \mathcal{P}'$ alors le nombre de processus de $\mathcal{P}$ et $\mathcal{P}'$ sont identiques. En revanche, le nombre de " + " est inférieur dans $\mathcal{P}'$, donc $w(\mathcal{P}') \prec w(\mathcal{P})$;
- si $\mathcal{P} \rightarrow_s \mathcal{P}'$ alors le nombre de processus est inférieur dans $\mathcal{P}'$, donc $w(\mathcal{P}') \prec w(\mathcal{P})$.

On en déduit que $\mathfrak{R}_{\mathbb{C}}$ est noetherien. $\square$

Corollaire 5.1 $\mathfrak{R}_{\mathbb{R}^+}$ et $\mathfrak{R}$ sont noetheriens.

Lemme 5.3 (Non confluence) *Aucun des systèmes $\mathfrak{R}$, $\mathfrak{R}_{\mathbb{R}^+}$ et $\mathfrak{R}_{\mathbb{C}}$ n'est confluent.*

Preuve : Considérons, par exemple, le $\mathbf{q}$ -terme $\mathcal{P}_6 = (\{\mathbf{i}, \mathbf{q}, \mathbf{p}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$ avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathbb{C}$, $\mathcal{H}_{\mathbf{q}} = \mathcal{H}_{\mathbf{p}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle|].\mathbf{q} \\ \mathbf{q} &= [|0\rangle\langle 0|, |1\rangle\langle 1|].\mathbf{p} \\ \mathbf{p} &= [\frac{1}{\sqrt{2}}(|0\rangle\langle 0| + |0\rangle\langle 1|)].\mathbf{q} + [\frac{1}{\sqrt{2}}(|\rangle\langle 0| - |\rangle\langle 1|)].\mathbf{f} \end{aligned}$$

La règle de substitution peut être appliquée à la définition de $\mathbf{p}$ ou de $\mathbf{q}$, ainsi on obtient les termes $\mathcal{P}_7$ et $\mathcal{P}_8$:

– $\mathcal{P}_7 = (\{\mathbf{i}, \mathbf{p}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$ avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathbb{C}$, $\mathcal{H}_{\mathbf{p}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle|, 0].\mathbf{p} \\ \mathbf{p} &= [\frac{1}{\sqrt{2}}(|0\rangle\langle 0| + |0\rangle\langle 1|), 0].\mathbf{p} + [\frac{1}{\sqrt{2}}(|\rangle\langle 0| - |\rangle\langle 1|)].\mathbf{f} \end{aligned}$$

La règle de regroupement peut être appliquée deux fois sur le terme $\mathcal{P}_7$ afin de supprimer les actions 0 : $\mathcal{P}_7 \xrightarrow{0}_r \xrightarrow{0}_r \mathcal{P}'_7$

$\mathcal{P}'_7 = (\{\mathbf{i}, \mathbf{p}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$ avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathbb{C}$, $\mathcal{H}_{\mathbf{p}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle|].\mathbf{p} \\ \mathbf{p} &= [\frac{1}{\sqrt{2}}(|0\rangle\langle 0| + |0\rangle\langle 1|)].\mathbf{p} + [\frac{1}{\sqrt{2}}(|\rangle\langle 0| - |\rangle\langle 1|)].\mathbf{f} \end{aligned}$$

– $\mathcal{P}_8 = (\{\mathbf{i}, \mathbf{q}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$ avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{\mathbf{f}} = \mathbb{C}$, $\mathcal{H}_{\mathbf{q}} = \mathcal{H}_{\{0,1\}}$ et R :

$$\begin{aligned} \mathbf{i} &= [|0\rangle\langle|].\mathbf{q} \\ \mathbf{q} &= [\frac{1}{\sqrt{2}}|0\rangle\langle 0|, \frac{1}{\sqrt{2}}|0\rangle\langle 1|].\mathbf{q} + [\frac{1}{\sqrt{2}}|\rangle\langle 0|, \frac{1}{\sqrt{2}}|\rangle\langle 1|].\mathbf{f} \end{aligned}$$

Les termes $\mathcal{P}_7$ et $\mathcal{P}_8$ sont irréductibles dans $\mathcal{R}$, et pourtant ils sont différents, donc $\mathcal{R}$ n'est pas confluent. De même, les termes $\mathcal{P}'_7$ et $\mathcal{P}_8$ sont irréductibles dans $\mathcal{R}_{\mathbb{R}^+}$ et dans $\mathcal{R}_{\mathbb{C}}$, donc $\mathcal{R}_{\mathbb{R}^+}$ et $\mathcal{R}_{\mathbb{C}}$ ne sont pas confluents. $\square$

5.4.2 Perspectives

Les règles de réécriture introduites dans cette section préservent les sémantiques du $\mathbf{q}$ -calcul, terminent, mais ne confluent pas. Une raison importante de la non-confluence est la présence de définitions récursives de processus ; en effet, aucune règle ne permet de dérouler les boucles.

Une perspective est donc de proposer des règles permettant de dérouler les boucles, soit en calculant localement le point fixe du processus défini récursivement,

soit dans une perspective plus opérationnelle, en autorisant une approximation au niveau des déroulement des boucles. Dans le cadre de l'approximation, une telle règle ne préserve pas les sémantiques du q -terme. Une perspective est donc aussi de mettre en place un (ou plusieurs) domaine abstrait dans lequel les points fixes peuvent être obtenus en un nombre fini d'étapes permettant un déroulement des définitions récursives.

5.5 Un rôle unificateur

Dans cette section, le rôle unificateur du q -calcul est mis en évidence. En effet, le q -calcul est un modèle permettant de décrire une évolution quantique contrôlée classiquement. En établissant des restrictions au pouvoir expressif du q -calcul (par exemple en n'autorisant que certaines transformations admissibles, ou en interdisant les définitions récursives de processus ...), on obtient dans certains cas un modèle équivalent à un modèle déjà existant.

Ainsi un fragment du q -calcul où les seules actions possibles sont des transformations unitaires et où les définitions de processus sont de la forme $q = [a].p$ (sans opérateur $+$) avec $p \neq q$, permet de décrire le fragment réversible du calcul quantique, tout comme les circuits quantiques ou les machines de Turing quantiques.

Un modèle équivalent au m -calcul est obtenu en utilisant un fragment du q -calcul où les seules actions possibles sont les mesures sur 1-qubit, la transformation unitaire ΛZ et les opérateurs de Pauli, et où les définitions de processus ne doivent pas être récursives.

Le modèle des circuits quantiques généralisés [Aha99] est un modèle où les super-opérateurs sont autorisés, ainsi ce modèle est équivalent au fragment du q -calcul où le contrôle classique est équivalent à celui des circuits quantiques, i.e toutes les définitions de processus sont de la forme $q = [a].p$ avec $p \neq q$. En revanche toutes les transformations admissibles sont autorisées.

Enfin, le fragment du q -calcul où seules les mesures projectives (i.e. transformations admissibles où les opérateurs linéaires sont des projecteurs) s'avère être un modèle formel pour le modèle introduit par Nielsen [Nie03], où seules les mesures projectives sont autorisées, mais où un contrôle récursif est nécessaire. Dans le chapitre 8, ce fragment du q -calcul est utilisé pour démontrer des propriétés de simulations et d'universalité du modèle introduit par Nielsen. Le q -calcul et les machines de Turing quantiques contrôlées classiquement (introduites dans le chapitre 6) sont les seuls modèles formels permettant de représenter le calcul quantique par mesures projectives.

5.6 Conclusion

Nous avons introduit le $\mathbf{q}$ -calcul, un modèle formel du calcul quantique contrôlé classiquement. Chaque $\mathbf{q}$ -terme possède trois sémantiques différentes. Des relations d'interprétation exacte et d'abstraction exacte ont été mises en évidence entre ces sémantiques admissible, pure et observable.

Des règles de transformations sur les $\mathbf{q}$ -termes ont été introduites, elles forment le $\mathbf{q}$ -calcul. Les règles de factorisation et de substitution préservent les trois sémantiques des $\mathbf{q}$ -termes, en revanche la règle de regroupement ne préserve pas la sémantique admissible, ne préserve la sémantique pure que si le paramètre α de cette règle est un réel positif et préserve la sémantique observable dans tous les cas.

L'objectif de ces règles est de permettre une transformation des $\mathbf{q}$ -termes afin d'atteindre un $\mathbf{q}$ -terme "plus simple". Pourtant, même si le $\mathbf{q}$ -calcul termine, il n'est pas confluent, donc non convergent.

Un $\mathbf{q}$ -terme représente une exécution quantique, qui est décrite par la sémantique du $\mathbf{q}$ -terme, mais la notion de complexité n'est pas apparente dans ce modèle de calcul quantique contrôlé classiquement. Afin de définir de façon rigoureuse le coût d'une évolution quantique contrôlée classiquement, et de comparer la puissance du calcul quantique contrôlé classiquement par rapport au calcul quantique réversible, ou au calcul classique, nous introduisons dans le prochain chapitre une machine de Turing quantique contrôlée classiquement.

Chapitre 6

Machine de Turing quantique contrôlée classiquement

6.1 Introduction

L'un des principaux modèles abstraits du calcul quantique est la machine de Turing quantique (MTQ) définie par Deutsch [Deu85], qui est un analogue quantique de la machine de Turing *classique* (MT). Ce modèle a été largement étudié par Bernstein et Vazirani [BV97] : une machine de Turing quantique est un modèle abstrait d'ordinateurs quantiques, qui étend le modèle classique des machines de Turing en introduisant une fonction de transition *quantique*. Pour une MTQ, des superpositions et interférences de configurations sont possibles, mais la partie classique organisant l'enchaînement des opérations n'est pas formalisée et les entrées et sorties de la machine sont toujours traitées classiquement.

Tandis que des modèles traitant d'états quantiques tels que des circuits quantiques et QRAM sont largement utilisés pour décrire des algorithmes spécifiques, le développement des classes de complexité, tel que QMA [Wat00], qui traite des états quantiques, montre la nécessité des modèles abstraits pour les calculs agissant sur des données quantiques.

La machine de Turing quantique linéaire (MTQL), introduite par S. Iriyama, M. Ohya et I. Volovich [IOV04] est une généralisation de la machine de Turing quantique opérant sur des états mixtes et des fonctions de transition irréversibles, et qui permet la représentation de mesures quantiques sans résultat classique. Une conséquence de cette absence de résultat classique est l'impossibilité de formaliser un contrôle classique. Le protocole de téléportation quantique, parmi d'autres, ne peut être exprimé dans le modèle des machines de Turing quantiques linéaires. De plus, comme avec les MTQ, ces machines n'opèrent que sur des entrées/sorties classiques.

Nous introduisons ici une machine de Turing quantique à contrôle classique (MTQC) qui est une machine de Turing avec un ruban quantique pour manipuler des données quantiques et une fonction de transition classique pour formaliser le contrôle classique. Dans le formalisme des MTQC, les transformations unitaires mais aussi les mesures sont autorisées. Plus généralement, les transformations admissibles peuvent être représentées. Ce formalisme des MTQC et le q -calcul (cf chapitre 5) sont complémentaires. Là où le q -calcul permet de définir formellement la sémantique d'un calcul quantique contrôlé classiquement, une MTQC représente le calcul dans un plus bas niveau, plus proche de l'architecture. Le nombre de transitions permet de définir le temps d'exécution d'un calcul, et ainsi de comparer la puissance de calcul des MTQC avec d'autres modèles de calcul comme les machines de Turing classique et quantique, ou les circuits quantiques.

Le théorème 6.1 montrera que toute MT peut être simulée par une MTQC sans perte d'efficacité. Dans la section 6.4, une MTQC avec plusieurs rubans est introduite. Le théorème 6.2 démontre que toute MTQC à k rubans peut être simulée par une MTQC à deux rubans avec un ralentissement quadratique.

De plus, une séparation entre classique et quantique est mise en évidence. En effet, une MTQC à un ruban peut simuler efficacement toute machine de Turing Classique, en revanche certaines MTQC à deux rubans ne peuvent pas être simulées efficacement au moyen de MTQC à un ruban.

Dans la section 6.5.2, le modèle des MTQC est comparé au modèle des machines de Turing quantiques.

Dans la section 6.5, le modèle des MTQC est comparé à deux modèles existant de calcul quantique : le modèle des circuits quantiques [Yao93] et celui du m -calculus [DKP04a]. Les deux modèles sont simulés efficacement par des MTQC.

Afin de souligner la similarité entre la programmation d'une MT et d'une MTQC, des exemples de plusieurs MTQC sont présentés pour résoudre des problèmes tels que la reconnaissance de palindromes quantiques ou l'insertion d'un symbole neutre dans une donnée quantique. Un des objectifs est de faire du modèle des MTQC non seulement un modèle théorique du calcul quantique à contrôle classique, mais aussi un pont vers les modèles concrets de calcul quantique comme la QRAM, en s'appuyant sur le fait que les modèles naturels du calcul quantique sont contrôlés classiquement.

6.2 Machines de Turing quantiques contrôlées classiquement

La définition d'une machine de Turing déterministe est rappelée dans la définition 6.1. Une machine de Turing quantique à contrôle classique (définition 6.2)

est composée d'un ruban de cellules quantiques, d'un ensemble fini d'états internes classiques et d'une tête pour l'application de transformations admissibles sur les cellules du ruban. Le rôle de la tête est crucial puisqu'il implémente l'interaction entre la partie classique et la partie quantique de la machine.

Définition 6.1 Une machine de Turing (classique) est définie par un triplet $M = (K, \Sigma, \delta)$, dans lequel K est un ensemble fini d'états incluant un état initial s , Σ est un alphabet fini incluant un symbole neutre $\#$ et δ est une fonction de transition :

$$\delta : K \times \Sigma \rightarrow (K \cup \{\text{"yes"}, \text{"no"}, h\}) \times \Sigma \times \{\leftarrow, \rightarrow, -\}.$$

On supposera que h (l'état d'arrêt), "yes" (l'état acceptant) et "no" (l'état rejetant) ne sont pas dans K .

Définition 6.2 Une Machine de Turing Quantique à contrôle Classique (MTQC) est un quintuplet $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$. Ici, K est un ensemble fini d'états classiques incluant un état initial s , Σ_Q est un alphabet fini représentant les états de bases de chaque cellule quantique, Σ_C est un alphabet fini de résultats classiques, $\mathcal{A}$ est un ensemble fini de transformations admissibles agissant sur une cellule et δ est une fonction de transition classique :

$$\delta : K \times \Sigma_C \rightarrow (K \cup \{\text{"yes"}, \text{"no"}, h\}) \times \{\leftarrow, \rightarrow, -\} \times \mathcal{A}.$$

On supposera que h (l'état d'arrêt), "yes" (l'état acceptant) et "no" (l'état rejetant) ne sont pas dans K et que tous les résultats classiques possibles de chaque transformation admissible de $\mathcal{A}$ sont dans Σ_C . De plus, on supposera que Σ_Q contient toujours un symbole neutre $\#$, que Σ_C contient toujours un symbole neutre $\#$ et son complémentaire $\overline{\#}$, et que $\mathcal{A}$ contient toujours la transformation admissible $\mathcal{T}_{\#}$ de test du symbole neutre (voir figure 6.1).

La fonction δ est une formalisation du contrôle classique régissant le calcul quantique et peut également être vue comme le "programme" de la machine. Elle détermine, pour chaque combinaison d'état interne classique $q \in K$ et du dernier résultat classique obtenu $\tau \in \Sigma_C$, le triplet $\delta(q, \tau) = (p, D, A)$, où p est le prochain état interne de la machine, $D \in \{\leftarrow, \rightarrow, -\}$ est la direction selon laquelle la tête de la machine va se déplacer et $A \in \mathcal{A}$ est la prochaine transformation admissible appliquée. La transformation admissible $\mathcal{T}_{\#} = \{M_{\#}, M_{\overline{\#}}\}$ testant le symbole neutre établit la correspondance entre le symbole quantique $\# \in \Sigma_Q$ et les symboles classiques $\#, \overline{\#} \in \Sigma_C$: si l'état $|\varphi\rangle$ de la cellule sur laquelle la transformation opère est $|\#\rangle$, alors le résultat classique est $\#$, alors que si $|\varphi\rangle$ est orthogonal à $|\#\rangle$ ($\langle\varphi|\#\rangle = 0$) alors le résultat est $\overline{\#}$.

Comment le programme commence-t-il ? L'entrée quantique du calcul généralement inconnue $|\varphi\rangle = \sum_{\tau \in (\Sigma_Q - \{\#\})^n} \alpha_{\tau} |\tau\rangle$ est placée sur n cellules adjacentes sur

le ruban, alors que l'état de toutes les autres cellules du ruban est $|\#\rangle$. La tête lit la cellule neutre immédiatement à la gauche de l'entrée. Initialement, l'état interne de la machine est s tandis que $\#$ est considéré comme le dernier résultat classique, ainsi la première transition est toujours $\delta(s, \#)$.

Comment le programme s'arrête-t-il ? La fonction de transition δ est totale sur $K \times \Sigma_C$ (les transitions inapplicables sont écartées de cette description). La seule raison pour laquelle la machine ne peut continuer est la suivante : l'un des trois états d'arrêt h , “yes”, et “no” a été atteint. Si une machine M s'arrête sur une entrée $|\varphi_{in}\rangle$, alors la sortie $M(|\varphi_{in}\rangle)$ de la machine M sur $|\varphi_{in}\rangle$ est définie. Si l'un des états “yes” ou “no” est atteint, alors $M(|\varphi_{in}\rangle) = \text{“yes”}$ ou “no” respectivement. Autrement, si l'état d'arrêt h est atteint, alors la sortie est l'état $|\varphi_{out}\rangle$ du ruban de la machine M , au moment de l'arrêt. Puisque la machine a effectué un nombre fini d'étapes, seulement un nombre fini de cellules quantiques ne sont pas dans l'état $|\#\rangle$. L'état de sortie $|\varphi_{out}\rangle$ est l'état d'un registre fini composé de cellules quantiques, de la cellule la plus à gauche dans un état qui n'est pas $|\#\rangle$, vers la cellule la plus à droite dans un état qui n'est pas $|\#\rangle$. Bien entendu, il est possible que M ne s'arrête jamais sur l'entrée $|\varphi_{in}\rangle$. Si c'est le cas, on indique $M(|\varphi_{in}\rangle) = \nearrow$.

Etant donné un espace de Hilbert $\mathcal{H}_{\Sigma_Q}$, nous exhibons quelques transformations admissibles qui seront utilisées dans la suite, et dont les résultats classiques associés sont dans l'ensemble fini $\Sigma_C = \Sigma_Q \cup \overline{\Sigma_Q} \cup \{\lambda\}$, où $\overline{\Sigma_Q} = \{\bar{\tau} : \tau \in \Sigma_Q\}$ et $\lambda \notin \Sigma_Q$:

- $Std = (M_\tau)_{\tau \in \Sigma_Q}$ est une mesure projective dans la base standard : $\forall \tau \in \Sigma_Q, M_\tau = |\tau\rangle \langle \tau|$,
- $\mathcal{T}_\tau = (M_\tau, M_{\bar{\tau}})$ est un test de neutralité $\tau : M_\tau = |\tau\rangle \langle \tau|$ et $M_{\bar{\tau}} = I - |\tau\rangle \langle \tau|$,
- $\mathcal{P}_{[\tau_a, \tau_b]} = (M_\lambda)$ est une transformation unitaire dont le résultat classique est λ , et $M_\lambda = (\sum_{\tau \in \Sigma_Q - \{\tau_a, \tau_b\}} |\tau\rangle \langle \tau| + |\tau_a\rangle \langle \tau_b| + |\tau_b\rangle \langle \tau_a|)$ est une permutation des symboles τ_a and τ_b .
- $Swap = (M_\lambda)$ est une transformation unitaire sur deux cellules $M_\lambda = (\sum_{\tau, \sigma \in \Sigma_Q} |\tau\sigma\rangle \langle \sigma\tau|)$ et dont le résultat classique est λ .
- $\mathcal{U}_V = (M_\lambda)$ est la transformation unitaire $M_\lambda = V$ et dont le résultat classique est λ .
- $\mathcal{O}_O = (P_k)_k$, est une mesure projective selon l'observable $O = \sum_k \lambda_k P_k$.

FIG. 6.1 – Exemples de transformations admissibles

Puisque l'informatique quantique est probabiliste, pour un état d'entrée donné $|\varphi_{in}\rangle$, une MTQC ne produit pas toujours, en général, la même sortie. donc, il existe une distribution probabiliste de toutes les sorties possibles. De plus, le temps d'arrêt d'une MTQC sur une entrée $|\varphi_{in}\rangle$ est également probabiliste. Ainsi, deux

classes spéciales des MTQC peuvent être distinguées : *Monte Carlo* et *Las Vegas*. Pour un MTQC donnée, si pour une entrée donnée $|\varphi_{in}\rangle$, il existe une limite finie et non-probabiliste pour le temps d'exécution de M , alors M est *Monte Carlo*. Si la sortie $M(|\varphi_{in}\rangle)$ est non probabiliste, alors M est *Las Vegas*. Un exemple de *Monte Carlo* est donné dans l'exemple 6.1 : cette MTQC reconnaît un langage composé de "palindromes quantiques", c'est-à-dire d'états quantiques qui sont des superpositions de palindromes. Dans la section 6.3, on utilise une MTQC qui est à la fois *Las Vegas* et *Monte Carlo* pour simuler une MT classique.

La *configuration* d'une MTQC M est une description complète de l'état courant de la machine. Formellement, une configuration est un quadruplet $(q, \tau, |\psi\rangle, x)$, où $q \in K \cup \{h, \text{"yes"}, \text{"no"}\}$ est l'état interne de M , $\tau \in \Sigma_C$ est le dernier résultat obtenu, $|\psi\rangle \in \mathcal{H}_{\Sigma_Q}$ représente l'état du ruban, et $x \in \mathbb{Z}$ représente la position de la tête avec par convention $x = 0$ pour la configuration initiale de la machine.

Exemple 6.1 (Palindromes quantiques) *Considérons la MTQC $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, avec $K = \{s, q, q_0, q_1, q'_0, q'_1, \tilde{q}\}$, $\Sigma_C = \{\#, \overline{\#}, 0, 1, \lambda\}$, $\Sigma_Q = \{\#, 0, 1\}$ et $\mathcal{A} = \{\mathcal{T}_{\#}, Std, \mathcal{P}_{[0, \#]}, \mathcal{P}_{[1, \#]}\}$ (ces transformations admissibles sont décrites dans la figure 6.1), et δ comme décrit dans la figure 6.2.*

Le but de cette machine est de dire si son entrée est un palindrome quantique, c'est-à-dire une superposition d'états de base tels que chaque état de base de la superposition est un palindrome. Par exemple, les états : $|00\rangle$, $\frac{1}{\sqrt{2}}(|010\rangle + i|111\rangle)$, $\frac{1}{\sqrt{2}}(|00000\rangle + |11111\rangle)$ sont des palindromes quantiques. La machine fonctionne comme suit : la première cellule de l'entrée est mesurée dans la base standard et remplacée par $|\#\rangle$, le résultat est mémorisé au moyen des états internes q_0 et q_1 , puis, M se déplace sur la droite, jusqu'à la fin de l'entrée. La dernière cellule est ensuite mesurée dans la base standard : si le résultat correspond à celui mémorisé, il est remplacé par $|\#\rangle$. M revient alors sur la gauche au début de l'entrée restante et le procédé se répète. La fonction de transition est décrite à la figure 6.2. Par exemple, si l'état interne est q_0 et le dernier résultat obtenu est $\#$, alors l'état interne devient q'_0 , la tête se déplace sur la gauche et ainsi la cellule concernée est mesurée dans la base standard.

Cette machine est une MTQC Monte Carlo opérant en un temps $O(n^2)$, où n est la taille de l'entrée. Soit $L \subseteq \mathcal{H}_{\Sigma_Q}$ le langage composé de palindromes quantiques, si $|\varphi_{in}\rangle \in L$, alors la probabilité que M accepte $|\varphi_{in}\rangle$ est $\Pr[M(|\varphi_{in}\rangle) = \text{"yes"}] = 1$: si l'entrée est un palindrome quantique, alors, dans tous les cas, la machine reconnaît $|\varphi_{in}\rangle$, mais M pourra accepter les états qui ne sont pas des palindromes avec une forte probabilité, par exemple $\forall \epsilon > 0, \Pr[M(\sqrt{1-\epsilon}|00\rangle + \sqrt{\epsilon}|10\rangle) = \text{"yes"}] = 1 - \epsilon$.

$p \in K,$	$\tau \in \Sigma_C$	$\delta(p, \tau)$	$p \in K$	$\tau \in \Sigma_C$	$\delta(p, \tau)$
s	$\#$	$(q, \rightarrow, Std)$	q'_0	$\#$	$(\text{"yes"}, -, -)$
q	$\#$	$(\text{"yes"}, -, -)$	q'_0	0	$(\tilde{q}, -, \mathcal{P}_{[0, \#]})$
q	0	$(q_0, -, \mathcal{P}_{[0, \#]})$	q'_0	1	$(\text{"no"}, -, -)$
q	1	$(q_1, -, \mathcal{P}_{[1, \#]})$	q'_1	$\#$	$(\text{"yes"}, -, -)$
q_0	λ	$(q_0, \rightarrow, \mathcal{T}_\#)$	q'_1	0	$(\text{"no"}, -, -)$
q_0	$\overline{\#}$	$(q_0, \rightarrow, \mathcal{T}_\#)$	q'_1	1	$(\tilde{q}, -, \mathcal{P}_{[1, \#]})$
q_0	$\#$	$(q'_0, \leftarrow, Std)$	$\tilde{q}$	λ	$(\tilde{q}, \leftarrow, \mathcal{T}_\#)$
q_1	λ	$(q_1, \rightarrow, \mathcal{T}_\#)$	$\tilde{q}$	$\overline{\#}$	$(\tilde{q}, \leftarrow, \mathcal{T}_\#)$
q_1	$\overline{\#}$	$(q_1, \rightarrow, \mathcal{T}_\#)$	$\tilde{q}$	$\#$	$(q, \rightarrow, Std)$
q_1	$\#$	$(q'_1, \leftarrow, Std)$			

FIG. 6.2 – MTQC pour des palindromes quantiques. Le symbole “ $-$ ” utilisé pour une transformation admissible, signifie $\mathcal{U}_I$, i.e. la transformation identité avec λ comme résultat classique.

6.3 MTQC et MT

Le théorème suivant montre que toute MT est simulée par une MTQC sans perte d'efficacité.

Théorème 6.1 *Pour toute MT M_C donnée, opérant en un temps $f(n)$, où n est la taille de l'entrée, il existe une MTQC M_Q opérant en un temps $O(f(n))$ et telle que pour toute entrée x , $M_C(x) = M_Q(|x\rangle)$.¹*

Preuve : Pour une MT $M_C = (K, \Sigma, \delta_C)$ donnée, on décrit une MTQC M_Q qui simule M_C . Une façon de réaliser cela est de simuler un ruban classique de M_C en utilisant uniquement les états de base du ruban quantique de M_C .

Formellement, on considère la MTQC $M_Q = (K \cup K_\Sigma \cup \{s'\}, \Sigma \cup \{\overline{\#}, \lambda\}, \Sigma, \mathcal{A}, \delta_Q)$. Ici, $K_\Sigma = \{q_\tau : q \in K, \tau \in \Sigma\}$, $\mathcal{A} = \{Std\} \cup \{\mathcal{P}_{[\tau_1, \tau_2]}\}_{\tau_1, \tau_2 \in \Sigma}$. L'état initial de M_Q est s' et la première transition est $\delta_Q(s', \#) = (s, -, Std)$, où s est l'état initial de M_C . Pour tout $(q, \tau) \in K \times \Sigma$, la transition $\delta_C(q, \tau) = (q', \tau', D)$ est décomposée en deux transitions : $\delta_Q(q, \tau) = (q_\tau, -, \mathcal{P}_{[\tau, \tau']})$ et $\delta_Q(q_\tau, \lambda) = (q', D, Std)$.

Puisque chaque transition de M_C est simulée avec une probabilité de 1 par deux transitions de M_Q , si M_C opère en un temps $f(n)$, M_Q opère en un temps $2f(n)$, où n est la taille de l'entrée. $\square$

Chaque MT est donc simulée par une MTQC sans perte d'efficacité. Cependant, comme il sera montré dans le lemme 6.1, une MTQC avec un ruban ne peut pas

¹Si l'état d'arrêt h est atteint, $M_Q(|x\rangle)$ indique l'état final du ruban. Donc, si h est atteint, $M_C(x) = M_Q(|x\rangle)$ doit être remplacé par $M_Q(|x\rangle) = |M_C(x)\rangle$.

simuler certains modèles de calcul quantique, tels que les circuits quantiques, parce que seules les transformations admissibles sur une cellule sont permises. Dans le but d'autoriser les transformations sur plus d'une cellule, des MTQC multi-rubans sont introduites. Avec k têtes, des transformations admissibles sur k cellules peuvent être réalisées.

6.4 MTQC multi-rubans

On montre que toute MTQC avec k rubans est simulée par une MTQC avec deux rubans avec une perte d'efficacité inconséquente. De plus, en montrant que les MTQC avec 1 ruban et 2 rubans ne sont pas équivalentes, on met en évidence une *séparation* entre classique et quantique.

Définition 6.3 (k -MTQC) *Une machine de Turing quantique contrôlée classiquement à k rubans, où $k > 0$, est un quintuplet $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, où K est un ensemble fini d'états classiques avec un état initial identifié s , Σ_Q est un alphabet fini qui représente les états de base de chaque cellule quantique, $\mathcal{A}$ est un ensemble fini de transformations admissibles agissant sur k cellules, Σ_C est un alphabet fini de résultats classiques de transformations admissibles sur k cellules et δ est une fonction de transition classique :*

$$\delta : K \times \Sigma_C \rightarrow (K \cup \{"yes", "no", h\}) \times (\{\leftarrow, \rightarrow, -\})^k \times \mathcal{A}.$$

On supposera que tous les résultats possibles de chaque mesure de $\mathcal{A}$ sont dans Σ_C et que $\mathcal{A}$ contient toujours les transformations admissibles testant le symbole neutre sur un des k rubans.

Intuitivement, $\delta(q, \tau) = (q', (D_1, \dots, D_k), A)$ signifie que, si M est dans l'état q et le dernier résultat est τ , alors le prochain état sera q' , les k têtes de la machine se déplaceront suivant $D_1, \dots, D_k$ et la prochaine transformation admissible sur k cellule quantique sera A . Cette transformation admissible sera réalisée sur les k cellules quantiques pointées par les têtes de la machine après leur déplacement. Une transformation admissible A agissant sur k cellules peut être définie directement, par exemple en utilisant une transformation unitaire V agissant sur k cellules ($A = \mathcal{U}_V$). A peut aussi être définie comme une composition de deux transformations admissibles A_1, A_2 respectivement sur les cellules j et l tel que $j + l = k$, alors, $A = [A_1, A_2]$ signifie que les j premières têtes appliquent A_1 et, simultanément, les l dernières têtes appliquent A_2 . Le résultat classique est la concaténation des résultats de A_1 et A_2 , où λ est l'élément neutre de la concaténation (i.e. $\tau.\lambda = \tau$).

Une k -MTQC commence par un état initial $|\varphi\rangle$ sur un ruban spécifique T_1 , tous les autres rubans étant dans l'état $|\#\rangle$, et si l'état d'arrêt h est atteint, la machine s'arrête et la sortie est l'état du ruban spécifique T_1 .

$p \in K,$	$\tau \in \Sigma_C$	$\delta(p, \tau)$
s	$\#$	$(q_0, (\leftarrow, -), \text{Swap})$
q_0	λ	$(q_1, (\rightarrow, -), \text{Swap})$
q_1	λ	$(h, (\rightarrow, -), -)$

FIG. 6.3 – Une MTQC à 2 rubans pour l'insertion d'un symbole neutre

Exemple 6.2 (Insérer un symbole neutre) *Considérons le problème de l'insertion d'un symbole neutre entre la première et la seconde cellule d'un état quantique $|\psi_{in}\rangle$ qui réside sur l'un des rubans. Par exemple, $|abba\rangle$ se transforme en $|a\#bba\rangle$, et $\frac{1}{\sqrt{2}}(|aa\rangle + |bb\rangle)$ en $\frac{1}{\sqrt{2}}(|a\#a\rangle + |b\#b\rangle)$. Considérons la 2-MTQC $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, avec $K = \{s, q_0, q_1\}$, $\Sigma_C = \{\#, \overline{\#}, \lambda\}$ et $\mathcal{A} = \{\mathcal{T}_\#, \text{Swap}\}$. δ est décrit à la figure 6.3.*

L'état d'entrée est sur le premier ruban. Soit c_0 la première cellule sur la gauche de l'entrée. Dans le but d'insérer un symbole neutre en seconde position de l'état d'entrée, l'état de c_0 est échangé avec celui d'une cellule du second ruban. Alors, l'état de cette cellule sur le second ruban est échangé avec celui de la cellule immédiatement sur la gauche de c_0 .

Théorème 6.2 *Pour toute k -MTQC M opérant en un temps $f(n)$, où n est la taille de l'entrée, il existe une 2-MTQC M' opérant en un temps $O(f(n)^2)$ et telle que pour toute entrée $|\psi\rangle$, $M(|\psi\rangle) = M'(|\psi\rangle)$.*

Preuve : Supposons que $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$ possède k rubans, on décrit $M' = (K', \Sigma'_C, \Sigma'_Q, \mathcal{A}', \delta')$ possédant seulement deux rubans. M' doit "simuler" les k rubans de M . Une façon de réaliser cela est de maintenir sur un ruban T_1 de M' la *concaténation* du contenu des rubans de M . La position de chaque tête doit également être mémorisée.

Pour accomplir cela, $\Sigma'_Q = \Sigma_Q \cup \underline{\Sigma}_Q \cup \{\triangleright, \triangleleft\}$, où $\underline{\Sigma}_Q = \{\underline{\tau} : \tau \in \Sigma_Q\}$ est un ensemble de versions pointées de symboles dans Σ_Q , et où $\triangleright$ ($\triangleleft$) signale la fin à gauche (à droite) de chaque ruban simulé. Intuitivement, à chaque étape du calcul, si $|\varphi_j\rangle$ est l'état de chaque ruban j de M , l'état du ruban T_1 de M' est $|\triangleright\rangle|\varphi_1\rangle|\triangleleft\rangle|\varphi_2\rangle|\triangleleft\rangle\ldots|\triangleleft\rangle|\varphi_k\rangle|\triangleleft\rangle$. Dans le but de mémoriser les positions des k têtes, une transformation unitaire est appliquée sur les cellules de M' correspondant à celles de M sur lesquelles se déplacent les têtes de M . Cette transformation unitaire remplace les symboles de Σ_Q par leur versions correspondantes dans $\underline{\Sigma}_Q$.

Puisque chaque transformation admissible agissant sur k cellules provenant de $\mathcal{A}$ peut être décomposée en $l_{\mathcal{A}}$ transformations admissibles agissant sur 2 cellules, (voir [MS00]), alors $\mathcal{A}'$, qui est composé de transformations admissibles agissant sur 1 ou 2 cellules, est défini tel que toute transformation provenant de $\mathcal{A}$ puisse

être simulée par un nombre l_A de transformations de $\mathcal{A}'$. l_A est exponentiel en le nombre k de rubans, or k est fixé, donc l_A est une constante.

Pour que la simulation commence, M' insère $\triangleright\triangleright$ sur la gauche et $\triangleleft(\triangleright\triangleleft)^k\triangleleft$ sur la droite de l'entrée, puisque l'entrée de M est localisée sur son premier ruban. Pour simuler la transition $\delta(q, \tau) = (q', D, A)$ de M , les cellules pointées par les têtes changent en premier selon D . Notons que si une tête rencontre le symbole $\triangleright$, alors un symbole neutre est inséré à la droite de cette cellule (voir l'exemple 6.2) pour simuler l'infinité des rubans, et de la même façon pour le symbole $\triangleleft$. A est simulé via une séquence de transformations agissant sur 2 cellules. Puisque de telles transformations sont possibles uniquement sur des cellules localisées sur des rubans différents, l'état de l'une des cellules est transféré (au moyen de *Swap*, voir exemple 6.2) du ruban T_1 à l'autre ruban T_2 . Alors, la transformation sur 2 cellules peut avoir lieu et l'état localisé sur T_2 est re-transféré sur T_1 , etc.

Dans le but de reconstruire le résultat classique de la transformation simulée A , M' doit utiliser de nouveaux états internes, afin de garder en mémoire les résultats classiques des différentes transformations sur 1 et 2 cellules.

La simulation se poursuit jusqu'à ce que M s'arrête. Combien de temps le calcul d'une entrée $|\varphi\rangle$ de taille n prend-il ? Puisque M s'arrête en un temps $f(n)$, pas plus de $k.f(n)$ cellules de M sont des cellules remplies. Ainsi la longueur totale des cellules remplies de M' est $k.(f(n) + 2) + 3$ (pour tenir compte de $\triangleleft, \triangleright$ et les cellules de T_2 utilisées pour l'application de transformations de 2 cellules quantiques). Simuler le mouvement des têtes prend tout au plus deux traversées des cellules non neutres de T_1 . Chaque simulation d'une transformation admissible de $\mathcal{A}$ requiert un nombre constant l_A de transformations de $\mathcal{A}'$ (l_A est indépendant de la taille de l'entrée), de plus, la simulation de chaque transformation dans $\mathcal{A}'$ requiert deux traversées. Par conséquent, la simulation de chaque transition de M demande $O(f(n))$ transitions de M' , ainsi, le temps d'exécution total de M' est $O(f(n)^2)$. $\square$

Le lemme suivant montre que certaines MTQC à 2 rubans ne peuvent pas être simulées par une MTQC à 1 ruban.

Lemme 6.1 *Il existe une MTQC à 2 rubans M telle qu'aucune MTQC à 1 ruban ne peut simuler M .*

Preuve : Soit $M = (\{s\}, \{\lambda, \#, \overline{\#}\}, \{\#, 0\}, \{\mathcal{U}_V\}, \delta)$ une 2-MTQC, où $\delta(s, \#) = (h, -, \mathcal{U}_V)$ et $V = \frac{1}{2\sqrt{2}}(|\#\# \rangle + |00\rangle) \langle \#\#| + (|\#\# \rangle - |00\rangle) \langle \#0| + (|\#0 \rangle + |0\# \rangle) \langle 0\#| + (|\#0 \rangle - |0\# \rangle) \langle 00|$. Si l'entrée est $|0\rangle$, alors la machine s'arrête, l'état des cellules pointées par les têtes est intriqué : $\frac{1}{\sqrt{2}}(|\#0 \rangle + |0\# \rangle)$. Ainsi, il n'existe pas de MTQC à 1 ruban qui simule M , puisque l'intrication ne peut pas être créée au moyen de transformations admissibles sur une cellule. $\square$

6.5 MTQC et les modèles de calcul quantique réversible

Dans cette section, le modèle des MTQC est comparé aux modèles de calcul quantiques existants et en particulier au modèle des circuits quantiques et à celui des machines de Turing quantiques. Ces deux modèles sont équivalents. Nous montrons dans un premier temps que toute famille de circuits quantiques peut être simulée efficacement par une MTQC à 2 rubans, puis que toute MTQC multirubans peut être simulée efficacement par une machine de Turing quantique. Ainsi, les trois modèles ont une puissance de calcul équivalente.

6.5.1 MTQC et Circuits quantiques

Le théorème 6.2 est une preuve forte de la puissance et de la stabilité des MTQC : l'ajout d'un nombre fini de rubans à une MTQC à 2 rubans n'augmente pas ses capacités de calcul et a un impact sur son efficacité uniquement d'une façon polynomiale. Cette stabilité fait de la MTQC à 2 rubans une bonne candidate à l'universalité quantique, c'est-à-dire la capacité de simuler tout calcul quantique. Cette capacité est prouvée dans le théorème suivant par la simulation d'une famille semi-uniforme quelconque de circuits quantiques [NO02a]. Dans ce paragraphe, certaines notions de base et des propriétés des circuits quantiques sont présentées, on se référera au chapitre 3 et à [Yao93, KSV02] pour une présentation complète des circuits quantiques.

Une famille uniforme de circuits quantiques est un ensemble de circuits $F = \{C_n\}$ tel qu'une machine de Turing classique M_F peut produire une description de C_n sur l'entrée n en un temps $\text{poly}(s(n))$, où $s(n)$ est la taille de l'entrée de C_n . Une famille semi-uniforme de circuits quantiques $\{C_n\}$ est une famille uniforme définie sur un ensemble fini $\mathcal{G}$ d'opérateurs.

Théorème 6.3 *Pour toute famille semi-uniforme de circuits quantiques $F = \{C_n\}$ de taille s , il existe une 2-MTQC M , opérant en un temps $\text{poly}(n, s(n))$, et telle que pour tout état $|\psi\rangle$ des n qubits, $C_n |\psi\rangle = M(|\psi\rangle)$.*

Preuve : L'application de $F = \{C_n\}$ sur un état de n qubits $|\psi\rangle$ consiste à appliquer le circuit quantique C_n avec $|\psi\rangle$ comme entrée. Soit $\mathcal{G}$ une base de F , i.e. l'ensemble de toutes les portes utilisées dans F . Puisque $\mathcal{G}$ est fini, $\mathcal{G}$ possède une arité finie w , i.e. pour toute $G \in \mathcal{G}$, l'arité de G est inférieure à w , où l'arité de la porte est le nombre de qubits sur lesquels elle opère.

La description de C_n produite par M_F est de la forme $(G_1, R_1) \cdots (G_{s(n)}, R_{s(n)})$, signifiant que $G_i \in \mathcal{G}$ est appliqué sur l'ensemble $\{R_i^{(j)}\}_j$ des qubits dans R_i , puis G_{i+1} est appliqué, etc.

Soit M une MTQC à $(w + 1)$ rubans. Les transformations admissibles de M incluent la transformation unitaire $\mathcal{U}_G$, pour tout $G \in \mathcal{G}$.

Une description générale de l'évolution de M est :

- La taille de n de l'entrée $|\psi\rangle$ localisée sur le ruban 1 est calculée, en utilisant la transformation admissible $\mathcal{T}_\#$ du test neutre, et stockée sur le ruban T_{w+1} . La phase d'initialisation peut être réalisée en un nombre d'étapes linéaire en n .
- M_F est simulée (théorème 6.1) et produit une description classique de C_n sur le ruban T_{w+1} . La complexité de cette phase est $p(s(n))$, pour un polynôme p .
- Pour chaque (G_i, R_i) , si $R_i = \{R_i^{(j)}\}_j$, pour chaque j , le qubit $R_i^{(j)}$ du ruban 1 est transféré au ruban j , alors, $\mathcal{U}_{G_i}$ est appliqué, et les qubits de R_i sont retransférés au ruban 1. Cette phase peut être réalisée en $O(ns(n))$ étapes.

On peut montrer que l'état résultant sur le ruban 1 est l'état produit par $C_n |\psi\rangle$. De plus, cette simulation peut être réalisée en $ns(n) + p(s(n))$ étapes.

Au final, la MTQC M à $(w + 1)$ rubans est simulée par une MTQC M' à 2 rubans (théorème 6.2 : d'abord les deux premières étapes de M , qui consistent à calculer la taille de l'entrée et à simuler la machine de Turing classique, puis M_F peut être simulée sans ralentissement sur M' , puisque seulement deux rubans sont nécessaires pour ces étapes).

Ainsi, la MTQC M' à 2 rubans opère en un temps $O(n^2 s(n)^2 + p(s(n)))$. $\square$

Corollaire 6.1 *Pour toute famille semi-uniforme de circuits quantiques de taille polynomiale, $F = \{C_n\}$, il existe, M , une MTQC polynomiale à 2 rubans telle que pour tout n et pour toute entrée $|\psi\rangle$ sur n qubits, $C_n |\psi\rangle = M(|\psi\rangle)$.*

Les familles semi-uniformes de circuits quantiques peuvent être simulées en un temps polynomial au moyen d'une MTQC à 2 rubans. Contrairement aux familles semi-uniformes, les familles uniformes de circuits quantiques possèdent une base dénombrable mais pas forcément finie de portes. Puisque toute MTQC est fondée sur un ensemble fini de transformations admissibles, on conjecture que certaines familles de circuits quantiques ne peuvent pas être simulées au moyen d'une MTQC. Cependant, toutes les familles uniformes de circuits quantiques peuvent être approchées :

Théorème 6.4 *Pour tout $\epsilon > 0$, et pour toute famille uniforme de circuits quantiques $F = \{C_n\}$ de taille s , il existe une MTQC M à 2 rubans opérant en un temps $\text{poly}(n, s(n), 1/\epsilon, 2^w)$ (où w est l'arité des portes de C_n) et telle que pour tout n et pour toute entrée $|\psi\rangle$ sur n qubits, $\|C_n |\psi\rangle - M(|\psi\rangle)\| < \epsilon$.²*

²où $\|\cdot\|$ est la norme Euclidienne.

Preuve : La preuve consiste à combiner le théorème 6.3 et l'approximation de toute transformation unitaire qui utilise un ensemble fini de transformations unitaires. Il existe un ensemble fini de transformations unitaires $\mathcal{U}$ sur au plus 2 qubits, tel que l'approximation à $\epsilon > 0$ près d'une transformation unitaire U sur w qubits utilise $\text{poly}(1/\epsilon, 2^w)$ portes de $\mathcal{U}$ [Aha98]. De plus, il existe un algorithme opérant en un temps $\text{poly}(1/\epsilon, 2^w)$ qui construit une description du circuit réalisant l'approximation de U [BV97].

Pour tout $F = \{C_n\}$ et $\epsilon > 0$ donnés, soit M une MTQC à 2 rubans. Une description générale de l'évolution de M est :

- La taille n de l'entrée $|\psi\rangle$ localisée sur le ruban 1 est calculée, en utilisant la transformation admissible testant le symbole blanc $\mathcal{T}_\#$, puis elle est stockée sur le ruban 2. Cette phase initiale peut être réalisée en un nombre d'étapes linéaire en n .
- M_F est simulé (voir le théorème 6.1) et produit une description classique de C_n sur 2 rubans. Cette phase est réalisée en un temps $\text{poly}(s(n))$.
- Dans la description de C_n , chaque porte est remplacée par une approximation à $\epsilon/s(n)$ près en un temps L , où L est un polynôme en $s(n)/\epsilon$ et 2^w . Cette phase est réalisée en un temps $Ls(n)$ et la description du circuit est maintenant composée de $\text{poly}(s(n), 1/\epsilon, 2^w)$ portes.
- Pour chaque (G_i, R_i) , si $R_i = \{p_i^{(j)}\}_j$, pour chaque j , le qubit $R_i^{(j)}$ du ruban 1 est transféré au ruban j , alors $\mathcal{U}_{G_i}$ est appliqué, et les qubits de R_i sont retransférés vers le ruban 1.

Ainsi M opère en un temps $\text{poly}(n, s(n), 1/\epsilon, 2^w)$. □

Corollaire 6.2 *Pour tout $\epsilon > 0$, et pour toute famille uniforme $F = \{C_n\}$ de circuits quantiques de taille polynomiale, il existe une MTQC M à 2 rubans opérant en un temps $\text{poly}(n, 1/\epsilon, 2^w)$ (où w est l'arité des portes de C_n) et telle que pour tout n et pour toute entrée $|\psi\rangle$ sur n qubits, $\|C_n |\psi\rangle - M(|\psi\rangle)\| < \epsilon$.*

6.5.2 MTQC et MTQ

Dans la section précédente, nous avons prouvé que le formalisme des MTQC est suffisamment puissant pour simuler efficacement toute famille (semi-) uniforme de circuits quantiques. À l'inverse, le théorème 6.5 montre que toute machine de Turing quantique contrôlée classiquement peut être simulée par une machine de Turing quantique, sans ralentissement. Ce résultat démontre que malgré l'utilisation du contrôle classique et la possibilité d'utiliser des transformations admissibles, la machine de Turing quantique contrôlée classiquement est un modèle réaliste puisque sa puissance de calcul n'est pas supérieure à celle d'une machine de Turing quantique.

Les définitions des machines de Turing quantiques avec un ou plusieurs rubans sont rappelées, ainsi que les conditions de bonne formation des MTQ. Une étude plus complète des machines de Turing quantiques avec un ou plusieurs rubans est donnée dans [BV97, Yam99].

Définition 6.4 (MTQ) Une machine de Turing quantique M est définie par un triplet (K, Σ, δ) , où Σ est un alphabet fini avec un symbole neutre $\#$, K est un ensemble fini d'états avec un état initial q_0 identifié et un état final $q_f \neq q_0$, et δ la fonction de transition quantique est une fonction :

$$\delta : K \times \Sigma \rightarrow \mathcal{H}_{K \times \Sigma \times \{\leftarrow, -, \rightarrow\}}$$

L'expression $\delta(p, \sigma, q, \tau, d)$ sera utilisée pour représenter l'amplitude de $\delta(p, \sigma)$ dans $|q\rangle |\tau\rangle |d\rangle$.

Définition 6.5 (k-MTQ) Une machine de Turing quantique à k rubans est un triplet $(K, \Sigma_1 \times \Sigma_2 \times \dots \times \Sigma_k, \delta)$, où chaque Σ_i est un alphabet fini avec un symbole neutre $\#$, K est un ensemble fini d'états avec un état initial q_0 identifié et un état final $q_f \neq q_0$, et δ la fonction de transition quantique est une fonction de $K \times \Sigma_1 \times \Sigma_2 \times \dots \times \Sigma_k$ dans $\mathcal{H}_{K \times \Sigma_1 \times \Sigma_2 \times \dots \times \Sigma_k \times \{\leftarrow, -, \rightarrow\}^k}$.

La notation $\{-1, 0, 1\}$ sera utilisée pour représenter les mouvements $\{\leftarrow, -, \rightarrow\}$.

Une configuration est un élément de $\mathcal{H}_{K \times (\Sigma_1^* \times \Sigma_2^* \times \dots \times \Sigma_k^*) \times \mathbb{Z}^k}$. $T \in \Sigma_1^* \times \Sigma_2^* \times \dots \times \Sigma_k^*$ est un état de base des k rubans de la machine, dont seule une partie finie est utilisée. $x \in \mathbb{Z}^k$, représente la position des têtes et $T_x \in \Sigma_1 \times \Sigma_2 \times \dots \times \Sigma_k$ représente l'état de base pointé par les têtes. L'opérateur d'évolution d'une MTQ M est un opérateur linéaire U_M sur $\mathcal{H}_{K \times (\Sigma_1^* \times \Sigma_2^* \times \dots \times \Sigma_k^*) \times \mathbb{Z}^k}$ tel que :

$$U_M |p, T, x\rangle = \sum_{q \in K, \tau \in \Sigma_1 \times \Sigma_2 \times \dots \times \Sigma_k, d \in \{-1, 0, 1\}^{(k)}} \delta(p, T_x, q, \tau, d) |q, T_x^\tau, x + d\rangle$$

où $T_x^\tau \in \Sigma_1^* \times \Sigma_2^* \times \dots \times \Sigma_k^*$ représente T où les éléments indexés par x_i sont remplacés par τ_i .

Une MTQ est *bien-formée* si et seulement si elle a une évolution unitaire, i.e. U_M est unitaire. Le lemme suivant donne une caractérisation de la bonne formation d'une k -MTQ [Yam99] :

Lemme 6.2 (Bonne formation) Une k -MTQ $M = (K, \Sigma_1 \times \dots \times \Sigma_k, \delta)$ est bien-formée si et seulement si les trois conditions suivantes sont vérifiées :

1. (unité) $\|\delta(p, \sigma)\| = 1$ pour tout $(p, \sigma) \in K \times \Sigma_1 \times \dots \times \Sigma_k$.
2. (orthogonalité) $\delta(p_1, \sigma_1) \cdot \delta(p_2, \sigma_2) = 0$ pour toutes paires distinctes $(p_1, \sigma_1), (p_2, \sigma_2) \in K \times \Sigma_1 \times \dots \times \Sigma_k$.

3. (séparabilité) $\delta[p_1, \sigma_1, \tau_1 | \epsilon] \cdot \delta[p_2, \sigma_2, \tau_2 | \epsilon'] = 0$ pour toutes paires distinctes $\epsilon, \epsilon' \in E^k$ et pour toute paire $(p_1, \sigma_1, \tau_1), (p_2, \sigma_2, \tau_2) \in K \times (\Sigma_1 \times \dots \times \Sigma_k)^2$.

Si l'exécution d'une MTQ M débute avec une configuration $|\varphi\rangle$, alors la configuration finale est notée $M(|\varphi\rangle)$.

Le théorème 6.5 démontre que toute k -MTQC M peut être simulée sans ralentissement par une $k + 3$ -MTQ M' . Ici la simulation de M par M' signifie que l'exécution de M doit être équivalente à l'exécution de M' . L'entrée $|\varphi\rangle$ détermine la configuration initiale de M' . La configuration finale de M' étant formée de l'état interne, l'état de tous les rubans et la position des têtes, une partie de cette configuration finale est mesurée dans la base standard, seuls k rubans ne sont pas mesurés. Il y a simulation de M par M' si la mesure de la configuration finale de M' vérifie les propriétés suivantes :

- La probabilité de mesurer “yes” (“no”) pour l'état interne est $Pr[M(|\varphi\rangle) = \text{“yes”}] (Pr[M(|\varphi\rangle) = \text{“no”}])$.
- La probabilité de mesurer h pour l'état interne et que l'état des k rubans après la mesure soit $|\psi\rangle$ est $Pr[M(|\varphi\rangle) = |\psi\rangle]$.

Théorème 6.5 Pour toute k -MTQC $M = (K, \Sigma_C, \Sigma_C, \mathcal{A}, \delta)$, opérant en temps $f(n)$, où n est la taille de l'entrée, il existe une $k + 3$ -MTQ $M' = ((K \cup \{\text{“yes”}, \text{“no”}, h\}) \times \Sigma, \Sigma_Q^k \times K \times \Sigma_C^2, \delta')$ opérant en temps $f(n)$ telle que pour toute entrée $|\varphi\rangle$, la configuration initiale de M' est $|\varphi'\rangle = |s\rangle |\varphi\rangle |\#, \#, \#, 0\rangle$ et :

- $Pr[M(|\varphi\rangle) = \text{“yes”}] = \sum_{T \in \Sigma_Q^{*k} \times K^* \times \Sigma_C^{*2} \times \mathbb{Z}^k} |\langle \text{“yes”}, T | M'(|\varphi'\rangle) \rangle|^2$
- $Pr[M(|\varphi\rangle) = \text{“no”}] = \sum_{T \in \Sigma_Q^{*k} \times K^* \times \Sigma_C^{*2} \times \mathbb{Z}^k} |\langle \text{“no”}, T | M'(|\varphi'\rangle) \rangle|^2$
- $Pr[M(|\varphi\rangle) = |\psi\rangle] = \sum_{T \in K^* \times \Sigma_C^{*2} \times \mathbb{Z}^k} |\langle h | \langle \psi | \langle T | M'(|\varphi'\rangle) \rangle|^2$

Preuve : Pour toute k -MTQC $M = (K, \Sigma_C, \Sigma_C, \mathcal{A}, \delta)$, soit $M' = ((K \cup \{\text{“yes”}, \text{“no”}, h\}) \times \Sigma, \Sigma_Q^k \times K \times \Sigma_C^2, \delta')$ une $k + 3$ -MTQ telle que pour tout $(p, \tau, \mu) \in K \times \Sigma_C \times \Sigma_Q^k$:

$$\delta'((p, \tau), (\mu, \#, \#, \#)) = \sum_{\sigma \in \Sigma_C} |q^{p, \tau}, \sigma\rangle (M_\sigma^{p, \tau} |\mu\rangle) |p, \tau, \sigma, d^{p, \tau}, \rightarrow, \rightarrow, \rightarrow\rangle$$

avec $\delta(p, \tau) = (q^{p, \tau}, d^{p, \tau}, A)$ et $A = (M_\sigma^{p, \tau})_{\sigma \in \Sigma_C}$.

Seul le cas où le symbole des rubans auxiliaires est $\#$ est pris en compte. En effet, les têtes positionnées sur ces rubans se déplacent toujours vers la droite découvrant nécessairement un symbole $\#$ à chaque transition. Les rubans auxiliaires sont utilisés comme un historique, permettant de mémoriser l'évolution de l'état interne et des résultats classiques de la machine simulée. Nous écrirons $\delta'(p, \tau, \mu)$ pour $\delta'((p, \tau), (\mu, \#, \#, \#))$, de plus les têtes des trois rubans auxiliaires ayant

exactement les mêmes déplacements, seule une position sera représentée pour ces trois têtes.

La fonction de transition sera complétée, via un lemme de complétion [Yam99] lors de la vérification de la bonne formation de M' .

En effet, nous vérifions dans un premier temps que M' simule M , puis les conditions de bonne formation de M' seront vérifiées dans un second temps.

Une configuration de M' est de la forme

$$\left(\sum_{p \in K, \tau \in \Sigma_C, T \in \Sigma_Q^*, w_1 \in K^*, w_2 \in \Sigma_C^*, w_3 \in \Sigma_C^*, x \in \mathbb{Z}^k} \alpha_{p, \tau, T, w_1, w_2, w_3, x} |(p, \tau), T, w_1, w_2, w_3, x\rangle \right) |n\rangle$$

En effet dans la configuration initiale la position des têtes des rubans auxiliaires est séparable et le restera tout au long de l'exécution. De plus, w_1, w_2, w_3 sont des états de base des rubans auxiliaires et chacun d'eux est de longueur n .

L'évolution associée à M' est :

$$U_{M'} |(p, \tau), T, w_1, w_2, w_3, x, n\rangle = \sum_{\sigma \in \Sigma_C, \rho \in \Sigma_Q^k} \langle \rho | M_{\sigma}^{p, \tau} | T_x \rangle |(q^{p, \tau}, \sigma), T_x^{\rho}, w_1 p, w_2 \tau, w_3 \sigma, x + d^{p, \tau}, n + 1\rangle$$

On remarque que si après chaque évolution de M' , la dernière cellule du ruban auxiliaire $k + 3$ est mesurée dans la base standard, alors le résultat classique σ_0 est obtenu avec probabilité $p_0 = \sqrt{\langle T_x | M_{\sigma_0}^{p, \tau \dagger} M_{\sigma_0}^{p, \tau} | T_x \rangle}$, et la configuration de la machine devient :

$$\frac{1}{\sqrt{p_0}} |(q, \sigma_0)\rangle \left(\sum_{\rho \in \Sigma_Q^k} \langle \rho | M_{\sigma_0}^{p, \tau} | T_x \rangle | T_x^{\rho} \rangle \right) |w_1, w_2, w_3 \sigma_0, x + d^{p, \tau}, n + 1\rangle$$

Ainsi, si ces mesures sont effectuées, l'état interne, la position des têtes et l'historique sont des états de base à toute étape. De plus, les k premiers rubans ont une évolution correspondant à l'application de $M_{\sigma_0}^{p, \tau}$, avec renormalisation, sur les cellules pointées par les têtes. On obtient donc une évolution similaire à celle de M .

Or la mesure effectuée après chaque transition de M' est appliquée à chaque étape sur une cellule différente, cellule qui de plus, n'est plus modifiée par la suite. Donc, chaque mesure commute avec toutes les autres mesures et les transitions qui lui succèdent. Ainsi ces mesures peuvent être reportées en fin d'exécution.

Autrement dit, $\langle w_3 | M'(|\varphi'\rangle) \rangle$ est, à une renormalisation près, le résultat $M(|\varphi\rangle)$ obtenu par la machine M si l'historique des résultats classiques est w_3 . La probabilité associée est $(M'(|\varphi'\rangle))^{\dagger} |w_3\rangle \langle w_3| M'(|\varphi'\rangle) = ||\langle w_3 | M'(|\varphi'\rangle) ||^2$.

La probabilité $Pr[M(|\varphi\rangle) = \text{"yes"}]$ est la probabilité que $M(|\varphi\rangle) = \text{"yes"}$ quelle que soit l'historique des résultats classiques, donc $Pr[M(|\varphi\rangle) = \text{"yes"}] = \sum_{w_3 \in \Sigma_C^*} ||\langle \text{"yes"}, w_3 | M'(|\varphi\rangle) \rangle||^2 = ||\langle \text{"yes"} | M'(|\varphi\rangle) \rangle||^2$.

Vérifions maintenant que M' est bien formée. La fonction de transition δ' est partielle puisqu'elle n'est définie sur le ruban auxiliaire que pour le symbole $\#$. Le lemme de complétion montre que si les trois conditions d'unité, d'orthogonalité et de séparabilité sont vérifiées sur le domaine de définition des δ' alors δ' peut être complétée pour que M' soit une MTQ bien formée.

– **Unité :**

$$\forall (p, \tau, \mu) \in K \times \Sigma_C \times \Sigma_1 \times \cdots \times \Sigma_k$$

$$\begin{aligned} ||\delta'(p, \tau, \mu)|| &= \left(\sum_{\sigma \in \Sigma_C} |q^{p, \tau}, \sigma\rangle (M_{\sigma}^{p, \tau} |\mu\rangle) |p, \tau, \sigma, d^{p, \tau}, \rightarrow\rangle \right)^{\dagger} \\ &\quad \left(\sum_{\sigma' \in \Sigma_C} |q^{p, \tau}, \sigma'\rangle (M_{\sigma'}^{p, \tau} |\mu\rangle) |p, \tau, \sigma', d^{p, \tau}, \rightarrow\rangle \right) \\ &= \sum_{\sigma \in \Sigma_C} \langle \mu | M_{\sigma}^{p, \tau \dagger} M_{\sigma}^{p, \tau} | \mu \rangle \\ &= 1 \end{aligned}$$

– **Orthogonalité :**

Pour tous triplets distincts $(p, \tau, \mu), (p', \tau', \mu') \in K \times \Sigma_C \times \Sigma_1 \times \cdots \times \Sigma_k$.

$$\begin{aligned} \delta'(p, \tau, \mu)^{\dagger} \delta'(p', \tau', \mu') &= \left(\sum_{\sigma \in \Sigma_C} |q^{p, \tau}, \sigma\rangle (M_{\sigma}^{p, \tau} |\mu\rangle) |p, \tau, \sigma, d^{p, \tau}, \rightarrow\rangle \right)^{\dagger} \\ &\quad \left(\sum_{\sigma' \in \Sigma_C} |q^{p', \tau'}, \sigma'\rangle (M_{\sigma'}^{p', \tau'} |\mu'\rangle) |p', \tau', \sigma', d^{p', \tau'}, \rightarrow\rangle \right) \end{aligned}$$

Si $(p, \tau) \neq (p', \tau')$, alors $\delta'((p, \tau, \mu)^{\dagger} \delta'(p', \tau', \mu')) = 0$. Sinon $\mu \neq \mu'$, d'où :

$$\begin{aligned} \delta'(p, \tau, \mu)^{\dagger} \delta'(p, \tau, \mu') &= \sum_{\sigma \in \Sigma_C} \langle \mu | M_{\sigma}^{p, \tau \dagger} M_{\sigma}^{p, \tau} | \mu' \rangle \\ &= \langle \mu | \mu' \rangle \\ &= 0 \end{aligned}$$

– **Séparabilité :**

La séparabilité est satisfaite si pour tout $(p, \tau, \mu, \rho), (p', \tau', \mu', \rho') \in K \times \Sigma_C \times (\Sigma_Q^{(k)})^2$, et pour tout $\epsilon, \epsilon' \in \{0, \pm 1, \pm 2\}^k$ distincts, $\delta'[p, \tau, \mu, \rho | \epsilon]^{\dagger} \delta'[p', \tau', \mu', \rho' | \epsilon'] = 0$, avec

$$\delta'[p, \tau, \mu, \rho | \epsilon] = \sum_{\sigma \in \Sigma_C} \langle \rho | M_{\sigma}^{p, \tau} | \mu \rangle |q^{p, \tau}, \sigma, \rho, p, \tau, \sigma, d^{p, \tau}, \rightarrow\rangle |h_{(d^{p, \tau}, \rightarrow), \epsilon}\rangle$$

où $h_{d, \epsilon} = 2d - \epsilon$ if $\epsilon \neq 0$ et $h_{d, \epsilon} = \#$ sinon. On remarque que si $\epsilon \neq \epsilon'$ alors $h_{d, \epsilon} \neq h_{d, \epsilon'}$.

$$\begin{aligned} \delta'[p, \tau, \mu, \rho | \epsilon]^{\dagger} \delta'[p', \tau', \mu', \rho' | \epsilon'] &= \\ &\quad \left(\sum_{\sigma \in \Sigma_C} \langle \rho | M_{\sigma}^{p, \tau} | \mu \rangle |q^{p, \tau}, \sigma, \rho, p, \tau, \sigma, d^{p, \tau}, \rightarrow\rangle |h_{(d^{p, \tau}, \rightarrow), \epsilon}\rangle \right)^{\dagger} \end{aligned}$$

$$\left(\sum_{\sigma' \in \Sigma_C} \langle \rho' | M_{\sigma'}^{p', \tau'} | \mu' \rangle | (q^{p', \tau'}, \sigma'), \rho', p', \tau', \sigma', d^{p', \tau'}, \rightarrow \rangle \left| h_{(d^{p', \tau'}, \rightarrow), \epsilon'} \right\rangle \right)$$

Si $(p, \tau) \neq (p', \tau')$ alors la séparabilité est satisfaite. Sinon :

$$\delta'[p, \tau, \mu, \rho | \epsilon]^\dagger \delta'[p, \tau, \mu', \rho' | \epsilon'] = \sum_{\sigma \in \Sigma_C} \langle \mu | (M_{\sigma}^{p, \tau})^\dagger | \rho \rangle \langle \rho' | M_{\sigma}^{p, \tau} | \mu' \rangle \langle h_{(d^{p, \tau}, \rightarrow), \epsilon} | h_{(d^{p, \tau}, \rightarrow), \epsilon'} \rangle$$

Puisque $\epsilon \neq \epsilon'$, $\langle h_{(d^{p, \tau}, \rightarrow), \epsilon} | h_{(d^{p, \tau}, \rightarrow), \epsilon'} \rangle = 0$, donc la séparabilité est satisfaite.

L'unité, l'orthogonalité et la séparabilité sont satisfaites sur le domaine de définition de δ' donc, d'après le lemme de complétion, δ' peut être complétée afin de faire de M' une MTQ bien formée. $\square$

6.5.3 Circuits quantiques et MTQ

Nous évoquons très brièvement quelques résultats existants sur les simulations entre machine de Turing quantiques et circuits quantiques. Tout d'abord, les théorèmes 6.3 et 6.5 montrent que toute famille semi-uniforme de circuits quantiques peut être simulée efficacement par une MTQ. De même, le théorème 6.4 permet de conclure que toute famille uniforme de circuits quantiques peut être approximée efficacement par une MTQ.

Ces résultats ont déjà été prouvés, en utilisant une preuve directe, par Nishimura et Ozawa [NO02b].

Nishimura et Ozawa prouvent également la simulation inverse, c'est-à-dire que toute machine de Turing quantique peut être simulée efficacement par une famille semi-uniforme de circuits quantiques [NO02b]. Cette simulation complète les résultats précédents et permet de conclure que les trois modèles (MTQ, MTQC et circuits quantiques) sont équivalents.

6.6 MTQC à 1 ruban et MTQC à 2 rubans

Pour résumer, deux rubans sont suffisants pour le calcul quantique (théorème 6.3), alors qu'un ruban est suffisant pour le calcul classique (théorème 6.1) mais pas pour le calcul quantique (lemme 6.1). Ainsi, une séparation apparaît entre classique et quantique. Notons que ce résultat ne contredit pas l'équivalence, en terme de décidabilité, entre le classique et le quantique : la séparation apparaît si et seulement si on considère les données quantiques.

On peut se demander pourquoi les MTQC à 1 ruban ne sont pas quantiquement universelles alors que Briegel et Raussendorf ont prouvé, dans le cadre du calcul quantique par consommation d'intrication, que les mesures sur 1 qubit sont

universelles [RB02a]. La preuve de Briegel et Raussendorf est donnée avec une hypothèse forte qui est qu'il existe une grille de qubits auxiliaires qui a été préparée initialement, par un certain procédé externe non spécifié, dans un état globalement intriqué appelé état graphe (voir chapitre 9). Cette hypothèse n'étant pas vérifiée avec une MTQC à 1 ruban, les résultats précédents ne sont pas contradictoires avec les résultats de Briegel et Raussendorf.

6.7 Conclusion

Ce chapitre introduit les machines de Turing quantiques contrôlées classiquement (MTQC), un nouveau modèle pour le calcul quantique. Ce modèle permet une formalisation rigoureuse des interactions inhérentes entre le classique et le quantique pendant un calcul quantique. L'étude de ce modèle montre que toute machine de Turing classique est simulée par une MTQC sans perte d'efficacité, de plus, toute MTQC à k rubans est simulée par une MTQC à 2 rubans avec un ralentissement polynomial.

Nous avons également comparé le pouvoir de ce nouveau modèle de calcul quantique à celui des modèles de calcul quantique existants : le modèle des circuits quantiques et celui des machines de Turing quantiques. Le modèle des MTQC à 2 rubans permet de simuler efficacement le modèle des circuits quantiques et est simulé efficacement par les machines de Turing quantiques. Ces résultats font de la MTQC un modèle réaliste de calcul quantique, ils permettent également de prouver que l'utilisation du contrôle classique ne permet pas d'augmenter le pouvoir de calcul des modèles de calcul quantique réversibles, en leur ajoutant un contrôle classique.

De plus, une séparation entre le calcul classique et le calcul quantique est mise en évidence. En effet, l'utilisation des modèles de calcul quantique contrôlé classiquement a aussi un intérêt théorique, en permettant d'identifier une séparation entre calcul classique et calcul quantique : les MTQC à 1 ruban suffisent à la simulation efficace des modèles de calcul classiques mais pas à celle des modèles quantiques.

L'équivalence entre les modèles contrôlés classiquement et les modèles de calcul quantique réversibles montre que les puissances de calcul des deux modèles sont équivalentes, ce qui élimine par exemple la perspective de trouver un algorithme quantique contrôlé classiquement qui serait plus efficace qu'un algorithme quantique réversible.

L'introduction de modèles formels pour le calcul quantique contrôlé classiquement est, en revanche, fortement motivée par l'avènement de nouveaux modèles de calcul quantique, comme le calcul par consommation d'intrication ou par mesures projectives. Ces modèles ouvrent de nouvelles possibilités de réalisations physiques

d'un ordinateur quantique. Le développement de ces modèles nécessite des outils formels permettant une meilleure compréhension des ressources utilisées, comme les états graphes (voir chapitre 9), ou les mesures projectives. La mise en place de tels outils formels peuvent permettre :

- de minimiser les ressources nécessaires : par exemple en terme de nombre de mesures projectives et de qubits auxiliaires dans le cadre du calcul par mesures projectives ;
- de décider si un état graphe peut être utilisé comme support pour un calcul quantique par consommation d'intrication (voir les conditions de *flot*, chapitre 9).

Les MTQC et le $\mathbf{q}$ -calcul introduits dans cette partie sont des modèles formels pour le calcul quantique contrôlé classiquement, dont le calcul quantique réversible s'avère être une sous famille. Nous allons étudier dans les parties suivantes deux autres sous familles du calcul quantique contrôlé classiquement : le calcul par mesures projectives (partie III) et le calcul par consommation d'intrication (partie IV).

Troisième partie

Calcul par mesures projectives

Chapitre 7

Vers les ressources minimales du calcul par mesures projectives

7.1 Introduction

Dans les modèles de calcul réversible, comme le modèle des circuits quantiques, ou celui des machines de Turing quantiques, seul un fragment des évolutions quantiques est utilisé, celui des transformations unitaires. En effet, le théorème 6.5 montre que le fragment des transformations unitaires suffit au calcul quantique. Nielsen [Nie03] a prouvé qu'un autre fragment, celui des mesures projectives, associé à un contrôle classique récursif, suffit également au calcul quantique. Nous présentons dans ce chapitre le modèle introduit par Nielsen. Le caractère universel de ce modèle de calcul quantique par mesures projectives est prouvé par la simulation de toute transformation unitaire. Suite à l'introduction de ce modèle, plusieurs améliorations se sont succédées afin de diminuer les ressources nécessaires à ce type de calcul. Nous introduisons dans ce but un modèle à base de transfert d'état, où les ressources minimales sont en partie atteintes.

Les mesures projectives sont donc l'ingrédient de base du calcul introduit par Nielsen. Une mesure projective F est une transformation admissible, donc une famille d'opérateurs linéaires $(P_i)_{i \in A}$, telle que ces opérateurs sont des projecteurs orthogonaux, i.e. $\forall i, j, P_i P_j = \delta_{i,j} P_i$. Une mesure projective peut également être décrite par un *observable*. Un observable est une matrice $\mathcal{O}$ hermitienne ($\mathcal{O}^\dagger = \mathcal{O}$). Une telle matrice admet une décomposition spectrale $\mathcal{O} = \sum_{i \in A} \lambda_i P_i$, où les P_i sont des projecteurs. Un observable $\mathcal{O} = \sum_{i \in A} \lambda_i P_i$ décrit la mesure projective $(P_i)_{i \in A}$.

Ainsi, une mesure selon un observable $\mathcal{O} = \sum_i \lambda_i P_i$ d'un état quantique $|\Phi\rangle$ produit avec probabilité $p_i = \langle \Phi | P_i | \Phi \rangle$ le résultat classique λ_i . L'état quantique du système devient alors $\frac{1}{\sqrt{p_i}} P_i |\Phi\rangle$.

Des exemples d'observables sont les matrices de Pauli X, Y et Z :

$$X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

Les compositions tensorielles de matrices de Pauli sur plusieurs qubits sont également des observables. On utilisera des lettres capitales pour dénoter un opérateur de Pauli quand celui-ci est utilisé comme observable, on préférera la notation σ_x, σ_y ou σ_z quand celui-ci est utilisé comme transformation unitaire.

Comme pour le modèle des circuits quantiques, nous ne travaillerons dans ce chapitre qu'avec des registres de qubits, les espaces de Hilbert étant alors de la forme $\mathcal{H}_{\{0,1\}^n}$, où n est la taille du registre. Dans le chapitre suivant, nous verrons que le formalisme du $\mathbf{q}$ -calcul peut être utilisé pour représenter le calcul par mesures projectives dans un cadre plus général que celui des qubits.

7.2 Calcul par mesures projectives à base de téléportation

La démarche suivie par Nielsen dans [Nie03] consiste dans un premier temps à remarquer que le protocole de téléportation [BBC⁺93] qui est habituellement réalisé à l'aide des transformations unitaires H et ΛX peut être réalisé en utilisant uniquement des mesures projectives sur deux qubits, appelées mesures de Bell.

Une mesure de Bell est une mesure projective $(|B_{ij}\rangle\langle B_{ij}|)_{i,j \in \{0,1\}}$, où les $|B_{ij}\rangle$ sont les états de Bell :

$$\begin{aligned} |B_{00}\rangle &= \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle) \\ |B_{01}\rangle &= \frac{1}{\sqrt{2}}(|00\rangle - |11\rangle) \\ |B_{10}\rangle &= \frac{1}{\sqrt{2}}(|01\rangle + |10\rangle) \\ |B_{11}\rangle &= \frac{1}{\sqrt{2}}(|01\rangle - |10\rangle) \end{aligned}$$

On remarque que $|B_{ij}\rangle = \sigma_z^j \otimes \sigma_x^i |B_{00}\rangle$.

La téléportation de l'état $|\Phi\rangle$ d'un qubit a peut être obtenue en appliquant une mesure de Bell sur deux qubits auxiliaires b et c puis une mesure de Bell sur les qubits a et b . L'état $|\Phi\rangle$ se trouve alors téléporté, à un opérateur de Pauli près sur le qubit c , i.e. l'état de c est $\sigma |\Phi\rangle$, avec $\sigma \in \{Id, \sigma_x, \sigma_y, \sigma_z\}$ (figure 7.1).

En effet si (i, j) est le résultat de la première mesure, et (k, l) celui de la seconde, alors le qubit c est dans l'état $\sigma_x^{j+k} \otimes \sigma_z^{i+l} |\Phi\rangle$.

Si une transformation unitaire U est appliquée juste avant la téléportation, alors l'état du qubit c après la téléportation est $U |\Phi\rangle$ à un opérateur de Pauli près (figure 7.2, à gauche). D'après la loi de composition des transformations admissibles

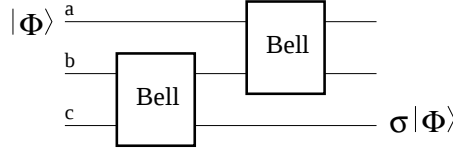
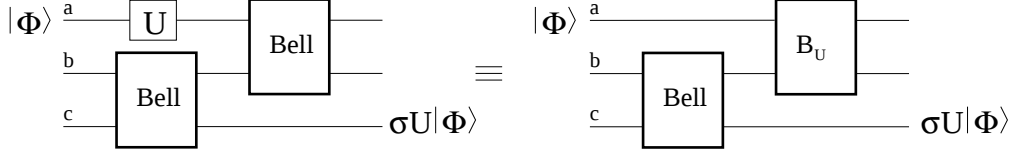


FIG. 7.1 – Téléportation à base de mesures

FIG. 7.2 – A gauche : Téléportation de $U|\Phi\rangle$; A droite : un pas de simulation de U

(définition 1.6, chapitre 1) :

$$\begin{aligned} (|B_{ij}\rangle\langle B_{ij}|)_{i,j\in\{0,1\}} \circ (U \otimes Id) &= (|B_{ij}\rangle\langle B_{ij}| (U \otimes Id))_{i,j\in\{0,1\}} \\ &= (U \otimes Id) \circ ((U^\dagger \otimes Id) |B_{ij}\rangle\langle B_{ij}| (U \otimes Id))_{i,j\in\{0,1\}} \end{aligned}$$

On remarque que pour tout U , $B_U = ((U^\dagger \otimes Id) |B_{ij}\rangle\langle B_{ij}| (U \otimes Id))_{i,j\in\{0,1\}}$ est une mesure projective. Ainsi, plutôt que d'appliquer U sur le qubit a puis d'appliquer une mesure de Bell sur a, b , on peut, de façon équivalente, appliquer une mesure selon B_U sur a, b puis appliquer U sur a . Or seul l'état du qubit c est important après la téléportation, donc l'application de U sur a après la mesure selon B_U est facultative.

Ainsi un schéma permettant de simuler, à un opérateur de Pauli près, toute transformation unitaire U sur un qubit est obtenu (figure 7.2, à droite).

Ce schéma à base de téléportation constitue un premier pas dans la simulation d'une transformation unitaire U sur un qubit. En effet, la seconde étape consiste à faire disparaître l'opérateur de Pauli dépendant des résultats des mesures.

Une simulation complète est obtenue en composant de façon conditionnelle un pas de simulation de U avec un pas de simulation de chacun des opérateurs de Pauli. Afin de décrire la stratégie utilisée pour obtenir une simulation complète de U , une étape de simulation est représentée par une boîte noire dont l'entrée est un état quantique $|\Phi\rangle$ et dont les sorties sont les quatre résultats possibles : $U|\Phi\rangle, \sigma_x U|\Phi\rangle, \sigma_y U|\Phi\rangle, \sigma_z U|\Phi\rangle$ (figure 7.3).

La stratégie est la suivante : après un premier pas de simulation de U sur $|\Phi\rangle$, le résultat $\sigma U|\Phi\rangle$ est obtenu. Si $\sigma = Id$ alors la simulation est terminée, sinon l'opérateur de Pauli σ est lui-même simulé produisant l'état $\sigma'\sigma U|\Phi\rangle = \sigma'U|\Phi\rangle$, où σ' est un opérateur de Pauli, etc.

Cette stratégie peut être décrite à l'aide d'un automate (figure 7.3).

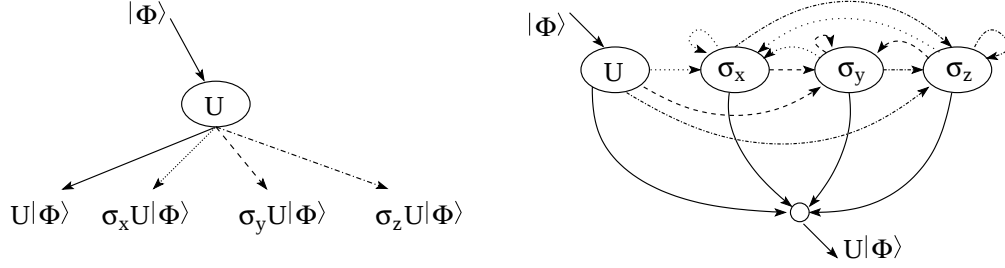


FIG. 7.3 – A gauche : Boîte noire pour un pas de simulation ; A droite : Simulation complète de U .

Un schéma similaire permet de simuler des transformations unitaires sur deux qubits et en particulier ΛX [Nie03].

Remarque 7.1 *La simulation des transformations unitaires sur un qubit présentée ici est adaptée de celle proposée par Nielsen [Nie03]. Dans son article original, Nielsen propose d'appliquer une transformation unitaire U après la téléportation et sur le qubit c . Cette transformation unitaire peut être intégrée à la première mesure de Bell sur les qubits b, c . Bien que semblable à celui décrit en figure 7.2, ce schéma originel possède une stratégie de correction plus complexe.*

Tout circuit quantique ayant pour base ΛX et des transformations unitaires sur un qubit peut être simulé en n'utilisant que des mesures projectives. Afin de déterminer le coût de cette simulation, remarquons premièrement que les résultats classiques des différentes mesures effectuées sont indépendants de U et de $|\Phi\rangle$. Plus précisément, dans la stratégie décrite en figure 7.3, chaque transition est effectuée avec probabilité $1/4$. On en déduit que, pour tout $\epsilon > 0$, $O(k \cdot \log(k/\epsilon))$ mesures projectives suffisent à simuler un circuit composés de k portes avec une probabilité d'échec inférieur à ϵ .

7.2.1 Ressources

Une des ressources du calcul réversible est l'ensemble des transformations unitaires de base qui permettent de décomposer, de façon exacte ou approchée, toute transformation unitaire. Dans le cadre du calcul par mesure projective, l'ensemble des actions de base est composé de mesures projectives.

Définition 7.1 (Universalité exacte) *Un ensemble E de mesures projectives est universel si et seulement si pour toute transformation unitaire U et tout $\epsilon > 0$, il existe une stratégie permettant de simuler U avec une probabilité d'échec inférieure à ϵ .*

Remarque 7.2 *La notion de stratégie est formalisée dans le chapitre 8.*

Même si la simulation est probabiliste, car la probabilité d'échec est a priori non nulle, l'universalité est qualifiée d'exacte. L'universalité approchée est définie comme suit :

Définition 7.2 (Universalité approchée) *Un ensemble E de mesures projectives est approximativement universel si et seulement si pour toute transformation unitaire U , tout $\epsilon > 0$ et tout $\delta > 0$, il existe une stratégie permettant de simuler $\tilde{U}$ avec une probabilité d'échec inférieure à ϵ , avec $\|U - \tilde{U}\| \leq \delta$.*

Ainsi le modèle introduit par Nielsen prouve que l'ensemble de mesures projectives $\{B_U \mid U \in \mathcal{U}_1\} \cup \{B_{\Lambda X}\}$ est universel, où $\mathcal{U}_1$ est l'ensemble des transformations unitaires sur 1 qubit et $B_{\Lambda X}$ est une mesure sur 4 qubits permettant la simulation de ΛX .

Les observables constituent donc naturellement une ressource du calcul par mesures projectives. Les qubits auxiliaires en sont une autre ressource. En effet, le schéma de simulation proposé par Nielsen nécessite des qubits auxiliaires : 2 pour la simulation d'une transformation unitaire d'ordre 1 et 4 pour la simulation de ΛX . Ainsi la simulation d'un circuit composé de portes d'ordre 1 et 2, utilise 4 qubits auxiliaires.

Donc 4 qubits auxiliaires associés à l'ensemble $\{B_U \mid U \in \mathcal{U}_1\} \cup \{B_{\Lambda X}\}$ de mesures projectives constituent des ressources suffisantes au calcul quantique.

La minimisation de ces ressources est un point essentiel, notamment afin de rendre réalisable l'implémentation physique de ce modèle de calcul quantique. La minimisation a été entreprise par Fenner et Zhang [FZ01] puis par Leung [Leu04]. Fenner et Zhang ont montré qu'un certain ensemble formé uniquement de mesures sur deux et trois qubits, associé à 4 qubits auxiliaires, est universel. Leung a démontré à son tour que des mesures sur 2 qubits associées à 4 qubits auxiliaires sont universelles. Citons le résultat de Leung dans le cas de l'universalité approchée :

Théorème 7.1 ([Leu04]) *L'ensemble de 5 observables $\{Z, X \otimes X, Z \otimes Z, X \otimes Z, \frac{1}{\sqrt{2}}(X - Y) \otimes X\}$ est approximativement universel en utilisant 4 qubits auxiliaires.*

Afin de réduire les ressources nécessaires au calcul par mesures projectives, nous introduisons un nouveau schéma de simulation, non plus fondé sur la téléportation, mais sur le *transfert d'état*.

7.3 Transfert d'état

Alors que le schéma de simulation introduit par Nielsen puis amélioré par Leung est fondé sur la téléportation, nous introduisons un nouveau schéma de calcul par

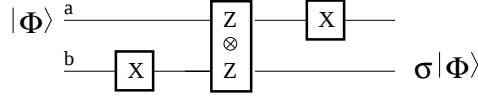


FIG. 7.4 – Transfert d'état

mesure projective : le *transfert d'état*. Le transfert d'état est une alternative à la téléportation dans le cadre du calcul par mesures projectives permettant de diminuer les ressources utilisées. En revanche, le transfert d'état ne peut pas remplacer la téléportation dans des traitements en situation de non localité.

Lemme 7.1 (Transfert d'état) *Etant donné un qubit a et un qubit auxiliaire b , la suite de mesures selon X_b , puis $Z_a \otimes Z_b$ et enfin X_a (voir figure 7.4) transfère l'état $|\Phi\rangle$ du qubit a au qubit b , à un opérateur de Pauli près.*

Preuve : Soit $|\Phi\rangle = \alpha|0\rangle + \beta|1\rangle$ l'état du qubit a . Après la première mesure, le qubit b est dans l'état $\frac{|0\rangle+|1\rangle}{\sqrt{2}}$ si le résultat de la mesure est $j = 1$, ou $\frac{|0\rangle-|1\rangle}{\sqrt{2}}$ si le résultat de la mesure est $j = -1$. Donc l'état du registre a, b après la première mesure est

$$|\psi_1\rangle = Id \otimes \sigma_z^{(1-j)/2} (\alpha|0\rangle + \beta|1\rangle) \otimes \left(\frac{|0\rangle+|1\rangle}{\sqrt{2}}\right)$$

$Z \otimes Z = P_1 - P_{-1}$ est la décomposition spectrale de $Z \otimes Z$ avec $P_k = \frac{1}{2}(Id + kZ \otimes Z)$. On remarque que $\langle\psi_1| Z \otimes Z |\psi_1\rangle = 0$ donc pour $k = -1$ et $k = 1$, $\langle\psi_1| P_k |\psi_1\rangle = 1/2$. Ainsi l'état $|\psi_2\rangle$ du registre a, b après la deuxième mesure, est

$$\begin{aligned} |\psi_2\rangle &= \sqrt{2}P_k |\psi_1\rangle \\ &= (Id \otimes \sigma_z^{(1-j)/2} \cdot \sigma_x^{(1-k)/2}) (\alpha|00\rangle + \beta|11\rangle) \end{aligned}$$

où $k \in \{-1, 1\}$ est le résultat de la deuxième mesure.

Enfin, $X \otimes I = P'_1 - P'_{-1}$ avec $P'_l = \frac{1}{2}(Id + lX \otimes I)$. On remarque également que $\langle\psi_2| P'_l |\psi_2\rangle = 1/2$. Ainsi l'état du registre a, b après la dernière mesure, est

$$\begin{aligned} |\psi_2\rangle &= \sqrt{2}P'_l |\psi_2\rangle \\ &= (\sigma_z^{(1-l)/2} \otimes \sigma_z^{(1-j)/2} \cdot \sigma_x^{(1-k)/2} \cdot \sigma_z^{(1-l)/2}) \left(\frac{|0\rangle+|1\rangle}{\sqrt{2}}\right) (\alpha|0\rangle + \beta|1\rangle) \end{aligned}$$

Ainsi, l'état du qubit b après les trois mesures est $\sigma_z^{(1-j)/2} \cdot \sigma_x^{(1-k)/2} \cdot \sigma_z^{(1-l)/2} |\Phi\rangle$, où $j, k, l \in \{-1, 1\}$ sont les résultats classiques des mesures. $\square$

Le transfert d'état est au cœur du nouveau schéma que nous présentons. Ainsi, si une transformation unitaire U d'ordre 1 est appliquée avant le transfert d'un état $|\Phi\rangle$, et une transformation unitaire V est appliquée juste après le transfert d'état, alors le résultat obtenu est $V\sigma U |\Phi\rangle$ (figure 7.5).

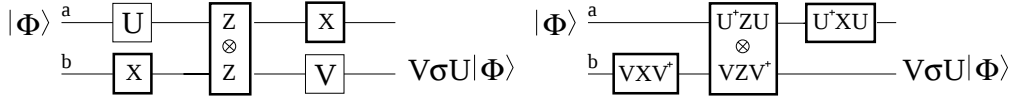
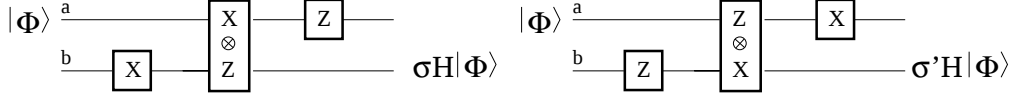


FIG. 7.5 – Transfert d'état généralisé

FIG. 7.6 – Pas de simulation de H : à gauche $U = H$, $V = Id$; à droite : $U = Id$ et $V = H$ (on note que pour tout σ , il existe σ' tel que $H\sigma = \sigma'H$)

Afin d'intégrer les transformations unitaires U et V aux mesures projectives du transfert d'état, pour obtenir un schéma n'utilisant que des mesures projectives, le lemme 7.2 est introduit. Ce lemme traite de la composition et la commutativité d'une transformation unitaire avec une mesure projective.

Lemme 7.2 *Pour toute transformation unitaire U et tout observable $\mathcal{O}$, appliquer U puis mesurer selon $\mathcal{O}$ est équivalent à mesurer selon $U^\dagger \mathcal{O} U$ puis appliquer U .*

Preuve : Si $\mathcal{O} = \sum_{i \in A} \lambda_i P_i$ est la décomposition spectrale de $\mathcal{O}$, alors la loi de composition des transformations admissibles implique :

$$(P_i)_{i \in A} \circ U = (P_i U)_{i \in A} = U \circ (U^\dagger P_i U)_{i \in A}$$

$$\text{Or } \mathcal{O}' = \sum_{i \in A} \lambda_i U^\dagger P_i U \text{ est un observable avec } \mathcal{O}' = U^\dagger \mathcal{O} U \quad \square$$

On déduit du lemme 7.2 que la suite de mesures décrite dans la figure 7.5 à droite, permet d'obtenir l'état $V\sigma U |\Phi\rangle$ à partir de l'état $|\Phi\rangle$ en n'appliquant que des mesures projectives. Une telle suite de mesures est appelée *transfert d'état généralisé*.

Ainsi, on peut dériver du transfert d'état généralisé des pas de simulation des transformations unitaires H , T , et $H.T$ (figures 7.6 et 7.7), en choisissant correctement les transformations U et V .

$$H = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}, T = \begin{pmatrix} 1 & 0 \\ 0 & e^{\frac{i\pi}{4}} \end{pmatrix}, \Lambda X = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

Alors que le schéma proposé par Nielsen puis amélioré par Leung propose un pas de simulation de l'opérateur ΛX utilisant 4 qubits auxiliaires, nous introduisons un nouveau schéma qui n'utilise qu'un seul qubit auxiliaire :

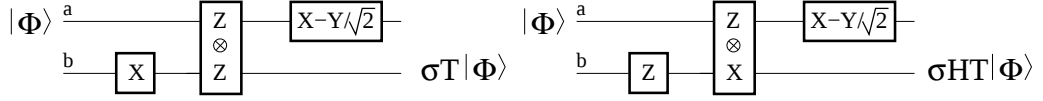


FIG. 7.7 – À gauche : pas de simulation de $T : U = T, V = Id$; à droite : pas de simulation de $HT : U = T$ et $V = H$.

Lemme 7.3 *Etant donné un registre de deux qubits a, b et un qubit auxiliaire c , la suite de mesures selon $Z_c, Z_a \otimes X_c, Z_c \otimes X_b$ et enfin X_c (voir figure 7.8) simule l'application de ΛX sur l'état $|\Phi\rangle$ des qubits a, b , à un opérateur de Pauli près.*

Preuve : Les notations $|+\rangle = \frac{|0\rangle+|1\rangle}{\sqrt{2}}$ et $|-\rangle = \frac{|0\rangle-|1\rangle}{\sqrt{2}}$ sont utilisées dans cette preuve.

Si l'état $|\Phi\rangle$ du registre a, b est $\alpha|00\rangle + \beta|01\rangle + \gamma|10\rangle + \delta|11\rangle$, alors l'état $|\psi_1\rangle$ du registre a, b, c après la première mesure est :

$$|\psi_1\rangle = (I \otimes I \otimes \sigma_x^{\frac{1-j}{2}})[(\alpha|000\rangle + \beta|010\rangle + \gamma|100\rangle + \delta|110\rangle)]$$

où $j \in \{-1, 1\}$ est le résultat classique de la mesure selon Z_c . On remarque que $\langle\psi_1|Z_a \otimes X_c|\psi_1\rangle = 0$. Si P_k pour $k \in \{-1, 1\}$ sont les projecteurs associés à cette mesure, alors l'état $|\psi_2\rangle$ du registre après la deuxième mesure est :

$$\begin{aligned} |\psi_2\rangle &= \sqrt{2}P_k|\psi_1\rangle \\ &= (\sigma_z^{\frac{1-j}{2}} \otimes I \otimes \sigma_z^{\frac{1-k}{2}})[\alpha|00+\rangle + \beta|01+\rangle + \gamma|10-\rangle + \delta|11-\rangle] \end{aligned}$$

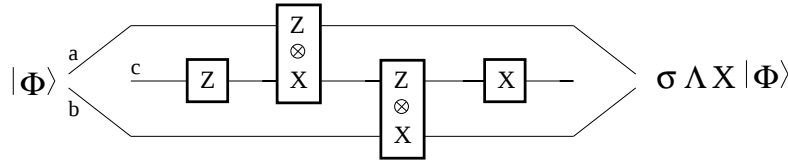
De même, $\langle\psi_2|Z_c \otimes X_b|\psi_2\rangle = 0$, donc si P'_l , pour $l \in \{-1, 1\}$, sont les projecteurs associés à cette troisième mesure, alors :

$$\begin{aligned} |\psi_3\rangle &= \sqrt{2}P'_l|\psi_2\rangle \\ &= \frac{1}{\sqrt{2}}(\sigma_z^{\frac{1-j}{2}} \otimes \sigma_x^{\frac{1-j,l}{2}} \otimes \sigma_x^{\frac{1-k}{2}} \otimes \sigma_x^{\frac{1-l}{2}})[\alpha(|00+\rangle + |01-\rangle) + \beta(|01+\rangle + |00-\rangle) \\ &\quad + \gamma(|10-\rangle + |11+\rangle) + \delta(|11-\rangle + |10+\rangle)] \end{aligned}$$

Enfin, la dernière mesure selon X_c vérifie $\langle\psi_3|X_c|\psi_3\rangle = 0$, donc si $m \in \{-1, 1\}$ est le résultat classique de cette dernière mesure, l'état $|\psi_4\rangle$ après cette mesure est :

$$\begin{aligned} |\psi_4\rangle &= (\sigma_z^{\frac{1-j,l}{2}} \otimes \sigma_x^{\frac{1-k,m}{2}} \otimes \sigma_z^{\frac{1-m}{2}})[(\alpha|00\rangle + \beta|01\rangle + \gamma|11\rangle + \delta|10\rangle) \otimes |+\rangle] \\ &= (\sigma_z^{\frac{1-j,l}{2}} \otimes \sigma_x^{\frac{1-k,m}{2}})(\Lambda X|\Phi\rangle) \otimes (\sigma_z^{\frac{1-m}{2}}|+\rangle) \end{aligned}$$

□

FIG. 7.8 – Simulation de ΛX

7.3.1 Ressources

Dans cette section, de nouveaux ensembles universels de mesures projectives améliorant les résultats de Leung sont présentés. Les cas de l'universalité exacte et de l'universalité approchée sont pris en compte.

Théorème 7.2 *L'ensemble $E_0 = \{Z \otimes X, X, Z, \frac{X-Y}{\sqrt{2}}\}$ est approximativement universel, en utilisant un seul qubit auxiliaire.*

Preuve : La base $\{H, T, \Lambda X\}$ est approximativement universelle [NC00], donc pour tout $\delta > 0$ et toute transformation unitaire U , il existe un circuit C ayant pour base $\{H, T, \Lambda X\}$ tel que l'action $\tilde{U}$ de C vérifie $\|U - \tilde{U}\| < \delta$. Soit k la taille de C . Chacune des portes H , T et ΛX peut être simulée en utilisant les transferts d'état généralisés décrits dans les figures 7.6, 7.7 et 7.8.

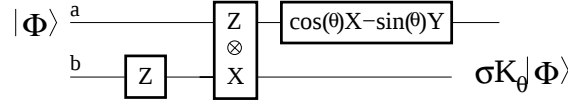
Afin d'obtenir une simulation complète de chacune des portes, la stratégie décrite dans la figure 7.3 peut être appliquée. Cette stratégie requiert l'utilisation d'un pas de simulation de σ_x , σ_y et σ_z , or l'application directe du schéma de transfert d'état généralisé pour les opérateurs de Pauli implique l'utilisation de mesures selon l'observable $Z \otimes Z$, non présent dans l'ensemble E_0 des observables disponibles.

En revanche, l'application successive de deux pas de simulation de H permet d'obtenir la simulation de l'identité à un opérateur de Pauli près. En effet, pour tous les opérateurs de Pauli $\sigma, \sigma', \sigma' H \sigma H = \sigma'' Id$ où σ'' est un opérateur de Pauli. Cette simulation peut être utilisée pour simuler chacun des opérateurs de Pauli, en effet une simulation de l'identité à σ près, peut être interprétée comme une simulation de $\tilde{\sigma}$ à $\sigma \tilde{\sigma}$ près¹, où $\tilde{\sigma}$ est l'opérateur de Pauli à simuler.

Ainsi, pour tout $\epsilon > 0$, chaque porte peut être simulée avec une probabilité d'échec inférieure à ϵ/k . La probabilité de succès de la simulation totale du circuit est donc supérieure à $(1 - \epsilon/k)^k \geq 1 - \epsilon$, garantissant que la probabilité d'échec de $\tilde{U}$ soit inférieure à ϵ . $\square$

Théorème 7.3 *L'ensemble $E_1 = \{Z \otimes X, Z, \cos(\theta)X + \sin(\theta)Y, \theta \in [0, 2\pi]\}$ est universel, en utilisant un seul qubit auxiliaire.*

¹Les opérateurs de Pauli sont auto-adjoints donc $\sigma \tilde{\sigma} \sigma = \sigma Id$

FIG. 7.9 – Pas de simulation de $HR_{\hat{z}}(\theta)$: $V_1 = R_{\hat{z}}(\theta)$ et $V_2 = H$.)

Preuve : Etant donné un vecteur $\mathbf{n} = (a, b, c)$, une rotation d'angle α autour de $\mathbf{n}$ est définie par $R_{\mathbf{n}}(\alpha) = \cos(\alpha/2)I - i \sin(\alpha/2)(a\sigma_x + b\sigma_y + c\sigma_z)$.

Pour toute transformation unitaire U d'ordre 1, il existe $\alpha, \beta, \gamma, \delta$ tels que : $U = e^{i\alpha}R_{\hat{x}}(\beta)R_{\hat{z}}(\gamma)R_{\hat{x}}(\delta)$, avec $\hat{x} = (1, 0, 0)$ et $\hat{z} = (0, 0, 1)$ (voir [NC00]). Ainsi toute transformation unitaire U d'ordre 1 peut être décomposée en trois rotations élémentaires, à une phase globale près.

De plus, pour tout θ , $R_{\hat{x}}(\theta) = HR_{\hat{z}}(\theta)H$ donc $U = e^{i\alpha}HR_{\hat{z}}(\beta)HR_{\hat{z}}(\gamma)HR_{\hat{z}}(\delta)H$ d'où $U = e^{i\alpha}K_{\beta}K_{\gamma}K_{\delta}K_0$, avec $K_{\theta} = HR_{\hat{z}}(\theta)$. La phase globale $e^{i\alpha}$ n'étant pas importante dans la simulation de U (voir [NC00] pour une discussion sur la phase globale), on en déduit que la base $\{\Lambda X, K_{\theta}, \theta \in [0, 2\pi]\}$ est universelle. Une preuve de l'universalité de cette base a été donnée indépendamment par [DKP04a].

De plus, pour tout θ , la transformation $HR_{\hat{z}}(\theta)$ peut être simulée à l'aide du schéma de transfert d'état généralisé présenté dans la figure 7.9. On en déduit l'universalité exacte de l'ensemble $E_1 = \{Z \otimes X, Z, \cos(\theta)X + \sin(\theta)Y, \theta \in [0, 2\pi]\}$. $\square$

Nous avons introduit un nouveau schéma, le transfert d'état, permettant de diminuer les ressources utilisées pour le calcul par mesures projectives. Ces ressources sont-elles minimales ? Tout d'abord, en terme de qubits auxiliaires, la simulation de transformations unitaires, donc réversibles, par des opérations non réversibles comme les mesures projectives, nécessite un espace auxiliaire supplémentaire jouant un rôle de mémoire afin de rendre globalement l'évolution réversible. Ainsi, un minimum de un qubit auxiliaire est nécessaire à la simulation d'une transformation unitaire par des mesures projectives.

Les théorèmes 7.2 et 7.3 montrent qu'aussi bien dans le cadre de l'universalité exacte qu'approchée, un seul qubit auxiliaire est suffisant au calcul par mesure projective. Les ressources minimales, en terme de qubits auxiliaires, sont donc atteintes.

En terme d'ensemble d'observables, remarquons premièrement que la création d'intrication nécessite des mesures projectives multi-qubits, or un seul observable multiqubit, $Z \otimes X$ est nécessaire à l'universalité exacte et approchée. Ainsi, en termes d'observables multi-qubits, les ressources minimales sont également atteintes.

Enfin, pour les observables sur 1 qubit, les cas des universalités exacte et approchée sont à distinguer.

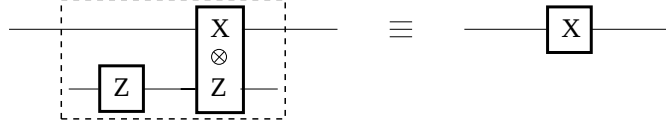


FIG. 7.10 – Simulation d’une mesure selon X , par une mesure selon Z puis selon $Z \otimes X$

- Pour l’universalité exacte, un simple argument de dénombrement montre qu’un ensemble universel de façon exacte doit nécessairement être non dénombrable. Ainsi, l’ensemble E_1 est satisfaisant, même s’il n’est pas minimal. En effet, pour tout $\gamma > 0$, l’ensemble $E_\gamma = \{Z \otimes X, Z, \cos(\theta)X + \sin(\theta)Y, \theta \in [0, \gamma]\}$ est universel et $E_\gamma \subset E_1$.
- Pour l’universalité approchée, aucun argument ne nous permet actuellement de conclure sur l’aspect nécessaire des observables de E_0 , au contraire, dans la section suivante, un ensemble d’observables strictement inclus dans E_0 est montré approximativement universel.

7.4 Compromis entre observables et qubits auxiliaires ?

L’ensemble approximativement universel E_0 inclut les observables X, Z et $Z \otimes X$, or le lemme suivant, illustré par la figure 7.10, montre qu’une mesure selon X peut être simulée en utilisant les observables Z et $Z \otimes X$:

Lemme 7.4 *Etant donné un registre a et un qubit auxiliaire b , une mesure de selon X du qubit a est équivalente à une mesure selon Z de b suivie d’une mesure selon $X \otimes Z$ du registre a, b .*

Preuve : Soit $|\Phi\rangle$ l’état du qubit a .

- Si a est mesuré selon X , la probabilité que le résultat classique soit $j \in \{-1, 1\}$ est $p(j) = \langle \Phi | (Id + jX) | \Phi \rangle / 2$, l’état de a après la mesure est $|\Phi'(j)\rangle = \frac{1}{2\sqrt{p(j)}}((Id + jX) | \Phi \rangle)$
- Si b est dans l’état $\alpha |0\rangle + \beta |1\rangle$ avant la mesure selon Z , la probabilité que le résultat classique de la mesure de b selon Z soit k est $p'(k)$, avec $p'(-1) = |\beta|^2$ et $p'(1) = |\alpha|^2$. Après cette mesure le registre a, b est dans l’état $|\psi(k)\rangle = |\Phi\rangle \otimes (\sigma_x^{(1-k)/2} |0\rangle)$. Après la mesure du registre a, b selon $X \otimes Z$, la probabilité que le résultat classique soit $l \in \{-1, 1\}$ est : $p''(k, l) = \langle \psi(k) | (Id + lX \otimes Z) | \psi(k) \rangle / 2 = p(k.l)$. Après cette mesure, l’état du registre est $|\psi'(k, l)\rangle = |\Phi'(k.l)\rangle \otimes \sigma_x^{(1-k)/2} |0\rangle$.

Ainsi une mesure selon X peut être simulée par deux mesures, selon Z puis selon $X \otimes Z$ (figure 7.10) : la probabilité que la mesure X produise le résultat j est la probabilité que les deux mesures produisent des résultats k et l tels que $k.l = j$. $\square$

Ainsi, une mesure selon X peut être simulée par une mesure selon X et une mesure selon $X \otimes Z$, en utilisant un qubit auxiliaire. On en déduit l'ensemble approximativement universel suivant :

Lemme 7.5 *L'ensemble $E_2 = \{Z \otimes X, Z, \frac{X-Y}{\sqrt{2}}\}$ est approximativement universel, en utilisant deux qubits auxiliaires.*

Ainsi un compromis entre qubits auxiliaires et observables apparaît : afin de diminuer le nombre d'observables, un qubit auxiliaire supplémentaire est utilisé. D'un autre côté, le lemme 7.5 ne montre pas que pour que l'ensemble E_2 soit universel, deux qubits auxiliaires sont nécessaires.

7.5 Vers les ressources minimales

L'ensemble $E_2 = \{Z \otimes X, Z, \frac{X-Y}{\sqrt{2}}\}$ est approximativement universel en utilisant 2 qubits. La preuve d'universalité de E_2 repose sur l'universalité de $E_3 = E_2 \cup \{X\}$ et sur la capacité de simuler la mesure selon X en utilisant des mesures de l'ensemble E_2 . Une approche différente, qui est adoptée dans cette section, permet de montrer que E_2 est approximativement universel en utilisant un seul qubit auxiliaire.

Alors que l'universalité de l'ensemble E_3 repose sur l'universalité de la base $B = \{H, T, \Lambda X\}$, nous introduisons une base différente de transformations unitaires et nous prouvons son universalité. Dans un second temps, nous montrerons comment simuler les éléments de cette nouvelle base à l'aide des observables de E_2 .

La preuve de l'universalité de B est adaptée de la preuve de d'universalité de $\{H, T, \Lambda X\}$. La démonstration s'appuie sur les deux propriétés suivantes (voir [NC00]).

Propriété 7.1 *Etant donné un vecteur $\mathbf{n}$, si θ est un multiple irrationnel de π , alors pour tout α et tout $\epsilon > 0$, il existe k tel que*

$$||R_{\mathbf{n}}(\alpha) - R_{\mathbf{n}}(\theta)^k|| < \epsilon/3$$

Propriété 7.2 (Włodarski [Wlo69]) *Si α n'est pas un multiple entier de $\pi/4$ et $\cos \beta = \cos^2 \alpha$, alors soit α , soit β est un multiple irrationnel de π .*

Lemme 7.6 *$\{HT, \sigma_y, \Lambda Z\}$ est une base approximativement universelle.*

Preuve : Tout d'abord, nous prouvons que toute transformation unitaire d'ordre 1 qubit peut être approchée par HT et σ_y . Considérons les transformations $U_1 = T$, $U_2 = HTH$ et $U_3 = \sigma_y HTH \sigma_y$. On note que T est, à une phase globale non pertinente près, une rotation d'angle $\pi/4$ autour de l'axe $\hat{z} = (0, 0, 1)$:

$$\begin{aligned} U_1 &= T &= e^{-i\pi/8}(\cos(\pi/8)I - i\sin(\pi/8)\sigma_z) \\ U_2 &= HTH &= e^{-i\pi/8}(\cos(\pi/8)I - i\sin(\pi/8)\sigma_x) \\ U_3 &= \sigma_y HTH \sigma_y &= e^{-i\pi/8}(\cos(\pi/8)I + i\sin(\pi/8)\sigma_x) \end{aligned}$$

En composant U_1 et U_2 on obtient, à une phase globale près :

$$\begin{aligned} U_2 U_1 &= (\cos(\pi/8)I - i\sin(\pi/8)\sigma_x)(\cos(\pi/8)I - i\sin(\pi/8)\sigma_z) \\ &= \cos^2(\pi/8)I - i[\cos(\pi/8)(\sigma_x + \sigma_z) - \sin(\pi/8)\sigma_y]\sin(\pi/8) \end{aligned}$$

$U_2 U_1$ est donc une rotation autour de l'axe $\mathbf{n} = (\cos(\pi/8), -\sin(\pi/8), \cos(\pi/8))$, d'angle θ défini comme la solution de $\cos(\theta/2) = \cos^2(\pi/8)$. Puisque $\pi/8$ n'est pas un multiple entier de $\pi/4$, mais un multiple rationnel de π , d'après la propriété 7.2, un tel θ est un multiple irrationnel de π .

Cette irrationnalité implique que pour tout angle α , la rotation d'angle α autour de $\mathbf{n}$ peut être approchée avec une précision arbitraire en répétant des rotations autour de $\mathbf{n}$ d'angle θ (voir propriété 7.2). Pour tout α et tout $\epsilon > 0$, il existe k tel que

$$\|R_{\mathbf{n}}(\alpha) - R_{\mathbf{n}}(\theta)^k\| < \epsilon/3$$

De plus, en composant U_1 et U_3 , on obtient, à une phase globale près :

$$\begin{aligned} U_3 U_1 &= (\cos(\pi/8)I + i\sin(\pi/8)\sigma_x)(\cos(\pi/8)I - i\sin(\pi/8)\sigma_z) \\ &= \cos^2(\pi/8)I - i[\cos(\pi/8)(-\sigma_x + \sigma_z) + \sin(\pi/8)\sigma_y]\sin(\pi/8) \end{aligned}$$

$U_3 U_1$ est une rotation autour de l'axe $\mathbf{m} = (-\cos(\pi/8), \sin(\pi/8), \cos(\pi/8))$ et d'angle θ .

Pour tout α et tout $\epsilon > 0$, il existe k tel que

$$\|R_{\mathbf{m}}(\alpha) - R_{\mathbf{m}}(\theta)^k\| < \epsilon/3$$

Puisque les axes $\mathbf{n}$ et $\mathbf{m}$ ne sont pas parallèles, toute transformation unitaire U d'ordre 1 peut être décomposée en rotation autour de $\mathbf{n}$ et $\mathbf{m}$. Il existe $\alpha, \beta, \gamma, \delta$ tels que :

$$U = e^{i\alpha} R_{\mathbf{n}}(\beta) R_{\mathbf{m}}(\gamma) R_{\mathbf{n}}(\delta)$$

Enfin, pour tout U et $\epsilon > 0$, il existe k_1, k_2, k_3 tel que

$$\|U - R_{\mathbf{n}}(\theta)^{k_1} R_{\mathbf{m}}(\theta)^{k_2} R_{\mathbf{n}}(\theta)^{k_3}\| < \epsilon$$

Ainsi, toute transformation unitaire d'ordre 1 peut être approchée en utilisant $U_2 U_1$ et $U_3 U_1$. Puisque $U_2 U_1 = (HT)(HT)$ et $U_3 U_1 = \sigma_y HTH \sigma_y T = -(\sigma_y HT)(\sigma_y HT)$, l'ensemble $\{HT, \sigma_y\}$ approxime toute transformation unitaire d'ordre 1.

En ajoutant ΛZ , l'ensemble B est approximativement universel. $\square$

Théorème 7.4 $E_2 = \{Z \otimes X, Z, \frac{X-Y}{\sqrt{2}}\}$ est approximativement universel, en utilisant un seul qubit auxiliaire.

Tout d'abord, on remarque que le pas de simulation de HT présenté en figure 7.7 peut être réalisé en utilisant les observables de E_2 et un seul qubit auxiliaire. Il reste donc à simuler ΛZ et les opérateurs de Pauli, utiles pour obtenir σ_y , mais également pour réaliser une simulation complète à partir d'un pas de simulation, selon la stratégie décrite en figure 7.3. On introduit quelques lemmes préalables avant de prouver ce théorème.

Lemme 7.7 Etant donné un registre de deux qubits a, b et un qubit auxiliaire c , la suite de mesures selon Z_c , $Z_a \otimes X_c$, $Z_c \otimes X_b$ et enfin Z_b (voir figure 7.11) simule l'application de ΛZ sur l'état $|\Phi\rangle$ des qubits a, b , à un opérateur de Pauli près. L'état final est sur les qubits a, c .

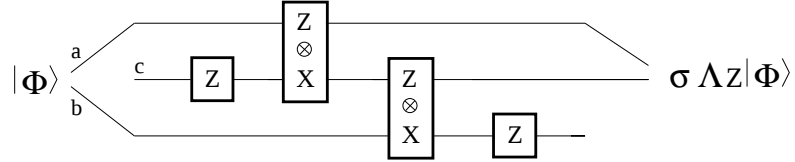


FIG. 7.11 – Simulation de ΛZ

La preuve du lemme 7.7 est similaire à celle du lemme 7.3.

Afin de simuler les opérateurs de Pauli σ_x et σ_z , un nouveau schéma différent du transfert d'état est introduit :

Lemme 7.8 Etant donné un qubit b et un qubit auxiliaire a , la suite de mesures selon Z_a , $X_a \otimes Z_b$, et Z_a (voir figure 7.12) simule l'application, sur le qubit b , de σ_z avec probabilité $1/2$ et l'identité avec probabilité $1/2$.

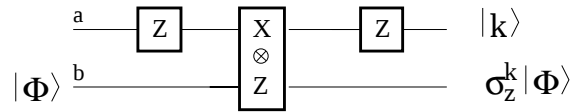


FIG. 7.12 – Simulation de σ_z

Preuve : Soit $|\Phi\rangle = \alpha|0\rangle + \beta|1\rangle$ l'état du qubit b . Après la première mesure dont le résultat classique est $j \in \{-1, 1\}$, l'état du registre a, b est :

$$|\psi_1\rangle = (\sigma_x^{(1-j)/2} \otimes Id) |0\rangle \otimes |\Phi\rangle$$

On remarque que $\langle \psi_1 | X \otimes Z | \psi_1 \rangle = 0$, donc l'état $|\psi_2\rangle$ après la deuxième mesure, dont le résultat classique est $k \in \{-1, 1\}$, est :

$$\begin{aligned} |\psi_2\rangle &= \frac{\sqrt{2}}{2}(\sigma_x^{(1-j)/2} \otimes Id)(Id + kX \otimes Z) |0\rangle \otimes |\Phi\rangle \\ &= \frac{\sqrt{2}}{2}(\sigma_x^{(1-j)/2} \otimes Id)(|0\rangle \otimes |\Phi\rangle + |1\rangle \otimes (\sigma_z |\Phi\rangle)) \end{aligned}$$

Ainsi, si la mesure du qubit a selon Z produit le résultat classique $l \in \{-1, 1\}$, alors si $k = j$ alors le registre a, b sera dans l'état $|0\rangle |\Phi\rangle$, sinon le registre sera dans l'état $|1\rangle \sigma_z |\Phi\rangle$. $\square$

Lemme 7.9 *Etant donné un qubit b et un qubit auxiliaire a , la suite de mesures selon $\left(\frac{X-Y}{\sqrt{2}}\right)_a$, $Z_a \otimes X_b$, et $\left(\frac{X-Y}{\sqrt{2}}\right)_a$ (voir figure 7.13) simule l'application, sur le qubit b , de σ_x avec probabilité $1/2$ et l'identité avec probabilité $1/2$.*

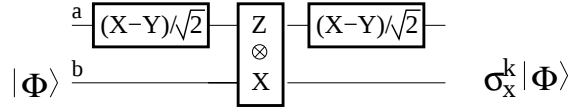


FIG. 7.13 – Simulation de σ_x

Preuve : Une preuve directe, comme pour le lemme 7.8 peut être donnée, on préfère ici utiliser le lemme sur la commutativité des transformations unitaires et des observables pour déduire le lemme 7.9 du lemme 7.8.

En effet, d'après le lemme 7.2 sur la commutativité des transformations unitaires et des observables, la suite de mesures $\left(\frac{X-Y}{\sqrt{2}}\right)_a$, $Z_a \otimes X_b$, et $\left(\frac{X-Y}{\sqrt{2}}\right)_a$ équivaut à appliquer HT sur a , H sur b , puis les mesures Z_a , $X_a \otimes Z_b$, et Z_a et enfin à appliquer $T^\dagger H$ sur a et H sur b . En effet on peut vérifier que $T^\dagger H Z H T = \frac{X-Y}{\sqrt{2}}$, $T^\dagger H X H T = Z$ et $H Z H = X$.

Il reste à interpréter l'application de ces transformations unitaires avant et après la simulation de σ_z sur le qubit b . Les transformations unitaires agissant sur le qubit a ne modifient pas l'action globale sur b car il n'y a pas de condition sur l'état du qubit auxiliaire avant et après le pas de simulation. En revanche, l'application de H avant et après la simulation sur le qubit b produit une simulation de $H\sigma_z H = \sigma_x$. $\square$

Preuve du théorème 7.4 : L'ensemble universel de transformations unitaires $B = \{HT, \sigma_y, \Lambda Z\}$ peut être simulé en utilisant des mesures de l'ensemble $E_2 = \{Z \otimes X, Z, \frac{X-Y}{\sqrt{2}}\}$.

Tout d'abord ΛZ et HT peuvent être simulés, à un opérateur de Pauli près, en utilisant des observables de E_2 (Figures 7.11 et 7.7).

La preuve se réduit donc à la simulation des opérateurs de Pauli permettant de simuler σ_y et d'obtenir une simulation de HT et ΛZ , à partir d'un pas de simulation de ces transformations unitaires. Les lemmes 7.8 et 7.9 proposent un schéma de simulation de σ_z et σ_x . Ces simulations ne sont pas obtenues à un opérateur de Pauli près, mais réussissent avec probabilité $1/2$, la stratégie décrite en 7.3 est donc modifiée :

- pour simuler σ_x , on applique le schéma 7.13 tant que la simulation de σ_x n'est pas obtenue ;
- pour simuler σ_z , on applique le schéma 7.12 tant que la simulation de σ_z n'est pas obtenue ;
- pour simuler σ_y , on utilise le fait que $\sigma_y = i\sigma_z\sigma_x$, en simulant σ_x puis σ_z . On obtient une simulation de σ_y à une phase globale près.

On remarque que même si la stratégie de simulation est légèrement modifiée, le coût en temps de la simulation est inchangé : le nombre de mesures à effectuer pour obtenir une simulation, avec probabilité d'erreur $\epsilon > 0$, d'un circuit composé de k portes et ayant pour base B , est toujours $O(k \cdot \log(k/\epsilon))$.

7.6 Conclusion

Nous avons montré que les ressources minimales nécessaires à la simulation exacte de transformations unitaires ont été atteintes : une mesure sur deux qubits et un ensemble infini de mesures sur un qubits sont universelles, en utilisant un seul qubit auxiliaire.

Nous avons également montré que trois mesures projectives (dont une seule sur deux qubits) et un qubit auxiliaire forment un ensemble approximativement universel. Ce résultat améliore les résultats de Nielsen et de Leung.

Une perspective est de trouver une borne inférieure pour ces ressources. Un qubit auxiliaire, ainsi qu'une mesure sur deux qubits sont nécessaires, la diminution des ressources porte donc sur le nombre de mesures sur un qubit. Afin de montrer qu'un ensemble d'observables est minimal, les phénomènes permettant de simuler des transformations unitaires à partir de mesures projectives doivent être mieux compris. Une étude préliminaire montre, par exemple, que deux observables appliqués successivement doivent satisfaire une propriété d'anticommutation nécessaire à la simulation d'une transformation réversible. Une perspective est donc d'exploiter ce type de propriétés afin de caractériser les ressources nécessaires au calcul par mesures projectives.

Le chapitre suivant traite de l'utilisation, dans le cadre du calcul par mesures projectives, de modèles formels comme le $\mathbf{q}$ -calcul, ou les machines de Turing contrôlées classiquement.

Chapitre 8

Modèles formels du calcul par mesures projectives

L'universalité du calcul par mesures projectives a été prouvée dans le chapitre précédent. L'élaboration de modèles formels pour décrire le calcul par mesures projectives est nécessaire. Tout d'abord, la notion de *stratégie* permettant de composer des *pas de simulation* doit être formalisée et n'est encore décrite que de façon *ad hoc*. Un modèle formel permet de représenter de façon homogène les mesures projectives qui sont appliquées et leur composition conditionnelle.

De plus, nous n'avons, pour l'instant, vu le calcul par mesure que par sa capacité à simuler une évolution unitaire. Ceci suggère un schéma où les algorithmes seraient d'abord décrits sous forme de circuits quantiques (composés uniquement de transformations unitaires) pour être ensuite "compilés" en une succession de mesures projectives. Or, comme le montre le $\mathbf{q}$ -calcul, ceci n'est qu'une extrémité du spectre des possibilités : un algorithme est décrit par une composition, contrôlée classiquement, de transformations admissibles, dont les transformations unitaires et les mesures projectives sont des cas particuliers. Des modèles formels, comme le $\mathbf{q}$ -calcul ou les MTQC, permettent en effet de décrire tout type de calculs quantiques, y compris un calcul composé exclusivement de mesures projectives.

L'utilisation d'un modèle formel comme le $\mathbf{q}$ -calcul pour décrire un calcul par mesures projectives permet d'associer à chaque terme sa sémantique. Alors que les transformations unitaires sont stables par composition, les mesures projectives ne le sont pas. Il apparaît donc important d'utiliser un formalisme capable de représenter les mesures projectives, mais qui soit de plus stable par composition. Le $\mathbf{q}$ -calcul satisfait cette propriété.

8.1 Fragment observable du $\mathbf{q}$ -calcul

Tout d'abord, nous introduisons le fragment du $\mathbf{q}$ -calcul dans lequel les seules transformations admissibles autorisées sont les mesures projectives :

Définition 8.1 *Un $\mathbf{o}$ -terme est un $\mathbf{q}$ -terme tel que tous les opérateurs linéaires utilisés comme actions, sont des projecteurs ($P^2 = P$).*

Exemple 8.1 *Les $\mathbf{q}$ -termes $\mathcal{P}_2, \mathcal{P}_3$ (Exemples 5.8, dans la chapitre 5) sont des $\mathbf{o}$ -termes, en revanche $\mathcal{P}_4$ ou $\mathcal{P}_5$ ne sont pas des $\mathbf{o}$ -termes.*

Une utilisation efficace des mesures projectives nécessite l'utilisation d'espace auxiliaire. C'est pourquoi nous définissons un fragment du $\mathbf{q}$ -calcul incluant en plus des mesures projectives, la possibilité d'initialiser un système dans un état standard $|\tau_0\rangle$ mais également les mesures destructrices :

- Une initialisation *Init* d'un espace auxiliaire est une transformation admissible de $\mathcal{H}_{B_1}$ dans $\mathcal{H}_{B_2}$, avec $B_1 \subseteq B_2$, telle que $Init = (|\tau_0\rangle\langle| \otimes Id_{\mathcal{H}_{B_1}})$
- Une mesure destructrice *D* est une transformation admissible de $\mathcal{H}_{B_2}$ dans $\mathcal{H}_{B_1}$, avec $B_1 \subseteq B_2$, telle que $D = (|\langle\tau| \otimes Id_{\mathcal{H}_{B^\tau}})_{\tau \in A}$, avec $B^\tau = \{x \mid (\tau, x) \in B_1\}$ et $\cup_{\tau \in A} B^\tau = B_1$.

Définition 8.2 *Un $\mathbf{o}'$ -terme est un $\mathbf{q}$ -terme $\mathcal{P} = (K, I, F, R)$ tel que pour tout processus $\mathbf{q} \in K \setminus F$, les opérateurs linéaires utilisés dans la définition de $\mathbf{q}$ dans R , forment soit :*

- une mesure projective ;
- une initialisation ;
- une mesure destructrice.

Afin d'établir la possibilité de simuler toute transformation unitaire à l'aide de mesures projectives, nous introduisons, pour toute transformation unitaire U , deux termes $\mathcal{P}_U$ et $\tilde{\mathcal{P}}_U$. $\tilde{\mathcal{P}}_U$ est un $\mathbf{q}$ -terme servant de référence qui consiste à appliquer U , alors que $\mathcal{P}_U$ est un $\mathbf{o}'$ -terme. Le lemme 8.1 montre l'équivalence observable de $\tilde{\mathcal{P}}_U$ et $\mathcal{P}_U$.

Définition 8.3 *Pour toute transformation unitaire $U \in \mathbf{L}(\mathcal{H}_{B_1}, \mathcal{H}_{B_2})$, $\tilde{\mathcal{P}}_U$ est un $\mathbf{q}$ -terme $(\{\mathbf{i}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, \tilde{R})$, avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{B_1}$, $\mathcal{H}_{\mathbf{f}} = \mathcal{H}_{B_2}$ et $\tilde{R}$:*

$$\mathbf{i} = U.\mathbf{f}$$

Définition 8.4 *Pour toute transformation unitaire $U \in \mathbf{L}(\mathcal{H}_{B_1}, \mathcal{H}_{B_2})$, $\mathcal{P}_U$ est un $\mathbf{o}'$ -terme $(\{\mathbf{i}, \mathbf{q}, \mathbf{p}, \mathbf{f}\}, \{\mathbf{i}\}, \{\mathbf{f}\}, R)$ avec $\mathcal{H}_{\mathbf{i}} = \mathcal{H}_{B_1}$, $\mathcal{H}_{\mathbf{f}} = \mathcal{H}_{B_2}$, $\mathcal{H}_{\mathbf{p}} = \mathcal{H}_{\mathbf{q}} = \mathcal{H}_{\{\top\} \times B_2 \cup \{\perp\} \times B_1}$ et R :*

$$\begin{aligned}
\mathbf{i} &= [|\perp\rangle\langle\perp| \otimes Id_{\mathcal{H}_{B_1}}].\mathbf{p} \\
\mathbf{p} &= [P_U, P_{-U}].\mathbf{q} \\
\mathbf{q} &= [|\perp\rangle\langle\perp| \otimes Id_{\mathcal{H}_{B_1}}].\mathbf{i} + [|\top\rangle\langle\top| \otimes Id_{\mathcal{H}_{B_2}}].\mathbf{f}
\end{aligned}$$

avec $P_U = \frac{1}{2}(|\top\rangle\langle\top| \otimes Id_{\mathcal{H}_{B_2}} + |\perp\rangle\langle\perp| \otimes Id_{\mathcal{H}_{B_1}} + |\top\rangle\langle\perp| \otimes U + |\perp\rangle\langle\top| \otimes U^\dagger)$.

$\mathcal{P}_U$ est un σ' -terme, car en particulier P_U et P_{-U} sont des projecteurs orthogonaux tels que $P_U + P_{-U} = Id_{\{\top\} \times B_2 \cup \{\perp\} \times B_1}$. Le σ' -terme $\mathcal{P}_U$ consiste tout d'abord à initialiser un espace auxiliaire formé des symboles $\perp$ et $\top$. Plus précisément B_1 représente une base de l'espace de départ, B_2 une base de l'espace d'arrivée et $B = \{\perp\} \times B_1 \cup \{\top\} \times B_2$. Le symbole $\perp$ signifie que la simulation n'est pas réalisée, alors que le symbole $\top$ signifie que la simulation est terminée, c'est pourquoi $\perp$ est associé à l'espace de départ et $\top$ à l'espace d'arrivée. Le système est initialisé dans l'état $|\perp\rangle|\varphi\rangle$. Après l'application du projecteur P_U ou P_{-U} , le système est dans une superposition uniforme de $\frac{1}{\sqrt{2}}(|\perp\rangle|\varphi\rangle \pm |\top\rangle U|\varphi\rangle)$. En mesurant l'espace auxiliaire, on obtient avec probabilité $1/2$ l'état $|\varphi\rangle$ et avec probabilité $1/2$, $U|\varphi\rangle$.

Nous établissons dans le lemme suivant la possibilité de simuler toute transformation unitaire en utilisant uniquement des mesures projectives, en prouvant que $\mathcal{P}_U$ et $\tilde{\mathcal{P}}_U$ ont la même sémantique observable.

Lemme 8.1 *Pour toute transformation unitaire $U \in L(\mathcal{H}_{B_1}, \mathcal{H}_{B_2})$,*

$$\mathcal{P}_U \equiv_{obs} \tilde{\mathcal{P}}_U$$

Preuve : Plutôt que de calculer directement la sémantique de $\mathcal{P}_U$, on remarque que la règle de substitution peut être appliquée au processus $\mathbf{q}$ et $\mathbf{p}$, l'ensemble des définitions de processus se réduit alors à :

$$\mathbf{i} = [\frac{1}{2}Id_{\mathcal{H}_{B_1}}, \frac{1}{2}Id_{\mathcal{H}_{B_1}}].\mathbf{i} + [\frac{1}{2}U, \frac{-1}{2}U].\mathbf{f}$$

La règle de regroupement peut être appliquée deux fois, avec comme paramètre 1 et -1 . On obtient un terme $\mathcal{P}'_U$ tel que $\mathcal{P}'_U \equiv_{obs} \mathcal{P}_U$ dont l'ensemble des définitions est :

$$\mathbf{i} = [\frac{1}{\sqrt{2}}Id_{\mathcal{H}_{B_1}}].\mathbf{i} + [\frac{1}{\sqrt{2}}U].\mathbf{f}$$

On peut alors calculer la sémantique admissible de $\mathcal{P}'_U$:

$$\begin{aligned}
\llbracket \mathcal{P}'_U \rrbracket^b &= \llbracket \frac{1}{\sqrt{2}}U \rrbracket^b_{\mathbf{i}, \mathbf{f}} \circ \cup_{n \in \mathbb{N}} (\llbracket \frac{1}{\sqrt{2}}Id_{\mathcal{H}_{B_1}} \rrbracket^b_{\mathbf{i}, \mathbf{i}})^n \\
&= \cup_{n \in \mathbb{N}^*} \left((\mathbf{i}, \mathbf{f}), \frac{1}{\sqrt{2}}U \right)
\end{aligned}$$

D'où $\llbracket \mathcal{P}'_U \rrbracket = \mathcal{X}(\llbracket \mathcal{P}'_U \rrbracket^b) = \lambda(d, |\Phi\rangle) \cdot \eta_{(\mathfrak{f}, U|\Phi)}.$

La sémantique admissible de $\tilde{\mathcal{P}}_U$ est $\llbracket \tilde{\mathcal{P}}_U \rrbracket^b = ((\mathfrak{f}, \mathfrak{f}), U)$ donc $\llbracket \tilde{\mathcal{P}}_U \rrbracket = \lambda(d, |\Phi\rangle) \cdot \eta_{(\mathfrak{f}, U|\Phi)}.$

Ainsi $\llbracket \tilde{\mathcal{P}}_U \rrbracket = \llbracket \mathcal{P}'_U \rrbracket$ or $\mathcal{P}'_U \equiv_{obs} \mathcal{P}_U$, donc $\mathcal{P}_U \equiv_{obs} \tilde{\mathcal{P}}_U$ $\square$

Le lemme 8.1 permet de généraliser l'universalité des mesures projectives au cas où les données ne sont pas composées exclusivement de qubits. On remarque ici une technique de preuve différente de celle utilisée dans le chapitre précédent. En effet, plutôt que de chercher à reconstruire une transformation unitaire à partir de pas de simulation, puis de stratégie de correction, la manipulation du σ' -terme $\mathcal{P}_U$ est basée sur les règles de transformations permettant de simplifier un $\mathfrak{q}$ -terme, et sur le calcul de la sémantique admissible.

L'utilisation du $\mathfrak{q}$ -calcul permet donc de considérer le calcul par mesures projectives, non pas seulement comme une méthode permettant de simuler le calcul quantique réversible, mais comme un véritable modèle de calcul indépendant.

8.2 Machine de Turing quantique fondée sur la mesure

Dans cette section, une restriction des machines de Turing quantiques contrôlées classiquement est introduite, il s'agit des machines de Turing quantiques fondées sur la mesure (MTQM) pour lesquelles seules les mesures projectives sont autorisées.

L'objectif de cette section est de montrer que ce nouveau modèle permet de simuler efficacement le modèle des MTQC. Nous montrons dans un premier temps que toute machine de Turing (classique) peut être simulée efficacement par une MTQM.

Dans le modèle des MTQC, l'ensemble des transformations admissibles $\mathcal{A}$ représente les ressources quantiques qui sont autorisées pendant le calcul. En limitant l'ensemble $\mathcal{A}$, on peut obtenir des modèles particuliers de calcul quantique tel que le calcul quantique fondé sur la mesure.

Définition 8.5 (Machine de Turing quantique fondée sur la mesure) *Une machine de Turing quantique fondée sur la mesure (MTQM) est une MTQC $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, où $\mathcal{A}$ est uniquement composée de mesures projectives.*

Le théorème suivant montre que toute MT est simulée par une MTQM *Las Vegas*, c'est-à-dire une MTQM avec un résultat non probabiliste, même si le temps d'exécution est non borné.

Théorème 8.1 *Pour toute MT M donnée opérant en un temps $f(n)$, où n est la taille de l'entrée, il existe une MTQM M' opérant en un temps espéré $O(f(n))$, et telle que, pour toute entrée x , $M(x) = M'(|x\rangle)$.*

Preuve : La preuve est similaire à la simulation de toute MT au moyen d'une MTQC (voir théorème 6.1), mais la permutation $\mathcal{P}_{[\tau,\sigma]}$ de deux symboles, qui est une transformation unitaire, doit être simulée à l'aide de mesures projectives. La simulation consiste à appliquer une mesure projective $\mathcal{O}_{[\tau,\sigma]}$ dans la base diagonale, transformant l'état $|\tau\rangle$ en une superposition uniforme $(|\tau\rangle \pm |\sigma\rangle)/\sqrt{2}$. Puis la cellule est mesurée dans la base standard produisant $|\sigma\rangle$ avec une probabilité $1/2$ et $|\tau\rangle$ avec une probabilité $1/2$. Si la simulation de la permutation échoue, la séquence des deux mesures $\mathcal{O}_{[\tau,\sigma]}$ et Std est appliquée à nouveau, jusqu'à la réussite.

Formellement, si $M = (K, \Sigma, \delta)$, soit $M' = (K', (\Sigma \cup \{\#, \lambda, \top, \perp\}), \Sigma, \{Std\} \cup \{\mathcal{O}_{[\tau,\tau']}\}_{\tau,\tau' \in \Sigma}, \delta')$. Pour tout $(q, \tau) \in K \times \Sigma$, if $\delta(q, \tau) = (p, \sigma, D)$, alors :

$$\begin{aligned} \delta'(q, \tau) &= (q_{\tau,\sigma}, -, \mathcal{O}_{[\tau,\sigma]}) \\ \delta'(q_{\tau,\sigma}, -) &= (q'_\tau, -, Std) \\ \delta'(q'_\tau, \sigma) &= (p, D, Std) \\ \delta'(q'_\tau, \tau) &= (q_{\tau,\sigma}, -, \mathcal{O}_{[\tau,\sigma]}) \end{aligned}$$

Chaque transition est simulée avec probabilité $1/2$ par deux mesures ($\mathcal{O}_{[\tau,\sigma]}$ et Std). Le nombre de transitions de M est $f(n)$, le nombre de transitions X de M' vérifie :

$$\begin{aligned} Pr(X = 2n + 1) &= 0 \\ Pr(X = 2n) &= \binom{f(n) - 1}{2n - 1} \frac{1}{2^{2n}} \end{aligned}$$

On en déduit que le temps espéré est $4f(n)$. □

D'après le théorème 8.1, les transformations unitaires ne sont donc pas nécessaires pour simuler le calcul classique. On peut se demander quelle est la puissance des MTQM comparées aux MTQC. Dans l'exemple 8.2, une MTQM qui simule l'application de la transformation d'Hadamard est décrite. La simulation est fondée sur le transfert d'état 7.3. D'une façon plus générale, le théorème 8.2 prouve que toute MTQC peut être simulée efficacement par une MTQM.

Exemple 8.2 (Simulation fondée sur la mesure de Hadamard)

Considérons le problème de la simulation d'une transformation unitaire donnée à l'aide de mesures. On choisit de simuler la transformation d'Hadamard H , sur une entrée à 1 qubit $|\Phi\rangle$, en utilisant uniquement des mesures projectives. Considérons une MTQM à 2 rubans $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, avec $K = \{s, s', q', q''\} \cup \{q_{j(a,b)} \mid a, b \in \{-1, 1\} \wedge j = 1 \dots 5\}$, $\Sigma_Q = \{\#, 0, 1\}$, $\Sigma_C = \{\#, \overline{\#}, -1, 0, 1, \top, \perp\}$, $\mathcal{A} = \{\mathcal{T}_\#, Std, \mathcal{O}_{[\#,0]}, \mathcal{O}_Z, \mathcal{O}_X, \mathcal{O}_{X \otimes Z}, \mathcal{O}_{Z \otimes Z}\}$ et δ est décrite dans la figure 8.1.

	$p \in K,$	$\tau \in \Sigma_C$	$\delta(p, \tau)$
	s	$\#$	$(s', (\rightarrow, -), [-, \mathcal{O}_{[\#,0]}])$
$\forall t \in \{-1, 0\}$	s'	t	$(q', (-, -), [-, Std])$
	q'	$\#$	$(s', (-, -), [-, \mathcal{O}_{[\#,0]}])$
	q'	0	$(q'', (-, -), [-, \mathcal{O}_X])$
$\forall i \in \{-1, 1\}$	q''	i	$(q_{1(i,1)}, (-, -), \mathcal{O}_{X \otimes Z})$
$\forall i, j \in \{-1, 1\}$	$q_{1(i,1)}$	j	$(q_{2(i,j)}, (-, -), [\mathcal{O}_Z, -])$
$\forall i, j, k \in \{-1, 1\}$	$q_{2(i,j)}$	k	$(q_{3(i,k,j)}, (-, -), [\mathcal{O}_X, -])$
$\forall k', j, l \in \{-1, 1\}$	$q_{3(k',j)}$	k	$(q_{4(k',l,j)}, (-, -), \mathcal{O}_{Z \otimes Z})$
$\forall l', j, m \in \{-1, 1\}$	$q_{4(l',j)}$	k	$(q_{5(l',j,m)}, (-, -), [-, \mathcal{O}_X])$
$\forall l', n \in \{-1, 1\} t.q.l' = n$	$q_{5(l',1)}$	n	$(h, (-, -), -)$
$\forall l', m', n \in \{-1, 1\} t.q.m' \neq 1 \vee l' \neq n$	$q_{5(l',m')}$	n	$(q_{2(l',n,m')}, (-, -), \mathcal{O}_{Z \otimes Z})$

FIG. 8.1 – MTQM à 2 rubans avec uniquement des mesures projectives pour la simulation de H .

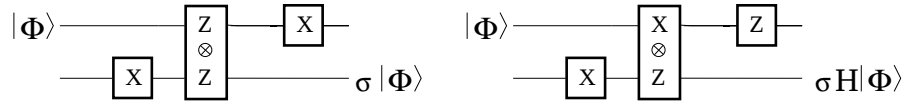


FIG. 8.2 – à gauche : transfert d'état - à droite : transfert d'état généralisé pour la simulation de H .

L'état d'entrée est placé sur le ruban T_1 . Les quatre premières transitions sont utilisées pour transformer l'état désigné par la seconde tête, de l'état $|\#\rangle$ vers $|0\rangle$. Puis les mesures projectives sont appliquées d'après le transfert d'état 7.3 (voir figure 8.2) : après ces mesures, $H|\Phi\rangle$, à un opérateur de Pauli près, est placé sur T_2 . Puisque le résultat du calcul doit être placé sur T_1 , les trois transitions suivantes transfèrent le résultat de la simulation de T_2 vers T_1 .

Le transfert d'état est obtenu à un opérateur de Pauli près. Cet opérateur dépend des résultats classiques des différentes mesures, c'est pourquoi les résultats classiques sont mémorisés à l'aide des états internes de la machines : $q_{j(a,b)}$ signifie que le transfert d'état est obtenu à $\sigma_z^{\frac{1-a}{2}} \sigma_x^{\frac{1-b}{2}}$ près. Dans le but de corriger cet opérateur de Pauli, l'état est transféré deux fois : de T_1 vers T_2 , puis à l'inverse de T_2 vers T_1 . Quand M s'arrête, l'état de la cellule pointée par la première tête est $H|\Phi\rangle$, donc $M(|\Phi\rangle) = H|\Phi\rangle$.

Notons que le nombre de transitions espéré est constant.

Théorème 8.2 *Pour toute k -MTQC M opérant en un temps $f(n)$, il existe une k -MQTM M' opérant en un temps $O(f(n))$, telle que pour toute entrée $|\psi\rangle$ de taille n , $M(|\psi\rangle) = M'(|\psi\rangle)$.*

Preuve : La démonstration est composée de deux étapes : chaque transformation admissible utilisée dans la fonction de transition de M est transformée en une transformation unitaire immédiatement suivie d'une mesure projective, puis dans un second temps, l'opération unitaire précédemment obtenue est simulée à l'aide de mesures projectives.

Si $M = (K, \Sigma_C, \Sigma_Q, \mathcal{A}, \delta)$, alors $M' = (K', (\Sigma_C^2 \times \{\top, \perp\}), (\Sigma_Q \times \Sigma_C \times \{\top, \perp\}), \mathcal{A}', \delta')$. L'alphabet des cellules quantiques de M' est composé de triplets (Φ, c, r) : $\Phi \in \Sigma_Q$ est utilisé pour simuler la cellule quantique correspondante de M , $c \in \Sigma_C$ est utilisé dans la première étape transformant les transformations admissibles en unitaires suivie d'une mesure, et $r \in \{\top, \perp\}$ est un élément de l'espace de travail additionnel nécessaire à la simulation des transformations unitaires à l'aide de mesures projectives. $\mathcal{A}'$ est un ensemble de transformations admissibles agissant sur les états quantiques dans $\mathcal{H}_{(\Sigma_Q \times \Sigma_C \times \{\top, \perp\})^k}$. Dans la suite, nous utilisons implicitement le fait que $\mathcal{H}_{(\Sigma_Q \times \Sigma_C \times \{\top, \perp\})^k}$ est isomorphe à $\mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k \times \{\top, \perp\}^k}$.

- Pour tout $(q, c) \in K \times \Sigma_C$, si $\delta(q, c) = (p, D, A)$, alors $\tilde{\delta}(q, c) = (q_c, -, \mathcal{U}_V)$ et $\tilde{\delta}(q_c, \lambda) = (p, D, \mathcal{O})$ tel que si $A = \{M_c, c \in \Sigma_C\}$, alors $V : \mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k} \rightarrow \mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k}$ est une transformation unitaire telle que $V|\psi\rangle|\#^k\rangle = \sum_{c \in \Sigma_C} (M_c|\psi\rangle)|c\#^{k-1}\rangle$ (V peut être étendue au cas où le second registre n'est pas dans l'état $|\#^k\rangle$), et $\mathcal{O}$ est une mesure projective du second registre dans la base $\{|c\#^{k-1}\rangle, c \in \Sigma_C\}$. Si δ est remplacé par $\tilde{\delta}$, une simulation

exacte de M est alors obtenue.¹

- Afin de simuler une transition consistant à appliquer une transformation V , reprenons la technique utilisée dans le théorème 8.1. Soit $R = V \otimes |\perp\rangle \langle \top| \otimes Id_{\mathcal{H}_{\{\top, \perp\}^{k-1}}} + V^\dagger \otimes |\top\rangle \langle \perp| \otimes Id_{\mathcal{H}_{\{\top, \perp\}^{k-1}}}$, $P_\top = (Id_{\mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k \times \{\top, \perp\}^k}} + R)/2$, et $P_\perp = (Id_{\mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k \times \{\top, \perp\}^k}} - R)/2$. $\mathcal{V} = \{P_\top, P_\perp\}$ est une mesure projective. On peut montrer que pour tout $|\Phi\rangle \in \mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k}$, une mesure de $|\Phi\rangle \otimes |\top^k\rangle$ selon $\mathcal{V}$, suivie par une mesure projective selon $L = \{Id_{\mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k}} \otimes |\top\rangle \langle \top| \otimes Id_{\mathcal{H}_{\{\top, \perp\}^{k-1}}}, Id_{\mathcal{H}_{\Sigma_Q^k \times \Sigma_C^k}} \otimes |\perp\rangle \langle \perp| \otimes Id_{\mathcal{H}_{\{\top, \perp\}^{k-1}}}\}$, produit $(V|\Phi\rangle) \otimes |\perp \top^{k-1}\rangle$ avec une probabilité $1/2$, et $|\Phi\rangle \otimes |\top^k\rangle$ avec une probabilité $1/2$. Alors, pour tout $(q, c) \in K \times \Sigma_C$, si $\tilde{\delta}(q, c) = (q_c, -, \mathcal{U}_V)$ et $\tilde{\delta}(q_c, \lambda) = (p, D, \mathcal{O}_O)$, soit δ' tel que :

$$\begin{aligned} \delta'(q, c) &= (q'_c, -, \mathcal{O}_V) \\ \delta'(q'_c, -) &= (q''_c, -, \mathcal{O}_L) \\ \delta'(q''_c, \top) &= (q'_c, -, \mathcal{O}_V) \\ \delta'(q''_c, \perp) &= (p, D, \mathcal{O}_O) \end{aligned}$$

Ainsi, toute transition de la MTQC M peut être simulée sur la MTQM M' avec un nombre espéré constant de transitions. Par conséquent, M' simule M avec un temps d'exécution espéré $O(f(n))$. $\square$

Puisque toute MTQC à k rubans est simulée efficacement par une MTQM à k rubans, les résultats prouvés pour les MTQC dans le chapitre 6 sont aussi valables pour les MTQM. Ainsi, bien que le calcul fondé sur la mesure nécessite un espace auxiliaire pour simuler des transformations réversibles, toute MTQC à k rubans peut être simulée avec un ralentissement polynomial par une MTQM à 2 rubans. On en déduit également que le formalisme des MTQM est aussi puissant que celui des circuits quantiques et des machines de Turing quantiques.

8.3 Conclusion

Les outils formels développés pour le calcul quantique contrôlé classiquement peuvent être utilisés dans le cadre plus restreint du calcul par mesures projectives. Les outils formels obtenus sont les machines de Turing fondées sur la mesure (MTQM), ainsi qu'une restriction du $\mathbf{q}$ -calcul aux $\mathbf{o}'$ -termes.

Le théorème 8.2 montre que toute MTQC à k rubans peut être simulée efficacement par une MTQM à k rubans. Ce résultat généralise celui de Nielsen, puisque

¹La simulation d'une transformation admissible à l'aide de transformations unitaires suivies de mesures projectives, est une variante, adaptée aux MTQM, du lemme 1.1.

les opérations simulées ici ne sont pas des transformations unitaires, mais des transformations admissibles. Cependant, le point le plus important de ce théorème est sans doute que cette simulation ne nécessite pas de ruban supplémentaire. En effet, le fait que le calcul par mesures projectives requiert un espace auxiliaire, se traduit par l'utilisation de cellules quantiques de plus grande dimension, et non par l'utilisation d'un ruban supplémentaire. Ainsi les résultats sur les MTQC, dépendant notamment du nombre de rubans, sont également valables pour les MTQM.

L'utilisation du $\mathbf{q}$ -calcul dans le cadre du calcul par mesures projectives, permet de montrer que ce type de calcul peut être utilisé en dehors de la simulation des transformations unitaires dans un schéma très formaté de type téléportation. En effet, le modèle de Nielsen [Nie03], ainsi que les améliorations qui en ont été faites, sont fondés sur un même schéma de simulation : tout d'abord un pas de simulation produisant la simulation d'une transformation unitaire à un opérateur de Pauli près, puis une stratégie de correction. Le $\mathbf{q}$ -calcul offre une possibilité de donner une sémantique à un calcul par mesures projectives, sans pour autant que ce calcul soit un mimétisme du calcul réversible.

Ainsi les règles de transformations (substitution, factorisation et regroupement, 5) permettent de simplifier un σ' -terme en un $\mathbf{q}$ -terme dont la sémantique admissible est plus facilement calculable.

Quatrième partie

Représentation de l'intrication et calcul par consommation d'intrication

Chapitre 9

Etats Graphes

9.1 Introduction

Choisir une représentation abstraite des états quantiques qui soit adéquate pour formuler simplement les propriétés de ces états et les traitements qu'on leur applique est évidemment une question qui se pose en informatique quantique. L'état d'un système quantique composé de n qubits est un vecteur composé de 2^n coordonnées complexes. Par exemple, le vecteur état d'un système de 10 qubits est composé de 1024 nombres complexes. Face à cette explosion exponentielle, des formalismes plus *compacts* ont été introduits : les états stabilisables [NC00] et les états graphes [HEB04, VdNDD04, HDE⁺05] sont les plus utilisés. Fondés sur un compromis entre compacité et expressivité, ces formalismes ne permettent de représenter que certaines classes d'états quantiques, mais la taille de la représentation de l'état est polynomiale en la taille du système.

Notons que le formalisme des stabiliseurs¹ est strictement plus expressif que celui des états graphes. Le formalisme des états graphes est fondé sur une représentation graphique des états quantiques, au moyen de graphes simples non orientés. Cette représentation permet une caractérisation combinatoire de certaines propriétés quantiques. Par exemple, certaines transformations unitaires peuvent être interprétées comme des transformations sur le graphe correspondant. Afin d'augmenter l'expressivité du formalisme et de pouvoir interpréter certaines mesures quantiques de façon combinatoire, nous introduisons une extension du formalisme, les états graphes signés.

Les états graphes ont été largement étudiés ces cinq dernières années. L'étude récente [HDE⁺05] par Hein *et al.* donne une excellente introduction au domaine et

¹Le formalisme des stabiliseurs est largement utilisé, par exemple pour les codes correcteurs quantique. Le stabiliseur d'un état $|\varphi\rangle$ est un ensemble $\{P_i\}$ d'opérateurs de Pauli, dont $|\varphi\rangle$ est l'unique point fixe : $\forall i P_i |\varphi\rangle = |\varphi\rangle$.

inclut plus de 200 références. Ces travaux ont mis en évidence des liens entre des concepts de base du calcul quantique et la théorie des graphes. Ils ont également établi plusieurs résultats fondamentaux sur les implémentations d'états graphes, sur l'intrication représentée par des états graphes et sur l'universalité du calcul quantique fondé sur des états graphes.

Les états graphes sont également utilisés dans le calcul par consommation d'intrication, où les états graphes représentent la ressource initiale qui est consommée pas à pas par des mesures locales effectuées pour conduire un calcul. Ce modèle est présenté et développé dans le chapitre 10, où notamment, de nouveaux critères permettant de savoir si un état graphe peut être utilisé pour mener un calcul quantique sont introduits.

La question de la préparation des états graphes est également importante. En effet, les états graphes étant par exemple la ressource initiale dans un calcul par consommation d'intrication, quelles sont les ressources plus traditionnelles (nombre d'opérations, nombre de qubits auxiliaires) nécessaires pour préparer un état graphe donné ? Ce problème est discuté dans le chapitre 11.

Dans ce chapitre, les états graphes sont présentés et une extension aux états graphes signés est introduite. Nous associons également une caractérisation combinatoire à différentes transformations quantiques.

9.2 Etats graphes et états graphes signés

Pour tout graphe non orienté G à n sommets, l'état graphe $|G\rangle$ est un état quantique sur n qubits :

Définition 9.1 *Pour tout graphe $G = (V, E)$, $|G\rangle$ est l'état graphe sur $|V|$ qubits tel que :*

$$|G\rangle = \frac{1}{\sqrt{2^n}} \sum_{x \in \{0,1\}^n} (-1)^{q_\Gamma(x)} |x\rangle,$$

où Γ est la matrice d'adjacence² de G et $q_\Gamma(x) = \sum_{i < j: (i,j) \in E} x_i x_j$.

La forme quadratique q_Γ vérifie $q_\Gamma(x) = x^T \Gamma^{\text{upper}} x$ où x est vu comme un vecteur colonne de $\{0,1\}^n$, x^T est le vecteur transposé de x et Γ^{upper} est une matrice triangulaire supérieure obtenue à partir de Γ telle que $\forall i < j, \Gamma_{i,j}^{\text{sup}} = \Gamma_{i,j}$.

Les états graphes vérifient les propriétés suivantes [HDE⁺05] :

Propriété 9.1 *Pour tout graphe $G = (V, E)$, $|G\rangle$ est l'unique point fixe de l'ensemble d'opérateurs de Pauli $\{X_u Z_{N_G(u)}, u \in V\}$, où X_u signifie l'application de*

² $\Gamma_{i,j} = 1$ si i, j est une arête du graphe et $\Gamma_{i,j} = 0$ sinon.

l'opérateur de Pauli X sur le qubit u , $N_G(u)$ est l'ensemble des sommets adjacents à u dans le graphe G et $Z_S = \prod_{v \in S} Z_v$.

On en déduit que pour tout $u \in V$,

$$|G\rangle = X_u Z_{N_G(u)} |G\rangle$$

Le formalisme des états graphes est un formalisme compact. En effet, il permet de représenter un état quantique sur n qubit en utilisant seulement $n(n-1)/2$ bits au lieu de 2^n nombres complexes pour la représentation standard des états quantiques, sous forme de vecteur. La contre partie de cette compacité est que seule une classe particulière d'états quantiques peuvent être représentés par un graphe. Par exemple, on remarque que tous les états graphes sont des superpositions uniformes, au signe près, de tous les vecteurs de la base standard.

Afin d'en augmenter le pouvoir expressif, nous introduisons une généralisation des états graphes, les états graphes signés :

Définition 9.2 Pour tout graphe $G = (V, E)$ et tout $S \subseteq V$, l'état graphe signé $|G; S\rangle$ est l'état

$$|G; S\rangle = Z_S |G\rangle$$

S est le *signe* de l'état graphe signé $|G; S\rangle$. L'application d'un opérateur de Pauli à un état graphe signé peut être intégré au signe de l'état graphe :

Propriété 9.2 Pour tout graphe $G = (V, E)$, tout signe $S \subseteq V$, et tout opérateur de Pauli P sur $|V|$ qubits, il existe $S' \subseteq V$ et $k \in \mathbb{N}$ tels que

$$P |G; S\rangle = i^k |G; S'\rangle$$

où $i = \sqrt{-1}$.

Preuve : P peut être décomposé, à une phase globale près, en un produit d'opérateurs X et Z : il existe $k \in \mathbb{N}$, et $S_1, S_2 \subseteq V$ tels que $P = i^k X_{S_1} Z_{S_2}$, Donc

$$\begin{aligned} P |G; S\rangle &= i^k X_{S_1} Z_{S_2 \Delta S} |G\rangle \\ &= i^{k'} Z_{S_2 \Delta S} X_{S_1} \prod_{u \in S_1} (X_u Z_{N_G(u)}) |G\rangle \\ &= i^{k'} Z_{S_2 \Delta S} \prod_{u \in S_1} Z_{N_G(u)} |G\rangle \\ &= i^{k'} |G; S'\rangle \end{aligned}$$

où Δ est la différence symétrique et $S' = (S_2 \Delta S) \Delta (\Delta_{u \in S_1} N_G(u))$. □

De plus, pour tout graphe fixé G , l'ensemble des états graphes signés $\{|G; S\rangle, S \subseteq V\}$ forme une base orthonormée de $\mathcal{H}_{\{0,1\}^n}$:

Proposition 9.1 *Pour tout graphe $G = (V, E)$ et tout sous-ensemble non vide $S \subseteq V$,*

$$\langle G; S | G \rangle = 0,$$

et ainsi, $\langle G; S | G; S' \rangle = 0$ pour tous les sous-ensembles distincts $S, S' \subseteq V$, et les $2^{|V|}$ états $|G; S\rangle$, pour G fixé et $S \subseteq V$ forment une base orthonormale.

Preuve : Pour tout G et tout $S \neq \emptyset$, on remarque que pour $u \in S$, $X_u Z_S = -Z_S X_u$, donc $\langle G; S | G \rangle = \langle G | Z_S | G \rangle = \langle G | Z_S X_u Z_{N_G(u)} | G \rangle = -\langle G | Z_S | G \rangle$. Ainsi $\langle G; S | G \rangle = 0$. $\square$

Avec cette représentation unique des états graphes signés, le pouvoir expressif des états graphes est augmenté. En effet, chaque paire formée d'un graphe et d'un signe représente un état quantique différent :

Propriété 9.3 *Pour tout graphe $G = (V, E)$, $G' = (V, E')$, et tout signe $S, S' \subseteq V$,*

$$|G; S\rangle = |G'; S'\rangle \implies G = G' \wedge S = S'$$

Preuve : Si $G = G'$, alors la proposition 9.1 permet de conclure que $S = S'$.

Par contradiction, supposons que $G \neq G'$. Il existe alors un sommet $u \in V$ tel que $N_G(u) \neq N_{G'}(u)$. De plus, $|G\rangle = |G'; S \Delta S'\rangle$, donc $X_u Z_{N_G(u)} |G\rangle = \beta X_u Z_{N_{G'}(u)} |G'; S \Delta S'\rangle$ avec $\beta = 1$ si $u \notin S \Delta S'$ et $\beta = -1$ sinon.

Ainsi $|G\rangle = \beta |G; N_G(u) \Delta N_G(u)\rangle$, si $\beta = 1$, alors la proposition 9.1 permet de conclure que $N_{G'}(u) \Delta N_G(u) = \emptyset$ donc $N_{G'}(u) = N_G(u)$.

Si $\beta = -1$, on considère $v \in N_{G'}(u) \Delta N_G(u)$, donc $X_v Z_{N_G(v)} |G\rangle = X_v Z_{N_{G'}(v)} |G; N_G(u) \Delta N_G(u)\rangle$ ainsi $|G\rangle = |G; N_G(u) \Delta N_G(u)\rangle$, ce qui permet également de conclure que $N_{G'}(u) = N_G(u)$. $\square$

9.3 Propriétés combinatoires des états graphes

L'utilisation de graphes pour représenter des états quantiques établit une connexion entre les états quantiques et la théorie de graphes. Dans cette section, nous abordons quelques propriétés combinatoires des états graphes. Par exemple, l'application d'une transformation unitaire sur un état graphe peut être interprétée formellement sur le graphe. L'utilisation des états graphes signés permet de donner des interprétations formelles à des transformations comme les mesures projectives. Nous montrons également les limites de l'interprétation combinatoire des transformations quantiques. En effet, certaines transformations, comme la mesure d'un seul qubit selon l'observable X , ne peuvent être interprétées directement comme des transformations combinatoires. Cette impossibilité souligne l'importance du formalisme des états graphes signés.

La complémentation locale et le pivot sont des transformations combinatoires déterminantes dans le cas des transformations unitaires locales. Certaines transformations unitaires non locales et certaines mesures locales ou non, sont également caractérisées de façon combinatoire.

9.3.1 Complémentation locale et pivot

Les opérations unitaires locales, sur un qubit, ne modifient pas l'intrication de l'état quantique global auquel ce qubit participe. On peut ainsi définir des classes d'équivalence d'états quantiques ayant la même intrication :

Définition 9.3 (LU-équivalence) *Pour tout $|\varphi\rangle, |\psi\rangle \in \mathcal{H}_{\{0,1\}^n}$, $|\varphi\rangle \equiv_{LU} |\psi\rangle$ si et seulement si il existe une transformation unitaire U telle que $|\varphi\rangle = U|\psi\rangle$ et $U = U_1 \otimes \dots \otimes U_n$, où U_i est l'application d'une transformation unitaire locale au $i^{\text{ème}}$ qubit.*

Une version plus restrictive de la LU -équivalence est la LC -équivalence, dans laquelle la transformation locale est une opération de Clifford :

Définition 9.4 (LC-équivalence) *Pour tout $|\varphi\rangle, |\psi\rangle \in \mathcal{H}_{\{0,1\}^n}$, $|\varphi\rangle \equiv_{LC} |\psi\rangle$ si et seulement si $|\varphi\rangle = C_1 \otimes \dots \otimes C_n |\psi\rangle$ où chaque C_i est un élément du groupe engendré par $\langle \sqrt{X}, \sqrt{Z} \rangle$, avec $\sqrt{\sigma} = \frac{1}{\sqrt{2}}(Id - i\sigma)$, pour $\sigma \in \{X, Y, Z\}$.*

Tous les états ne peuvent être représentés par des états graphes signés. Mais existe-t-il au moins un état graphe signé par classe d'équivalence induite par la LU -équivalence ? Existe-t-il des états graphes différents et pourtant LU -équivalents ?

Tout d'abord, certains états n'ont pas d'état graphe pour équivalent. Par exemple, aucun état graphe n'est LU équivalent à l'état $|W\rangle = \frac{1}{\sqrt{3}}(|001\rangle + |010\rangle + |100\rangle)$. Cela signifie que les états graphes ne permettent pas de représenter tous les types d'intrication. De plus, il existe des états graphes différents et pourtant équivalents. Par exemple, tous les graphes connexes à 3 sommets représentent des états graphes LC -équivalents.

Le problème consistant à savoir si deux graphes représentent des états graphes LU -équivalents a été en partie résolu par Van den Nest [VdN05] en utilisant la *complémentation locale*, une transformation sur les graphes étudiée par Bouchet [Bou89], qui consiste à compléter le voisinage d'un sommet :

Définition 9.5 *La complémentation locale de $G = (V, E)$ selon $u \in V$ est le graphe $G \star u = G \Delta K_{N_G(u)}$ où $K_{N_G(u)}$ est le graphe complet sur les voisins de u dans G .*

On dit que deux graphes G et G' sont localement équivalents si et seulement si il existe une suite de complémentations locales transformant G en G' , on écrit alors $G \approx_{\text{loc}} G'$.

Van den Nest a montré que si deux graphes G, G' sont localement équivalents, alors $|G\rangle$ et $|G'\rangle$ sont LC -équivalents :

Théorème 9.1 (Van den Nest [VdN05]) *Deux graphes G et G' sont localement équivalents si et seulement si $|G\rangle \equiv_{LC} |G'\rangle$*

Par exemple, pour tout graphe $G = (V, E)$ et tout $u \in V$, $|G\rangle$ et $|G \star u\rangle$ sont LC -équivalents, plus précisément :

$$|G \star u\rangle = \sqrt{X_u}^\dagger \sqrt{Z}_{N_G(u)} |G\rangle$$

Ainsi l'application de $\sqrt{X_u}^\dagger \sqrt{Z}_{N_G(u)}$ peut être interprétée comme une complémentation locale sur le sommet u . Nous étendons cette interprétation aux états graphes signés :

Propriété 9.4 *Pour tout graphe $G = (V, E)$ et tout signe $S \subseteq V$, il existe $S' \subseteq V$ et $k \in \mathbb{N}$ tels que :*

$$i^k |G \star u; S'\rangle = \sqrt{X_u}^\dagger \sqrt{Z}_{N_G(u)} |G; S\rangle$$

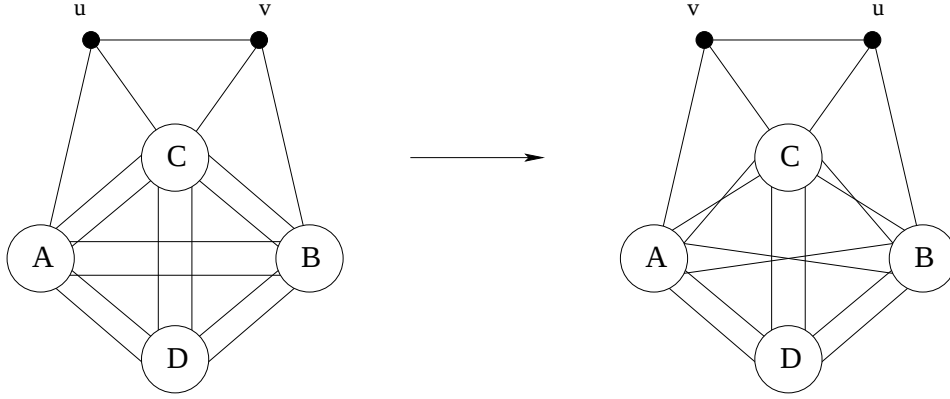
Preuve : On remarque que $\sqrt{X}Z = -Z\sqrt{X}$ et $\sqrt{Z}Z = Z\sqrt{Z}$, on en déduit que pour tout S , il existe un opérateur de Pauli P tel que $\sqrt{X_u}^\dagger \sqrt{Z}_{N_G(u)} Z_S |G\rangle = P |G \star u\rangle$. La propriété 9.2 permet de conclure. $\square$

Etant donné un graphe $G = (V, E)$, un *pivot* selon l'arête $uv \in E$ transforme le graphe G en $G \wedge uv = G \star u \star v \star u$ (voir figure 9.1). On vérifie facilement que $G \star u \star v \star u = G \star v \star u \star v$ pour tout $uv \in E$. Un pivot correspond à une complémentation du sous graphe tripartite A, B, C (c.f. figure) et un échange des sommets u et v .

Propriété 9.5 *Pour tout graphe $G = (V, E)$, tout signe $S \subseteq V$ et toute arête $(u, v) \in E$, il existe $S' \subseteq V$ et $k \in \mathbb{N}$ tel que :*

$$i^k |G \wedge uv; S'\rangle = \sqrt{Y_{u,v}}^\dagger |G; S\rangle$$

Preuve : Dans un premier temps, supposons $S = \emptyset$. On remarque, en utilisant les notations de la figure 9.1, que $|G \star u\rangle = \sqrt{X_u}^\dagger \sqrt{Z}_{A \cup C} |G\rangle$, $|G \star u \star v\rangle =$

FIG. 9.1 – Pivot selon uv

$\sqrt{X}_v^\dagger \sqrt{Z}_{A \cup B} |G \star v\rangle$ et $|G \star u \star v \star u\rangle = \sqrt{X}_u^\dagger \sqrt{Z}_{B \cup C} |G \star u \star v\rangle$. On en déduit que :

$$\begin{aligned}
 |G \wedge uv\rangle &= \sqrt{X}_u^\dagger \sqrt{Z}_u \sqrt{X}_u^\dagger \sqrt{Z}_v \sqrt{X}_v^\dagger \sqrt{Z}_v Z_{A \cup B \cup C} |G\rangle \\
 &= -\sqrt{Y}_{u,v}^\dagger X_u X_v Z_{A \cup B \cup C} |G\rangle \\
 &= \sqrt{Y}_{u,v}^\dagger Z_{C \cup \{u,v\}} |G\rangle
 \end{aligned}$$

Si $S \neq \emptyset$, on a $\sqrt{Y}^\dagger Z = -X \sqrt{Y}^\dagger$, donc il existe un opérateur de Pauli P tel que $\sqrt{Y}_{u,v}^\dagger Z_S |G\rangle = P \sqrt{Y}_{u,v}^\dagger Z_{C \cup \{u,v\}} |G\rangle = P |G \wedge uv\rangle$. La propriété 9.2 permet de conclure.

□

La caractérisation de la LU -équivalence par la complémentation locale est une conjecture :

Conjecture 9.1 *G et G' sont localement équivalents si et seulement si $|G\rangle$ et $|G'\rangle$ sont LU -équivalents.*

9.3.2 Transformations

Dans cette section, nous nous intéressons aux transformations qui peuvent être appliquées à un état graphe ou à un état graphe signé. Les transformations préservant l'intrication ont été étudiées dans la section précédente. L'objet de cette section est l'étude d'autres transformations quantiques, comme les transformations unitaires non locales, ou les mesures sur un ou plusieurs qubits selon un observable de Pauli.

La transformation unitaire, non locale, ΛZ , peut facilement être interprétée de façon combinatoire : l'application de ΛZ sur deux qubits u, v d'un état graphe ajoute l'arête (u, v) si celle-ci n'existait pas et la supprime sinon.

Proposition 9.2 *Pour tout graphe (V, E) , tout $S \subseteq V$ et tout $u, v \in V$,*

$$\Lambda Z_{u,v} |(V, E); S\rangle = |(V, E \Delta(u, v)); S\rangle$$

Ainsi l'application de l'opération ΛZ sur un état graphe signé, crée ou supprime une arête entre les sommets correspondants. On en déduit une méthode de construction des états graphes : en partant d'un état graphe sans arête, où tous les qubits sont dans l'état $|+\rangle = \frac{|0\rangle + |1\rangle}{\sqrt{2}}$, on obtient l'état graphe $|G\rangle$ en appliquant ΛZ sur chaque paire de qubits correspondant à des sommets voisins dans G .

Certaines mesures de Pauli peuvent être interprétées de façon combinatoire. Nous supposons, dans un premier temps, que les mesures sont destructrices, c'est-à-dire que le ou les qubits mesurés sont consommés par la mesure. Nous considérons les transformations suivantes sur les graphes, afin de donner une caractérisation combinatoire de certaines mesures de Pauli :

Pour tout graphe $G = (V, E)$:

- La *suppression d'un sommet* $v \in V$ dans G est le graphe $G \setminus v$ obtenu à partir de G en retirant v et toutes les arêtes incidentes à v .
- La *suppression d'une arête* $e \in E$ dans G est le graphe $G \setminus e$ obtenu à partir de G en retirant simplement l'arête e .
- La *contraction d'une arête* $(u, v) \in E$ dans G est le graphe $G/(u, v)$ obtenu à partir de G en introduisant des arêtes entre v et chaque sommet de $N_G(u) \setminus N_G(v)$, et ensuite en supprimant u et toutes les arêtes qui lui sont incidentes.

Un graphe G est un *sous graphe induit* de G' , si G est isomorphe à un graphe qui peut être obtenu à partir de G' par suppressions de sommets.

Un graphe G est un *sous graphe* de G' si on peut obtenir un graphe isomorphe à G par suppressions d'arêtes dans G' . C'est un *mineur* de G' s'il est isomorphe à un graphe qui peut être obtenu à partir de G' par suppressions ou contractions d'arêtes. Il s'agit d'un *mineur induit* de G' s'il est isomorphe à un graphe qui peut être obtenu à partir de G' par suppressions de sommets et contractions d'arêtes. Et G est un *pivot mineur* de G' s'il est isomorphe à un graphe qui peut être obtenu à partir de G' par suppression de sommets et par pivots.

Une des transformations les plus simples est la mesure selon Z , c'est-à-dire dans la base standard $\{|0\rangle, |1\rangle\}$. En effet, une telle mesure peut être interprétée comme la suppression du sommet correspondant :

Lemme 9.1 (Suppression d'un sommet) *Soient $G = (V, E)$ un graphe, $S \subseteq V$ et $v \in V$ un sommet. Appliquer une mesure selon Z sur un qubit v transforme $|G; S\rangle$ en $|G \setminus v; S'\rangle$ et ainsi correspond à la suppression du sommet v .*

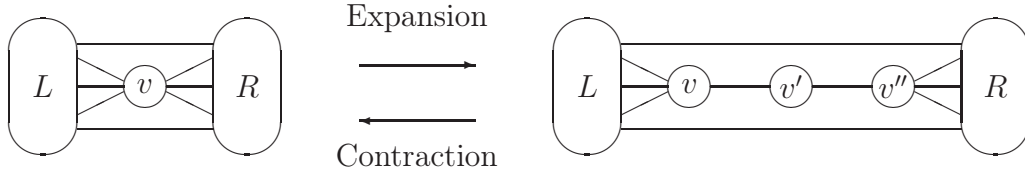


FIG. 9.2 – Appliquer des mesures selon X sur des qubits v' et v'' adjacents contracte les arêtes (v, v') et (v', v'') .

Preuve : On remarque que $|G\rangle$ peut être obtenu à partir de $|G \setminus v\rangle$, en ajoutant le sommet v et les arêtes entre (u, v) pour tout $u \in N_G(v)$. On en déduit, d'après la propriété 9.2 que :

$$\begin{aligned}
 |G; S\rangle &= Z_S \Pi_{u \in N_G(v)} \Lambda Z_{u,v} |G \setminus v\rangle |+\rangle_v \\
 &= \frac{Z_S}{\sqrt{2}} (\Pi_{u \in N_G(v)} \Lambda Z_{u,v} |G \setminus v\rangle |0\rangle_v + \Pi_{u \in N_G(v)} \Lambda Z_{u,v} |G \setminus v\rangle |1\rangle_v) \\
 &= \frac{Z_S}{\sqrt{2}} (|G \setminus v\rangle |0\rangle_v + Z_{N_G(v)} |G \setminus v\rangle |1\rangle_v) \\
 &= \frac{1}{\sqrt{2}} (|G \setminus v; S \setminus \{v\}\rangle |0\rangle_v + |G \setminus v; (S \setminus \{v\}) \Delta N_G(v)\rangle (Z_{S \cap \{v\}} |1\rangle_v))
 \end{aligned}$$

Ainsi, si le qubit v est mesuré selon Z , l'état résultant est soit $|G \setminus v; S \setminus \{v\}\rangle$, soit $|G \setminus v; (S \setminus \{v\}) \Delta N_G(v)\rangle$, et par conséquent égal à $|G \setminus v\rangle$, à un signe (connu) près. $\square$

Les mesures selon X , dans certaines conditions, peuvent également être interprétées de façon combinatoire. La mesure de deux sommets voisins selon X peut être interprétée, sous certaines conditions énoncées dans le lemme 9.2, comme une *contraction* du graphe (voir figure 9.2). A noter que cette transformation est réversible : la transformation inverse est appelée *expansion* du graphe.

Lemme 9.2 (Contraction d'arêtes) *Etant donné un graphe $G = (V, E)$ et un signe $S \subseteq V$, soient $v, v', v'' \in V$ tels que $N_G(v') = \{v, v''\}$, $N_G(v) \cap N_G(v'') = \{v'\}$ et $(v, v'') \notin E$ (voir figure 9.2). Appliquer une mesure selon X sur chacun des deux qubits v' et v'' transforme $|G; S\rangle$ en $|G/(v'', v')/(v', v); S'\rangle$ et correspond ainsi à la contraction successive des arêtes (v, v') puis (v', v'') .*

Preuve : Soient $v, v', v'' \in V$ tels que dans l'énoncé. Pour tout sous-ensemble $S \subseteq V$,

$$|G; S\rangle = \frac{Z_S}{2} \sum_{k, l \in \{0, 1\}} \left(Z_{v''}^k Z_{N_G(v) \setminus \{v'\}}^l |G/v''v/v'v\rangle \right) (|0\rangle + Z_{v'}^l |1\rangle)_{v'} (|0\rangle + Z_{v''}^k |1\rangle)_{v''}$$

et ainsi, si les qubits v' et v'' sont mesurés selon X , l'état résultant de la mesure est $|G/(v'', v')/(v', v)\rangle$ à un signe (connu) près. $\square$

Ainsi, une mesure selon Z peut être interprétée comme une suppression de sommet, et deux mesures selon X , sous certaines conditions (voir lemme 9.2) comme la contraction de deux arêtes. La mesure d'un seul qubit ne transforme pas, en général, un état graphe signé en un état graphe signé. Une mesure selon X seule ne peut donc pas être interprétée de façon combinatoire. En revanche, la mesure de deux qubits adjacents selon X peut toujours être interprétée de façon combinatoire, même si les conditions du lemme 9.2 ne sont pas vérifiées. En effet, la mesure de deux sommets adjacents selon X peut être interprétée comme un pivot sur ces deux sommets, suivie de la suppression des deux sommets :

Lemme 9.3 *Etant donné un graphe $G = (V, E)$ et un signe $S \subseteq V$, pour tout arête $(u, v) \in E$, appliquer une mesure selon X sur chacun des deux qubits u et v transforme $|G; S\rangle$ en $|G \wedge uv \setminus u \setminus v; S'\rangle$, où S' est entièrement déterminé à partir de S et du résultat classique des mesures.*

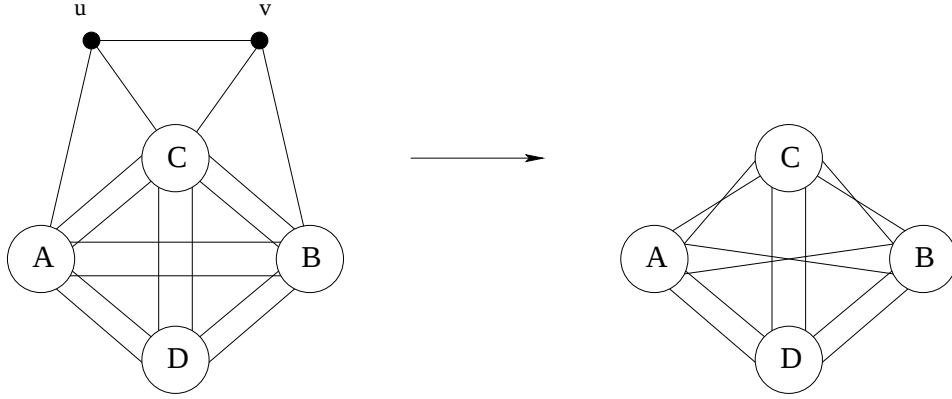


FIG. 9.3 – Pivot selon uv puis suppression des sommets u et v .

Preuve : D'après le lemme 9.3 et le lemme 9.1, l'application de $\sqrt{Y}^\dagger$ sur u et v suivie d'une mesure selon Z de u et v transforme $|G; S\rangle$ en $|G \wedge uv \setminus u \setminus v\rangle$ à un signe près. Or le lemme 7.2 de commutation entre observable et transformation unitaire implique que l'application $\sqrt{Y}^\dagger$ suivie d'une mesure selon Z est équivalente à la mesure directement selon $\sqrt{Y}Z\sqrt{Y}^\dagger = X$. $\square$

Force est de constater que l'état graphe obtenu après la mesure de qubits selon Z et de qubits adjacents selon X est représenté par un graphe qui est un pivot mineur du graphe initial. La caractérisation des mesures selon X étant conditionnée par l'adjacence des qubits mesurés, on peut se demander si l'application de

mesures selon Z et/ou X permet d'atteindre un état graphe qui n'est pas un pivot mineur du graphe initial. Nous montrons dans le théorème suivant que si une suite de mesures selon X et/ou Z transforme un état graphe signé en un état graphe signé, alors le second graphe est un pivot mineur du premier :

Théorème 9.2 *Pour tous graphes G_1, G_2 , et tous signes S_1, S_2 , si $|G_2; S_2\rangle$ est obtenu à partir de $|G_1; S_1\rangle$ en effectuant des mesures uniquement selon X ou Z , alors G_2 est un pivot mineur de G_1 .*

Preuve : Tout d'abord, on remarque que pour tout graphe $G = (V, E)$, et tout $K \subseteq V$ si $K \neq \emptyset$, alors $Z_K |G\rangle \neq |G\rangle$ et $Z_K |G\rangle \neq -|G\rangle$. Par contradiction, supposons que $Z_K |G\rangle = |G\rangle$. Pour $u \in K$, on a alors $Z_K |G\rangle = Z_K X_u Z_{N_G(u)} |G\rangle = -X_u Z_{N_G(u)} |G\rangle = -|G\rangle$. De même, si $Z_K |G\rangle = -|G\rangle$, pour $u \in K$, on a $Z_K |G\rangle = Z_K X_u Z_{N_G(u)} |G\rangle = -X_u Z_{N_G(u)} |G\rangle = |G\rangle$.

On remarque que les mesures effectuées sur G commutent car elles agissent sur des qubits différents. Toutes les mesures selon Z peuvent être interprétées comme des suppressions de sommets. Ainsi l'état obtenu après toutes les mesures selon Z est un état graphe $|G'\rangle$.

Nous montrons que pour tout sommet de G' mesuré selon X , soit ce sommet est isolé, et la mesure peut alors être interprétée comme une suppression de sommet, soit il existe un sommet adjacent dans G' qui est également mesuré selon X et les deux mesures peuvent alors être interprétées comme un pivot suivi d'une suppression des deux sommets. Une récurrence permet alors de conclure.

Soit u un sommet de G' , non isolé, tel que u est mesuré selon X . Il existe un état $|\varphi\rangle$, tel que la mesure transforme l'état $|G'; S\rangle$ en $|\varphi\rangle \otimes (Z_u^l |+\rangle_u)$, où $l \in \{-1, 1\}$ est le résultat classique de la mesure. $X_u Z_{N_{G'}(u)}$ vérifie $X_u Z_{N_{G'}(u)} |G'; S'\rangle = Z_{S' \cap \{u\}} |G'; S'\rangle$. D'après le lemme 7.2 de commutation entre observable et transformation unitaire, appliquer $X_u Z_{N_{G'}(u)}$ puis la mesure selon X_u (ce qui produit l'état $Z_{S \cap \{u\}} |\varphi\rangle \otimes (Z_u^l |+\rangle_u)$) est équivalent à mesurer selon $(X_u Z_{N_{\bar{G}}(u)})^\dagger X_u (X_u Z_{N_{\bar{G}}(u)}) = X_u$ puis appliquer $X_u Z_{N_{G'}(u)}$. On en déduit que $X_u Z_{N_{G'}(u)} |\varphi\rangle \otimes (Z_u^l |+\rangle_u) = Z_{S \cap \{u\}} |\varphi\rangle \otimes (Z_u^l |+\rangle_u)$. En particulier, $Z_{N_{G'}(u)} |\varphi\rangle = |\varphi\rangle$, donc si toutes les autres mesures selon X sont sur des sommets qui ne sont pas adjacents, alors l'opérateur $Z_{N_{G'}(u)}$ commute avec les mesures, ce qui contredit la remarque initiale. Un des sommets adjacents à u dans G' est donc également mesuré selon X . Ces deux mesures transforment $|G'; S'\rangle$ en $|G''; S''\rangle$, où G'' est un pivot mineur de G . On conclut par récurrence que l'état graphe final $|G_2; S_2\rangle$ est tel que G_2 est un pivot mineur de G_1 . $\square$

Le théorème précédent renforce l'interprétation des transformations engendrées par des mesures selon X et Z , puisque l'état graphe obtenu est nécessairement un pivot mineur du graphe initial. En revanche, étant donné deux graphes G et G'

tels que G' est un pivot mineur de G , il n'est pas toujours possible de transformer l'état graphe $|G\rangle$ en $|G'\rangle$, à un signe près, en utilisant uniquement des mesures selon X et Z . Par exemple, $G \wedge uv$, pour (u, v) une arête de G , est un pivot mineur de G et n'est pas a priori isomorphe à G , et il n'est pas possible de transformer $|G; S\rangle$ en $|G \wedge uv; S'\rangle$ par des mesures selon X ou Z car les deux graphes ont le même nombre de sommets.

Alors que les mesures destructrices selon X et Z consomment les qubits mesurés, leur interprétation combinatoire est donc naturellement liée au concept de mineur, et plus particulièrement au pivot mineur. En revanche, les mesures non destructrices et non locales peuvent être utilisées pour créer de l'intrication :

Lemme 9.4 (Ajout d'un sommet) *Etant donné un graphe $G = (V, E)$, un signe S , un qubit $v \notin V$ et un sous ensemble de sommets $K \subseteq V$, une mesure selon Z du qubit v suivie d'une mesure selon $X_v Z_K$ transforme l'état $|G; S\rangle$ en $|G'; S'\rangle$, où $G' = (V \cup \{v\}, E \cup \{(u, v), u \in K\})$ est obtenu à partir du graphe G en ajoutant le sommet v et les arêtes entre v et K .*

Preuve : L'état graphe $|G'\rangle$ peut être obtenu, à un signe près, à partir de $|G; S\rangle$ en mesurant un nouveau qubit v selon X , puis en ajoutant les arêtes entre v et les sommets de K en appliquant ΛZ entre les qubits correspondants. En effet, après la mesure selon X , le qubit v est dans l'état $Z_v^l |+\rangle$, où l est le résultat de la mesure. Or, $|G; S\rangle \otimes (Z_v^l |+\rangle) = |G_1; S_1\rangle$, où $G_1 = (V \cup \{v\}, E)$. Les applications ΛZ permettent d'ajouter les arêtes entre v et K (voir la propriété 9.2).

On remarque qu'une mesure préliminaire du qubit v selon Z peut être ajoutée. D'après le lemme 7.2 de commutation entre observable et transformation unitaire, une mesure selon Z_v puis selon X_v suivie de $\Pi_{u \in K} \Lambda Z_{v,u}$ est équivalente à mesurer selon Z , puis appliquer $\Pi_{u \in K} \Lambda Z_{v,u}$, et mesurer selon $(\Pi_{u \in K} \Lambda Z_{v,u}) X_v (\Pi_{u \in K} \Lambda Z_{v,u})^\dagger = X_v Z_K$. Or l'état du qubit v après la mesure selon Z étant $|0\rangle$ ou $|1\rangle$, l'application de $\Pi_{u \in K} \Lambda Z_{v,u}$ ne modifie que le signe de $|G; S\rangle$, et peut donc être supprimée.

Ainsi une mesure de v selon Z suivie d'une mesure selon $X_v Z_K$ transforme $|G; S\rangle$ en $|G'; S'\rangle$, où S' dépend de S et du résultat des mesures. $\square$

9.4 Conclusion

Nous avons présenté le formalisme des états graphes, et introduit une généralisation de ce formalisme : les *états graphes signés*. Un état graphe signé est une représentation compacte de certains états quantiques, permettant une caractérisation combinatoire d'un ensemble de transformations quantiques. L'introduction du *signe* permet une interprétation formelle, non seulement de certaines transformations unitaires comme pour les états graphes, mais aussi de certaines mesures sur un ou plusieurs qubits.

Les états graphes ont différentes applications en informatique quantique, incluant la représentation de l'intrication et le calcul quantique par consommation d'intrication. Pour ce dernier modèle de calcul quantique, où un état graphe est transformé au cours de l'exécution par des mesures sur un qubit, le chapitre 10 présente un modèle formel : le m -calcul.

Chapitre 10

Calcul quantique par consommation d'intrication

10.1 Introduction

Le modèle par consommation d'intrication a été introduit par Briegel et Rausendorf [RB00]. Une exécution dans ce modèle consiste à appliquer des mesures sur un qubit à un état graphe. Un mécanisme de correction à base d'opérateurs de Pauli permet de rendre une telle exécution déterministe, alors que les mesures quantiques sont non déterministes. La possibilité d'appliquer les corrections pour rendre le calcul déterministe dépend de la géométrie de l'état graphe utilisé et des mesures appliquées.

Un modèle formel pour le calcul par consommation d'intrication, le m -calcul¹ a été introduit par Danos, Kashefi et Panangaden [DKP04a, DKP04b]. Afin de savoir si un terme du m -calcul (un m -terme) admet ou non une stratégie de correction le rendant déterministe, un critère de flot a été développé par Danos et Kashefi [DK05]. Si un terme admet un flot, alors il existe une stratégie rendant l'exécution déterministe.

Dans ce chapitre, nous introduisons le m -calcul, ainsi que la condition de flot, suffisante pour que soit déterministe le calcul par consommation d'intrication. Puis, nous introduisons un nouveau critère de flot, plus général, qui permet de montrer que la condition de flot n'est pas une condition nécessaire. Ce nouveau critère permet également dans certains cas une diminution du temps d'exécution d'un m -terme.

Dans un second temps, nous proposons une généralisation du m -calcul autorisant des mesures dans plusieurs plans, au lieu d'un plan et un axe dans le m -calcul. Nous établissons également une condition de déterminisme pour ce nouveau mo-

¹*measurement calculus*

dèle.

Enfin, les utilisations de différents états graphes initiaux sont comparées. Traditionnellement, et suivant en cela Briegel et Raussendorf, une grille rectangulaire est utilisée comme état graphe initial. Nous montrons que l'utilisation d'une grille triangulaire permet de diminuer les ressources nécessaires au calcul par consommation d'intrication, en terme d'observables. Ce résultat permet également de démontrer une propriété de théorie des graphes : tout graphe est pivot mineur d'une grille triangulaire. Ce résultat est le premier résultat de théorie des graphes, à la connaissance de l'auteur, prouvé en utilisant des résultats d'informatique quantique.

10.2 m -calcul

Le m -calcul, [DKP04a], est un langage formel pour le calcul par consommation d'intrication introduit par Briegel et Raussendorf [RB00, RB02b, RBB02]. Un tel calcul consiste à mesurer individuellement les qubits d'un état graphe.

10.2.1 Syntaxe

Etant donné un ensemble V de qubits, quatre opérations, appelées *commandes* sont définies :

- N_i est une préparation du qubit i dans l'état $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$;
- E_{ij} est la transformation unitaire ΛZ sur les qubits i, j : $E_{ij} = \Lambda Z_{i,j}$;
- M_i^α est une mesure destructrice du qubit i dans la base $\{\frac{1}{\sqrt{2}}(|0\rangle + e^{i\alpha}|1\rangle), \frac{1}{\sqrt{2}}(|0\rangle - e^{i\alpha}|1\rangle)\}^2$;
- X_i et Z_i sont les opérateurs de Pauli X et Z appliqués au qubit i .

Dans le modèle par consommation d'intrication, les commandes N_i et $E_{i,j}$ permettent la création de l'état graphe, alors que les mesures individuelles le consomment. Les opérations de corrections X_i et Z_i sont utilisées pour rendre globalement déterministe une évolution composée de mesures qui sont par nature non déterministes.

Puisque les qubits sont mesurés de façon destructrice, le résultat de la mesure d'un qubit j peut être représenté de façon non ambiguë par $s_j \in \{0, 1\}$. Les opérations de correction X_i et Z_i peuvent dépendre des résultats de mesures précédentes, ainsi on écrit $X_i^{s_j}$ et $Z_i^{s_j}$ avec $X_i^0 = Z_i^0 = Id_i$, $X_i^1 = X_i$ et $Z_i^1 = Z_i$.

Un terme du m -calcul (un m -terme) est un quadruplet (V, I, F, A) , avec V un ensemble de qubits, $I \subseteq V$ et $F \subseteq V$ représentent respectivement les qubits

² M_i^α est une mesure décrite par la transformation admissible $(\frac{1}{\sqrt{2}}(|\rangle\langle 0| + e^{i\alpha}|\rangle\langle 1|), \frac{1}{\sqrt{2}}(|\rangle\langle 0| - e^{i\alpha}|\rangle\langle 1|))$.

d'entrée et de sortie, et A est une suite finie de commandes, lue de droite à gauche. Un m -terme est exécutable si et seulement si :

- Aucune commande ne dépend du résultat classique associé à un qubit qui n'a pas encore été mesuré ;
- Aucune commande n'agit sur un qubit déjà mesuré ou pas encore préparé ;
- Un qubit i est préparé (mesuré) si et seulement si i n'est pas un qubit d'entrée (sortie).

Par exemple, $(\{1, 2\}, \{1\}, \{2\}, X_2^{s_1} M_1^0 E_{12} N_2)$ est un terme du m -calcul. Ce terme consiste à initialiser le qubit 2 dans l'état $|+\rangle$, à appliquer ΛZ sur les qubits 1 et 2, puis à mesurer le qubit 1 selon X . Enfin, une correction est appliquée sur le qubit 2, dépendant du résultat de la mesure. Ce schéma, appelé parfois téléportation à 1 qubit est proche du transfert d'état (chapitre 7), sauf qu'ici une transformation unitaire ΛZ est utilisée pour créer de l'intrication entre les qubits, alors qu'une mesure sur deux qubits est utilisée dans le transfert d'état.

10.2.2 Sémantique

L'exécution d'un terme du m -calcul (V, I, F, A) produit $|V \setminus F|$ résultats classiques. Ces résultats classiques peuvent être représentés par un mot binaire $s \in \{0, 1\}^{|V \setminus F|}$. Alors que l'article original de Danos, Kashefi et Panangaden associe à chaque terme un super-opérateur, nous proposons une adaptation de cette sémantique dénotationnelle, associant à chaque terme une transformation admissible :

Définition 10.1 *Pour tout m -terme $\mathfrak{P} = (V, I, F, A)$, soit $n = |V \setminus F|$. Pour tout $K \subseteq V$, $\mathcal{H}_{\{0,1\}^{|K|}}$ est noté $\mathcal{H}^K$. Pour chaque suite de résultats s classiques, l'opération effectuée par la suite de commandes A est une opération linéaire $\llbracket A \rrbracket(s)$. La sémantique $\llbracket \mathfrak{P} \rrbracket$ du terme est la transformation admissible formée par toutes ces évolutions possibles.*

Pour toute suite de commandes t et tout $\Gamma \subseteq V$, $\llbracket t \rrbracket_\Gamma : \{0, 1\}^n \rightarrow \mathcal{L}(\mathcal{H}^\Gamma, \mathcal{H}^F)$:

$$\llbracket tN_j \rrbracket_\Gamma = \lambda s. (\llbracket t \rrbracket_{\Gamma \cup \{j\}}(s) (|+\rangle_j \langle +| \otimes Id_{\mathcal{H}^\Gamma}))$$

$$\llbracket tE_{ij} \rrbracket_\Gamma = \lambda s. (\llbracket t \rrbracket_\Gamma(s) (\Lambda Z_{i,j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i,j\}}}))$$

$$\llbracket tM_j^{\alpha_j} \rrbracket_\Gamma = \lambda s. \left(\llbracket t \rrbracket_{\Gamma \setminus \{j\}}(s) \left(\frac{|\langle 0_j | + (-1)^{s_j} e^{i\alpha_j} | \rangle \langle 1_j |}{\sqrt{2}} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{j\}}} \right) \right)$$

$$\llbracket tX_i^{s_j} \rrbracket_\Gamma = \lambda s. (\llbracket t \rrbracket_\Gamma(s) (X_i^{s_j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i\}}}))$$

$$\llbracket tZ_i^{s_j} \rrbracket_\Gamma = \lambda s. (\llbracket t \rrbracket_\Gamma(s) (Z_i^{s_j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i\}}}))$$

$$\llbracket \mathfrak{P} \rrbracket \in T(\mathcal{H}^I, \mathcal{H}^F) :$$

$$\llbracket \mathfrak{P} \rrbracket = (\llbracket A \rrbracket_I(s))_{s \in \{0,1\}^n}$$

La sémantique définie par Danos, Kashefi et Panangaden dans [DKP04a] est en fait le super-opérateur $\mathcal{X}^{\natural}(\llbracket \mathfrak{P} \rrbracket) = \lambda \rho. \sum_{s \in \{0,1\}^n} \llbracket A \rrbracket_I(s) \rho \llbracket A \rrbracket_I^{\dagger}(s)$, qui est une interprétation exacte de $\llbracket \mathfrak{P} \rrbracket$.

Un terme du m -calcul est dit *déterministe* si l'état des qubits de sortie après l'exécution successive de toutes les commandes dépend de l'état des qubits d'entrée avant l'exécution, mais est indépendant des résultats classiques des mesures effectuées.

Définition 10.2 (Déterminisme) *Un terme $\mathfrak{P} = (V, I, F, A)$ est déterministe si et seulement si il existe une transformation admissible, composée d'un seul opérateur U , tel que $\llbracket \mathfrak{P} \rrbracket \equiv (U)$.*

Un terme est dit *uniformément déterministe* si et seulement si il est déterministe pour toutes les valeurs possibles des angles α des mesures M_i^{α} utilisées dans A .

10.2.3 Forme standard

Le m -calcul est un ensemble de règles de réécriture permettant de transformer un m -terme, en préservant sa sémantique. L'objectif de ce calcul est d'atteindre une forme normale, appelée forme standard. Nous présentons ici les règles de réécriture définies dans [DKP04a]. Ce système de réécriture a été prouvé terminant et confluent. La forme standard est un terme où les commandes sont regroupées par type : toutes les commandes N_i sont au début de la suite, puis les commandes $E_{i,j}$, puis M_i^{α} et les commandes de corrections X_i^{sj} et Z_i^{sj} .

Les règles de réécritures sont utilisées pour repousser en fin de calcul les facteurs correctifs.

$$\begin{aligned} E_{ij} X_i^s &\rightarrow X_i^s Z_j^s E_{i,j} \\ E_{ij} Z_i^s &\rightarrow Z_i^s E_{i,j} \\ {}^t[M_i^{\alpha}]^s X_i^r &\rightarrow {}^t[M_i^{\alpha}]^{s+r} \\ {}^t[M_i^{\alpha}]^s Z_i^r &\rightarrow {}^{t+r}[M_i^{\alpha}]^s \end{aligned}$$

${}^t[M_i^{\alpha}]^s$ est une notation pour représenter la commande $M_i^{(-1)^s \alpha + t\pi}$. Ces quatre règles préservent la sémantique des m -termes. Les deux premières règles découlent directement des propriétés des transformations unitaires X, Z et ΛZ . Les deux règles suivantes consistent à modifier l'observable selon lequel la mesure est effectuée. La préservation de la sémantique par ces deux dernières règles peut être prouvée en utilisant le lemme 7.2 (chapitre 7) sur la commutation observable / transformation unitaire.

A ces règles s'ajoutent des règles permettant de commuter des commandes agissant sur des qubits différents. Le système de réécriture obtenu est le m -calcul. Pour tout terme $\mathfrak{P}$, il existe un terme sous forme standard $\mathfrak{P}'$ tel que $\mathfrak{P} \rightarrow^* \mathfrak{P}'$.

Comme les opérations $E_{i,j}$ commutent entre elles, leur ordre n'importe pas. Ainsi, pour tout terme sous forme standard $\mathfrak{P}$, il existe un graphe G tel que :

$$\mathfrak{P} \rightarrow^* (V, I, F, C_F(\Pi_{i \in V \setminus F}^> M_i^{\alpha_i}) E_G N_{V \setminus I})$$

où C_F est composé de commandes de correction, $\Pi^>$ signifie que le produit n'est pas commutatif et $E_{(V,D)} = \Pi_{u,v \in D} E_{u,v}$. (G, I, F) est appelé *géométrie* de $\mathfrak{P}$.

Les règles de réécriture permettant d'atteindre la forme standard préservent le déterminisme, le déterminisme uniforme et l'exécutabilité d'un m -terme.

10.3 Condition de flots

10.3.1 Condition de flot simple

Danos et Kashefi [DK05] ont introduit une condition suffisante pour qu'un terme soit uniformément déterministe : un terme est uniformément déterministe si la *géométrie* associée admet un *flot*.

Définition 10.3 *Un flot $(f, \prec)$ pour une géométrie (G, I, F) est composé d'une fonction $f : V \setminus F \rightarrow V \setminus I$ et d'un ordre partiel $\prec$ sur V tel que pour tout $x \in V \setminus F$:*

- x et $f(x)$ sont adjacents dans G ;
- $x \prec f(x)$;
- pour tout $y \in N_G(f(x))$, $x \prec y$.

Théorème 10.1 [DK05] *Un terme du m -calcul est uniformément déterministe si et seulement si la géométrie associée admet un flot.*

Ainsi, si une géométrie admet un flot, alors un calcul déterministe peut être mené. De plus, la relation d'ordre partiel associée au flot donne le nombre d'étapes permettant d'effectuer le calcul. En effet, en supposant que l'état graphe est préparé, certaines mesures sur un qubit peuvent être effectuées en parallèle. Ce parallélisme est limité par la stratégie de correction. Danos et Kashefi ont montré que si une géométrie admet un flot $(f, \prec)$, alors k étapes de mesure suffisent pour effectuer toutes les mesures, où k est la longueur de la plus longue suite strictement décroissante selon $\prec$. k est appelé la profondeur du flot.

La première étape de mesure consiste à mesurer en parallèle tous les qubits qui correspondent à des éléments minimaux pour $\prec$: $W_1 = \sqcup V$. La seconde étape consiste à mesurer, parmi les qubits restants, les plus petits, et cœtera. Ainsi, les qubits mesurés à l'étape l sont $W_l = \sqcup (V \setminus (\cup_{k < l} W_k))$.

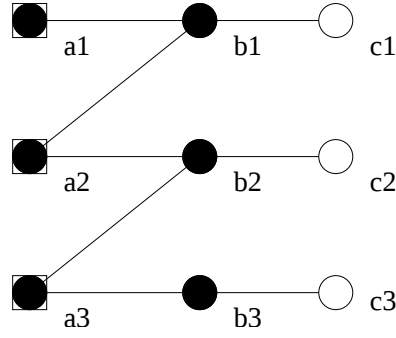


FIG. 10.1 – Géométrie admettant un flot de profondeur 5.

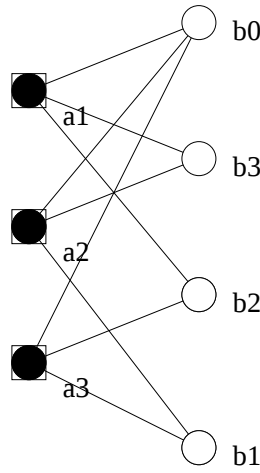


FIG. 10.2 – Géométrie n'admettant pas de flot.

Exemple 10.1 La géométrie $(G_1, \{a_1, a_2, a_3\}, \{c_1, c_2, c_3\})$ avec G_1 décrit dans la figure 10.1 admet un flot : $f(a_i) = b_i$ et $f(b_i) = c_i$ avec $a_1 \prec \{b_1, a_2\} \prec \{c_1, b_2, a_3\} \prec \{c_2, b_3\} \prec c_3$. La profondeur est 5. La structure du graphe permet de conclure que tous les flots de ce graphe sont de profondeur au moins 5. En effet, la seule fonction f possible est $f(a_i) = b_i$ et $f(b_i) = c_i$, de plus chaque b_i est strictement plus grand que a_i , chaque a_i est plus grand que b_{i-1} . On en déduit que la profondeur minimum est 5.

La géométrie $(G_2, \{a_1, a_2, a_3\}, \{b_0, b_1, b_2, b_3\})$ avec G_2 décrit dans la figure 10.2 n'admet pas de flot.

En effet, par contradiction, l'image par f d'au moins deux des sommets d'entrée est un sommet de degré 2, supposons par exemple que $f(a_1) = b_3$ et $f(a_2) = b_1$. On en déduit que $a_1 \prec \{b_3, a_2\} \prec \{b_1, a_3\}$. Or, $f(a_3) = b_0$ ou $f(a_3) = b_2$, dans les deux cas on a $a_3 \prec a_1$ ce qui contredit la relation d'ordre.

10.3.2 Condition de flot généralisé

Est ce que tout terme uniformément déterministe admet un flot ? Est ce que pour tout terme uniformément déterministe, il existe un flot dont la profondeur est le nombre minimum d'étapes nécessaires ?

Nous proposons une généralisation de la condition de flot, qui permet de répondre par la négative à ces deux questions.

Etant donnée une fonction $f : V \setminus F \rightarrow \mathbb{N}$, pour tout $u \in V$, soient :

- $A(u) = \{v \in V \setminus F, v \neq u \wedge f(v) \leq f(u)\}$
- $S(u) = \{v \in V \setminus F, f(v) > f(u)\}$

La fonction f induit un ordre sur les éléments de V , $A(u)$ désigne les antécédents de u , et $S(u)$ désigne les successeurs.

On remarque que pour tout $v \in V \setminus F$, $\{v\}$, $A(v)$ et $S(v)$ forment une partition de $V \setminus F$.

Définition 10.4 (Flot généralisé) Une géométrie (G, I, F) admet un flot généralisé si et seulement si il existe $f : V \setminus F \rightarrow \mathbb{N}$ telle que pour tout $v \in V \setminus F$, il existe un ensemble correctif $C(v) \subseteq (S(v) \cup F) \setminus I$ satisfaisant :

- $v \in \text{Imp}_G(C(v))$,
- $A(v) \subseteq \text{Pair}_G(C(v))$.

où $\text{Imp}_G(X) = \{u \in V \setminus X : N_G(v) \cap X = 1 \pmod{2}\}$ et $\text{Pair}_G(X) = \{u \in V \setminus X : N_G(v) \cap X = 0 \pmod{2}\}$. $\text{Imp}_G(X)$ ($\text{Pair}_G(X)$) est l'ensemble des voisins impairs (pairs) de X .

Théorème 10.2 Un terme est uniformément déterministe si la géométrie associée admet un flot généralisé.

La preuve, comme dans le cas de la condition de flot simple, est fondée sur l'idée selon laquelle l'application d'un opérateur de correction $Z_i^{s_i}$ juste avant la mesure du qubit i permettrait d'obtenir une évolution déterministe. Un tel terme, bien qu'uniformément déterministe, n'est pas exécutable. En revanche, ce terme non exécutable peut être transformé en un terme exécutable, en préservant le déterminisme uniforme. Cette transformation s'appuie sur une des propriétés combinatoires des états graphes : pour tout graphe $G = (V, E)$ et tout $u \in V$, $|G\rangle = X_u Z_{N_G(u)} |G\rangle$. Le lemme suivant donne une version équivalente de cette propriété dans le cadre du m -calcul :

Lemme 10.1 Pour toute géométrie (G, I, F) , et tout $u \in V \setminus I$,

$$E_G N_{V \setminus I} = X_u Z_{N_G(u)} E_G N_{V \setminus I}$$

Preuve : $E_G N_{V \setminus I}$ est une opération linéaire de $\mathcal{H}_I = \mathcal{H}_{\{0,1\}^{|I|}}$ dans $\mathcal{H}_V$. L'ensemble $\{Z_K N_I, K \subseteq I\}$ est une base de $\mathcal{H}_I$. Soit $K \subseteq V$,

$$\begin{aligned} E_G N_{V \setminus I} Z_K N_I &= Z_K E_G N_I \\ &= Z_K |G\rangle \\ &= Z_K X_u Z_{N_G(u)} |G\rangle \\ &= X_u Z_{N_G(u)} E_G N_{V \setminus I} Z_K N_I \end{aligned}$$

□

Preuve du théorème 10.2 : Soit (G, I, F) une géométrie et $f : V \setminus F \rightarrow \mathbb{N}$ telle que pour tout $v \in V \setminus F$, il existe un ensemble correctif $C(v) \subseteq S(v) \setminus I$ avec :

- $v \in \text{Imp}_G(C(v))$,
- $A(v) \subseteq \text{Pair}_G(C(v))$.

Le terme $\mathfrak{P} = (V, I, F, A)$, avec $A = (\Pi_{i \in V \setminus F}^> M_i^{\alpha_i} Z_i^{s_i}) E_G N_{V \setminus I}$, est uniformément déterministe mais non exécutable. De plus, la géométrie associée à $\mathfrak{P}$ est (G, V, I) . La stratégie permettant de transformer le terme $\mathfrak{P}$ en un terme exécutable consiste à utiliser, pour chaque sommet mesuré v , le lemme 10.1 pour chacun des sommets de l'ensemble correctif $C(v)$. En effet, $E_G N_{V \setminus I} = \Pi_{u \in C(v)} X_u Z_{N_G(u)} E_G N_{V \setminus I}$. Puisque $Z^2 = Id$, on a, à une phase globale près non importante :

$$E_G N_{V \setminus I} = X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s_v} E_G N_{V \setminus I}$$

La définition du flot généralisé garantit que l'opérateur $X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s(v)}$ commute avec les mesures effectuées avant celle de v . De plus, $v \in \text{Imp}_G(C(v))$ permet de compenser le facteur correctif $Z_v^{s_v}$ introduit pour obtenir le déterminisme uniforme. Enfin, en utilisant les règles de réécriture, l'opérateur $X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus \{v\}}^{s(v)}$ peut être intégré aux mesures qui suivent celles de v , ou à un facteur correctif global agissant sur les qubits de sortie.

Ainsi, par récurrence sur la position de la mesure $M_v^{\alpha_v}$ dans le terme, on montre que ce terme peut être transformé en un terme exécutable :

$$\begin{aligned} A &= C_F(\Pi_{i \in S(v)}^> M_i^{\alpha_i} Z_i^{s_i}) M_v^{\alpha_v} Z_v^{s_v} (\Pi_{j \in A(v)}^> M_j^{\alpha_j}) E_G N_{V \setminus I} \\ &= C_F(\Pi_{i \in S(v)}^> M_i^{\alpha_i} Z_i^{s_i}) M_v^{\alpha_v} Z_v^{s_v} (\Pi_{j \in A(v)}^> M_j^{\alpha_j}) X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s_v} E_G N_{V \setminus I} \\ &= C_F(\Pi_{i \in S(v)}^> M_i^{\alpha_i} Z_i^{s_i}) X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus \{v\}}^{s_v} M_v^{\alpha_v} (\Pi_{j \in A(v)}^> M_j^{\alpha_j}) E_G N_{V \setminus I} \\ &\rightarrow^* C'_F(\Pi_{i \in S(v)}^> M_i^{\alpha'_i} Z_i^{s'_i}) M_v^{\alpha_v} (\Pi_{j \in A(v)}^> M_j^{\alpha_j}) E_G N_{V \setminus I} \end{aligned}$$

□

La condition de flot simple, introduite par Danos et Kashefi, peut être vue comme un cas particulier du flot généralisé, dans le cas où l'ensemble correctif est composé d'un seul élément.

Alors que la relation d'ordre partiel $\prec$ du flot donne un nombre suffisant d'étapes pour réaliser le calcul par consommation d'intrication, ce nombre d'étapes, appelé profondeur, est donné par la taille de l'image de f dans le cas du flot généralisé. En effet, tous les éléments de $V \setminus F$ ayant la même image par f peuvent être mesurés en parallèle, et si $f(v) < f(u)$, alors v peut être mesuré avant u .

Le théorème 10.2 permet de répondre à deux questions ouvertes : existe-t-il des termes uniformément déterministes n'admettant pas de flot simple ? De plus la profondeur donnée par le flot simple est-elle optimale ?

La géométrie de la figure 10.2 n'admet pas de flot simple. En revanche, elle admet un flot généralisé : pour tout i , $f(a_i) = 1$. De plus, pour tout i , $C(a_i) = \{b_0, b_i\}$. La profondeur de ce flot généralisé est 2. L'existence d'un flot simple n'est donc pas nécessaire au déterminisme uniforme.

La géométrie de la figure 10.1 admet un flot simple et un flot généralisé. Le flot généralisé est le suivant : pour tout i , $f(a_i) = 1$ et $f(b_i) = 2$. De plus, $C(a_1) = \{b_1\}$, $C(b_1) = \{c_1\}$ et pour tout $i > 1$, $C(a_i) = \{b_{i-1}, b_i\}$, $C(b_i) = \{c_i\}$. La profondeur de ce flot généralisé est 2, alors que la profondeur minimale du flot simple est 5.

Plus généralement, soit $H_n = (V_n, E_n)$ tel que $V_n = \{a_i, b_i, c_i\}_{i \in 1 \dots n}$, et $E_n = \{(a_i, b_i), (b_i, c_i)\}_{i \in 1 \dots n} \cup \{(a_{i+1}, b_i)\}_{i \in 1 \dots n-1}$. On remarque que l'exemple de la figure 10.1 est H_3 . Pour tout $n > 0$, H_n a une profondeur de flot simple minimale de $n+1$, alors qu'il existe un flot généralisé de profondeur 2. L'utilisation du flot généralisé peut donc permettre d'augmenter le parallélisme du calcul par consommation d'intrication.

La condition de flot généralisé permet donc de montrer que la condition de flot simple n'est pas nécessaire au déterminisme uniforme. La nécessité de la condition de flot généralisé est une conjecture :

Conjecture 10.1 *Un terme est uniformément déterministe si et seulement si la géométrie associée admet un flot généralisé.*

Un autre problème ouvert concerne la profondeur du flot généralisé. La profondeur minimale d'un flot généralisé correspond-elle au nombre minimum d'étapes de mesure nécessaires à un calcul uniformément déterministe ?

10.4 *m-calcul 3P*

Dans le modèle du calcul par consommation d'intrication, deux types de mesures sont autorisés, d'une part des mesures selon l'observable Z , c'est-à-dire dans la base standard $\{|0\rangle, |1\rangle\}$, d'autre part des mesures selon un observable de la forme $\cos(\alpha)X + \sin(\alpha)Y$, ces dernières étant appelées mesures dans le plan (X, Y) .

Comme les mesures selon Z peuvent être interprétées directement en termes de transformations de l'état graphe initial, seules les mesures dans le plan (X, Y) sont présentes dans le m -calcul.

A partir de ce modèle initial, limité aux mesures selon Z et dans le plan (X, Y) , de nouvelles propositions, composées de suite de mesures selon des plans (X, Z) et (Y, Z) sont récemment apparues [BB06], qui étendent les possibilités offertes par le calcul par consommation d'intrication.

Le modèle de calcul par consommation d'intrication peut être ainsi étendu au cas où les mesures sont selon un observable de l'un des trois plans : (X, Y) , (X, Z) et (Y, Z) . Le m -calcul peut alors être généralisé en un m -calcul $\mathcal{3P}$ permettant l'utilisation des mesures selon les trois plans. La syntaxe, la sémantique et la standardisation du m -calcul sont alors aussi généralisées. Nous nous intéressons également à la généralisation de la condition de flot dans le cadre du m -calcul $\mathcal{3P}$.

10.4.1 Définitions

La syntaxe du m -calcul $\mathcal{3P}$ est la même que celle du m -calcul, à l'exception des commandes de mesure, où la mesure M_i^α dans le plan (X, Y) est remplacée par trois commandes :

- $M_i^{Z,\alpha}$ est une mesure dans la base $\{|+^{Z,\alpha}\rangle, |-^{Z,\alpha}\rangle\}$, avec $|+^{Z,\alpha}\rangle$ et $|-^{Z,\alpha}\rangle$ tels que $(\cos(\alpha)X + \sin(\alpha)Y)|+^{Z,\alpha}\rangle = |+^{Z,\alpha}\rangle$ et $(\cos(\alpha)X + \sin(\alpha)Y)|-^{Z,\alpha}\rangle = -|-^{Z,\alpha}\rangle$.
- $M_i^{X,\alpha}$ est une mesure dans la base $\{|+^{X,\alpha}\rangle, |-^{X,\alpha}\rangle\}$, avec $|+^{X,\alpha}\rangle$ et $|-^{X,\alpha}\rangle$ tels que $(\cos(\alpha)Y + \sin(\alpha)Z)|+^{X,\alpha}\rangle = |+^{X,\alpha}\rangle$ et $(\cos(\alpha)Y + \sin(\alpha)Z)|-^{X,\alpha}\rangle = -|-^{X,\alpha}\rangle$.
- $M_i^{Y,\alpha}$ est une mesure dans la base $\{|+^{Y,\alpha}\rangle, |-^{Y,\alpha}\rangle\}$, avec $|+^{Y,\alpha}\rangle$ et $|-^{Y,\alpha}\rangle$ tels que $(\cos(\alpha)Z + \sin(\alpha)X)|+^{Y,\alpha}\rangle = |+^{Y,\alpha}\rangle$ et $(\cos(\alpha)Z + \sin(\alpha)X)|-^{Y,\alpha}\rangle = -|-^{Y,\alpha}\rangle$.

Sur le modèle de la sémantique dénotationnelle du m -calcul, la sémantique dénotationnelle du m -calcul $\mathcal{3P}$ est la suivante :

Définition 10.5 *Pour tout terme $\mathfrak{P} = (V, I, F, A)$, soit $n = |V \setminus F|$.*

Pour toute suite de commandes t et tout $\Gamma \subseteq V$, $\llbracket t \rrbracket_\Gamma : \{0, 1\}^n \rightarrow \mathbf{L}(\mathcal{H}^\Gamma, \mathcal{H}^F)$:

$$\begin{aligned}
\llbracket tN_j \rrbracket_\Gamma &= \lambda s. (\llbracket t \rrbracket_{\Gamma \cup \{j\}}(s)(|+_j\rangle\langle| \otimes Id_{\mathcal{H}^\Gamma})) \\
\llbracket tE_{ij} \rrbracket_\Gamma &= \lambda s. (\llbracket t \rrbracket_\Gamma(s)(\Lambda Z_{i,j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i,j\}}})) \\
\llbracket tM_j^{X,\alpha_j} \rrbracket_\Gamma &= \lambda s. \left(\llbracket t \rrbracket_{\Gamma \setminus \{j\}}(s) \left(\left((1-s_j)|\rangle\langle+_j^{X,\alpha_j}| + s_j|\rangle\langle-_-^{X,\alpha_j}| \right) \otimes Id_{\mathcal{H}^{\Gamma \setminus \{j\}}} \right) \right) \\
\llbracket tM_j^{Y,\alpha_j} \rrbracket_\Gamma &= \lambda s. \left(\llbracket t \rrbracket_{\Gamma \setminus \{j\}}(s) \left(\left((1-s_j)|\rangle\langle+_j^{Y,\alpha_j}| + s_j|\rangle\langle-_-^{Y,\alpha_j}| \right) \otimes Id_{\mathcal{H}^{\Gamma \setminus \{j\}}} \right) \right) \\
\llbracket tM_j^{Z,\alpha_j} \rrbracket_\Gamma &= \lambda s. \left(\llbracket t \rrbracket_{\Gamma \setminus \{j\}}(s) \left(\left((1-s_j)|\rangle\langle+_j^{Z,\alpha_j}| + s_j|\rangle\langle-_-^{Z,\alpha_j}| \right) \otimes Id_{\mathcal{H}^{\Gamma \setminus \{j\}}} \right) \right) \\
\llbracket tX_i^{s_j} \rrbracket_\Gamma &= \lambda s. (\llbracket t \rrbracket_\Gamma(s)(X_i^{s_j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i\}}})) \\
\llbracket tZ_i^{s_j} \rrbracket_\Gamma &= \lambda s. (\llbracket t \rrbracket_\Gamma(s)(Z_i^{s_j} \otimes Id_{\mathcal{H}^{\Gamma \setminus \{i\}}})) \\
\llbracket \mathfrak{P} \rrbracket &\in T(\mathcal{H}^I, \mathcal{H}^F) : \\
\llbracket \mathfrak{P} \rrbracket &= (\llbracket A \rrbracket_I(s))_{s \in \{0,1\}^n}
\end{aligned}$$

Les règles de standardisation peuvent être généralisées au cas du *m*-calcul $\mathcal{3P}$:

$$\begin{array}{ll}
E_{ij}X_i^s &\rightarrow X_i^sZ_j^sE_{i,j} \\
E_{ij}Z_i^s &\rightarrow Z_i^sE_{i,j} \\
t[M_i^{X,\alpha}]^sX_i^r &\rightarrow t+r[M_i^{X,\alpha}]^s \\
t[M_i^{X,\alpha}]^sZ_i^r &\rightarrow t+r[M_i^{X,\alpha}]^{s+r} \\
t[M_i^{Y,\alpha}]^sX_i^r &\rightarrow t+r[M_i^{Y,\alpha}]^{s+r} \\
t[M_i^{Y,\alpha}]^sZ_i^r &\rightarrow t[M_i^{Y,\alpha}]^{s+r} \\
t[M_i^{Z,\alpha}]^sX_i^r &\rightarrow t[M_i^{Z,\alpha}]^{s+r} \\
t[M_i^{Z,\alpha}]^sZ_i^r &\rightarrow t+r[M_i^{Z,\alpha}]^s
\end{array}$$

Ce système de réécriture préserve la sémantique, est confluent et terminant. De plus, les termes irréductibles sont en forme standard : pour tout terme $\mathfrak{P}$, il existe un terme sous forme standard $\mathfrak{P}'$ tel que $\mathfrak{P} \rightarrow^* \mathfrak{P}'$ et

$$\mathfrak{P}' = (V, I, F, C_F(\Pi_{i \in V \setminus F}^> M_i^{\lambda(i), \alpha_i}) E_G N_{V \setminus I})$$

où $\lambda(i) \in \{X, Y, Z\}$ indique le plan de la mesure avec $\forall i \in I, \lambda(i) = Z$.

La géométrie associée est le quadruplet (G, λ, I, F) où $\lambda : V \setminus F \rightarrow \{X, Y, Z\}$ est une fonction partielle d'étiquetage indiquant le plan selon lequel le qubit correspondant est mesuré.

Définition 10.6 (Flot 3P) Une géométrie (G, λ, I, F) , admet un flot 3P si et seulement si il existe $f : V \rightarrow \mathbb{N}$ telle que pour tout $v \in V$, il existe un ensemble correctif $C(v) \subseteq (S(v) \cup F) \setminus I$ tel que :

- 1 - si $\lambda(v) = X$ alors
 - $v \in \text{Pair}_G(C(v))$,
 - $A(v) \cap N_G(v) \subseteq \text{Imp}_G(C(v))$,
 - $A(v) \setminus N_G(v) \subseteq \text{Pair}_G(C(v))$,
- 2 - si $\lambda(v) = Y$ alors
 - $v \in \text{Imp}_G(C(v))$,
 - $A(v) \cap N_G(v) \subseteq \text{Imp}_G(C(v))$,
 - $A(v) \setminus N_G(v) \subseteq \text{Pair}_G(C(v))$,
- 3 - si $\lambda(v) = Z$ alors
 - $v \in \text{Imp}_G(C(v))$,
 - $A(v) \subseteq \text{Pair}_G(C(v))$.

Théorème 10.3 Un terme est uniformément déterministe si la géométrie associée admet un flot généralisé.

Preuve : Le terme $\mathfrak{P} = (V, I, F, A)$, avec $A = (\prod_{i \in V \setminus F}^> M_i^{\lambda(i), \alpha_i} R(i)_i^{s_i}) E_G N_{V \setminus I}$, est uniformément déterministe mais non exécutable, avec $R(i)$ un facteur correctif adapté au plan $\lambda(i)$: si $\lambda(i) \in \{X, Z\}$ alors $R(i) = \lambda(i)$, sinon $R(i) = Z(i)X(i)$.

La preuve est donc similaire à celle du théorème 10.2, sauf qu'ici trois cas sont à considérer, suivant la valeur de l'opérateur correctif $R(v)$. Par récurrence sur la position de la mesure $M_v^{\lambda(i), \alpha_v}$ dans le terme, on montre que ce terme peut être transformé en un terme exécutable :

- Si $R(v) = X$, puisque $(\{v\} \cup C(v)) \cap I = \emptyset$, le lemme 10.1 est utilisé pour les éléments de $C(v)$ et pour v lui même, on obtient, à une phase globale près :

$$\begin{aligned} E_G N_{V \setminus I} &= X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s_v} X_v Z_{N_G(v)} E_G N_{V \setminus I} \\ &= X_{C(v) \cup \{v\}}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus A(v)}^{s_v} Z_{N_G(v) \setminus A(v)}^{s_v} E_G N_{V \setminus I} \end{aligned}$$

Ce terme correctif $X_{C(v) \cup \{v\}}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus A(v)}^{s_v} Z_{N_G(v) \setminus A(v)}^{s_v}$ commute avec les mesures effectuées avant v , compense la correction $X_v^{s_v}$, et le reste de l'opérateur, $X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus A(v)}^{s_v} Z_{N_G(v) \setminus A(v)}^{s_v}$ peut être intégré, en utilisant les règles de réécriture, aux mesures qui suivent celles de v , ou à un facteur correctif global agissant sur les qubits de sortie.

- Si $R(v) = Z$, puisque $C(v) \cap I = \emptyset$, le lemme 10.1 est utilisé pour les éléments de $C(v)$, et on obtient, à une phase globale près :

$$E_G N_{V \setminus I} = X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s_v} E_G N_{V \setminus I}$$

Ceci permet de conclure, de façon analogue au cas $R(i) = X$.

- Si $R(i) = Y$, puisque $(\{v\} \cup C(v)) \cap I = \emptyset$, le lemme 10.1 est utilisé pour les éléments de $C(v)$ et pour v lui-même, on obtient, à une phase globale près :

$$\begin{aligned} E_G N_{V \setminus I} &= X_{C(v)}^{s_v} Z_{\text{Imp}_G(C(v))}^{s_v} X_v Z_{N_G(v)} E_G N_{V \setminus I} \\ &= X_{C(v) \cup \{v\}}^{s_v} Z_{\text{Imp}_G(C(v)) \setminus A(v)}^{s_v} Z_{N_G(v) \setminus A(v)}^{s_v} E_G N_{V \setminus I} \end{aligned}$$

Ceci permet de conclure, de façon analogue au cas $R(i) = X$.

□

Tout comme le flot généralisé du m -calcul, la condition de flot du m -calcul $3P$ est une condition suffisante. La nécessité du flot est une conjecture :

Conjecture 10.2 *Un terme du m -calcul $3P$ est uniformément déterministe si et seulement si la géométrie associée admet un flot généralisé.*

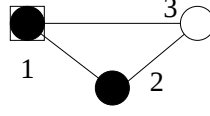
Ici aussi, l'autre problème ouvert concerne la profondeur du flot : la profondeur minimale d'un flot détermine-t-elle le nombre minimum d'étapes de mesure nécessaires à un calcul uniformément déterministe ?

10.5 Calcul par mesures dans le plan (X, Z) sur une grille

Nous avons introduit une généralisation du m -calcul, le m -calcul $3P$, en autorisant des mesures dans trois plans différents afin d'augmenter le pouvoir expressif du langage. Un point de vue inverse et complémentaire peut être adopté : quelles restrictions peuvent être apportées au modèle par consommation d'intrication, tout en ayant un modèle de calcul universel ? Autrement dit, quelles sont les ressources minimales du calcul par consommation d'intrication ? Un premier constat est que les ressources sont de deux types : l'état graphe initial, et les mesures effectuées.

La complexité de préparation d'un état graphe donné est étudiée dans le chapitre 11. Pour des raisons de réalisation physique, il peut être intéressant d'utiliser comme ressource un état graphe provenant d'un graphe régulier, comme une grille. Ainsi dans leur modèle, Briegel et Raussendorf proposent l'utilisation d'une grille rectangulaire comme état graphe initial.

Afin d'exécuter un terme $\mathfrak{P}$ du m -calcul sur une grille rectangulaire, une première phase consiste à transformer la grille en l'état graphe donné par la géométrie de $\mathfrak{P}$. Une telle transformation peut être réalisée en utilisant des mesures de Pauli selon X , Y ou Z . Ainsi, la grille rectangulaire initiale et les mesures selon Z et dans le plan (X, Y) constituent des ressources *universelles* pour le calcul par consommation d'intrication. Van den Nest *et al.* [VdNMDB06, dNDVB07] étudient l'universalité de certains états graphes, en établissant un lien entre la *largeur*

FIG. 10.3 – Géométrie de $\mathfrak{D}(\alpha)$.

*d'intrication*³ fondée sur la largeur de rang⁴ d'un graphe et l'universalité de l'état graphe correspondant. Cette mesure permet de prouver notamment que les grilles triangulaires et hexagonales sont universelles pour les mesures selon Z et dans le plan (X, Y) .

Dans cette section, nous montrons que les grilles triangulaires associées aux mesures dans le plan (X, Z) sont universelles. Ce résultat offre un degré de liberté supplémentaire au niveau de la réalisation physique, puisqu'aucune opération selon l'axe Y n'est effectuée. Le développement du m -calcul $\mathcal{3P}$ permet de représenter de façon formelle un calcul par consommation d'intrication où les mesures sont effectuées dans le plan (X, Z) .

Afin de montrer l'universalité de la grille triangulaire associée aux mesures dans le plan (X, Z) , nous montrons d'abord que toute transformation unitaire réelle peut être simulée par le m -calcul $\mathcal{3P}$ en utilisant uniquement des mesures dans le plan (X, Z) . L'universalité des grilles triangulaires associées aux mesures dans le plan (X, Z) se réduit alors à la capacité à obtenir tous les graphes qui correspondent à la géométrie d'un m -terme, à partir d'une grille triangulaire.

10.5.1 Universalité des mesures dans le plan (X, Z)

On considère les termes suivants du m -calcul $\mathcal{3P}$:

$$\begin{aligned} \mathfrak{H} &= (\{1, 2\}, \{1\}, \{2\}, Z_2^{s_1} M_1^{Z,0} E_{1,2}) \\ \mathfrak{D}(\alpha) &= (\{1, 2, 3\}, \{1\}, \{3\}, Z_3^{1-s_2} X_3^{s_2} M_2^{Y,\alpha} Z_2^{s_1} M_1^{Z,0} E_{1,2} E_{1,3} E_{2,3}) \\ \Lambda \mathfrak{Z} &= (\{1, 2\}, \{1, 2\}, \{1, 2\}, E_{1,2}) \end{aligned}$$

La géométrie associée à $\mathfrak{D}(\alpha)$ est donnée dans la figure 10.3. Cette géométrie admet le flot $\mathcal{3P}$ suivant : $f(1) = 1, f(2) = 2, C(1) = \{2\}$ et $C(2) = \{3\}$. $\mathfrak{H}$ admet également un flot $\mathcal{3P}$. On remarque que toutes les mesures effectuées dans les termes $\mathfrak{D}(\alpha)$ et $\mathfrak{H}$ sont dans le plan (X, Z) . En effet, une mesure de la forme $M^{Y,\alpha}$ est selon l'observable $\cos(\alpha)Z + \sin(\alpha)X$, et une mesure $M^{Z,0}$ est selon X .

³Entanglement-width.

⁴Rank-width.

Théorème 10.4 (Universalité) *Toute transformation unitaire réelle peut être simulée par un terme du m -calcul $3P$ dans lequel toutes les mesures sont dans le plan (X, Z) .*

Preuve : La famille de transformations unitaires $\{\Lambda Z, H, \tilde{R}_x(\alpha), \alpha \in [0, 2\pi]\}$ est universelle pour les transformations unitaires réelles, où $\tilde{R}_x(\alpha) = \cos(\alpha/2)Id + \sin(\alpha/2)X$. En effet, toute transformation unitaire réelle U sur un qubit peut être décomposée en trois rotations élémentaires :

$$U = \tilde{R}_x(\alpha)\tilde{R}_z(\alpha)\tilde{R}_x(\alpha) = \tilde{R}_x(\alpha)H\tilde{R}_x(\alpha)H\tilde{R}_x(\alpha)$$

Or, les sémantiques de $\mathfrak{H}$, $\mathfrak{D}(\alpha)$ et $\Lambda\mathfrak{Z}$ sont :

$$\begin{aligned} \llbracket \mathfrak{H} \rrbracket &= (H) \\ \llbracket \mathfrak{D}(\alpha) \rrbracket &= (\tilde{R}_x(\alpha)) \\ \llbracket \Lambda\mathfrak{Z} \rrbracket &= (\Lambda Z) \end{aligned}$$

Puisque les mesures effectuées dans les termes $\mathfrak{H}$, $\mathfrak{D}(\alpha)$ et $\Lambda\mathfrak{Z}$ sont toutes dans le plan (X, Z) , on en déduit l'universalité réelle des mesures dans le plan (X, Z) . $\square$

De plus Bernstein et Vazirani ont montré que tout circuit quantique peut être transformé en un circuit quantique utilisant uniquement des transformations unitaires réelles [BV97].

Les mesures dans le plan (X, Z) sont donc suffisantes pour le calcul quantique, à condition d'être capable de préparer tous les états graphes associés aux géométries des termes du m -calcul $3P$ restreint aux mesures selon (X, Z) .

La complexité de préparation d'un état graphe correspondant à un graphe donné est l'objet du chapitre 11. Certaines propositions de réalisation physique du calcul par consommation d'intrication proposent l'utilisation d'un graphe régulier comme état graphe de départ. Cet état graphe régulier est alors transformé, à l'aide de mesure de Pauli selon X, Y ou Z , en l'état graphe correspondant à la géométrie du terme que l'on veut exécuter. Les états graphes réguliers initiaux sont par exemple des grilles (rectangulaires, triangulaires ou hexagonales) plus facilement réalisables expérimentalement qu'un état graphe sans régularité.

Les seules mesures de Pauli du plan (X, Z) sont les mesures X, Z . Quels états graphes peuvent être obtenus en appliquant uniquement des mesures selon X et selon Z si l'état initial est une grille rectangulaire, triangulaire ou hexagonale ?

10.5.2 Pivot mineur

Les grilles rectangulaires et hexagonales ne sont pas adaptées au calcul utilisant uniquement des mesures dans le plan (X, Z) . En effet, certains états graphes ne

peuvent pas être obtenus en appliquant des mesures selon X ou Z à un état graphe associé à une grille rectangulaire ou hexagonale, notamment les graphes ayant des cycles composés d'un nombre impair de sommets, comme un triangle par exemple.

Théorème 10.5 *Pour tout graphe G et toute arête uv de G , si G n'a pas de cycle impair, alors $G \wedge uv$ n'a pas de cycle impair.*

Preuve : Si G n'a pas de cycle impair alors pour tout arête uv :

1. Les voisinages propres $A = N_G(u) \setminus N_G(v)$ et $B = N_G(v) \setminus N_G(u)$ sont des stables (il n'y a pas d'arête entre deux sommets de A , ni entre deux sommets de B).
2. Le voisinage commun $C = N_G(u) \cap N_G(v)$ de u et v est vide.
3. Tout chemin $u_1 \dots u_n$ entre deux sommets de A (de B), a un nombre pair d'arêtes.
4. Tout chemin $u_1 \dots u_n$ entre un sommet de A et un sommet de B a un nombre impair d'arêtes.

Par l'absurde, supposons qu'après le pivot sur l'arête uv , un cycle impair ait été créé, il existe alors un chemin avec un nombre pair d'arêtes entre A et B (en effet le pivot ne fait que créer et retirer des arêtes entre A et B).

Considérons un plus court chemin wPw' dans $G' = G \wedge uv$ qui viole 3 ou 4. P doit intersecter $A \cup B$ sinon le chemin contredirait 3 ou 4 pour G . Il est donc de la forme $wP_1w''P_2w'$ et sa minimalité amène alors à une contradiction. □

Corollaire 10.1 :

- *Il existe un graphe qui n'est pas pivot mineur de la grille rectangulaire.*
- *Il existe un graphe qui n'est pas pivot mineur de la grille hexagonale.*

Par exemple, un triangle n'est pivot mineur ni de la grille rectangulaire, ni de la grille hexagonale.

D'après le lemme 9.2, si un état graphe $|G\rangle$ est obtenu à partir d'un autre état graphe $|H\rangle$ en utilisant des mesures selon X ou selon Z , alors G est un pivot mineur de H . On en déduit que l'état graphe représenté par un triangle ne peut pas être obtenu à partir d'un état graphe associé à une grille rectangulaire ou hexagonale en appliquant uniquement des mesures selon X ou Z . Or, la géométrie associée au terme $\mathfrak{D}(\alpha)$ est un triangle (figure 10.3). On en déduit que les grilles rectangulaire et hexagonale ne sont pas adaptées au calcul par mesures dans le plan (X, Z) .

D'après le théorème 10.5, il est important d'utiliser une grille initiale possédant des cycles impairs. La grille triangulaire possède cette propriété. Le théorème suivant montre que tout état graphe peut être obtenu en appliquant des mesures selon X ou Z sur une grille triangulaire.

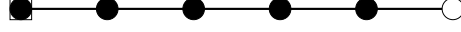


FIG. 10.4 – Géométrie admettant un flot et dont la sémantique est H si tous les qubits noirs sont mesurés selon X .

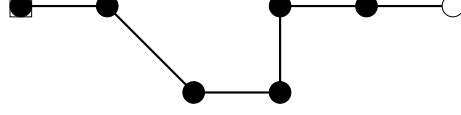


FIG. 10.5 – Géométrie admettant un flot et dont la sémantique est l'identité si tous les qubits noirs sont mesurés selon X .

Théorème 10.6 *Pour tout graphe G , l'état graphe $|G\rangle$ peut être obtenu à un signe près, en appliquant des mesures selon Z ou X sur un état graphe représenté par une grille triangulaire.*

Preuve : Soit $G = (V, E)$ un graphe, $|G\rangle = \Pi_{(u,v) \in E} \Lambda Z_{u,v} |+\rangle_V$. Le circuit quantique $C_G = (\Lambda Z[u, v])_{(u,v) \in E}$ (voir chapitre 11) permet de produire l'état $|G\rangle$ à partir de l'état $|+\rangle_V$. C_G peut être transformé en un circuit planaire C'_G en utilisant pour base $\{H, \Lambda Z\}$. En effet C_G a pour base $\{\Lambda Z\}$ et l'opération *Swap*, suffisante pour transformer C_G en C'_G peut être décomposée comme suit :

$$\text{Swap}_{u,v} = H_{u,v} \Lambda Z_{u,v} H_{u,v} \Lambda Z_{u,v} H_{u,v} \Lambda Z_{u,v}$$

Les graphes décrits dans les figures 10.4, 10.5 et 10.6 admettent un flot simple, si tous les qubits noirs sont mesurés dans le plan (X, Y) . A l'aide de la sémantique du m -calcul, on montre facilement que si tous les qubits, sauf les qubits blancs, sont mesurés selon l'observable X , alors les transformations induites des qubits encadrés vers les qubits blancs sont respectivement H , Id et ΛZ .

On remarque que les graphes décrits dans les figures 10.4, 10.5 et 10.6 sont des sous graphes induits (donc des pivots mineurs) de grilles triangulaires, comme illustré dans la figure 10.7.

D'après les lois de composition des m -termes, les termes simulant H , Id et ΛZ peuvent être composés pour obtenir une simulation du circuit C'_G . Etant donné que les simulations de H , Id et ΛZ sont réalisées par des géométries ayant la même "largeur" une fois plongées dans la grille triangulaire, la simulation du circuit C'_G peut être faite couche par couche. On en déduit qu'il existe une géométrie (H, I, F) associée à la simulation du circuit C'_G , telle que H est un sous graphe de la grille triangulaire. De plus, comme les qubits d'entrée de C'_G sont dans l'état $|+\rangle$, si tous les qubits de $|H\rangle$, sauf les qubits de sortie situés à droite du graphe, sont mesurés selon X , alors l'état des qubits non mesurés est $|G\rangle$ à un signe près.

□

Ainsi les grilles triangulaires associées aux mesures selon (X, Z) sont universelles pour les transformations réelles. Etant donné qu'un état graphe ne possède pas de coordonnée complexe, et que les mesures dans le plan (X, Z) n'introduisent pas de nombre complexe, les mesures selon (X, Z) sur un état graphe sont donc nécessairement limitées à une universalité réelle. Même si toute transformation unitaire peut être simulée en utilisant uniquement des transformations unitaires réelles, la question de savoir si l'utilisation de mesures dans le plan (Y, Z) pourrait permettre d'atteindre une universalité complète, reste ouverte.

Le théorème 10.6 permet de prouver la propriété combinatoire suivante :

Propriété 10.1 *Tout graphe est pivot mineur d'une grille triangulaire.*

Preuve : Cette propriété est prouvée en utilisant les états graphes. Soit G un graphe. D'après le théorème 10.6, il existe une grille triangulaire H telle que $|G\rangle$ est obtenu en appliquant des mesures selon Z et X sur l'état graphe $|H\rangle$. D'après le théorème 9.2 (chapitre 9), G est un pivot mineur de H . $\square$

Il est intéressant de noter que la preuve de cette propriété combinatoire emprunte un chemin quantique.

Nous avons montré que les grilles triangulaires associées aux mesures selon (X, Z) sont universelles. En revanche, nous n'avons pas montré que les grilles triangulaires ou hexagonales ne le sont pas. En effet, aucun graphe ayant un cycle impair ne peut être atteint à partir d'une grille rectangulaire ou hexagonale. Un problème ouvert est de savoir si l'existence de cycle impair est une condition nécessaire à l'universalité pour les géométries d'un m -calcul $\mathcal{3P}$ qui n'effectuerait que des mesures dans le plan (X, Z) .

10.6 Unification des modèles par consommation d'intrication et par mesures projectives

10.6.1 Le secret du calcul par consommation d'intrication est caché dans la préparation de l'état graphe initial

Nous avons étudié deux modèles de calcul quantique fondés sur la mesure : le modèle de calcul quantique par mesures projectives (chapitre 7) et le modèle de calcul quantique par consommation d'intrication (ce chapitre). Il est naturel de comparer ces deux modèles. Le premier s'appuie sur une brique de base, le *transfert d'état* (chapitre 7, figure 7.4) qui permet de simuler des transformations unitaires, sans intrication préalable, alors que dans le second modèle, une ressource supplémentaire est l'état graphe initial qui est consommé durant l'exécution. Une

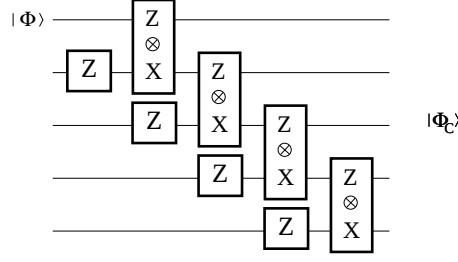


FIG. 10.8 – Préparation d'une ligne par mesures projectives.

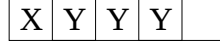


FIG. 10.9 – Calcul par consommation d'intrication.

comparaison rapide des deux modèles fait donc apparaître deux différences majeures en termes de ressources : la taille des opérations utilisées, et l'utilisation ou non d'intrication préalable. Dans le modèle par mesures projectives, des mesures multiqubits sont nécessaires, alors que seules des mesures locales sont utilisées dans le calcul par consommation d'intrication. En revanche, aucune intrication préalable n'est nécessaire dans le calcul par mesures projectives, contrairement au calcul par consommation d'intrication.

Afin de pouvoir comparer et unifier ces deux modèles de calcul quantique, la création de l'intrication doit être prise en compte dans le modèle par consommation d'intrication. Nous nous intéressons aux états graphes à une dimension. Un tel état graphe peut être préparé en utilisant uniquement des mesures projectives, comme décrit dans le lemme 9.4. Une telle préparation est décrite figure 10.8. Pour un graphe $G = (V, E)$ à une dimension, c'est-à-dire une ligne, cette préparation permet d'obtenir, à partir d'un état $|\Phi\rangle$ sur le qubit d'entrée $v_0 \in V$, l'état $|\Phi_C\rangle = \Pi_{(u,v) \in E} \Lambda Z_{u,v} |\Phi\rangle_{v_0} |+\rangle_{V \setminus \{v_0\}}$, à un opérateur de Pauli près.

Un calcul par consommation d'intrication consiste à mesurer les qubits de l'état graphe à une dimension, sauf un qubit qui est le qubit de sortie (figure 10.9). Si ces mesures sur un qubit sont représentées sur le même schéma que les mesures utilisées pour créer l'état graphe, alors une redécomposition naturelle des mesures peut être effectuée pour reconnaître une suite de transferts d'états (voir figure 10.10).

Au delà de l'unification des deux modèles, cette technique donne un outil sémantique pour le calcul par consommation d'intrication. En effet, ainsi décomposé en une suite de transferts d'états, une transformation unitaire peut être associée à chaque transfert d'état. L'action du calcul par consommation d'intrication est alors la composition de l'action de chaque transfert d'état. Ainsi, dans l'exemple donné dans la figure 10.9, la suite de mesures est X, Y, Y et Y . Or,

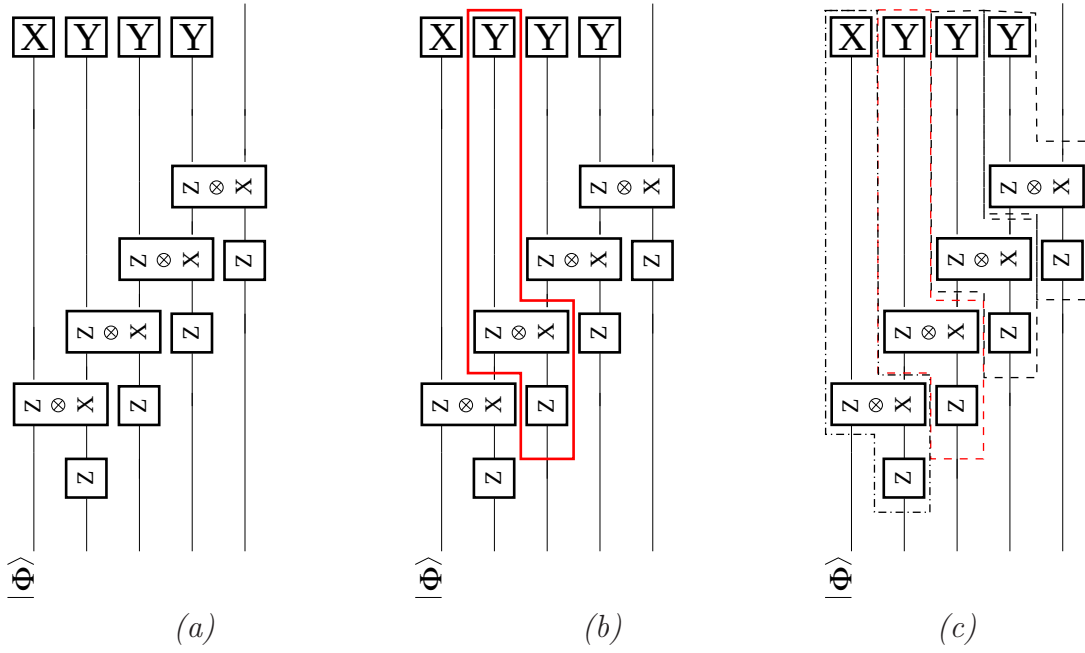


FIG. 10.10 – Unification

d'après le transfert d'état généralisé (chapitre 7, figure 7.5), de telles mesures correspondent à la simulation de $H\sqrt{Z}^\dagger$ pour une mesure selon Y et H pour une mesure selon X . La sémantique du calcul par consommation d'intrication est donc $H\sqrt{Z}^\dagger H\sqrt{Z}^\dagger H\sqrt{Z}^\dagger H = H$.

10.6.2 Unification et m -calcul

L'unification présentée dans la section précédente est la première [JP04c, JP05b], mais pas l'unique unification [CLN05, Gim05] entre les modèles de calcul par mesures projectives et par consommation d'intrication.

L'unification présentée n'est pas complète, puisqu'elle ne permet de traiter que le cas des graphes dont le degré est au plus 2, avec au moins un sommet de degré 1. Malgré son incomplétude, cette unification présente l'intérêt de n'utiliser que des mesures projectives, là où les autres utilisent des transformations unitaires, elles s'appuient sur une préparation des états graphes via la transformation unitaire ΛZ .

On remarque également que si une préparation par transformation unitaire est utilisée, alors le principe de l'unification présenté dans la figure 10.10 est similaire à la standardisation du m -calcul [DKP04a] (la standardisation transforme le terme de droite en celui de gauche sur la figure).

10.7 Conclusion

Le m -calcul, modèle formel pour le calcul quantique par consommation d'intrication a été présenté. La condition de flot a été généralisée. Cette condition permet de savoir si un état graphe peut être utilisé pour effectuer un calcul par consommation d'intrication. Un modèle de calcul plus général, le m -calcul $\mathcal{3P}$ a été introduit. Ce nouveau modèle étend le m -calcul en offrant la possibilité d'effectuer des mesures dans trois plans différents au lieu d'un seul pour le m -calcul. La condition de flot, mais aussi d'autres ingrédients du m -calcul, comme la standardisation, ont été généralisés au m -calcul $\mathcal{3P}$.

Une autre piste a également été explorée, celle de la recherche des ressources minimales du calcul quantique par consommation d'intrication. Alors que dans le modèle originel de calcul par consommation d'intrication, proposé par Briegel et Raussendorf [RB00], des mesures selon Z et dans le plan (X, Y) sur une grille rectangulaire sont prouvées universelles, nous montrons que des mesures dans le plan (X, Z) sur une grille triangulaire constituent également des ressources universelles. Une application de ce résultat est que tout graphe est pivot mineur d'une grille triangulaire. Ce résultat de théorie des graphes est le premier, à la connaissance de l'auteur, obtenu en utilisant des méthodes de l'informatique quantique.

Chapitre 11

Préparation des états graphes

11.1 Introduction

Quelles sont les ressources minimales nécessaires à un calcul quantique universel ? Cette simple question est l'une des plus fondamentales posée par la conception d'un ordinateur quantique, et c'est l'une des questions les plus étudiées au sein de l'informatique quantique. En 2000, Raussendorf and Briegel [RB00] ont proposé un nouveau modèle de calcul quantique. Ils ont montré que si certains états quantiques initiaux, appelés états graphes, sont fournis, alors la simple capacité à appliquer des mesures sur 1 qubit selon un observable dans le plan (X, Y) ou selon Z , suffit au calcul quantique. Un modèle de calcul dédié au calcul par consommation d'intrication, le m -calcul a été introduit par Dansos, Kashefi et Panangaden [DKP04a], ce modèle a été présenté et généralisé dans le chapitre 10. De plus, une caractérisation combinatoire de certaines transformations quantiques a été donnée dans le chapitre 9. Mais encore faut-il fournir, à un coût minimal, les états graphes initiaux.

Nous montrons d'abord et principalement que les états graphes peuvent être préparés en temps constant. C'est-à-dire, pour une description classique donnée d'un graphe $G = (V, E)$, on peut produire l'état graphe correspondant $|G\rangle$ par un circuit quantique de profondeur constante qui a une taille linéaire dans la taille de l'entrée $|V| + |E|$ et qui consiste seulement en des opérations sur 1 et 2 qubits. Cela implique que toutes les opérations sur 2 qubits jamais employées par un algorithme quantique peuvent être conduites au début de l'algorithme et en parallèle, le reste des opérations étant des opérations sur un qubit.

Ce regroupement des opérations multiqubits et la parallélisation de ces opérations nécessitent un espace auxiliaire. Un compromis entre le temps de la préparation et l'espace nécessaire à cette préparation est étudié, et il est montré qu'il est possible de fixer arbitrairement un temps de préparation, et d'en déduire l'espace

nécessaire pour y parvenir.

Pour déterminer la complexité de préparation des états graphes, nous introduisons une nouvelle quantité sur les graphes, le degré minimum local, noté δ_{loc} . Par exemple, nous l'utilisons pour prouver que toute préparation composée uniquement de mesures nécessite un espace auxiliaire, ou des mesures agissant simultanément sur au moins $\delta_{\text{loc}} + 1$ qubits. De plus, nous établissons également que le degré minimum local est lié à l'intrication dans les états graphes : une mesure agissant sur δ_{loc} qubits suffit à créer une séparation dans un état graphe.

Le degré minimum local apparaît comme une propriété importante des états graphes. Afin de montrer que le degré minimum local n'est pas borné, nous exhibons une famille de graphes pour laquelle le degré minimum local est grand.

11.2 Préparation d'un état graphe

Le problème que l'on considère est celui de la préparation d'un état graphe. Ce problème a la particularité de posséder une entrée classique, un graphe G , et une sortie quantique, l'état graphe $|G\rangle$. Nous considérons également le cas où le résultat de l'algorithme est un graphe signé $|G; S\rangle$, c'est-à-dire l'état graphe $|G\rangle$ à un opérateur de Pauli près.

Un algorithme de préparation peut en général être décomposé en deux étapes :

- Un algorithme classique, à partir de l'entrée classique, produit la description classique d'un circuit quantique C_G ;
- Le circuit C_G est exécuté, produisant l'état quantique $|G\rangle$.

Afin de mesurer la complexité de la préparation d'un état graphe, la profondeur du circuit C_G , sa taille, ainsi que le nombre de qubits sur lesquels il agit sont pris en compte. Ces quantités sont calculées en fonction de la taille $n + m$ de l'entrée $G = (V, E)$, où $n = |V|$ et $m = |E|$. La taille du circuit est le nombre de portes, elle représente donc la complexité en temps pour un déroulement séquentiel des opérations. La profondeur d'un circuit représente la complexité en temps dans le cadre d'une exécution des portes quantiques qui exploite au maximum le parallélisme possible entre portes agissant sur des qubits différents. Le nombre de qubits sur lesquels agit le circuit représente une complexité en *espace* de la préparation. Etant donné que le résultat du circuit est un état quantique sur n qubits, la taille est au moins n . Les éventuels qubits supplémentaires qui peuvent être nécessaires à la préparation, sont appelés qubits auxiliaires.

La complexité de la première étape de la préparation, c'est-à-dire la partie classique, n'est pas prise en compte dans la complexité de l'algorithme, tant que celle-ci est polynomiale (en temps).

Il existe un algorithme simple pour préparer l'état graphe $|G\rangle$ correspondant à un graphe quelconque $G = (V, E)$ sur $n = |V|$ sommets avec $m = |E|$ arêtes. On

prépare d'abord n qubits dans une superposition de tous les 2^n états de base,

$$|+\rangle_V = \frac{1}{\sqrt{2^n}} \sum_{x \in \{0,1\}^n} |x\rangle,$$

en appliquant par exemple l'opérateur d'Hadamard H sur chacun des n qubits dans l'état initial $|0\rangle$. Chaque qubit correspond à un sommet de G . Ensuite, on applique une séquence de m opérations sur 2 qubits, ce qui produit l'état graphe $|G\rangle$,

$$|G\rangle = \prod_{(u,v) \in E} \Lambda Z_{u,v} |+\rangle_V.$$

Pour chaque arête $(u, v) \in E$, on applique l'opérateur ΛZ , dont nous rappelons la définition :

$$\Lambda Z_{u,v} = |00\rangle\langle 00| + |01\rangle\langle 01| + |10\rangle\langle 10| - |11\rangle\langle 11|$$

Ces m opérations sur 2 qubits sont diagonales et ainsi commutent, ce qui nous permet de les appliquer dans l'ordre de notre choix. Pour résumer, il est possible de préparer un état graphe $|G\rangle$ en utilisant n opérations sur 1 qubit et m opérations ΛZ sur 2 qubits. En considérant ce simple algorithme comme étant un circuit quantique, on peut préparer un état graphe par un circuit sur n qubits de taille $n + m$ et de profondeur $m + 1$ en utilisant uniquement des opérations sur 1 et 2 qubits.

La profondeur du circuit peut être diminuée en effectuant des opérations sur 2 qubits en parallèle. Pour cela, on choisit d'abord une coloration des arêtes de G . Une *coloration des arêtes* utilisant χ' couleurs est une fonction de coloration $c : E \rightarrow \{1, 2, \dots, \chi'\}$ telle que pour deux arêtes distinctes e et e' partageant un sommet, $c(e) \neq c(e')$. Tout graphe possède une coloration des arêtes utilisant au plus $\Delta(G) + 1$ couleurs, où $\Delta(G)$ est le degré maximum des sommets dans G , et on peut trouver une coloration utilisant $O(\Delta)$ couleurs en un temps polynomial en n et m , par exemple par l'algorithme (classique) de Vizing [Viz64]. Cela implique que les m opérations sur 2 qubits peuvent être réorganisées dans le circuit de préparation de telle façon qu'il ait une profondeur d'au plus $\Delta(G) + 2$:

Proposition 11.1 *Tout état graphe $|G\rangle$ peut être préparé par un circuit quantique de profondeur $O(\Delta(G))$, qui consiste à appliquer $O(n + m)$ opérations sur 1 ou 2 qubits, où $\Delta(G)$ est le degré maximum de tout sommet dans G .*

L'algorithme ci-dessus implique que les graphes ayant un degré borné peuvent être préparés par des circuits quantiques de profondeur constante. En particulier, les états graphes associés à une grille peuvent être préparés par des circuits quantiques de profondeur constante. On montre à présent comment étendre cet algorithme à des graphes quelconques.

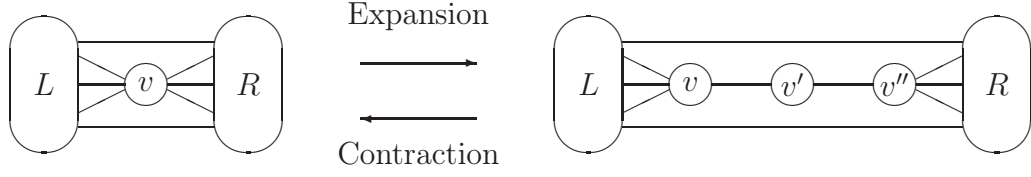


FIG. 11.1 – Appliquer des mesures selon X sur des qubits v' et v'' adjacents contracte les arêtes (v, v') et (v', v'') .

Théorème 11.1 (Préparation d'état graphe en profondeur constante)

Pour tout graphe G , $|G\rangle$ est préparé à un signe près¹ par un circuit quantique de profondeur constante qui consiste à appliquer sur $n + O(m)$ qubits, $O(n + m)$ opérations sur 1 ou 2 qubits.

Preuve : L'idée clé de cette preuve est d'assimiler G à un mineur induit² d'un graphe G' de degré borné, puis de montrer que G peut être obtenu à partir de G' en effectuant des mesures sur un qubit.

On commence par montrer que G est un mineur induit provenant d'un graphe plus grand G' de degré borné en répétant l'expansion d'un sommet v de degré $d \geq 4$ en trois sommets v, v', v'' de degrés strictement inférieurs, comme illustré dans la figure 11.1.

Formellement, le voisinage $N(v)$ de v est partitionné en deux ensembles L et R de taille $\lceil \frac{d}{2} \rceil$ et $\lfloor \frac{d}{2} \rfloor$, respectivement. Ensuite, on pose $\hat{V} = V \cup \{v', v''\}$ et $\hat{E} = E \setminus (\{v\} \times R) \cup \{(v, v'), (v', v'')\} \cup (\{v''\} \times R)$. On pose $\hat{G} = (\hat{V}, \hat{E})$, et l'expansion de tout sommet v de $\hat{G}$ est répétée récursivement jusqu'à ce que tous les sommets aient un degré d'au plus trois. Le graphe G' ainsi obtenu comportera $O(m + n)$ sommets et sera de degré 3 au plus.

L'état graphe $|G'\rangle$ est alors préparé par un circuit de profondeur constante sur $O(m + n)$ qubits en appliquant la proposition 11.1, puis en appliquant le lemme 9.2 pour contracter $|G'\rangle$ en $|G\rangle$ en appliquant des mesures selon X sur tous les sommets introduits pendant l'expansion de G en G' . Comme les mesures selon X commutent, elles peuvent toutes être effectuées en parallèle. $\square$

La proposition 11.1 et le théorème 11.1 fournissent deux circuits de taille linéaire pour préparer un état graphe quelconque, à un signe près. Le premier a une

¹Un état graphe signé $|G; S\rangle$ est l'état $Z_S |G\rangle$, où S est un sous ensemble des sommets de G appelé le signe (voir chapitre 9).

² G est un mineur induit de G' , si G peut être obtenu à partir de G' par contractions et suppressions d'arêtes (voir chapitre 9).

petite largeur et une grande profondeur, le second une grande largeur et une faible profondeur. On peut diminuer la largeur en augmentant la profondeur du circuit dans la construction ci-dessus, sans changer la taille globale du circuit, en arrêtant l'expansion dès que tous les sommets ont un degré d'au plus T .

Théorème 11.2 (Préparation d'état graphe de faible profondeur)

Pour tout graphe G et tout entier $T \geq 3$, $|G\rangle$ peut être préparé, à un signe près, par un circuit quantique agissant sur $n + O(m/T)$ qubits, de profondeur $O(T)$, composé de $O(n + m)$ opérations sur 1 ou 2 qubits.

Preuve : La preuve est similaire à celle du théorème 11.1 : les sommets dont le degré est supérieur à T subissent une expansion, diminuant leur degré de moitié. Cette opération est répétée tant que le degré du graphe est supérieur à T . Le graphe G' ainsi obtenu possède $n + O(m/T)$ sommets. G' est préparé par l'algorithme de la proposition 11.1, puis tous les qubits auxiliaires sont mesurés en parallèle selon X . $\square$

11.3 Préparation fondée sur la mesure

Alors que les états graphes constituent une ressource du calcul quantique par consommation d'intrication³, il est naturel de considérer la préparation des états graphes dans un cadre où seules les mesures projectives sont autorisées.

Dans un premier temps, les résultats sur la simulation des transformations unitaires par des mesures projectives (voir chapitre 7) permet de transformer les algorithmes précédents pour obtenir des algorithmes utilisant exclusivement des mesures. Une préparation *en place* (sans espace auxiliaire) utilisant exclusivement des mesures projectives est également proposée.

Théorème 11.3 (Préparation fondée sur la mesure) *Pour tout graphe G et pour tout entier $T \geq 3$, $|G\rangle$ est préparé à un opérateur de Pauli près, par une série de $O(n + m)$ mesures projectives sur 1 ou 2 qubits. Ces mesures agissent sur un espace composé de $O(n + m/T)$ qubits, et peuvent être exécutées en parallèle en $O(T)$ étapes.*

Preuve : La preuve est fondée sur l'utilisation du théorème 11.2 puis sur la simulation par des mesures projectives du circuit C_G ainsi obtenu. Le circuit C_G est composé de mesures selon X et de transformations unitaires ΛZ . Chaque ΛZ peut être simulé, à un opérateur de Pauli près par 3 mesures sur 1 et 2 qubits, comme

³Le calcul quantique par consommation d'intrication est un modèle de calcul fondé sur la mesure successive des qubits d'un état graphe, voir chapitre 10.

décrit dans la figure 7.11 du chapitre 7. Puisque le circuit est composé uniquement d'opérations ΛZ et que cet opérateur *normalise* les opérateurs de Pauli⁴, la stratégie permettant de corriger l'opérateur de Pauli produit aléatoirement par la simulation n'est pas nécessaire : les opérateurs de Pauli sont intégrés au signe de l'état graphe signé obtenu.

Afin de simuler les $O(n + m)$ opérations, $(n + m)/T$ qubits auxiliaires sont utilisés. Le parallélisme est donc limité à au plus $(n + m)/T$ opérations par étape. Malgré cette limitation, le nombre d'étapes est $O(T)$. $\square$

Etant donné que la simulation de chaque transformation unitaire par des mesures projectives nécessite un qubit auxiliaire, le cas où un seul qubit auxiliaire est disponible impose une séquentialité des opérations appliquées.

Proposition 11.2 (Préparation par mesures utilisant un qubit auxiliaire)

Pour tout graphe G , $|G\rangle$ est préparé à un opérateur de Pauli près, par une série de $O(m + n)$ mesures projectives sur 1 ou 2 qubits. Ces mesures agissent sur un espace composé de $n + 1$ qubits, et peuvent être exécutées en $O(m + n)$ étapes.

Si aucun qubit auxiliaire n'est disponible, alors la simulation des transformations unitaires n'est pas possible. Un état graphe peut cependant être préparé en place, sans qubit auxiliaire, en utilisant des mesures dont la taille dépend du degré minimum du graphe.

Lemme 11.1 (Préparation fondée sur la mesure, sans qubit auxiliaire)

Pour tout graphe G , $|G\rangle$ est préparé à un opérateur de Pauli près, par une suite de $O(n + m)$ mesures projectives sur au plus $\delta G + 1$ qubits, où $\delta(G) = \min\{\deg_G(v) : v \in V\}$ est le degré minimum de G , où $\deg_G(v)$ représente le degré de v dans G .

Preuve : La preuve consiste à choisir un sommet v de degré minimum dans G , donc $\deg_G(v) = \delta(G)$. L'état graphe $|G \setminus v\rangle$ peut être préparé, d'après la proposition 11.2, en utilisant le qubit v comme qubit auxiliaire. Puis, d'après le lemme 9.4, une mesure sur $\delta(G) + 1$ qubits selon l'observable $X_v Z_{N_G(v)}$ permet de transformer $|G \setminus v\rangle$ en $|G\rangle$ à un signe près. $\square$

11.4 Circuits et complémentation locale

Les constructions de circuits données ci-dessus pour préparer les états graphes prennent appui sur le concept de mineur induit. Pour améliorer les algorithmes

⁴Tout opérateur de Clifford C , dont ΛZ , normalisent les opérateurs de Pauli : pour tout opérateur de Pauli P , il existe P' tel que $CP = P'C$.

de préparation, le concept de complémentation locale de graphes est utilisé (voir définition 9.5 du chapitre 9).

En effet, si deux graphes sont localement équivalents⁵, alors $|G\rangle$ peut être transformé en $|G'\rangle$ par un circuit composé d'opérations sur 1 qubit, et donc de profondeur constante. Il est donc possible de préparer $|G\rangle$ en préparant dans un premier temps $|G'\rangle$, avec $G' \approx_{\text{loc}} G$.

Pour étudier les possibilités offertes par de telles préparations, nous introduisons deux propriétés locales : le degré minimum et le nombre minimal d'arêtes atteignables par complémentations locales.

Définition 11.1 (Degré minimum local) *Pour tout graphe $G = (V, E)$, $\delta_{\text{loc}}(G) = \min\{\delta(G') : G' \approx_{\text{loc}} G\}$ est le degré minimal des graphes atteignables par des complémentations locales. On se réfère à δ_{loc} comme étant le degré minimum local de G .*

Définition 11.2 (Nombre minimum local d'arêtes) *Pour tout graphe $G = (V, E)$, soit $m_{\text{loc}}(G) = \min\{|E'| : (V, E') \approx_{\text{loc}} G\}$. On se réfère à m_{loc} comme étant le nombre minimum local d'arêtes de G .*

Malheureusement, il n'existe pas d'algorithme polynomial connu qui calcule l'une ou l'autre des quantités $\delta_{\text{loc}}(G)$ et $m_{\text{loc}}(G)$, à partir d'un graphe G donné en entrée. Le résultat le plus avancé dans cette voie est celui de Bouchet [Bou87] affirmant que le problème consistant à décider si deux graphes sont localement équivalents est calculable en un temps polynomial. Van den Nest [VdN05] donne dans sa thèse une courte description de l'algorithme de Bouchet.

La quantité m_{loc} est reliée à la taille de tout circuit quantique préparant un état graphe. Supposons que l'on puisse trouver un algorithme qui, pour un graphe donné et en un temps polynomial, fournit un graphe localement équivalent à m_{loc} arêtes, alors, dans les théorèmes 11.1, 11.2 et 11.3, on pourrait remplacer m par m_{loc} , et toujours avoir des circuits quantiques constructibles en des temps polynomiaux. Toutefois, aucun algorithme de cette sorte n'est visible à l'horizon.

Une autre propriété locale, la largeur de rang⁶, introduite par Oum [Oum05], est utilisée par Van den Nest, Miyake, Dür et Briegel [VdNMDB06] comme mesure d'intrication sur les états graphes.

On montre maintenant que δ_{loc} est relié à l'usage de qubits auxiliaires dans la préparation d'états graphes. Pour prouver cela, commençons par donner trois définitions équivalentes de δ_{loc} , la première issue de la théorie des graphes, la seconde

⁵ G est localement équivalent à G' , si G peut être obtenu à partir de G' par application de complémentations locales.

⁶*rank-width*

de la combinatoire et la dernière de l'algèbre. Nous introduisons, ou rappelons les notations et concepts utiles suivants :

Pour tout sous-ensemble $X \subseteq V$ de sommets, $\text{Imp}_G(X) = \{u \in V \setminus X : N(u) \cap X = 1 \pmod 2\}$ est l'ensemble des sommets qui sont adjacents à un nombre impair de sommets de X dans G . De manière similaire, $\text{Pair}_G(X) = \{u \in V \setminus X : N(u) \cap X = 0 \pmod 2\}$ est l'ensemble des sommets qui sont adjacents à un nombre pair de sommets de X dans G . On dit que les sommets dans $\text{Imp}_G(X)$ sont *voisins impairs* de X dans G , et que les sommets dans $\text{Pair}_G(X)$ sont *voisins pairs* de X dans G .

La matrice de coupe d'un sous ensemble $X \subseteq V$ de sommets est la sous matrice $\Gamma_G(X, V \setminus X)$ indexée par $X \times (V \setminus X)$ de la matrice d'adjacence Γ_G de G . Le rang de coupe $\text{Cutrk}(X)$ de X est le rang de sa matrice de coupe, où on définit le rang sur $\text{GF}(2)$. Le rang de coupe de X est invariant par complémentation locale [Bou89], bien que le noyau de $\Gamma_G(X, V \setminus X)$ peut changer par complémentation locale. Cela a été utilisé par Bouchet [Bou87] et d'autres sous le terme de *fonction de connexion*, et renommé *rang de coupe* par Oum [Oum05].

Définition 11.3 *Etant donné un graphe $G = (V, E)$, on dit que $L \subseteq V$ est local si il existe $X \subseteq V$ tel que $L = X \cup \text{Imp}_G(X)$.*

On remarque que pour tout sommet $v \in V$, $\{v\} \cup N_G(v)$ est local. On remarque également qu'un ensemble local L ne peut pas être de rang de coupe plein, en effet le vecteur χ_X est tel que $\chi_X \Gamma_G[L, V \setminus L] = 0 \pmod 2$, où χ_X est la fonction indicatrice de X dans L . Le lien entre la localité d'un sous-ensemble de sommets et le rang de coupe de cet ensemble est donné dans le théorème 11.4.

Lemme 11.2 *Tout ensemble local L est invariant par complémentation locale. De plus, pour tout $y \in L$, il existe un graphe G' localement équivalent à G tel que $\{y\} \cup \text{Imp}_{G'}(\{y\}) \subseteq L$.*

Preuve : Supposons que $L = X \cup \text{Imp}_G(X)$. On considère la façon dont la partition tripartite $V = X \cup \text{Imp}_G(X) \cup \text{Pair}_G(X)$ change par complémentation locale sur un sommet $v \in V$. Soit $G' = G \star v$. Alors, la tripartition change comme suit :

	X'	$\text{Imp}_{G'}(X')$	$\text{Pair}_{G'}(X')$
$v \in \text{Pair}_G(X)$	X	$\text{Imp}_G(X)$	$\text{Pair}_G(X)$
$v \in \text{Imp}_G(X)$	$X \cup \{v\}$	$\text{Imp}_G(X) \setminus \{v\}$	$\text{Pair}_G(X)$
$v \in X$ et $ N_G(v) \cap X $ est impair	$X \setminus \{v\}$	$\text{Imp}_G(X) \cup \{v\}$	$\text{Pair}_G(X)$
$v \in X$ et $ N_G(v) \cap X $ est pair	X	$\text{Imp}_G(X)$	$\text{Pair}_G(X)$

En effet :

- Si $v \in \text{Pair}_G(X)$, alors pour tout $u \in \text{Imp}_G(X) \cup \text{Pair}_G(X)$, si $u \notin N_G(v)$, les arêtes issues de u ne sont pas modifiées par la complémentation locale sur le sommet v . Sinon, si $X \cap N_G(v) = A_1 \cup A_2$, avec $A_1 = X \cap N_G(v) \cap N_G(u)$ et $A_2 = (X \cap N_G(v)) \setminus N_G(u)$, A_1, A_2 forment une partition de $X \cap N_G(v)$, et $|A_1| + |A_2| \equiv 0 \pmod 2$. $X \cap N_{G'}(u) = ((X \cap N_G(u)) \setminus A_1) \cup A_2$, donc $|X \cap N_{G'}(u)| = |X \cap N_G(u)| - |A_1| + |A_2| \equiv |X \cap N_G(u)| \pmod 2$. Donc si $u \in \text{Imp}_G(X)$ alors $u \in \text{Imp}_{G'}(X)$ et si $u \in \text{Pair}_G(X)$ alors $u \in \text{Pair}_{G'}(X)$. Donc $L = X \cup \text{Imp}_{G'}(X)$.
- Si $v \in \text{Imp}_G(X)$, alors pour tout $u \in N_G(v)$, avec les mêmes notations que pour le cas précédent, on a $|A_1| + |A_2| \equiv 1 \pmod 2$, et donc $|(X \cup \{v\}) \cap N_{G'}(u)| = |X \cap N_G(u)| - |A_1| + |A_2| + |\{v\}| = |X \cap N_G(u)|$.
- Si $v \in X$, alors pour tout $u \in N_G(v)$, avec les mêmes notations que pour le cas précédent, on a $|A_1| + |A_2| \equiv 1 \pmod 2$, et donc $|(X \cup \{v\}) \cap N_{G'}(u)| = |X \cap N_G(u)| - |A_1| + |A_2| + |\{v\}| = |X \cap N_G(u)|$.
- Si $v \in X$, alors pour tout $u \in N_G(v)$, avec les mêmes notations que pour le cas précédent, on a $|A_1| + |A_2| \equiv |N_G(v) \cap X| \pmod 2$. Si $|N_G(v) \cap X| \equiv 0 \pmod 2$, alors $|X \cap N_{G'}(u)| = |X \cap N_G(u)| - |A_1| + |A_2| \equiv |X \cap N_G(u)| \pmod 2$. Sinon $(X \setminus v) \cap N_{G'}(u) = |X \cap N_G(u)| - |\{v\}| - |A_1| + |A_2| \equiv |X \cap N_G(u)| \pmod 2$. De plus si $|N_G(v) \cap X| \equiv 1 \pmod 2$, alors $v \in \text{Imp}_{G'}(X \setminus \{v\})$.

La quatrième et dernière colonne du tableau ci-dessus implique que tout ensemble local L est invariant par complémentation locale, ainsi seule la structure interne de L change. Grâce à la deuxième ligne, on peut déplacer le sommet y dans X , si $y \in \text{Imp}_G(X)$. Grâce à la troisième ligne, on peut déplacer des sommets hors de X tant qu'il existe un sommet dans X qui possède un nombre impair de voisins dans X . Si tous les sommets dans X ont un nombre pair de voisins dans X , et si tout sommet dans X a un voisin z dans $\text{Pair}_G(X)$, alors une complémentation locale appliquée à z crée au moins deux sommets dans X qui possèdent un nombre impair de voisins dans X . L'un d'entre eux doit être un sommet différent de y . Ainsi, par une suite de complémentations locales on peut transformer G en un certain graphe G' dans lequel il n'y a aucune arête entre X et $\text{Pair}_{G'}(X)$, et dans lequel $y \in X$. Par conséquent, $N_{G'}(y) \subseteq L$ et ainsi $\{y\} \cup \text{Imp}_{G'}(\{y\}) \subseteq L$. $\square$

Théorème 11.4 (Caractérisation du degré minimum local) *Pour tout graphe G , le degré minimum local $\delta_{\text{loc}}(G)$ est égal à :*

1. $\min \{ \delta(G') : G' \approx_{\text{loc}} G \}$.
2. $\min \{ |L| : L \text{ est non vide et local} \} - 1$.
3. $\min \{ |X| : \text{Cutrk}(X) < |X| \} - 1$.

Preuve : On montre d'abord que la quantité dans (1) est une borne supérieure de la quantité en (2). Soit $y \in V$ un sommet de degré $\delta_{\text{loc}}(G)$ dans $G' \approx_{\text{loc}} G$. Alors,

$\{y\} \cup N_{G'}(y)$ est local. De manière similaire, on montre que la quantité dans (2) est une borne supérieure de la quantité (3). Soit $L = X \cup \text{Imp}_G(X)$ local. Alors, $\chi_X \Gamma_G[L, V \setminus L] = 0$, où χ_X est la fonction indicatrice de X dans L , et ainsi, L n'est pas un rang de coupe plein. Finalement, on montre que la quantité dans (3) est une borne supérieure de la quantité dans (1). Soit $X \subset V$ un ensemble qui n'est pas un rang de coupe plein. Soit $Y \subseteq V \setminus X$ tel que $\chi_Y \Gamma[X, V \setminus X] = 0$. Alors, $\text{Imp}_G(Y) \subset X$. Grâce au lemme 11.2, pour tout $y \in Y$, il existe un graphe G' localement équivalent à G tel que $\deg_{G'}(y) \leq |X| - 1$. $\square$

D'après le théorème 11.4, pour tout entier fixé d , il existe un algorithme polynomial pour décider si $\delta_{\text{loc}} > d$. Si d fait partie de l'entrée, aucun algorithme polynomial n'est connu. Bien que cela soit plausible, on ne sait pas si le concept de rang de coupe est utile pour calculer δ_{loc} en un temps polynomial. Un résultat dans cette voie est dû à Oum [Oum05], qui fournit un algorithme polynomial qui, pour deux ensembles disjoints et non vides de sommets $A, B \subset V$ en entrée, calcule la valeur $\min\{\text{Cutrk}(Z) : X \subseteq Z \subseteq V \setminus B\}$ en recherchant des suites bloquantes telles qu'introduites par Geelen [Gee95].

La caractérisation ci-dessus est à présent utilisée pour montrer que si aucun qubit auxiliaire n'est disponible pour préparer un état graphe, alors des mesures projectives sur au moins $\delta_{\text{loc}}(G) + 1$ qubits sont nécessaires. Par conséquent, pour tous les graphes pour lesquels $\delta_{\text{loc}} > 1$, il n'existe pas de préparation fondée sur la mesure mettant en oeuvre uniquement des mesures projectives sur 1 et 2 qubits sans utiliser de qubits auxiliaires.

Théorème 11.5 (Borne inférieure de la préparation par mesure) *Soit $G = (V, E)$ un graphe. Toute préparation de $|G\rangle$ par un circuit quantique agissant sur $n = |V|$ qubits et qui consiste en des mesures projectives nécessite des mesures qui agissent sur $\delta_{\text{loc}}(G) + 1$ qubits.*

Preuve : Par contradiction. Supposons que la dernière mesure de la préparation agisse sur un ensemble R d'au plus $\delta_{\text{loc}}(G)$ qubits et que cela produise l'état graphe signé $|G; S\rangle$. Soit W l'observable décrivant cette mesure. Alors, $W|G; S\rangle = |G; S\rangle$, et ainsi $\langle G; S|W|G; S\rangle = 1$.

Soit U un sous ensemble de l'ensemble R des sommets mesurés. Puisque $|R| \leq \delta_{\text{loc}}(G)$, le sous ensemble R est un rang de coupe plein d'après le théorème 11.4, et ainsi, il existe un sous ensemble $T \subseteq V \setminus R$ tel que $\chi_U = \Gamma[R, V \setminus R]\chi_T$. L'opérateur X_T agit uniquement sur les qubits qui ne sont pas dans R , et ainsi commute avec W . De plus, $Z_{\text{Imp}_G(T)}X_T|G; S\rangle = \pm|G; S\rangle$, d'où :

$$\begin{aligned} 1 &= \langle G; S|W|G; S\rangle = \langle G; S|X_TX_TW|G; S\rangle = \langle G; S|Z_{\text{Imp}_G(T)}WZ_{\text{Imp}_G(T)}|G; S\rangle \\ &= \langle G; S|Z_UWZ_U|G; S\rangle = \langle G; S|\Delta U|W|\Delta U\rangle. \end{aligned}$$

Il en découle que W agit trivialement sur $|G; S\Delta U\rangle$ pour tous les sous ensembles $U \subseteq R$. Puisque ces $2^{|R|}$ états sont orthogonaux deux à deux, W est l'identité, ce qui est une contradiction. $\square$

Il est naturel de considérer des méthodes récursives pour préparer un état graphe $|G\rangle$, par exemple en découpant l'ensemble des sommets V en deux parties qui sont considérées individuellement. Le lemme suivant pose le fait que la suppression d'un sommet ou d'une arête quelconque peut faire diminuer le degré local d'au plus une unité.

Lemme 11.3 *Pour tout graphe $G = (V, E)$, tout sommet $u \in V$, et toute arête $e = (v, w) \in E$, $\delta_{\text{loc}}(G \setminus v) \geq \delta_{\text{loc}}(G) - 1$ et $\delta_{\text{loc}}(G \setminus e) \geq \delta_{\text{loc}}(G) - 1$.*

Preuve : Soit $R \subseteq V \setminus \{u\}$ un ensemble de sommets satisfaisant $\text{Cutrk}_{G \setminus u}(R) < |R|$. Alors, $\text{Cutrk}_G(R \cup \{u\}) \leq \text{Cutrk}_{G \setminus u}(R) + 1 < |R \cup \{u\}|$. A présent, soit $R \subseteq V$ un ensemble de sommets satisfaisant $\text{Cutrk}_{G \setminus e}(R) < |R|$, et considérons une arête $e = (v, w) \in E$. Premièrement, si $v, w \in R$ ou $v, w \notin R$, alors $\text{Cutrk}_G(R) = \text{Cutrk}_{G \setminus e}(R)$. Deuxièmement, si $v \in R$ et $w \notin R$, alors $\text{Cutrk}_G(R \cup \{w\}) = \text{Cutrk}_{G \setminus e}(R \cup \{w\}) \leq \text{Cutrk}_{G \setminus e}(R) + 1 < |R| + 1 = |R \cup \{w\}|$. $\square$

Supposons que l'on se donne un oracle $\mathcal{O}_\delta$ qui, pour tout graphe G , fournit $\delta_{\text{loc}}(G)$. Alors, il existe un algorithme déterministe qui, pour tout graphe G fournit un graphe G' localement équivalent à G avec $\delta(G') = \delta_{\text{loc}}(G)$. L'algorithme calcule en un temps polynomial en n et utilise au plus un nombre linéaire en n de questions à l'oracle.

Théorème 11.6 *Les deux problèmes de calcul suivants sont équivalents polynomialement : (1) calculer $\delta_{\text{loc}}(G)$ et (2) trouver un graphe G' avec $G' \equiv G$ et $\delta(G') = \delta_{\text{loc}}(G)$.*

Preuve : Supposons que l'on se donne un oracle qui, avec G en entrée, donne un graphe G' avec $G' \approx_{\text{loc}} G$ et $\delta(G') = \delta_{\text{loc}}(G)$. Alors, on peut trivialement calculer $\delta_{\text{loc}}(G)$ en calculant le degré de chaque sommet dans G' et en sortant le minimum.

Réciproquement, supposons que l'on se donne un oracle calculant δ_{loc} . Soit G un graphe et $d = \delta_{\text{loc}}(G)$. On désire construire un graphe G' qui soit localement équivalent à G et qui possède un sommet v_0 de degré d . Notons que $\{v_0\} \cup N_{G'}(v_0)$ est un ensemble local de sommets de taille $d + 1$.

Notre algorithme est récursif. Il y a trois cas dans chaque étape de la récursivité. D'abord, on calcule $\delta_{\text{loc}}(G \setminus w)$ pour chaque sommet $w \in V$.

- Si $\delta_{\text{loc}}(G \setminus w) < d$, on fait l'appel récursif sur $G \setminus w$.
- Si le premier cas ne se produit pas, on fait l'appel récursif sur $(G \star w) \setminus w$ si $\delta_{\text{loc}}((G \star w) \setminus w) < d$ pour un certain sommet $w \in V$.

- Enfin, si le premier cas ne se produit pas, on fait l'appel récursif sur $(G \star y \star w) \setminus w$ si $\delta_{\text{loc}}((G \star y \star w) \setminus w) < d$ pour une certaine paire de sommets $y, w \in V$.

On déduit alors de la preuve du lemme 11.2 qu'un des trois cas doit se présenter. Soit $L \subseteq V$ un ensemble local de sommets de taille $d+1$, et soit $L = X \cup \text{Imp}_G(X)$:

- Si $\text{Imp}_G(X)$ est non vide, alors le premier cas ci-dessus s'applique puisque $L \setminus w$ est local dans $G \setminus w$ pour chaque $w \in \text{Imp}_G(X)$.
- Si $\text{Imp}_G(X)$ est vide, mais s'il existe un sommet $w \in X$ tel que $|N_X(w)|$ soit impair, alors $\text{Imp}_{G \star w}(X \setminus w) = \{w\}$ d'après la preuve du lemme 11.2, et par conséquent, $L \setminus w$ est local dans $(G \star w) \setminus w$.
- Au final, si $\text{Imp}_G(X)$ est vide et $|N_X(w)|$ est pair pour tous les $w \in X$, alors, il existe un sommet $y \in \text{Pair}_G(X)$ tel qu'une complémentation locale sur y crée au moins deux sommets dans X ayant un nombre impair de voisins dans X et que le deuxième cas s'applique.

La récursivité s'arrête lorsqu'on trouve un graphe H contenant un sommet v_0 de degré 1. Soient $u_1, u_2, \dots, u_k$ les complémentations locales appliquées pendant la récursivité, dans cet ordre. Alors, $G' = G \star u_1 \star u_2 \star \dots \star u_k$ est un graphe dans lequel $\deg_{G'}(v_0) = d$. $\square$

11.5 Séparabilité et δ_{loc}

Un état sur n qubits $|\varphi\rangle$ est *séparable* s'il existe une partition A, B des qubits qui le compose telle que $|\varphi\rangle = |\varphi\rangle_A \otimes |\varphi\rangle_B$.

Théorème 11.7 *Pour tout graphe $G = (V, E)$, il existe $K \subseteq V$ de taille $\delta_{\text{loc}}(G)$, une mesure destructrice agissant sur K et une bipartition (A, B) de $V \setminus (\{v\} \cup K)$ tels que l'état obtenu après la mesure est séparable par rapport à la partition A, B , quel que soit le résultat classique de la mesure.*

Preuve : Soit G' un graphe localement équivalent à G qui contient un sommet v de degré $d = \delta_{\text{loc}}(G)$. Soit $K = N_{G'}(v)$. On a $|G'\rangle = \Pi_{u \in K} \Lambda Z_{v,u} |G' \setminus v\rangle |+\rangle_v$. Considérons la mesure de $|G'\rangle$ dans la base standard des qubits de K . Cette mesure peut être décrite par la mesure individuelle de chaque qubit de K selon Z . D'après le lemme 7.2, effectuer pour chaque $w \in K$ une mesure selon Z_w de l'état $|G'\rangle = \Pi_{u \in K} \Lambda Z_{v,u} |G' \setminus v\rangle |+\rangle_v$ est équivalent à effectuer pour chaque $w \in K$ une mesure selon $(\Pi_{u \in K} \Lambda Z_{v,u})^\dagger Z_w (\Pi_{u \in K} \Lambda Z_{v,u})$ de l'état $|G' \setminus v\rangle |+\rangle_v$, puis appliquer $\Pi_{u \in K} \Lambda Z_{v,u}$. Or, pour tout $w \in K$, $(\Pi_{u \in K} \Lambda Z_{v,u})^\dagger Z_w (\Pi_{u \in K} \Lambda Z_{v,u}) = Z_w$.

Les mesures selon Z_w , pour $w \in K$, transforment $|G' \setminus v\rangle |+\rangle_v$ en $|\psi\rangle = |\varphi\rangle_{V \setminus (\{v\} \cup K)} |1\rangle_S |0\rangle_{K \setminus S} |+\rangle_v$, avec $S \subseteq K$. $\Pi_{u \in K} \Lambda Z_{v,u}$ transforme $|\psi\rangle$ en $|\psi'\rangle = Z_v^{|S|} |\varphi\rangle |1\rangle_S |0\rangle_{K \setminus S} |+\rangle_v$. On remarque que l'état des qubits non mesurés est séparable selon la partition $(\{v\}, V \setminus (K \cup \{v\}))$.

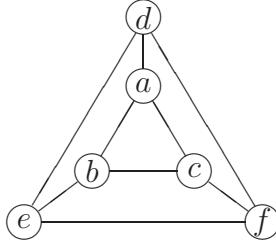


FIG. 11.2 – Le prisme P est un graphe à six sommets. Il existe une mesure des sommets b et e de $|G\rangle$ créant un état séparable, alors que $\delta_{\text{loc}}(P) = 3$.

Il existe C un opérateur de Clifford local tel que $|G'\rangle = C|G\rangle$. On remarque que si pour chaque $u \in K$, le graphe $|G\rangle$ est mesuré selon $C_u^\dagger Z_u C_u$, alors l'état produit est séparable selon la partition $(\{v\}, V \setminus (K \cup \{v\}))$, quel que soit le résultat des mesures.

□

La borne supérieure $\delta_{\text{loc}}(G)$ du nombre de qubits à mesurer pour créer un état séparable est également une borne inférieure pour certains graphes, mais pas pour tous. Un exemple où la borne inférieure n'est pas atteinte est le prisme P sur six sommets, illustré figure 11.2. Le degré minimum de P est 3 et tout ensemble local dans P a une taille d'au moins 4, et ainsi son degré minimum local est 3, d'après le théorème 11.4. En revanche, il existe une mesure de seulement deux qubits, b et e , dans $|P\rangle$, telle que l'état résultant est séparable par rapport à la coupe $(\{a, d\}, \{c, f\})$.

En effet, si l'arête (b, e) est supprimée, puis si une complémentation locale est appliquée sur b et sur e , suivie d'une suppression des sommets b et e , alors le graphe résultant n'est pas connexe. Or, ces transformations sur les graphes sont interprétées en transformations quantiques qui sont locales ou sur les qubits b, e . D'après le lemme 7.2 sur la commutation entre observable et transformation unitaire, il existe une mesure des qubits b et e telle que l'état résultant est séparable.

Il serait intéressant d'explorer davantage la relation entre δ_{loc} et la séparabilité. Une version faible de la caractérisation de la séparabilité par δ_{loc} peut être envisagée. En effet, la mesure effectuée dans le contre-exemple du prisme est non locale, c'est-à-dire qu'elle ne peut pas être décomposée en deux mesures sur un qubit. Il est donc naturel de se demander s'il est possible de créer une séparation en effectuant moins de δ_{loc} mesures *locales*. Cette version faible de la caractérisation de la séparabilité, qui reste un problème ouvert, est pertinente dans l'utilisation des états graphes dans le calcul par consommation d'intrication car seules des mesures locales sont autorisées.

11.6 Bornes inférieures sur δ_{loc}

Puisque que le degré minimum local d'un graphe intervient aussi bien dans la caractérisation de la séparabilité d'un état graphe, que dans la complexité de sa préparation, il est important de pouvoir borner cette quantité en fonction de la taille du graphe.

Les graphes ayant un degré minimum local égal à 1 sont des graphes ayant un sommet pendant (sommet de degré un), ou deux sommets jumeaux (c'est-à-dire deux sommets u et v tels que $N_G(u) \setminus \{v\} = N_G(v) \setminus \{u\}$). En effet, pour tout graphe G tel que $\delta_{\text{loc}}(G) = 1$, il existe un ensemble X de taille 2 tel que $X \cup \text{Imp}_G(X)$ est local. Si X est de taille 1, alors $\text{Imp}_G(X)$ est exactement le voisinage de X qui est également de taille 1, ce sommet est donc pendant. Si X est composé de deux sommets alors $\text{Imp}_G(X) = \emptyset$, les deux sommets ont alors le même voisinage à l'extérieur de X , ils sont jumeaux.

Notons que le prisme (voir figure 11.2) a un degré minimum local égal à 3.

À présent, nous montrons qu'il existe une famille naturelle de graphes pour laquelle δ_{loc} est polynomial en la taille du graphe. Nous montrons d'abord que $\delta_{\text{loc}}(B) \in \Theta(n)$ pour l'hypercube B sur 2^n sommets. Ce résultat soulève la question de savoir si l'état graphe correspondant $|B\rangle$ possède une intrication d'un nouveau type, pouvant être mise à profit dans des protocoles de communication.

Lemme 11.4 *Pour l'hypercube $B = (\{0, 1\}^n, E)$ où $E = \{(x, y) : |x \oplus y| = 1\}$ ⁷, $\delta_{\text{loc}}(B) \geq \frac{n}{2}$.*

Preuve : Soit $\emptyset \subset X \subseteq V$ un ensemble de sommets, et soit $L = X \cup \text{Imp}_G(X)$. On montre à présent que $|L| \geq \frac{n}{2} + 1$. Soit $z \in X$ un sommet de X . Pour chacun des n voisins z^i , pour $i \in \{1 \dots n\}$, de z dans l'hypercube B , si $z^i \notin L$ alors $z^i \in \text{Pair}_G(X)$, i.e., z^i a un nombre pair de voisins dans X , et ainsi $z^{ij} \in X$ pour un certain $j \in \{1 \dots n\} \setminus \{i\}$. Puisque z^{ij} possède seulement deux voisins en commun avec z , l'ensemble L contient au moins $\frac{n}{2}$ éléments parmi $\{z^i : i \in \{1 \dots n\}\} \cup \{z^{ij} : i, j \in \{1 \dots n\}\}$, et ainsi L contient au moins $\frac{n}{2} + 1$ éléments, l'un d'entre eux étant z . D'après le théorème 11.4, le degré minimum local est d'au moins $\frac{n}{2}$. $\square$

On considère également une généralisation naturelle de l'hypercube. Pour tout sous ensemble $H \subseteq \{1 \dots n\}$, soit $B_H = (\{0, 1\}^n, E_H)$ où $E_H = \{(x, y) : |x \oplus y| \in H\}$. Si $H = \{k\}$ est un singleton, on note parfois E_k pour $E_{\{k\}}$, et B_k pour $B_{\{k\}}$. L'hypercube est B_1 . Pour chaque $\ell \in \{0, 1, \dots, n\}$, soit $\mathcal{L}_\ell = \{x \in \{0, 1\}^n : |x| = \ell\}$ indiquant l'ensemble des mots de $\{0, 1\}^n$ de poids de Hamming ℓ .

⁷ $\oplus$ est la somme modulo deux et $|x|$ est le poids de Hamming de $x \in \{0, 1\}^n$.

Proposition 11.3 *Pour tout $\mathcal{H} \subseteq \mathcal{L}_\ell$, soit $\mathcal{K} = \{k \in \mathcal{L}_\ell : |(k \oplus \mathcal{L}_\ell) \cap \mathcal{H}| \text{ est pair}\}$. Alors, $|\mathcal{H} \cup \mathcal{K}| \geq \frac{1}{2}|\mathcal{L}_\ell|$ si ℓ est impair, $\ell = 0$, ou $\ell = 2$.*

Preuve : La proposition est trivialement vraie lorsque $\ell = 0$ ou ℓ est impair. Supposons $\ell = 2$. Soit $J_{\text{impair}} = \{x \in \mathcal{L}_1 : |(x \oplus \mathcal{L}_1) \cap \mathcal{H}| \text{ est impair}\}$, et soit $J_{\text{pair}} = \mathcal{L}_1 \setminus J_{\text{impair}}$. Pour au moins la moitié de toutes les paires distinctes $x, y \in \mathcal{L}_1$, on a que, soit les deux $x, y \in J_{\text{impair}}$, soit les deux $x, y \in J_{\text{pair}}$, dans lesquels cas $x \oplus y \in \mathcal{H} \cup \mathcal{K}$. $\square$

Théorème 11.8 *Le degré minimum local de l'hypercube généralisé B_ℓ est $\delta_{\text{loc}}(B_\ell) \geq \frac{1}{2} \binom{n}{\ell} / \binom{2\ell}{\ell}$ quand ℓ est impair, $\ell = 0$, ou $\ell = 2$.*

Preuve : Soit $\emptyset \subset X \subseteq V$ un ensemble de sommets, et soit $L = X \cup \text{Imp}_G(X)$. On montre à présent que $|L| \geq \frac{1}{2}|\mathcal{L}_\ell| / \binom{2\ell}{\ell}$. Soit $\mathcal{H} = X \cap \mathcal{L}_\ell$ et $\mathcal{K} = \{k \in \mathcal{L}_\ell : |(k \oplus \mathcal{L}_\ell) \cap \mathcal{H}| \text{ est pair}\}$. Alors, chaque $k \in \mathcal{K}$ a un nombre impair de voisins parmi $\mathcal{L}_0 \cup \mathcal{H}$ dans B_ℓ , et ainsi $k \in L$ où il existe un élément $m \in \mathcal{L}_{2\ell} \cap X$ tel que $k \oplus m \in \mathcal{L}_\ell$. Tout élément $m \in \mathcal{L}_{2\ell}$ peut être à une distance Hamming ℓ d'au plus $\binom{2\ell}{\ell}$ éléments de niveau ℓ . Ainsi, L a de manière cardinale au moins $|\mathcal{H} \cup \mathcal{K}| / \binom{2\ell}{\ell}$, qui est au moins $\frac{1}{2}|\mathcal{L}_\ell| / \binom{2\ell}{\ell}$ d'après la proposition 11.3. $\square$

Le corollaire 11.1 suit en posant $\ell = \lfloor n/3 \rfloor$ dans le théorème 11.8.

Corollaire 11.1 *Il existe une constante $c > 0$ et une famille de graphes G pour laquelle $\delta_{\text{loc}}(G) \in \Omega(|G|^c)$.*

Ainsi, il existe une famille naturelle de graphes pour laquelle δ_{loc} est polynomial en la taille du graphe. Cette famille de graphes semble posséder une intrication très robuste contre les opérations quantiques n'agissant que sur un nombre sous linéaire de qubits, et ainsi cela pourrait être utile par exemple en cryptographie quantique et en complexité de communication quantique.

11.7 Conclusion

Les états graphes sont de plus en plus utilisés en informatique quantique. L'étude de leur préparation est importante puisqu'ils constituent la ressource initiale du calcul quantique par consommation d'intrication. Nous avons proposé différents algorithmes de préparation, et montré notamment qu'une préparation de tout état graphe en temps constant est possible.

La complémentation locale joue un rôle important dans le formalisme des états graphes, notamment grâce aux travaux de Van den Nest [VdN05]. Nous avons introduit et étudié des propriétés locales (invariantes par complémentation locale),

comme le degré minimum local, afin de caractériser la complexité de la préparation des états graphes. Ces propriétés locales ont des applications qui dépassent le problème de la préparation des états graphes. Par exemple, nous avons montré le lien entre la robustesse de l'intrication d'un état graphe et le degré minimum local.

Cinquième partie

Conclusion et perspectives

Nous concluons cette thèse en donnant un rapide résumé de ses contributions et en ouvrant des perspectives d'utilisation de ses résultats ainsi que des directions pour des recherches futures.

Contributions

La minimisation des ressources nécessaires à un calcul est l'une des problématiques principales en informatique. Les ressources du calcul classique sont principalement le temps et l'espace nécessaires à la résolution d'un problème. En informatique quantique, cette question des ressources minimales est cruciale. En effet, afin de faciliter la réalisation d'un ordinateur quantique, la diminution des ressources est décisive. Le temps et l'espace ne sont plus les seules ressources à considérer, la taille des opérations utilisées et la quantité d'intrication sont également à prendre en compte.

Cette thèse contribue de plusieurs manières à la recherche de ressources minimales :

- Dans le calcul par mesures projectives, nous avons montré qu'un ensemble de trois observables associé à un qubit auxiliaire constituent des ressources universelles au calcul quantique (chapitre 7, théorème 7.4).
- Dans le cadre du calcul par consommation d'intrication, nous avons montré que les mesures dans le plan (X, Z) sur une grille triangulaire sont universelles (chapitre 10, théorème 10.4).
- Cette thèse a également permis de minimiser les ressources en temps et en espace nécessaires à la préparation d'un état graphe donné (chapitre 11).

Etudier la réduction des ressources nécessite l'abstraction et la formalisation des modèles de calcul quantique. Tout d'abord, il est important de pouvoir formaliser toutes les ressources offertes par la mécanique quantique, afin de dégager les propriétés générales du calcul quantique. Les propriétés ainsi dégagées s'appliqueront aussi bien au calcul par transformations unitaires que par mesures projectives par exemple. De plus, un cadre de travail unificateur permet de représenter dans un formalisme commun les différents modèles de calcul quantique, et par conséquent de pouvoir les comparer. Le $\mathbf{q}$ -calcul (chapitre 10) et les machines de Turing contrôlées classiquement (chapitre 6), introduits dans cette thèse, ont cet objectif. Des modèles plus spécifiques au calcul par consommation d'intrication (m -calcul et m -calcul $\mathcal{3P}$, chapitre 10), ou au calcul par mesures projective (fragment observable du $\mathbf{q}$ -calcul et machine de Turing quantique fondée sur la mesure, chapitre 8) sont alors considérés.

La démarche suivie dans cette thèse consiste à présenter tout d'abord les mo-

dèles les plus généraux du calcul quantique, généralisant les modèles standards, pour ensuite se focaliser sur l'étude de modèles émergents comme le calcul par mesures projectives et le calcul par consommation d'intrication.

Cette thèse contribue donc aussi à l'émergence et à l'étude de nouveaux modèles formels de calcul quantique comme le q -calcul (chapitre 5), les machines de Turing quantiques (chapitre 6) et le m -calcul (chapitre 10).

- Le q -calcul est un modèle formel pour le calcul quantique contrôlé classiquement, qui inclut aussi bien le calcul quantique réversible, que le calcul quantique par mesure. Ce modèle est né de la volonté de pouvoir représenter toutes les primitives offertes par la mécanique quantique, sans se limiter aux transformations unitaires ou aux mesures projectives. La représentation formelle de l'interaction nécessaire entre l'ordinateur quantique et son environnement classique est également un des objectifs du q -calcul. Trois domaines sémantiques ont été utilisés pour définir les sémantiques dénotationnelles des termes du q -calcul : le domaine des états purs, celui des matrices de densités, mais également celui des transformations admissibles, qui est utilisé pour la première fois comme domaine sémantique. Ces trois sémantiques ne sont pas équivalentes, mais une hiérarchie d'abstraction a été montrée : la sémantique admissible est la plus concrète, la sémantique pure en est une interprétation, elle est donc plus abstraite, enfin la sémantique observable (domaine des matrices de densité) est la plus abstraite.

Le q -calcul est un ensemble de règles agissant sur les q -termes. Nous avons montré que différentes règles naturelles de transformations des q -termes préservent tout ou partie de la hiérarchie sémantique. Ces règles de transformations facilitent le calcul de la sémantique des termes du q -calcul et permettent de montrer que deux termes sont équivalents. Par exemple, l'équivalence entre une transformation unitaire et sa simulation par des mesures projectives, a été montrée dans le cadre du q -calcul. Le développement ultérieur du q -calcul est un objectif qui est décrit dans la section *Perspective de recherche*.

- Les machines de Turing quantiques contrôlées classiquement ont, comme le q -calcul, l'objectif de permettre une formalisation de l'interaction entre le système quantique et son environnement. Complémentaires du q -calcul, les machines de Turing permettent de définir la notion de complexité. Par conséquent, elles sont utilisées comme une plateforme permettant de comparer le pouvoir du calcul quantique contrôlé classiquement aux modèles plus restreints de calcul quantique réversible, par mesures, et même de calcul classique. C'est ainsi, par exemple, qu'une séparation entre calcul classique et quantique a été mise en évidence.
- Le m -calcul est un modèle formel introduit par Danos, Kashefi et Panangaden [DKP04a] pour le calcul par consommation d'intrication. En se fondant sur

le m -calcul, Danos et Kashefi [DK05] ont montré que si un état graphe admet un flot simple, alors il peut être utilisé comme ressource pour le calcul par consommation d'intrication. De plus, la description de ce flot simple permet de majorer le temps d'exécution du calcul. Une généralisation de ce flot a été introduite dans cette thèse. Cette nouvelle condition de flot permet d'utiliser des états graphes pour le calcul par consommation d'intrication qui n'admettent pas de flot simple. Cette généralisation permet également de diminuer le temps d'exécution du calcul.

Une généralisation du m -calcul, le m -calcul $3P$ a été introduite, permettant la représentation de calculs par consommation d'intrication avec utilisation de mesures dans les trois plans possibles, alors que le m -calcul n'en utilise qu'un seul. Cette généralisation a servi à montrer l'universalité des mesures dans le plan (X, Z) sur une grille triangulaire. Cette universalité ouvre de nouvelles perspectives de réalisations physiques, en réduisant les ressources du calcul par consommation d'intrication.

Perspectives de recherches futures

Mes perspectives de recherches sont de deux types :

- Tout d'abord, poursuivre les pistes ouvertes pendant cette thèse : le q -calcul, les machines de Turing quantiques contrôlées classiquement, le m -calcul $3P$, mais aussi la caractérisation des ressources minimales du calcul par mesures projectives et par consommation d'intrication, sans oublier une meilleure compréhension des états graphes, pour imaginer de nouvelles applications de ceux-ci.
- Mais je tiens également à explorer de nouvelles pistes, comme l'analyse de propriétés quantiques, par exemple l'intrication, en utilisant des méthodes d'interprétations abstraites [CC77]. Des résultats préliminaires [Per07], non présentés dans cette thèse, montrent la faisabilité et le potentiel de l'analyse statique, et plus généralement de l'interprétation abstraite, pour l'étude de propriétés quantiques comme l'intrication.

Avant de rentrer dans une brève description de ces perspectives de recherches, il convient de s'interroger sur l'avenir du domaine de l'informatique quantique.

Le long terme est évidemment impossible à prévoir, et je ne me risquerai pas à formuler une réponse à la brûlante et épineuse question qui consiste à se demander si l'ordinateur quantique verra ou non le jour. En revanche, à plus court terme, je pense que les modèles alternatifs de calcul quantique, et en particulier le modèle de calcul par consommation d'intrication, vont devenir des candidats de plus en plus incontournables à la réalisation de l'ordinateur quantique. Mes perspectives

sont donc de contribuer au développement de ces modèles alternatifs, en particulier ceux fondés sur la mesure quantique. La réduction des ressources nécessaires à un calcul quantique ne fait qu'alléger le cahier des charges de la réalisation de l'ordinateur quantique, il est donc important de poursuivre la quête vers les ressources minimales. De plus, le pouvoir d'un ordinateur par consommation d'intrication est encore mal connu. Ainsi, une meilleure connaissance des états graphes, par la caractérisation combinatoire de leurs propriétés, constitue une étape importante vers une meilleure connaissance du modèle de calcul par consommation d'intrication.

La caractérisation combinatoire des propriétés quantiques s'avère donc être une piste à explorer. Une étude préliminaire montre que cette piste ne se limite pas à l'étude des états graphes, mais qu'une nouvelle classe d'états quantiques, les *états hyper-graphes* peut être considérée. En effet, il est pertinent d'introduire un nouveau formalisme, associant à chaque hyper-graphe H , un état quantique $|H\rangle$. Comme pour les états graphes, plusieurs définitions peuvent être considérées, notamment une définition constructive : en partant de n qubits dans l'état $|+\rangle$, $|H\rangle$ est obtenu en appliquant pour chaque hyper-arête $S \subseteq V$, une opération contrôlée $\Lambda^{|S|-1}Z_S$.

Ce nouveau formalisme est plus expressif que celui des états graphes, dépassant même l'expressivité des états stabilisables. En effet, l'état $|W\rangle = \frac{1}{\sqrt{3}}(|001\rangle + |010\rangle + |001\rangle)$, qui n'est pas stabilisable et possède une intrication particulièrement résistante aux mesures, peut être représenté par un hyper-graphe.

Une étude des propriétés combinatoires de ce formalisme permettra de savoir par exemple si la LC-équivalence peut être caractérisée par une transformation qui généraliserait la complémentation locale.

L'informatique quantique invite aussi à reconsidérer, avec un regard nouveau, les structures fondamentales de la théorie quantique. En effet, de la reformulation de la mécanique quantique via la théorie des catégories [AC02, Coe04, AC04a, AC04b, Sel05, Abr04b, AD06], aux développements de langages de programmation quantique [Gay06], de nouvelles approches prometteuses voient le jour et devraient permettre de mieux comprendre et exploiter les liens profonds qu'il y a entre la mécanique quantique et le traitement de l'information.

Citons deux problèmes, *a priori* inattendus, qui ont émergé récemment :

- La formalisation et l'abstraction du calcul quantique semblent se confronter au problème de l'*ordre supérieur*. Alors qu'en informatique *classique*, une fonction peut prendre comme argument une autre fonction (principe de base du λ -calcul), Selinger [Sel04c] a soulevé la difficulté de placer à un même niveau, les données quantiques et les fonctions quantiques.
- Une opération aussi simple que l'échange de deux qubits (*swap*) soulève des

problèmes intéressants en informatique quantique. En effet, Jozsa [Joz06] a montré que le gain du calcul quantique, par rapport au calcul classique se réduit à l'opération d'échange. En effet, si on est capable de simuler classiquement et efficacement l'échange de deux qubits, alors l'algorithme de Shor [Sho94] peut être simulé en temps polynomial sur un ordinateur classique.

Ce résultat surprenant trouve peut être une explication catégorique. En effet, une catégorie proche de celle des tresses pourrait être utilisée pour reformuler l'informatique quantique [Zha06]. Or, deux échanges successifs lors de la manipulation de tresses, ne reviennent pas toujours à ne rien faire, d'où le rôle singulier de l'échange.

Une perspective serait donc de tenter de caractériser la puissance de l'opération d'échange. En effet, une séparation devrait pouvoir être obtenue entre les machines à accès aléatoire, comme la QRAM [Kni96, BCS03] et les machines à accès séquentiel comme la machine de Turing quantique contrôlée classiquement, puisque l'adressage aléatoire peut être simulé par des échanges successifs sur une machine séquentielle. Cette séparation existe-t-elle ? Si oui, lequel de ces deux types de machines est un modèle réaliste pour l'informatique quantique ?

Une perspective supplémentaire est de développer le q -calcul et notamment tenter d'obtenir des règles de réécriture terminantes et confluantes permettant d'atteindre une forme normale. Une forme normale semble naturellement se dessiner, celle des termes sans récursion, composés uniquement de processus d'entrée et de sortie. Pour atteindre cette forme, les définitions récursives doivent être déroulées. Pour effectuer ce déroulement il est naturel de considérer des techniques d'interprétation abstraites utilisées en informatique classique [CC77]. Etablir un domaine abstrait et des règles de réécriture confluantes et terminantes agissant sur ce domaine abstrait est une des perspectives.

L'informatique quantique est une science jeune, née de la collaboration de physiciens, de mathématiciens et d'informaticiens. Son développement a été rapide, notamment grâce à des résultats majeurs en algorithmique [Sho94, Gro96a, DHHM06]. Son développement continue aujourd'hui avec de nouvelles perspectives pour le traitement de l'information quantique, qui dépassent les aspects algorithmiques, comme par exemple la reformulation de la mécanique quantique, ou l'utilisation des résultats d'informatique quantique en théorie des graphes, à laquelle nous avons modestement contribué dans cette thèse. Le traitement de l'information quantique évolue également de part les techniques auxquelles elle fait appel. En effet, la théorie des graphes, la théorie des catégories, et peut être demain l'interprétation abstraite, sont des domaines de recherche qui permettent

de faire progresser le domaine du traitement de l'information quantique. Le défi de l'informatique quantique est peut être de gérer cette multiplication des liens avec l'informatique théorique, qui sont extrêmement bénéfiques, tout en tirant le meilleur parti de l'une des richesses de l'informatique quantique, qui est de regrouper physiciens, mathématiciens et informaticiens, qui apportent des points de vue et des méthodes différentes, autour d'un projet commun, le traitement de l'information quantique.

Publications

Revues Internationales avec comité de lecture

- [Per05] Simon Perdrix. State transfer instead of teleportation in measurement-based quantum computation. *International Journal of Quantum Information*, 3(1) :219–223, 2005.
- [PJ06a] Simon Perdrix and Philippe Jorrand. Classically-controlled quantum computation. *Math. Struct. in Comp. Science*, 16 :601–620, 2006.

Conférences internationales avec actes et comité de sélection

- [Per04] Simon Perdrix. Qubit vs observable resource trade-offs in measurement-based quantum computation. In *proceedings of Quantum communication measurement and computing*, 2004.
- [JP05b] Philippe Jorrand and Simon Perdrix. Unifying quantum computation with projective measurements only and one-way quantum computation. In *Quantum Informatics 2004, Proceedings of the SPIE*, volume 5833, pages 44–51, 2005.
- [JP05a] Philippe Jorrand and Simon Perdrix. Resources for measurement-based quantum computation : A unifying view. In *proceedings of Quantum Information, Computation and Communication*, pages 111–120, 2005.
- [JP05b] Philippe Jorrand and Simon Perdrix. Unifying quantum computation with projective measurements only and one-way quantum computation. In *Quantum Informatics 2004, Proceedings of the SPIE*, volume 5833, pages 44–51, 2005.
- [Per07] Simon Perdrix. Quantum patterns and types for entanglement and separability. *To appear in Proceedings of the 3rd International Workshop on Quantum Programming Languages, ENTCS*, 2005.
- [HMP06a] Peter Hoyer, Mehdi Mhalla, and Simon Perdrix. Resources required for preparing graph states. In *Lecture Notes in Computer Science, Proceedings of ISAAC06 (To appear)*, 2006.
- [PJ06b] Simon Perdrix and Philippe Jorrand. Classically-controlled quantum computation. *Proceedings of DCM05, ENTCS*, 135(3) :119–128, 2006.

- [JP06] Philippe Jorrand and Simon Perdrix. Towards a quantum calculus. *To appear in Proceedings of the 4th International Workshop on Quantum Programming Languages, ENTCS*, 2006.

Conférences Nationales avec actes et comité de sélection

- [MP06] Mehdi Mhalla et Simon Perdrix, Etats graphes et calcul quantique. *Journées Graphes et Algorithmes, Orléans*, Novembre 2006.

Conférences Internationales avec comité de sélection

- Simon Perdrix. Measurement-based Quantum Computation. *FQI, Camerino (Italy)*, 2004.
- Simon Perdrix and Philippe Jorrand. Classically-Controlled Quantum Computation *EQIS Tokyo*, September 2004.
- Simon Perdrix. Optimizing the Resources for Measurement-based Preparation of Graph States *5th European QIPC Rome*, September 2004.
- Simon Perdrix and Mehdi Mhalla. Complexity of Graph State Preparation, *QIP Boston*, January 2005.
- Simon Perdrix, Minimal Resources of Measurement-based Quantum Computation. *Quantum Information, Computation and Logic : Exploring New Connections* PI Waterloo, Canada, July 2005.

Conférences Internationales

- Simon Perdrix. A stabilizer-based model for one-way quantum computation. *IQING II, London*, September 2002.
- Simon Perdrix. Around the measurement-based quantum computation. *IQING III, Munich*, December 2003.
- Simon Perdrix. Complexity of Graph State Preparation. *Q-Day I*, Paris, Decembre 2004.
- Simon Perdrix, Analysing and Handling Entanglement in Quantum Computation. *Q-Day II*, Paris, December 2005.

Colloques nationaux

- Simon Perdrix, Autour de l'universalité de la mesure quantique. *Colloque du GdR 2285 Information et communication Quantiques*, Les Houches, Mai 2003.
- Simon Perdrix. Peut-on "tout" faire avec le calcul quantique? *Journée "Co-hérence et information quantique"*, Grenoble, Février 2004.

- Simon Perdrix, Complexity of Graph State Preparation. *Colloque du GdR 2285 Information et communication Quantiques*, Orsay, Décembre 2004.
- Simon Perdrix, Graph States, *33ème École de printemps d’informatique théorique*, Montagnac-les-truffes, Juin 2005.
- Simon Perdrix, Graph State Preparation, *Journée quantique and co.*, ENS de Lyon, Novembre 2005.

Preprints et rapports de recherche

- Philippe Jorrand and Simon Perdrix. Non-probabilistic termination of measurement-based quantum computation. arXiv.org preprint quant-ph/0311142
- Philippe Jorrand and Simon Perdrix. Measurement-based quantum Turing machines and questions of universalities. arXiv.org preprint quant-ph/0402156.
- Simon Perdrix. State Transfer instead of Teleportation in Measurement-based Quantum Computation. arXiv.org preprint quant-ph/0402204.
- Philippe Jorrand and Simon Perdrix. Unifying Quantum Computation with Projective Measurements only and One-Way Quantum Computation. arXiv.org preprint quant-ph/0404125.
- Philippe Jorrand and Simon Perdrix. Measurement-based quantum Turing machines and their universality. arXiv.org preprint quant-ph/0404146.
- Simon Perdrix and Philippe Jorrand. Classically-Controlled Quantum Computation. arXiv.org preprint quant-ph/0407008.
- Mehdi Mhalla and Simon Perdrix. Complexity of graph state preparation. arXiv.org preprint quant-ph/0412071.
- Peter Hoyer, Mehdi Mhalla, and Simon Perdrix. Resources required for preparing graph states. *Cahier du Laboratoire Leibniz*, 151. 2006.

Bibliographie

- [Abr04a] S. Abramsky. A cook’s tour of a simple quantum programming language. *3rd International Symposium on Domain Theory, Xi’an, China*, May 2004.
- [Abr04b] S. Abramsky. High-level methods for quantum computation and information. In *Proceedings of the 19th Annual IEEE Symposium on Logic in Computer Science (LICS)*. IEEE Computer Society, 2004.
- [AC02] Samson Abramsky and Bob Coecke. Physical traces : Quantum vs. classical information processing. *Electr. Notes Theor. Comput. Sci.*, 69, 2002.
- [AC04a] S. Abramsky and B. Coecke. A categorical semantics of quantum protocols. In *Proceedings of the 19th Annual IEEE Symposium on Logic in Computer Science (LICS)*. IEEE Computer Society, 2004. Also arXiv :quant-ph/0402130.
- [AC04b] Samson Abramsky and Bob Coecke. A categorical semantics of quantum protocols. In *LICS*, pages 415–425, 2004.
- [AD04] P. Arrighi and G. Dowek. Operational semantics for formal tensorial calculus. In Selinger [Sel04a].
- [AD05] P. Arrighi and G. Dowek. Linear-algebraic λ -calculus. arXiv :quant-ph/0501150, 2005.
- [AD06] S. Abramsky and R. Duncan. A categorical quantum logic. *Mathematical Structures in Computer Science*, 16(3) :469–489, 2006.
- [AG05] T. Altenkirch and J. Grattage. A functional quantum programming language. In *Proceedings of the 20th Annual IEEE Symposium on Logic in Computer Science (LICS)*. IEEE Computer Society, 2005. Also arXiv :quant-ph/0409065.
- [AGR81] A. Aspect, P. Grangier, and G. Roger. Experimental tests of realistic local theories via Bell’s theorem. *Phys. Rev. Lett.*, 47 :460, 1981.

- [Aha98] D. Aharonov. Quantum computation. In Dietrich Stauffer, editor, *Annual Reviews of Computational Physics*, volume 6. World Scientific, 1998.
- [Aha99] D. Aharonov. *Noisy Quantum Computation*. PhD thesis, Hebrew University, 1999.
- [AJ94] S. Abramsky and A. Jung. Domain theory. In S. Abramsky, D. Gabbay, and T. S. E. Maibaum, editors, *Handbook of Logic in Computer Science Volume 3*, pages 1–168. Oxford University Press, 1994.
- [BB06] D. E. Browne and H. J. Briegel. One-way quantum computation - a tutorial introduction, 2006.
- [BBC⁺93] C. H. Bennett, G. Brassard, C. Crépeau, R. Jozsa, A. Peres, and W. K. Wootters. Teleporting an unknown quantum state via dual classical and Einstein-Podolsky-Rosen channels. *Phys. Rev. Lett.*, 70 :1895–1899, 1993.
- [BCS03] S. Bettelli, T. Calarco, and L. Serafini. Toward an architecture for quantum programming. *The European Physical Journal D*, 25 :181–200, 2003.
- [Ben80] P. Benioff. The computer as a physical system : A microscopic quantum mechanical hamiltonian model of computers as represented by Turing machines. *J. Stat. Phys.*, 22 :563–591, 1980.
- [Ben82] P. Benioff. Quantum mechanical Hamiltonian models of Turing machines. *J. Stat. Phys.*, 29 :515, 1982.
- [Boh52] D. Bohm. A suggested interpretation of the quantum theory in terms of "hidden" variables. i. *Physical Review*, 85 :166–179, January 1952.
- [Bou87] A. Bouchet. Digraph decompositions and eulerian systems. *SIAM Journal on Algebraic and Discrete Methods*, 8(3) :323–337, 1987.
- [Bou89] A. Bouchet. Connectivity of isotropic systems. In New York Academy of Sciences, editor, *Proceedings of the third international conference on Combinatorial mathematics*, pages 81–93, 1989.
- [BPM⁺97] D. Bouwmeester, J. W. Pan, K. Mattle, M. Eibl, H. Weinfurter, and A. Zeilinger. Experimental quantum teleportation. *Nature*, 390 :575–579, 1997.
- [BV93] E. Bernstein and U. Vazirani. Quantum complexity theory. In *Proceedings of the 25th Annual ACM Symposium on the Theory of Computation*, pages 11–20, New York, 1993. ACM press.
- [BV97] E. Bernstein and U. Vazirani. Quantum complexity theory. *SIAM J. Comput.*, 26 :1411–1478, 1997.

- [BvN36] Garrett Birkhoff and John von Neumann. The logic of quantum mechanics. *Annals of Mathematics*, 37 :823–843, 1936.
- [CC77] Patrick Cousot and Radhia Cousot. Abstract interpretation : A unified lattice model for static analysis of programs by construction or approximation of fixpoints. In *POPL*, pages 238–252, 1977.
- [CLN05] A. M. Childs, D. W. Leung, and M. A. Nielsen. Unified derivations of measurement-based schemes for quantum computation. *Phys. Rev. A*, 71 :032318, 2005.
- [Coe04] B. Coecke. Quantum information-flow, concretely, abstractly. In Selinger [Sel04a].
- [Cou97] P. Cousot. Types as abstract interpretations. In *POPL*, pages 316–331, 1997.
- [CVZ⁺98] I. L. Chuang, L. M. K. Vandersypen, X. Zhou, D. W. Leung, and S. Lloyd. Experimental realization of a quantum algorithm. *Nature*, 393 :143–146, 1998.
- [Deu85] D. Deutsch. Quantum theory, the Church-Turing principle and the universal quantum computer. *Proc. R. Soc. Lond. A*, 400 :97–117, 1985.
- [DHHM06] C. Dürr, M. Heiligman, P. Høyer, and M. Mhalla. Quantum query complexity of some graph problems. *SIAM J. Comput.*, 35(6) :1310–1328, 2006.
- [Dir47] P. A. M. Dirac. *The Principles of Quantum Mechanics*. The international series of monographs on physics. Clarendon Press, Oxford, 4 edition, 1947.
- [DK05] V. Danos and E. Kashefi. Determinism in the one-way model. arXiv :quant-ph/0506062, 2005.
- [DKP04a] V. Danos, E. Kashefi, and P. Panangaden. The measurement calculus. arXiv :quant-ph/0412135, 2004.
- [DKP04b] V. Danos, E. Kashefi, and P. Panangaden. Robust and parsimonious realisations of unitaries in the one-way model. arXiv :quant-ph/0411071, 2004.
- [dNDVB07] M. Van den Nest, W. Dur, G. Vidal, and H. J. Briegel. Classical simulation versus universality in measurement based quantum computation. *Physical Review A*, 75 :012337, 2007.
- [EPR35] A. Einstein, B. Podolsky, and N. Rosen. Can quantum-mechanical description of reality be considered complete? *Phys. Rev.*, 47(10) :777–780, May 1935.

- [FDJY06] Yuan Feng, Runyao Duan, Zhengfeng Ji, and Mingsheng Ying. Probabilistic bisimilarities between quantum processes. arXiv :cs/0601014, 2006.
- [Fey82] R. P. Feynman. Simulating physics with computers. *International Journal of Theoretical Physics*, 21 :467–488, 1982.
- [Fey84] R. P. Feynman. Quantum-mechanical computers. *J. Opt. Soc. Am. B*, 1 :464, 1984.
- [Fey86] R. P. Feynman. Quantum-mechanical computers. *Found. Phys.*, 16 :507–531, 1986.
- [FZ01] S. A. Fenner and Y. Zhang. Universal quantum computation with two- and three-qubit projective measurements, 2001.
- [Gay06] S. J. Gay. Quantum programming languages : Survey and bibliography. *Mathematical Structures in Computer Science*, 16(4), 2006.
- [Gee95] James F. Geelen. *Matchings, Matroids and Unimodular Matrices*. PhD thesis, University of Waterloo, 1995.
- [Gim05] S. Gimenez. Correspondance entre modèles de calcul quantique one-way et measurement-only. Master’s thesis, Université Paris 7, 2005.
- [Gle57] A. M. Gleason. Measures on the closed subspaces of a Hilbert space. *J. Math. and Mech.*, 6 :885–893, 1957.
- [GN05] S. J. Gay and R. Nagarajan. Communicating quantum processes. In *Proceedings of the 32nd ACM SIGACT-SIGPLAN Symposium on Principles of Programming Languages (POPL)*. ACM Press, 2005. Preliminary version in [Sel04a] ; also arXiv :quant-ph/0409052.
- [Gro96a] L. Grover. A fast quantum mechanical algorithm for database search. In *Proc. 28th Annual ACM Symposium on the Theory of Computation*, pages 212–219, New York, NY, 1996. ACM Press, New York.
- [Gro96b] Lov K. Grover. A fast quantum mechanical algorithm for database search. In *STOC*, pages 212–219, 1996.
- [Gru99] J. Gruska. *Quantum computing*. McGraw–Hill, New York, 1999.
- [HDE⁺05] M. Hein, W. Dür, J. Eisert, R. Raussendorf, M. Van den Nest, and H. J. Briegel. Entanglement in graph states and its applications. In *Proceedings of the International School of Physics “Enrico Fermi” on “Quantum Computers, Algorithms and Chaos”*, July 2005.
- [HEB04] M. Hein, J. Eisert, and H. J. Briegel. Multi-party entanglement in graph states. *Physical Review A*, 69 :062311, 2004.

- [HMP06a] P. Hoyer, M. Mhalla, and S. Perdrix. Resources required for preparing graph states. In *Proceedings 17th International Symposium on Algorithms and Computation (ISAAC 2006), Lecture Notes in Computer Science*, volume 4288, pages 638–649. Springer-Verlag, 2006.
- [HMP06b] P. Hoyer, M. Mhalla, and S. Perdrix. Resources required for preparing graph states. *Cahier du Laboratoire Leibniz 151*, 2006.
- [IOV04] S. Iriyama, M. Ohya, and I. Volovich. Generalized quantum Turing machine and its application to the sat chaos algorithm. arXiv quant-ph/0405191, 2004.
- [JL04] Ph. Jorrand and M. Lalire. Toward a quantum process algebra. In *Proceedings of the 1st ACM Conference on Computing Frontiers*. ACM Press, 2004. Also arXiv :quant-ph/0312067.
- [Joz06] Richard Jozsa. On the simulation of quantum circuits, 2006.
- [JP89] C. Jones and G. D. Plotkin. A probabilistic powerdomain of evaluations. In *LICS*, pages 186–195, 1989.
- [JP03] Ph. Jorrand and S. Perdrix. Non-probabilistic termination of measurement-based quantum computation. arXiv.org preprint quant-ph/0311142, 2003.
- [JP04a] Ph. Jorrand and S. Perdrix. Measurement-based quantum Turing machines and questions of universalities. arXiv.org preprint quant-ph/0402156, 2004.
- [JP04b] Ph. Jorrand and S. Perdrix. Measurement-based quantum Turing machines and their universality. arXiv.org preprint quant-ph/0404146, 2004.
- [JP04c] Ph. Jorrand and S. Perdrix. Unifying quantum computation with projective measurements only and one-way quantum computation. arXiv.org preprint quant-ph/0404125, 2004.
- [JP05a] Ph. Jorrand and S. Perdrix. Resources for measurement-based quantum computation : A unifying view. In *proceedings of Quantum Information, Computation and Communication*, pages 111–120, 2005.
- [JP05b] Ph. Jorrand and S. Perdrix. Unifying quantum computation with projective measurements only and one-way quantum computation. In *Quantum Informatics 2004, Proceedings of the SPIE*, volume 5833, pages 44–51, 2005.
- [JP06] Ph. Jorrand and S. Perdrix. Towards a quantum calculus. In *To appear in Proceedings of the 4th International Workshop on Quantum Programming Languages, ENTCS*, 2006.

- [Kas03] E. Kashefi. Quantum domain theory - definitions and applications. In *Proceedings of Computability and Complexity in Analysis (CCA03)*, 2003.
- [Kni96] E. Knill. Conventions for quantum pseudocode, 1996.
- [KSV02] A. Yu. Kitaev, A. H. Shen, and M. N. Vyalyi. *Classical and Quantum Computation*. American Mathematical Society, Boston, MA, USA, 2002.
- [Lal06] M. Lalire. Relations among quantum processes : Bisimilarity and congruence. *Mathematical Structures in Computer Science*, 16(3) :407–428, 2006.
- [Leu04] D. W. Leung. Quantum computation by measurements. *IJQI*, 2 :33, 2004.
- [MP04] M. Mhalla and S. Perdrix. Complexity of graph state preparation. arXiv.org preprint quant-ph/0412071, 2004.
- [MS00] A. Muthukrishnan and C. R. Stroud. Multivalued logic gates for quantum computation. *Phys. Rev. A*, 62(5) :052309, Oct 2000.
- [NC00] M. A. Nielsen and I. L. Chuang. *Quantum computation and quantum information*. Cambridge University Press, New York, NY, USA, 2000.
- [Nie03] M. A. Nielsen. Universal quantum computation using only projective measurement, quantum memory, and preparation of the 0 state. *Phys. Rev. A*, 308 :96–100, 2003.
- [NO02a] H. Nishimura and M. Ozawa. Computational complexity of uniform quantum circuit families and quantum Turing machines. *Theor. Comput. Sci.*, 276(1-2) :147–181, 2002.
- [NO02b] H. Nishimura and M. Ozawa. Computational complexity of uniform quantum circuit families and quantum Turing machines. *Theor. Comput. Sci.*, 276(1-2) :147–181, 2002.
- [Ome03] B. Omer. *Structured Quantum Programming*. PhD thesis, Technical University of Vienna, 2003.
- [Oum05] S.I. Oum. Approximating rank-width and clique-width quickly. In *WG*, pages 49–58, 2005.
- [Per04] S. Perdrix. Qubit vs observable resource trade-offs in measurement-based quantum computation. In *proceedings of Quantum communication measurement and computing*, 2004.
- [Per05] S. Perdrix. State transfer instead of teleportation in measurement-based quantum computation. *International Journal of Quantum Information*, 3(1) :219–223, 2005.

- [Per07] S. Perdrix. Quantum patterns and types for entanglement and separability. In *Proceedings of the 3rd International Workshop on Quantum Programming Languages (QPL 2005)*, *ENTCS*, volume 170, pages 125–138, 2007.
- [PJ06a] S. Perdrix and Ph. Jorrand. Classically-controlled quantum computation. *Math. Struct. in Comp. Science*, 16 :601–620, 2006.
- [PJ06b] S. Perdrix and Ph. Jorrand. Classically-controlled quantum computation. In *Proceedings of DCM05*, *ENTCS*, volume 135, pages 119–128, 2006.
- [RB00] R. Raussendorf and H. J. Briegel. Quantum computing via measurements only, arXiv.org quant-ph/0010033 2000.
- [RB02a] R. Raussendorf and H. J. Briegel. Computational model for the one-way quantum computer : Concepts and summary. arXiv.org quant-ph/0207183, 2002.
- [RB02b] R. Raussendorf and H. J. Briegel. Computational model underlying the one-way quantum computer. *QUANT.INF.COMP.*, 6 :433, 2002.
- [RBB02] R. Raussendorf, D. E. Browne, and H. J. Briegel. The one-way quantum computer - a non-network model of quantum computation. *Journal of Modern Optics*, 49 :1299, 2002.
- [Sel04a] P. Selinger, editor. *Proceedings of the 2nd International Workshop on Quantum Programming Languages*, number 33 in TUCS General Publications. Turku Centre for Computer Science, 2004.
- [Sel04b] P. Selinger. Towards a quantum programming language. *Mathematical Structures in Computer Science*, 14(4) :527–586, 2004.
- [Sel04c] P. Selinger. Towards a semantics for higher-order quantum computation. In *Proceedings of the 2nd International Workshop on Quantum Programming Languages* [Sel04a].
- [Sel05] P. Selinger. Dagger compact closed categories and completely positive maps. In P. Selinger, editor, *Proceedings of the 3rd International Workshop on Quantum Programming Languages*, Electronic Notes in Theoretical Computer Science. Elsevier Science, 2005.
- [Shi03] Y. Shi. Both toffoli and controlled-not need little help to do universal quantum computation. *Quantum Information and Computation*, 3(1) :84–92, 2003.
- [Sho94] P. Shor. Algorithms for quantum computation : Discrete logarithms and factoring. In IEEE Computer Society Press Shafi Goldwasser,

- editor, *Proceedings of the 35nd Annual Symposium on Foundations of Computer Science*, pages 124–134, 1994.
- [SV06] P. Selinger and B. Valiron. A lambda calculus for quantum computation with classical control. *Mathematical Structures in Computer Science*, 16(3) :527–552, 2006.
- [Tur36] A. Turing. On computable numbers with an application to the entscheidungsproblem. In *Proceedings of the London Mathematical Society*, volume 42, 1936.
- [VdN05] M. Van den Nest. *Local equivalence of stabilizer states and codes*. PhD thesis, Faculty of Engineering, K. U. Leuven, Belgium, May 2005.
- [VdNDD04] M. Van den Nest, J. Dehaene, and B. De Moor. Graphical description of the action of local clifford transformations on graph states. *Physical Review A*, 69 :022316, 2004.
- [VdNMDB06] M. Van den Nest, A. Miyake, W. Dür, and H. J. Briegel. Universal resources for measurement-based quantum computation, 2006.
- [Viz64] V. G. Vizing. On an estimate of the chromatic class of a p -graph. *Metody Diskret. Analiz.*, 3 :25–30, 1964.
- [VSB⁺01] L. M. K. Vandersypen, M. Steffen, G. Breyta, C. S. Yannoni, M. H. Sherwood, and I. L. Chuang. Experimental realization of shor’s quantum factoring algorithm using nuclear magnetic resonance. *Nature*, 414 :883–886, December 2001.
- [vT04] A. van Tonder. A lambda calculus for quantum computation. *SIAM Journal on Computing*, 33(5) :1109–1135, 2004. Also arXiv :quant-ph/0307150.
- [Wat00] J. Watrous. Succinct quantum proofs for properties of finite groups. In *Proc. 41st Annual Symposium on Foundation of Computer Science*, pages 537–546, 2000.
- [Wlo69] L. Wlodarski. On the equation $\cos \alpha_1 + \cos \alpha_2 \cos \alpha_3 + \cos \alpha_4 = 0$. *Ann. Univ. Sci. Budapest. Eötvös Sect. Math.*, 1969.
- [Yam99] T. Yamakami. A foundation of programming a multi-tape quantum Turing machine. In *MFCS ’99 : Proceedings of the 24th International Symposium on Mathematical Foundations of Computer Science*, pages 430–441, London, UK, 1999. Springer-Verlag.
- [Yao93] A. Yao. Quantum circuit complexity. In *Proc. 34th IEEE Symposium on Foundation of Computer Science*, 1993.
- [Zha06] Y. Zhang. Algebraic structures underlying quantum information protocols. *ArXiv.org quant-ph/0601050*, 2006.