



Turbo-codes quantiques

Mamdouh Abbara

► To cite this version:

Mamdouh Abbara. Turbo-codes quantiques. Théorie de l'information [cs.IT]. Ecole Polytechnique X, 2013. Français. NNT: . pastel-00842327

HAL Id: pastel-00842327

<https://pastel.hal.science/pastel-00842327>

Submitted on 8 Jul 2013

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.

Turbo-codes quantiques

Mamdouh ABBARA

Thèse réalisée à INRIA,
Sous la direction de Jean-Pierre TILLICH

En vue d'obtenir le diplôme de docteur de
l'Ecole Polytechnique

Soutenue le 9 avril 2013

Table des matières

Introduction	1
1 Codes correcteurs d'erreurs classiques	6
1.1 Canal, encodage et code	6
1.2 Distance minimale d'un code	7
1.3 Code correcteur d'un ensemble d'erreurs	9
1.4 Décodage au maximum de vraisemblance par bloc et par symbole	10
1.5 Encodeur et décodage par syndrome	12
1.6 Capacité d'un canal et conception d'encodages optimaux	14
1.7 L'encodage convolutif	15
1.8 Performances de décodage d'un code convolutif	19
1.9 Le turbo-encodage en parallèle	20
1.10 Le turbo-encodage en série	22
1.11 Distance minimale d'un turbo-encodage	23
1.12 Algorithme de décodage itératif d'un turbo-encodage en série . .	24
2 Codes correcteurs d'erreurs quantiques	30
2.1 Premier axiome : états d'un système	30
2.2 Deuxième axiome : mesure effectuée sur un système physique . .	33
2.3 Troisième axiome : évolution d'un système physique isolé	35
2.4 Erreur quantique et canal quantique	36
2.5 Encodage et code quantiques	40
2.6 Transformations, états distinguables	41
2.7 Code correcteur d'un ensemble d'erreurs	45
2.8 Distance minimale d'un code quantique	50
2.9 Erreurs de Pauli	51
2.10 Codes stabilisateurs	54
2.11 Transformations de Clifford	58
2.12 Décodage d'un encodage stabilisateur	65
2.13 Encodeurs quantiques	68

3	Le turbo-encodeur formel, construction dans un formalisme commun aux cadres classique et quantique	73
3.1	Un formalisme d'erreur et d'encodeur commun aux cadres classique et quantique	73
3.2	Turbo-encodeur formel	76
3.3	Propriétés recherchées pour un encodeur convolutif formel	85
3.4	Distance minimale d'un turbo-encodeur formel	87
3.5	Constructions expérimentales d'un turbo-encodeur quantique	88
3.5.1	Schéma de turbo-encodeur modifié	89
3.5.2	Propriétés et choix de l'encodeur externe	93
3.5.3	L'encodeur interne	97
3.5.4	Performances de réduction d'erreur des turbo-encodeurs modifiés	99
3.5.5	Turbo-encodeurs modifiés à deux étages	104
3.6	Algorithme de décodage itératif du turbo-encodeur quantique . .	108
4	Preuve de la distance minimale d'un turbo-encodeur formel	115
4.1	Schéma global de la preuve	115
4.2	Dénombrement d'erreurs au niveau de l'encodeur formel externe	122
4.3	L'encodeur convolutif interne : période de l'encodeur, et trace d'une erreur d'entrée	127
4.4	Dénombrement d'erreurs au niveau de l'encodeur formel interne .	139
4.5	Preuve des théorèmes	148

Remerciements

Au nom de Dieu.

Certains souvenirs résistent à l'épreuve du temps. Des années plus tard, ils reviennent à l'esprit, immuables et fidèles. Ainsi me submergent-ils en cet instant, alors qu'il s'agit de mettre la touche finale au manuscrit : une partie de baby-foot, un échange d'idées, une plaisanterie, des mots croisés, des moments fraternels... Le chemin pris par les membres de l'équipe, à 11h45, pour aller déjeuner, est si répétitif qu'il ne peut que laisser une marque. Les tours de capture de desserts, que nul n'a réussi avec brio autant que B.C., n'échappent pas à la règle. Ni l'appétit particulier de J-P.T. vis-à-vis de ces mêmes desserts, qui à partir de 16h, tombaient dans le domaine public s'ils n'étaient pas déjà consommés. Il serait également injuste de ne pas mentionner la contribution particulière de V.H., cet homme débordant d'imagination, au climat bon enfant de l'équipe.

Je remercie au préalable MM. Claude Berrou, Romain Alléaume, Nicolas Macris, David Poulin, Daniel Augot, Gilles Zémor, et Renato Renner, pour avoir, avec mon directeur de thèse M. Jean-Pierre Tillich, bien voulu m'honorer en acceptant d'être membres du jury de thèse.

Je souhaite ensuite remercier les membres et collaborateurs de l'équipe SE-CRET, au sein de laquelle j'ai effectué ma thèse, pour les bons moments connus avec eux : Christelle Guiziou, assistante de l'équipe, que j'apprécie vraiment et à qui je souhaite le plus grand bonheur ; Pascale Charpin, Anne Canteaut, Nicolas Sendrier et Daniel Augot, qui figurent parmi les piliers de l'équipe ; Jean-Pierre Tillich, directeur de thèse de grande valeur humaine et scientifique, Sumanta Sarkar et Baudoin Collard, qui ont occupé tour à tour le poste de voisin de bureau, et qui ont enrichi ma vision des Indiens et des Belges ; Anthony Leverrier, actuel titulaire de mon propre bureau ; Harold Ollivier, qui a passé sa thèse avant moi auprès de Jean-Pierre Tillich et qui m'a recommandé une telle expérience ; Vincent Herbert – qui a donné son nom à une unité de mesure d'une certaine activité – et l'inoubliable M. L'Ambassadeur de la ville de Lyon ; Matthieu Finiasz, expert des questions informatiques qui m'a bien souvent aidé ; Grégory Landais, dont j'abusais de la même manière, et grâce à qui je passais souvent de meilleures journées ; Denise Maurice et Anne Marin,

adeptes comme moi de mécanique quantique ; Stéphane Manuel, que je remercie pour le barbecue, et pour bien plus ; l'aimable Rafael Misoczki ; Valentin Suder, Benoît Gérard, Céline Blondeau, Maria Naya-Plasencia, Christina Boura, Sunandan Chakraborty, Stéphane Jacob, Maxime Côte, Joelle Roue, Marion Bellard, Dimistris Simos, et Audrey Tixier.

Merci également à Nicolas Delfosse, avec qui j'ai eu l'occasion d'échanger lors de divers événements, et à Mathieu Feuillet, qui, disait-on, alternait avec moi l'occupation de mon bureau.

Des salutations particulières, ainsi que des remerciements, sont adressés à mes frères : Ayoub Otmani, Samer Ammoun, Ahmed Rebai, Younes Bouchaala, Falou Ndoeye, Mohamed Marouf, Wajih Ouertani, Nebil Ben Marbouk. Rien n'égale le bien-être procuré par votre compagnie.

Merci à Iyad, d'abord pour sa satire de mes moments passés en thèse, mais aussi pour sa présence quand il le fallait à mes moments les plus difficiles. Merci à Ghazal et à Ahmed d'avoir suivi le même chemin, courageusement. Merci à l'association A.B.B., à l'ensemble de son Conseil d'Administration, ainsi qu'à M. Bouha, dont la fertilité des apports ne peut s'énumérer dans un espace aussi restreint.

Anfal, en n'apparaissant qu'une fois effectué l'essentiel de ma thèse, et en m'apportant ton soutien au moment sans doute le plus tendu, tu as, sans le vouloir, réalisé la combinaison optimale. Merci pour tes encouragements, et surtout, pour ta présence spéciale à mes côtés. A mes deux soeurs, merci pour votre éternelle gentillesse, et courage pour vos concours. Quant à mes parents, c'est à vous que je dédie ce manuscrit de thèse : je crois que cela suffit à exprimer la faveur que je vous dois.

A mes parents.

Résumé

On démontre l'existence d'hypothèses permettant à un turbo-code quantique basé sur les codes stabilisateurs de posséder une distance minimale non bornée. Cette distance minimale sera polynomiale ou sous-logarithmique en la longueur du code, selon la distance minimale du code externe. Les hypothèses, que sont le caractère récursif et anti-récursif de l'encodage interne, résolvent un obstacle théorique rencontré jusqu'à présent. En effet, un résultat principal des turbo-codes classiques est que ceux-ci possèdent une distance minimale polynomiale, pourvu, notamment, que l'encodage interne soit récursif et systématique. Or, il s'avère qu'un encodage convolutif quantique ne peut être à la fois récursif et systématique. On conçoit donc dans le cas quantique une nouvelle condition nommée anti-récursivité, pouvant s'allier au caractère récursif de l'encodage interne, qui permet d'établir un résultat théorique similaire au cas classique.

Dans un second temps, on propose une résolution du second obstacle majeur aux performances des turbo-codes quantiques, portant sur les performances de décodage de ces derniers par un algorithme itératif en temps linéaire. L'incompatibilité citée précédemment rend en effet également incompatibles les caractères récursif et non-catastrophique de l'encodage interne. Or, sans le caractère non catastrophique, l'algorithme de décodage itératif est voué à l'échec. On propose donc un turbo-encodage au schéma modifié, dans laquelle une portion des positions de sortie de l'encodage externe, au lieu d'alimenter l'encodage interne, est envoyée en entrée d'un second encodage permettant de réduire fortement le bruit en ces positions. Il s'avère que le second encodage réalisant cette tâche correspond également à un turbo-encodage modifié, dont les positions isolées de l'encodage externe sont toutefois envoyées vers le canal. Afin de garantir l'efficacité de la construction, on pose un ensemble de conditions devant être vérifiées par l'encodage externe. La démarche est guidée par la notion de distance épurée, distance pertinente pour une telle construction et correspondant à un modèle où le bruit est nul aux positions envoyées vers le second encodage. Le résultat des simulations montre alors que le turbo-code conçu est efficace pour transmettre de l'information quantique via un canal dépolarisant dont l'intensité de dépolarisation peut aller jusqu'à $p = 0,145$.

Abstract

We prove the existence of a set of hypothesis enabling a quantum turbo-code based on stabilizer codes to have an unbounded minimum distance. This minimum distance is either polynomial or sub-logarithmic, depending on the minimal distance of the outer code. The hypothesis are that the inner encoding is recursive and anti-recursive, and they tackle a theoretical issue previously encountered. Indeed, an essential result of classical turbo-codes is that they have a polynomial minimum distance provided that, mainly, the inner encoding is recursive and systematic. However, it turns out that a quantum convolutional encoding cannot be simultaneously recursive and systematic. We thus design a new condition in the quantum case called anti-recursiveness, which is combinable with the recursiveness of the inner encoding, and which enables to establish a theoretical result similar to the classical case.

We then propose a solution to the second main obstacle to the good performance of quantum turbo-codes, which is about their decoding performance under an iterative algorithm with linear complexity. Indeed, the previously mentioned incompatibility makes it also impossible for the inner encoding to be simultaneously recursive and non catastrophic. However, with a catastrophic inner encoding, the iterative decoding algorithm is doomed to failure. We hence propose to modify the turbo-encoding construction. Instead of feeding it into the inner encoding, a part of the output positions of the outer encoding is sent towards a second encoding which performs a high noise reduction at these positions. It turns out that such a second encoding is also a modified turbo-encoding in which, however, the isolated positions of the outer encoding are sent to the channel. In order to guarantee the efficiency of the setting, we choose to apply a set of conditions on the outer encoding. The approach is guided by the notion of purified distance, a notion pertinent to this setting, which corresponds to a model where the positions sent to the second encoding are noiseless. The simulations results show that the designed turbo-code is efficient in transmitting quantum information via a depolarizing channel up to a noise rate of $p = 0.145$.

Introduction

La théorie des codes correcteurs d'erreurs s'intéresse à la protection de l'information, par exemple une séquence de bits, contre les erreurs ; une telle protection se fait en appliquant à cette information un encodage bien choisi. Un défi central de la théorie des codes consiste à concevoir des familles de codes dont le rendement est linéaire, et encodables de façon telle qu'il existe un algorithme de décodage rapide et donnant des performances proches de la limite de Shannon. Les codes LDPC [18] [32] [31] et les turbo-codes en parallèle [4] [5] [8] [10] [27] [44] ou en série [3] [6] [7] [11] ont été largement étudiés, et figurent sous la coupole de la famille des codes à décodage itératif qui permet de répondre à ce défi. Plus récemment, les codes polaires [2] montrent qu'il est même possible d'atteindre la capacité pour un canal sans mémoire symétrique avec une procédure de décodage de complexité $n \log n$ où n est la longueur du code.

Le domaine de l'information quantique a introduit un nouveau paradigme où l'information correspond à un état quantique de la matière ou de la lumière. Dans ce domaine, l'information est régie par une structure mathématique plus riche que celle correspondant à l'information classique. Schématiquement, si l'information classique décrite par un bit peut se trouver dans deux états 0 ou 1, l'ensemble des valeurs que peut prendre l'information quantique correspondante est le $\mathbb{C}$ -espace vectoriel dont la paire d'états précédente est une base. De nombreuses expériences passées et présentes montrent que l'on se rapproche, en terme de progrès technique, de l'adoption de calculateurs et de réseaux de transmission exploitant l'information quantique. Le potentiel que représente ce progrès est très grand : citons, parmi les résultats pionniers du domaine, l'algorithme de factorisation des nombres entiers de Shor [46] qui permet de factoriser un nombre entier en temps polynomial, alors que les algorithmes classiques permettant de le faire requièrent un temps exponentiel, mettant ainsi en péril la sécurité du protocole de sécurité répandu RSS, ou encore le protocole de distribution de clé BB84 de Bennett et Brassard [9] basé sur la transmission de bits quantiques et révélant une sécurité inconditionnelle basée sur les lois de la physique. La puissance de calcul d'un ordinateur quantique peut également être utilisée afin de simuler de manière efficace les systèmes quantiques [13] tel que conjecturé par Feynman [17] en 1982 puis prouvé par Lloyd [30], en d'autres termes, avancer dans la compréhension des systèmes physiques dans lesquels les phénomènes quantiques jouent un rôle fondamental.

Fiabiliser la transmission et le stockage de l'information quantique représente

un problème central pour l'implémentation de systèmes d'information quantiques. Il existe des méthodes de protection passive de l'information quantique, où la nature même de l'évolution ou de l'interaction du système avec le milieu environnant garantit la protection de l'information, tels que le calcul quantique adiabatique imaginé par Farhi [16] et les sous-espaces sans décohérence ou *decoherence-free subspaces* [28]. La protection active de l'information quantique est quant à elle l'objet des codes correcteurs d'erreurs quantiques. Le défi qui se pose en théorie des codes quantiques est alors similaire à celui des codes classiques. On s'intéresse dans cette thèse aux codes correcteurs quantiques binaires, c'est-à-dire permettant d'encoder l'état quantique pris par un système de k bits. A priori, le problème est posé dans un cadre où l'ensemble des états quantiques est un espace vectoriel, donc un ensemble continu. Les codes stabilisateurs introduits par Gottesman [19] ramènent le problème à des variables discrètes : ces codes, ainsi que leurs propriétés, peuvent être entièrement étudiés grâce à un ensemble discret d'erreurs nommé *groupe de Pauli* et constitué dans le cas binaire de quatre opérateurs, $\mathcal{I}$, $\mathcal{X}$, $\mathcal{Y}$ et $\mathcal{Z}$ modulo une phase globale égale à ± 1 ou $\pm i$. Il s'agit du groupe d'erreurs engendré par les opérations d'inversion de bit et d'inversion de phase, la dernière étant propre au cadre quantique. Ces quatre opérateurs sont une base de l'espace des erreurs agissant sur un bit, et le produit cartésien de cet ensemble d'opérateurs n fois avec lui-même définit une base de l'espace des erreurs agissant sur un système de n bits. Le problème étant dorénavant exprimé en termes de ces quatre opérateurs, une manière naturelle de transposer les constructions de codes binaires classiques au cadre quantique est de construire des codes stabilisateurs dans lesquels, grossièrement, on substitue l'ensemble des quatre opérateurs $\mathcal{I}, \mathcal{X}, \mathcal{Y}, \mathcal{Z}$ à l'ensemble $\{0, 1\}$.

La limite théorique de capacité d'un canal quantique, à laquelle aspirent les codes quantiques optimaux en termes de rendement, n'a pas d'expression générale connue contrairement au théorème de Shannon classique. Une propriété exclusive au cadre quantique est que tout code est laissé invariant par un ensemble non trivial d'erreurs, et que par conséquent, plusieurs erreurs peuvent avoir le même syndrome lors du décodage. Ceci a une répercussion sur la capacité qui n'est pas bien connue ; cependant, on évalue en pratique les performances de décodage d'un code quantique grâce à la *capacité cohérente* ou *hashing bound*, capacité connue qui formule l'hypothèse que le décodage évalue l'erreur - et non la classe d'erreurs de même syndrome - la plus vraisemblable.

Il a récemment été montré que les codes polaires quantiques [50] [43] atteignent la capacité symétrique, soit la capacité obtenue sous contrainte que l'entrée suive une distribution uniforme. Cependant, la transposition au cadre des codes quantiques stabilisateurs des codes LDPC et des turbo-codes pose plusieurs problèmes.

La plupart des constructions de codes LDPC quantiques ont soit une distance minimale bornée, soit un rendement tendant vers 0 avec la longueur du code. On peut passer en revue quelques exceptions. La distance minimale de la construction [45] évolue comme la racine carrée de la longueur du code ; toutefois, ses performances sont clairement adaptées à un modèle de canal

asymétrique, c'est-à-dire présentant des probabilités d'inversion de bit et d'inversion de phase très différentes. Des codes LDPC quantiques basés sur des pavages de surfaces [24] [53] ont quant à eux une distance minimale logarithmique. Plus récemment, des codes LDPC quantiques ont été proposés avec une distance minimale évoluant comme la racine carrée de la longueur du code [49].

Toutes ces constructions souffrent néanmoins de deux défauts qui nuisent au décodage itératif. Le premier défaut est l'existence dans le graphe de Tanner de cycles de taille égale à 4. Le second défaut est le caractère hautement dégénéré des codes associés à ces constructions ; cette propriété consiste en ce que le code est laissé invariant par un grand nombre d'erreurs de poids inférieur à sa distance minimale. Si le caractère dégénéré est en théorie susceptible d'améliorer la transmission de l'information quantique, il se révèle en pratique qu'il dégrade les performances du décodage itératif. Une construction récente de code LDPC quantique inventée dans [23] et généralisée dans [1] fait exception ; elle se définit à partir d'une paire orthogonale de codes LDPC classiques binaires, transposée à un alphabet non binaire, puis, par une représentation binaire appropriée de cet alphabet, de nouveau à un alphabet binaire. Ce schéma ne s'applique toutefois qu'à une classe étroite de codes LDPC et ne permet pas d'approcher la capacité cohérente.

La branche des turbo-codes quantiques est quant à elle moins fournie que celle des LDPC quantiques. Le problème essentiel est que tout code convolutif quantique, composante de base d'un turbo-code quantique, ne peut concilier deux propriétés toutes deux nécessaires aux bonnes performances du turbo-code quantique. D'une part, un tel code doit être *récuratif*, condition nécessaire à ce que la distance minimale du turbo-code quantique soit non bornée. D'autre part, il ne doit *pas être catastrophique*, et ce afin que l'algorithme de décodage itératif permette de réduire le bruit à chaque itération et de converger vers une bonne estimation de l'erreur introduite par le canal. Ces deux propriétés ne peuvent être simultanément vérifiées [41]. Le turbo-code quantique proposée par [41] est basé sur un encodage convolutif interne non récuratif et non catastrophique ; sa distance minimale est bornée mais l'algorithme de décodage itératif converge. A cause du caractère borné de la distance minimale, les performances asymptotiques du décodage atteignent cependant un seuil limite d'efficacité. Une autre solution proposée par [52] est d'assister le turbo-code quantique par de l'*intrication* partagée entre l'émetteur et le récepteur. Cela consiste à supposer qu'émetteur et récepteur partagent préalablement une ressource quantique ; il s'agit de paires de bits quantiques dans un état de Bell [34], dont un bit quantique est détenu par l'émetteur et l'autre par le récepteur. Le schéma proposé fait alors intervenir lors du processus d'encodage un nombre linéaire de ces bits quantiques du côté de l'émetteur. Cela permet, lors du décodage, d'acquérir l'information complète sur l'erreur aux positions où les bits quantiques intriqués ont été encodés. Cela dope le processus de décodage itératif et contourne le caractère catastrophique de l'encodage convolutif interne, au prix de l'utilisation d'une ressource coûteuse en quantité linéaire avec la taille de l'encodage.

Une autre difficulté majeure qui se pose pour les turbo-codes quantiques est l'absence, jusqu'à cette thèse, d'un résultat théorique sur la distance minimale.

Il est en effet connu dans la théorie des turbo-codes classiques [22] qu'un turbo-code dont l'encodage convolutif interne est *récuratif* et *systématique* présente une distance minimale polynomiale, avec une puissance bornée par $(d^* - 2)/d^*$, où d^* est la distance libre de l'encodage externe. Dans la théorie quantique, encore une fois, les deux caractères *récuratif* et *systématique* sont incompatibles, car le caractère *systématique* implique en effet le caractère *non catastrophique* de l'encodage interne. Il est simplement connu que le caractère *récuratif* est nécessaire afin d'obtenir une distance minimale non bornée, mais aucune investigation connue n'a été menée afin de déterminer une condition suffisante.

Cette thèse apporte une double contribution au domaine des turbo-codes quantiques. D'abord, elle établit un résultat théorique concernant la distance minimale des turbo-codes quantiques. On y établit une nouvelle notion d'*anti-récurativité*, et on démontre qu'un turbo-code quantique dont l'encodage interne est *récuratif* et *anti-récuratif* possède une distance minimale dont le comportement est dicté par la distance *dégénérée* du code externe. La distance *dégénérée* correspond au plus petit poids d'une erreur non triviale laissant invariant le code, et s'identifie à la distance minimale si le code est non *dégénéré*.

La seconde contribution porte sur la construction d'un encodage quantique qui soit, d'une part, de rendement constant, et d'autre part, décodable par un algorithme de complexité linéaire performant, c'est-à-dire dont le taux de succès est proche de la capacité cohérente du canal. On propose pour cela une construction basée sur un turbo-encodage dont l'encodage interne est *récuratif* et *anti-récuratif*. Un tel turbo-encodage est par essence mauvais au décodage itératif à cause du caractère catastrophique de l'encodage interne. Toutefois, si l'on ajoute une contrainte supplémentaire sur l'encodage externe, on peut construire un turbo-encodage à deux niveaux, dans lequel une fraction des bits quantiques en sortie de l'encodage externe sont envoyées vers un second turbo-encodage. Le second turbo-encodage a pour premier effet de réduire considérablement le niveau de bruit sur les positions qui lui sont données en entrée, ce qui sert d'amorce pour décoder le premier turbo-encodage. Cette solution implémente pour la première fois un turbo-encodage de distance minimale polynomiale, sans qu'il n'y ait d'obstacle au décodage itératif, en l'assistant par un second étage de turbo-encodage.

Le premier chapitre reprend les principales notions de codes correcteurs d'erreurs classiques nécessaires au présent exposé. Le deuxième chapitre commence par une présentation des axiomes de la mécanique quantique sur lesquels la suite du chapitre se base de manière rigoureuse ; elle est suivie de l'introduction des notions d'erreur, de canal, d'encodage et de code quantiques. Le choix a ensuite été fait de démontrer une analogie entre la notion d'états *distinguable*s en mécanique quantique et la notion classique simple d'états *différents*, et de montrer que l'on peut en dériver ensuite naturellement les concepts de correction d'erreur et de distance minimale d'un code quantique. Sont enfin présentées les erreurs de Pauli, outils discrets de modélisation de l'encodage et du décodage quantiques dans le cadre des codes stabilisateurs. On aura pris soin, dans les

deux premiers chapitres, de montrer que l'étude de l'encodage binaire linéaire et de l'encodage stabilisateur peut se concentrer sur l'évolution des erreurs sous l'action de l'encodage ; cette évolution est un morphisme décrit par un *encodeur*. Ainsi, le couple erreurs-encodeur permet dans les deux cas de modéliser l'encodage par des variables discrètes.

Le chapitre 3 est dédié en premier lieu à l'énoncé du résultat théorique de la thèse sur la distance minimale d'un turbo-encodage quantique ; le choix a été fait de présenter ce résultat sous un formalisme plus large, qui englobe à la fois le cadre classique binaire linéaire et le cadre quantique stabilisateur, et d'énoncer dans ce formalisme un résultat qui rassemble à la fois la distance minimale d'un turbo-code classique et la distance minimale d'un turbo-code quantique. Pour ce faire, les deux modèles discrets d'erreurs et d'encodeur, dans le cadre binaire linéaire et le cadre stabilisateur, sont rassemblés sous un formalisme commun, qui saisit leurs similitudes essentielles et sert de point de départ pour déduire le modèle de turbo-encodage quantique à partir du modèle classique. On définit grâce à cela le *turbo-encodeur formel*, puis essentiellement les notions d'encodeur convolutif *récuratif*, *systématique* et *anti-récuratif*. Cette dernière notion d'anti-récurativité est une nouveauté proposée dans le cadre de la thèse, afin de pallier à l'absence du caractère systématique lorsque le turbo-encodeur correspond au cadre quantique. En second lieu, on présente différentes constructions basées sur le turbo-encodeur quantique et la simulation de leurs performances de décodage.

La preuve du résultat théorique est repoussée au chapitre 4. Elle repose sur la majoration de la probabilité que le poids d de la sortie du turbo-encodeur soit inférieure à la prétendue distance minimale. La majoration s'effectue par une sommation qui différencie les différentes valeurs du poids d de la sortie et le poids intermédiaire w entre l'encodeur externe et l'encodeur interne. Cette preuve inclut la procédure suivie par Kahale et Urbanke dans [22], à laquelle se greffent des éléments plus complexes afin de l'adapter au formalisme commun aux cadres classique et quantique. De surcroît, la preuve présentée traite deux nouveaux régimes de sommation couvrant l'aire définie par $w > d$. Cette aire est en effet de facto nulle dans le cas classique grâce au caractère *systématique* de l'encodeur convolutif interne ; elle doit cependant être traitée dans le cas quantique, et se majore par un procédé faisant intervenir le caractère *anti-récuratif* de l'encodeur convolutif interne.

Chapitre 1

Codes correcteurs d'erreurs classiques

1.1 Canal, encodage et code

Il est possible de définir l'information comme un état pris par un système physique. On peut ainsi représenter une quantité d'information finie par une séquence de n symboles appartenant à un alphabet fini A . Lorsque l'on cherche à conserver cette séquence de symboles sur un support physique ou à la transmettre d'un endroit à un autre, une erreur peut se produire et l'information qu'elle contient peut être endommagée, de sorte que l'on reçoit, après une telle transformation, une séquence de n symboles appartenant à un alphabet B .

Un exemple consisterait en l'envoi d'une séquence appartenant à $\{0, 1\}^n$ en n impulsions lumineuses séparées d'un intervalle de temps constant et pouvant prendre chacune deux états de polarisation différents. Chaque impulsion est soit reçue correctement, soit reçue avec la polarisation inverse, et soit perdue, de sorte que la séquence reçue est un élément de $\{0, 1, e\}^n$ où le symbole e désigne un effacement du symbole émis.

On modélise tout processus d'évolution subi par une séquence par un canal [14] qui transforme, pour tout n , une séquence en entrée $y \in A^n$ en une séquence de sortie $z \in B^n$ conformément à une certaine loi de probabilité $P_n(z|y)$. Les erreurs qui surviennent sur chaque symbole de la séquence peuvent être corrélées. Dans toute la suite, on ne considèrera cependant que le modèle simple du canal dit sans mémoire et pour lequel la loi de probabilité de l'erreur est identique et indépendante en chaque position.

Définition 1.1.1. Un canal est dit **sans mémoire** s'il existe une loi de probabilité $P(Z|Y)$, où les variables aléatoires Y et Z sont respectivement définies sur les alphabets d'entrée A et de sortie B du canal, telle que pour tout entier

n et pour toute paire de séquences $y = (y_1, \dots, y_n)$ et $z = (z_1, \dots, z_n)$:

$$P_n(z|y) = \prod_{i=1}^n P(z_i|y_i)$$

Deux exemples de canal sans mémoire méritent d'être cités.

Définition 1.1.2. Le **canal binaire symétrique** de probabilité d'erreur p est un canal qui transmet des bits en inversant chaque bit avec une probabilité p et en le laissant intact avec une probabilité $1 - p$. Ainsi pour ce canal, $A = B = \mathbb{F}_2$, $P(1|0) = P(0|1) = p$ et $P(0|0) = P(1|1) = 1 - p$.

Le **canal à effacement** de probabilité d'effacement p est un canal pour lequel $B = A \cup \{e\}$ où $e \notin A$ est le **symbole d'effacement**, qui transmet tout symbole de A correctement avec une probabilité $1 - p$ et le transforme en e avec probabilité p .

Par la suite, on se restreindra principalement au cas où l'alphabet de sortie B est égal à A , sauf lorsque l'on considérera le canal à effacement.

Afin de protéger l'information contenue dans une séquence $x = (x_1, \dots, x_k)$ de k symboles, on la transforme en une séquence $y = (y_1, \dots, y_n)$ de n symboles avec $n > k$ avant de l'envoyer via le canal.

Définition 1.1.3. On nomme **encodage** toute injection de A^k dans A^n , **message** un élément de A^k avant encodage, et **mot de code** la séquence de A^n obtenue après encodage. Le **code** est l'ensemble des $|A|^k$ mots de code, et n désigne la **taille du code**

La linéarité d'un encodage est une propriété élémentaire, dont l'utilité est multiple ; on peut ainsi citer la simplification de l'opération d'encodage, la simplification de la notion de distance minimale présentée à la section 1.2, et l'existence du concept de *syndrome* d'une erreur affectant un mot de code tel que présenté à la section 1.5. On se restreint ici aux encodages linéaires définis sur l'alphabet $A = \mathbb{F}_2$ muni de sa structure de corps.

Définition 1.1.4. Un **encodage linéaire** est une application $\mathbb{F}_2$ -linéaire injective ϕ de $\mathbb{F}_2^k$ dans $\mathbb{F}_2^n$ où $k \leq n$. Un **code linéaire** est un code $C = \text{Im } \phi$ engendré par un encodage linéaire.

1.2 Distance minimale d'un code

La distance naturelle définie sur l'ensemble des séquences de taille fixée d'un alphabet fini est la séquence de Hamming.

Définition 1.2.1. La **distance de Hamming** entre deux séquences $z = (z_1, \dots, z_n)$ et $z' = (z'_1, \dots, z'_n)$ de A^n est :

$$d(z, z') = |\{i \in \llbracket 1; n \rrbracket / z_i \neq z'_i\}|$$

Supposons que les alphabets d'entrée A et de sortie B sont identiques, et que la loi de probabilité du canal vérifie :

$$\begin{aligned} P(z|y) &= p && \text{si } z \neq y \\ P(z|y) &= 1 - (|A| - 1)p && \text{si } z = y \end{aligned}$$

où $p < 1/|A|$, de sorte à ce qu'il soit plus probable pour un symbole d'être transmis sans erreur que d'être transformé en un autre symbole donné. Dans ce cas, le décodage au maximum de vraisemblance équivaut à trouver un mot de code $(\hat{y}_1, \dots, \hat{y}_n)$ le plus proche de la séquence reçue au sens de la distance de Hamming. En effet, si l'on note t la distance de Hamming entre deux séquences y' et z , on peut écrire :

$$P_n(z|y') = p^t (1 - (|A| - 1)p)^{n-t}$$

C'est une probabilité qui ne dépend que de t et qui est d'autant plus importante que la distance t est petite. Il apparaît ainsi une propriété topologique importante d'un code, qui permet d'évaluer sa capacité à protéger l'information face à une erreur concentrée en un nombre borné de positions :

Définition 1.2.2. La **distance minimale** d'un code est la plus petite distance de Hamming d entre deux mots différents du code.

Une propriété immédiate permet de simplifier la formulation de la distance minimale lorsque le code est linéaire.

Propriété 1.2.1. La *distance minimale d'un code linéaire est le plus petit poids de Hamming d'un mot de code non nul.*

Démonstration. Il suffit de constater que la distance de Hamming entre deux mots de code différents y et y' correspond au poids de Hamming de $y + y'$, qui est à son tour un mot de code non nul. $\square$

Si un code est de distance minimale d , le décodage au maximum de vraisemblance permet en effet de retrouver sans ambiguïté la séquence initiale envoyée si l'on sait que le nombre de symboles erronés est inférieur ou égal à $t = (d-1)/2$. En effet, il ne peut exister deux mots de code différents situés tous deux à une distance de Hamming inférieure ou égale à t de la séquence reçue, car la distance entre ces deux mots de code serait inférieure ou égale à $2t = d-1$.

On peut de même étudier la capacité de correction d'un code de distance minimale d dans le cas d'un canal à effacement. On se place donc dans le cas où $B = A \cup \{e\}$ où $e \notin A$, et chaque symbole est transmis correctement avec une probabilité $1 - p$ et transformé en effacement e avec une probabilité p . Dans ce cas, la séquence initiale envoyée peut-être retrouvée sans ambiguïté si au plus $d-1$ symboles sont effacés. Supposons en effet que l'on émette un mot de code de A^n et que l'on reçoive une séquence de B^n contenant au plus $d-1$ symboles d'effacement e et concordant avec le mot de code émis sur toutes les autres positions. Si l'on peut trouver deux mots de code différents qui concordent avec la séquence reçue sur toutes les positions ne contenant pas un effacement, ces mots de code ne pourraient différer que sur les positions où la séquence reçue

contient le symbole e , et la distance de Hamming qui les sépare serait alors inférieure ou égale à $d - 1$.

Un défi majeur de la théorie des codes consiste à trouver des codes possédant une distance minimale non bornée, afin de pouvoir corriger un grand nombre d'erreurs, produits par un encodage possédant assez de structure pour permettre un décodage de complexité raisonnable. Le problème du décodage est exposé à la section 1.4. Une complexité raisonnable correspond dans l'idéal à une complexité linéaire en la taille du code. C'est le cas du décodage d'un code convolutif; cependant, un tel code est de distance minimale bornée comme l'atteste la propriété 1.8.1. A l'opposé, le problème du décodage d'un encodage aléatoire requiert a priori un temps exponentiel en la taille du code. Une des structures d'encodage permettant d'allier à la fois une distance minimale non bornée et un décodage de complexité raisonnable est le turbo-encodage, que l'on présente par la suite dans ce même chapitre. Pour un turbo-encodage, il existe des algorithmes de décodage en temps quasi-linéaire tels que l'algorithme de propagation de croyance ou *belief propagation* qui requiert une complexité en $n \log n$ où n est la taille du code.

1.3 Code correcteur d'un ensemble d'erreurs

On peut formaliser d'une manière générale la capacité d'un code à corriger un ensemble d'erreurs donné, indépendamment de la distribution de probabilité des erreurs induites par le canal. Il s'agit d'une propriété topologique du code, au même titre que la distance minimale. L'approche développée dans cette partie servira de point de repère lorsque l'on étudiera les propriétés des codes quantiques. En effet, et contrairement à la théorie des codes classiques, la question de savoir s'il existe un code quantique capable de corriger un ensemble d'erreurs donné n'est pas triviale, et impose d'aborder la théorie des codes quantiques directement sous cet aspect. Plaçons-nous dans le cas où l'alphabet d'entrée A est inclus dans l'alphabet de sortie B du canal. Commençons d'abord par définir une erreur et le poids d'une erreur.

Définition 1.3.1. On appelle **erreur** toute application de A^n dans B^n . Le **poids** d'une erreur E est la plus grande distance de Hamming d entre une séquence et son image, c'est-à-dire :

$$d = \max\{d(x, E(x)), x \in A^n\}$$

La capacité d'un code à corriger un ensemble d'erreurs se définit de la manière suivante.

Définition 1.3.2. Un code C **corrige un ensemble d'erreurs** $\{E_j, j \in J\}$ s'il existe une application de B^n dans A^n qui permet, pour tout mot de code $x \in C$ et pour tout $j \in J$, de transformer $E_j(x)$ en x .

Ainsi, si l'on suppose que l'erreur induite par le canal appartient à l'ensemble $\{E_j, j \in J\}$, on est capable de reconstituer le mot de code envoyé. La capacité

à effectuer cette même reconstitution sera exigée lorsque l'on définira un code quantique corrigeant un ensemble d'erreurs. La condition nécessaire et suffisante pour qu'un code corrige un ensemble d'erreurs est ici immédiate. Elle possède un équivalent dans le cadre quantique dont la preuve n'est pas directe. Sa formulation est très similaire à la condition que l'on présente ici, mais nécessitera au préalable l'introduction de la notion d'états distinguables, qui supprime la notion classique simple de séquences différentes.

Propriété 1.3.1. *Un code C corrige un ensemble d'erreurs $\{E_j, j \in J\}$ si et seulement si, pour tout $(i, j) \in J^2$ et pour tout $(x, x') \in C^2$ tel que $x \neq x'$, $E_j(x) \neq E_i(x')$.*

Comme on l'a vu précédemment, un code de distance minimale d permet de corriger l'ensemble des erreurs de poids inférieur ou égal à $t = (d - 1)/2$. Il permet également de corriger l'ensemble des effacements de poids inférieur ou égal à $d - 1$. La distance minimale d'un code possède la caractérisation suivante en termes d'erreurs, dont on utilisera une formulation similaire afin de définir la distance minimale d'un code quantique.

Propriété 1.3.2. *La distance minimale d'un code est le plus petit poids d d'une erreur E telle qu'il existe $(x, x') \in C^2$ tel que $x \neq x'$ et $E(x) = x'$.*

Démonstration. Soit d la distance minimale d'un code C . S'il existe une erreur E dont le poids est strictement inférieur à d , et pour laquelle il existe $(x, x') \in C^2$ tel que $x \neq x'$ et $E(x) = x'$, cela implique par définition du poids de l'erreur E que $d(x, x') < d$ ce qui est absurde.

Par ailleurs, si on note $(x, x') \in C^2$ une paire telle que $x \neq x'$ et $d(x, x') = d$, on peut poser E l'erreur qui envoie toute séquence de $A^n \setminus \{x\}$ sur elle-même et qui envoie x sur x' . Cette erreur est de poids d et il existe $(x, x') \in C^2$ tel que $x \neq x'$ et $E(x) = x'$. Cela prouve le résultat. $\square$

1.4 Décodage au maximum de vraisemblance par bloc et par symbole

Présentons à présent le problème du décodage d'une séquence reçue à la sortie d'un canal. Lorsqu'un mot de code $y = (y_1, \dots, y_n)$ est envoyé via le canal, certains de ses symboles subissent une erreur, et le récepteur obtient alors une séquence de n symboles $z = (z_1, \dots, z_n)$. Sachant qu'il a reçu la séquence z , le récepteur doit trouver un mot de code $\hat{y} = (\hat{y}_1, \dots, \hat{y}_n)$ ayant la plus grande probabilité d'être envoyé :

$$y \in \operatorname{argmax}_{y' \in C} p(y'|z)$$

Dans la configuration la plus courante, tous les mots de code y' sont envoyés avec la même probabilité $p(y') = 1/|C|$. On peut alors écrire la relation valable

pour tout mot de code y' :

$$p(y'|z) = \frac{p(y')P_n(z|y')}{p(z)} = \frac{P_n(z|y')}{|C|p(z)}$$

Par conséquent, pour une séquence reçue z donnée, l'estimation $\hat{y}$ produite par le décodage est la séquence maximisant la loi de probabilité conditionnelle du canal :

$$\hat{y} \in \operatorname{argmax}_{y' \in C} P_n(z|y')$$

Cela s'appelle le décodage au maximum de vraisemblance (par bloc), que l'on formalise donc par la définition suivante.

Définition 1.4.1. Soit C un code de longueur n dont on transmet un mot de code y via un canal de loi de probabilité P_n . Le **décodage au maximum de vraisemblance (par bloc)** consiste à estimer, pour une séquence reçue z , un mot de code $\hat{y}$ tel que :

$$\hat{y} \in \operatorname{argmax}_{y' \in C} P_n(z|y')$$

Dans l'algorithme de décodage itératif des turbo-encodages basé sur le *belief propagation* que l'on présente à la section 1.12, on s'intéresse plutôt à rechercher les symboles les plus probables en chaque position du mot de code y envoyé ou du message x qui lui correspond. Plus exactement, l'algorithme consiste à conjecturer une répartition de probabilité sur chaque symbole, puis par une succession d'itérations basées sur des échanges de messages entre les positions corrélées, à converger vers une répartition de probabilité finale que l'on estime proche de la répartition de probabilité réelle. Cette approche est plus pratique qu'une recherche brute du maximum de vraisemblance, dans la mesure où elle permet d'effectuer une recherche en temps quasi-linéaire en la taille n du code.

Ainsi, cette approche nommée décodage au maximum de vraisemblance par symbole s'énonce comme suit.

Définition 1.4.2. Soit C un code de longueur n dont on transmet un mot de code y via un canal de loi de probabilité P_n . Le **décodage au maximum de vraisemblance par symbole** appliqué au **mot de code** consiste à estimer, pour une séquence reçue z , un mot de code $\hat{y}$ dont chaque coordonnée $\hat{y}_i$ vérifie :

$$\hat{y}_i \in \operatorname{argmax}_{y'_i \in A} p(y'_i|z)$$

Appliqué au **message**, il consiste à estimer un message $\hat{x}$ dont chaque coordonnée $\hat{y}_i$ vérifie :

$$\hat{x}_i \in \operatorname{argmax}_{x'_i \in A} p(x'_i|z)$$

Dans cette définition, la probabilité p calculée pour chacun des cas peut être obtenue à partir de la loi de probabilité du canal. Par exemple, comme démontré en Annexe, on peut constater dans le cas d'un canal sans mémoire

que le décodage au maximum de vraisemblance par symbole appliqué au mot de code équivaut à rechercher la séquence $\hat{y}$ vérifiant pour tout $i \in [1; n]$ la propriété suivante :

$$\hat{y}_i \in \operatorname{argmax}_{y'_i \in A} (P(z_i | y'_i) \sum_{y' \in C_{i, y'_i}} \prod_{\substack{j=1 \\ j \neq i}}^n P(z_j | y'_j))$$

où C_{i, y'_i} est l'ensemble des mots du code qui contiennent en position i le symbole y'_i .

1.5 Encodeur et décodage par syndrome

Afin d'implémenter un encodage linéaire ϕ de $\mathbb{F}_2^k$ dans $\mathbb{F}_2^n$, on peut d'abord appliquer l'injection canonique $\iota_{k \rightarrow n}$ de $\mathbb{F}_2^k$ dans $\mathbb{F}_2^n$ qui consiste à ajouter $n - k$ bits initialisés à la valeur 0, puis appliquer un automorphisme $\bar{\phi}$ de $\mathbb{F}_2^n$ tel que $\bar{\phi} \circ \iota_{k \rightarrow n} = \phi$. Cette décomposition permet d'introduire une description simple du décodage par syndrome que l'on présente ci-dessous, et dont l'idée principale est que l'information pouvant être obtenue sur l'erreur ayant affecté un mot de code lors de son passage par un canal binaire symétrique se résume à la valeur des $n - k$ derniers bits de la séquence obtenue en appliquant $\bar{\phi}^{-1}$ à la séquence reçue. Une autre utilité de l'introduction de cette décomposition réside dans le fait qu'en mécanique quantique, toute action sur un système est une opération possédant le même ensemble de départ et d'arrivée. Cela contraint naturellement tout encodage en mécanique quantique à être défini comme une opération où le système physique de départ est plongé dans un système plus large dans un état préalablement fixé, afin de subir ensuite une transformation bijective sur l'ensemble des états du système global. La décomposition que l'on introduit ici permettra donc d'effectuer un rapprochement naturel des cadres classique et quantique sous un seul formalisme.

Néanmoins, lorsque l'on procèdera au rapprochement des cadres classique et quantique, on travaillera plus spécifiquement sur l'effet des encodages sur l'ensemble des erreurs affectant les séquences et les états plutôt que sur les séquences et les états eux-mêmes. Cette approche permettra de définir un modèle discret lié à ces erreurs, ce qui simplifiera grandement l'étude des encodages quantiques que l'on introduira. Dans le cas présent, le fait de travailler sur les séquences ou sur l'ensemble des erreurs affectant ces séquences ne demande aucun effort supplémentaire. En effet, toute erreur peut être vue comme une séquence de $\mathbb{F}_2^n$, où chaque coordonnée vaut 1 s'il y a inversion de bit en la position i et 0 sinon, de sorte qu'une erreur e transforme une séquence y en $y + e$. Par linéarité de l'encodage ϕ , l'image d'une telle séquence erronée correspond alors à $\phi(y) + \phi(e)$, ce qui signifie que l'encodage fait correspondre à l'erreur y l'erreur $\phi(y)$, et agit ainsi de la même manière sur l'ensemble des erreurs et sur l'ensemble des séquences. Cette similarité de l'effet prévaut également pour l'automorphisme $\bar{\phi}$ défini plus haut. Dans la suite de cette section, on portera donc simplement notre attention sur les séquences elles-mêmes.

Définition 1.5.1. Soit ϕ un encodage de $\mathbb{F}_2^k$ dans $\mathbb{F}_2^n$. Un automorphisme $\bar{\phi}$ de $\mathbb{F}_2^n$ dont la restriction à $\mathbb{F}_2^k \times \{0\}^{n-k}$ envoie toute séquence $(x_1, \dots, x_k, 0, \dots, 0)$ sur $\phi(x_1, \dots, x_k)$ sera appelé un **encodeur** correspondant à l'encodage ϕ .

Développons à présent l'idée du décodage par syndrome. Supposons que l'on cherche à transmettre un message $x = (x_1, \dots, x_k)$, choisi a priori suivant une loi uniforme sur $\mathbb{F}_2^k$, via un canal binaire symétrique de probabilité d'erreur par bit p . On applique pour cela à x un encodage linéaire ϕ définissant un code C , et le mot de code obtenu $y = (y_1, \dots, y_n)$ est envoyé via le canal. Soit $\bar{\phi}$ un encodeur associé à ϕ , de sorte que $y = \bar{\phi}(x, 0_{n-k})$ où 0_{n-k} est la séquence de $n - k$ bits égaux à 0. Lors du passage dans le canal, en chaque position i comprise entre 1 et n , le bit envoyé y_i est affecté d'une erreur modélisée par le bit e_i pour lequel le bit reçu en sortie s'écrit $y'_i = y_i + e_i$. D'après la loi de probabilité du canal, le bit d'erreur e_i vaut 1 avec probabilité p et 0 avec probabilité $1 - p$. La loi de probabilité de la séquence d'erreurs $e = (e_1, \dots, e_n)$ est donc égale à $p^d(1 - p)^{n-d}$, où d est le poids de Hamming de e . Cette loi de probabilité dépend uniquement de la séquence d'erreur et non de la séquence envoyée, si bien que $p(e|y) = p(e|y') = p(e)$. Ecrivons $e = \bar{\phi}(x_e, s)$, de sorte que par linéarité de $\bar{\phi}$, la séquence reçue $y' = y + e$ s'écrit $y' = \bar{\phi}(x + x_e, s)$. A la réception de la séquence y' en sortie du canal, on sait que la séquence d'erreurs e appartient à la classe d'équivalence de y' dans $\mathbb{F}_2^n/C$ soit $y' + C$. En d'autres termes, e appartient à l'ensemble $\{\bar{\phi}(t, s), t \in \mathbb{F}_2^k\}$. Cette information s'obtient donc en appliquant $\bar{\phi}^{-1}$ à la séquence reçue, et en mesurant la séquence s des $n - k$ derniers bits. La tâche du décodage au maximum de vraisemblance est alors de retrouver, dans cet ensemble, la séquence $\hat{e}$ qui maximise la probabilité $p(\hat{e}|y') = p(\hat{e})$. Cette tâche ne dépend que de la seule valeur de la séquence s appelée syndrome de la séquence d'erreurs. Une fois la séquence d'erreurs $\hat{e} = \bar{\phi}(\hat{x}_e, s)$ la plus vraisemblable trouvée, cela permet d'avoir l'estimation du message au maximum de vraisemblance $\hat{x} = x + \hat{x}_e$.

L'opération de décodage par syndrome se résume comme suit. Par linéarité, toute erreur s'appliquant à un mot de code peut être exprimée telle que dans la définition suivante.

Définition 1.5.2. Soit $\bar{\phi}$ un encodeur associé à un encodage ϕ de $\mathbb{F}_2^k$ dans $\mathbb{F}_2^n$. Une **erreur** est un élément $e \in \mathbb{F}_2^n$, dont l'effet est de transformer un mot de code $y \in \mathbb{F}_2^n$ en la séquence $e + y$. Le **syndrome** d'une erreur e est la séquence des $n - k$ dernières coordonnées de $\bar{\phi}^{-1}(e)$. Le **décodage par syndrome** consiste alors à trouver, pour une séquence z reçue, l'erreur e la plus vraisemblable connaissant son syndrome.

Les $n - k$ dernières positions de l'entrée de l'encodeur $\bar{\phi}$ jouent donc un rôle particulier aussi bien lors de l'encodage que lors du décodage. Elles sont initialisées à 0 lors de l'encodage, tandis que lors du décodage elles servent à mesurer le syndrome de la séquence d'erreurs.

1.6 Capacité d'un canal et conception d'encodages optimaux

En termes d'utilisation des ressources, un encodage est d'autant plus efficace qu'il permet de minimiser la probabilité d'erreur après décodage tout en gardant un taux k/n aussi grand que possible. Une question se pose alors. Peut-on transmettre de l'information via un canal avec un taux d'erreur par symbole arbitrairement petit, et si oui, au prix de quel taux d'encodage ? Intuitivement, la question revient à rechercher la quantité d'information pouvant être transmise pour chaque utilisation du canal.

La réponse à cette question est donnée par le théorème de codage en présence de bruit dû à Shannon [33] [14], donnant une propriété asymptotique sur tout canal sans mémoire. Ce théorème affirme les deux points clés suivants :

1. Il existe une quantité maximale C comprise entre 0 et 1 nommée capacité du canal, pour laquelle il existe une famille d'encodages dont le taux k/n tend vers C et dont la probabilité d'erreur après décodage tend vers 0.
2. C est l'information mutuelle maximale que l'on peut obtenir entre un symbole d'entrée et un symbole de sortie du canal. Formellement : $C = \max_{p_Y} I(Y; Z)$, où Y est la variable aléatoire désignant le symbole d'entrée et de distribution de probabilité p_Y , et Z est la variable aléatoire désignant le symbole de sortie correspondant du canal.

L'information mutuelle $I(Y; Z)$ des deux variables aléatoires Y et Z est définie comme

$$I(Y; Z) = H(Y) + H(Z) - H(Y, Z)$$

où la fonction H désigne l'entropie de Shannon définie pour une variable aléatoire X d'ensemble d'éventualités Ω et de loi p par :

$$H(X) = \sum_{x \in \Omega} -p(x) \log_2 p(x)$$

où $\log_2$ est le logarithme en base 2. Ce théorème très général, valable pour tout canal sans mémoire sur un alphabet fini quelconque, donne la limite théorique de l'utilisation d'un canal pour transmettre de l'information. La question pratique de la complexité du décodage peut toutefois constituer un frein à l'exploitation optimale du canal telle que promise par le théorème. Il est possible de montrer qu'une famille d'encodages atteignant la capacité peut s'obtenir par un procédé aléatoire comme suit. On construit une famille d'encodages dont le taux k/n tend vers C par valeurs inférieures, et où chacun des $|A|^k$ mots de code est obtenu en tirant chacun de ses n symboles selon une loi de probabilité p_Y qui maximise l'information mutuelle. Le défaut majeur de cette construction est la complexité exponentielle de l'algorithme de décodage qui, par défaut de l'existence de structure dans le code, revient à rechercher exhaustivement le mot de code le plus proche de la séquence reçue à la sortie du canal.

Dès lors, une question intéressante en théorie des codes correcteurs d'erreurs est de rechercher des schémas d'encodage réunissant les trois propriétés fondamentales suivantes : atteindre des taux proches de la capacité d'un canal donné, posséder un décodage rapide (idéalement en temps linéaire), et permettre une reconstitution de l'information après décodage avec une probabilité d'erreur tendant vers 0. Les familles d'encodages à décodage itératif tels que les turbo-encodages qui seront présentés au chapitre suivant possèdent ces propriétés. Même si la capacité théorique du canal n'est pas atteinte avec de tels codes, ces derniers réalisent un bon compromis en offrant une complexité de décodage linéaire avec la longueur du code.

1.7 L'encodage convolutif

Désormais, on se place dans le cas de l'alphabet $A = \mathbb{F}_2$ muni de sa structure de corps. Avant d'aborder le turbo-encodage, commençons par présenter l'encodage convolutif et certaines de ses propriétés. Cet encodage est le constituant de base des turbo-encodages.

Un encodage convolutif de paramètres n , k et m permet d'encoder une séquence de bits semi-infinie $(x_u)_{u \in \mathbb{N}}$ en une séquence de bits $(y_u)_{u \in \mathbb{N}}$ en mettant à jour continuellement le contenu d'un registre interne de m bits de mémoire initialisés à 0. L'entrée et la sortie sont échantillonnées en une suite de séquences de respectivement k bits

$$(x_{kt}, \dots, x_{k(t+1)-1})_{t \in \mathbb{N}}$$

et n bits

$$(y_{nt}, \dots, y_{n(t+1)-1})_{t \in \mathbb{N}}$$

L'indice t peut être nommé temps ou étape de l'encodage. Afin de simplifier les notations, désignons par

$$x^{(t)} = (x_{kt}, \dots, x_{k(t+1)-1})$$

et

$$y^{(t)} = (y_{nt}, \dots, y_{n(t+1)-1})$$

les séquences partielles d'entrée et de sortie au temps t , et par

$$a^{(t)} = (a_1(t), \dots, a_m(t))$$

le contenu du registre de mémoire. L'évolution de la mémoire et de la sortie est alors régie par les règles suivantes :

$$\begin{aligned} a^{(0)} &= (0, \dots, 0) \\ a^{(t+1)} &= \mu(x^{(t)}, a^{(t)}) \\ y^{(t)} &= \pi(x^{(t)}, a^{(t)}) \end{aligned}$$

où μ et π sont des opérateurs linéaires. De plus, afin de rendre l'opération d'encodage convolutif injective, l'endomorphisme constitué du couple (μ, π) doit être injectif.

Introduisons à présent une définition plus flexible d'un encodage convolutif qui généralise la description précédente et qui permet d'encoder des séquences de longueur finie. La définition que l'on introduit sera notamment utile afin de définir le turbo-encodage. Par ailleurs, *l'encodeur* correspondant sera également utile pour définir un encodage convolutif et un turbo-encodage dans le cadre quantique, ce qui sera présenté au chapitre 5.

L'encodage convolutif de taille N que l'on définit ci-dessous, et représenté dans la figure 1.1 sous la forme d'un circuit, intègre l'état initial du registre de mémoire, qui n'est plus nécessairement égal à la séquence nulle, dans la séquence d'entrée. Il opère ensuite comme précédemment par des relations d'évolution linéaires jusqu'au temps N , et termine en intégrant l'état du registre de mémoire au temps N dans la séquence de sortie.

On adopte la convention suivante : toute séquence s'écrivant sous la forme $x^{(t)}$, $a^{(t)}$ ou $y^{(t)}$ désignera respectivement une séquence de longueur k , m ou n et sera nommée état d'entrée, état de mémoire ou état de sortie. Une séquence formée par la concaténation de plusieurs états s'écrira comme une séquence d'états séparés par des virgules. Il faut noter que dans la figure 1.1, les lignes représentent non pas des bits, mais des états constitués potentiellement de plusieurs bits.

Définition 1.7.1. Un **encodage convolutif** ϕ_N de taille N et de paramètres $[n, k, m]$ est une opération linéaire injective de $\mathbb{F}_2^{m+Nk}$ dans $\mathbb{F}_2^{Nn+m}$, basée sur un encodage linéaire ϕ de $\mathbb{F}_2^{k+m}$ dans $\mathbb{F}_2^{n+m}$ appelé encodage de base. À une séquence d'entrée donnée elle associe une séquence de sortie selon le procédé ci-dessous.

Soit une séquence d'entrée de longueur $m+Nk$ écrite sous la forme suivante :

$$x = (a^{(0)}, x^{(0)}, x^{(1)}, \dots, x^{(N-1)})$$

Cette séquence permet alors de définir les deux séquences suivantes :

- la séquence d'états de la mémoire :

$$(a^{(0)}, a^{(1)}, \dots, a^{(N-1)}, a^{(N)})$$

- la séquence de sortie :

$$y = (y^{(0)}, y^{(1)}, \dots, y^{(N-1)}, a^{(N)})$$

Notons que l'on retrouve le dernier état de mémoire dans la séquence de sortie. Les relations permettant d'obtenir ces deux séquences sont données, pour tout $t \in \llbracket 0; N-1 \rrbracket$, par :

$$(y^{(t)}, a^{(t+1)}) = \phi(a^{(t)}, x^{(t)})$$

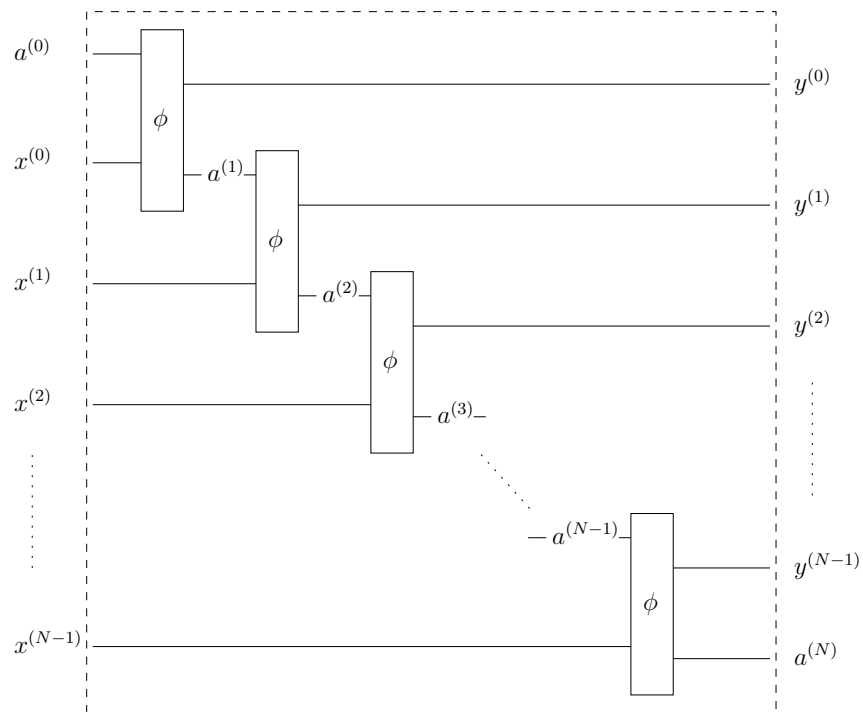


FIGURE 1.1 – Encodage convolutif ϕ_N de taille N basé sur l’encodage ϕ

Cette même définition peut être directement étendue à une séquence d'entrée semi-infinie et concorde avec la première définition d'un encodage convolutif :

Définition 1.7.2. Un encodage convolutif infini ϕ_∞ de paramètres $[n, k, m]$ basé sur un encodage de base ϕ associe à une séquence d'entrée infinie sous la forme $(a^{(0)}, x^{(0)}, x^{(1)}, \dots)$ la séquence de sortie infinie $(y^{(0)}, y^{(1)}, \dots)$ définie par les relations, pour tout $t \in \llbracket 0; \infty \rrbracket$:

$$(y^{(t)}, a^{(t+1)}) = \phi(a^{(t)}, x^{(t)})$$

Il sera utile de décomposer tout encodage de base sous la forme $\phi = (\mu, \pi)$ qui coïncide alors avec les définitions précédentes de μ et de π , de sorte à ce qu'à toute étape de l'encodage convolutif, on a :

$$y^{(t)} = \mu(a^{(t)}, x^{(t)}) \text{ et } a^{(t+1)} = \pi(a^{(t)}, x^{(t)})$$

Les propriétés recherchées chez un encodage convolutif afin de construire un turbo-encodage de distance minimale non bornée sont le caractère récursif et le caractère systématique. Afin de permettre par ailleurs le fonctionnement de l'algorithme de décodage itératif du turbo-encodage, il faut ajouter à cela le caractère non catastrophique. On expose ces propriétés dans les définitions suivantes.

Définition 1.7.3. Si l'encodage de base ϕ d'un encodage convolutif est tel que les k premiers bits y_t de la sortie au temps t sont une copie de l'état d'entrée x_t , l'encodage convolutif est dit **systématique**. Les positions en sortie d'un tel encodage convolutif qui correspondent à la copie des états d'entrée correspondent à la **sortie systématique**.

Définition 1.7.4. Si, pour un encodage convolutif infini, toute séquence d'entrée de poids de Hamming égal à 1 et dont l'état de mémoire $a^{(0)}$ est égal à 0 produit une séquence de sortie de poids de Hamming infini, l'encodage convolutif infini est dit **récursif**. Un encodage convolutif fini est également dit récursif si l'encodage convolutif infini correspondant est récursif.

Définition 1.7.5. Si, pour un encodage convolutif infini, il existe une séquence d'entrée de poids de Hamming infini pour laquelle la séquence de sortie est de poids de Hamming fini, l'encodage convolutif infini est dit **catastrophique**. De même, un encodage convolutif fini est également dit catastrophique si l'encodage convolutif infini correspondant est catastrophique.

Le caractère catastrophique d'un encodage convolutif est néfaste en ceci qu'une erreur du canal portant sur un nombre fini de positions peut, sans être détectée, engendrer une erreur sur une infinité de positions dans le message.

Le rôle du caractère systématique et le rôle du caractère récursif dans la distance minimale seront mis en évidence au chapitre 4.

1.8 Performances de décodage d'un code convolutif

Un encodage convolutif ne peut pas corriger une erreur de poids infini lorsque sa taille tend vers l'infini, ce qui constitue un frein à sa capacité de décodage comme l'affirme la propriété suivante.

Propriété 1.8.1. *La suite des distances minimales respectives des encodages convolutifs $(\phi_N)_{N \in \mathbb{N}^*}$ est bornée supérieurement par une constante qui dépend de l'encodage de base ϕ . Le maximum atteint par cette suite s'appelle la **distance libre** de l'encodage convolutif.*

Démonstration. Soit e le plus petit entier tel que $ke \geq m$. Considérons l'application de $\mathbb{F}_2^{ke}$ dans $\mathbb{F}_2^m$ qui associe à toute séquence d'états d'entrée $(x^{(0)}, \dots, x^{(e-1)})$ le dernier état de mémoire de la séquence $(0_m, x^{(0)}, \dots, x^{(e-1)})$, où 0_m désigne l'état de mémoire constitué de m 0. Comme $ke \geq m$, cette application ne peut être injective et il existe donc une séquence non nulle $(x^{(0)}, \dots, x^{(e-1)})$ telle que la séquence $\phi_e(0_m, x^{(0)}, \dots, x^{(e-1)})$ soit terminée par l'état de mémoire 0_m . Soit alors

$$(y^{(0)}, \dots, y^{(e-1)}, 0_m) = \phi_e(0_m, x^{(0)}, \dots, x^{(e-1)})$$

Pour tout $N \geq e$, on a alors :

$$\phi_N(0_m, x^{(0)}, \dots, x^{(e-1)}, 0_k, \dots, 0_k) = (y^{(0)}, \dots, y^{(e-1)}, 0_m, \dots, 0_m)$$

Comme l'encodage convolutif est injectif, la séquence $(y^{(0)}, \dots, y^{(e-1)})$ est non nulle. Ainsi ϕ_N a une distance minimale bornée par le poids de Hamming de $(y^{(0)}, \dots, y^{(e-1)})$ qui est indépendant de N . $\square$

Dans le cas simple du canal binaire symétrique, il est même possible d'aller plus loin en montrant l'échec certain du décodage au maximum de vraisemblance.

Propriété 1.8.2. *Le décodage au maximum de vraisemblance d'un encodage convolutif, dans le cas d'un canal binaire symétrique avec probabilité d'erreur $p > 0$, échoue avec une probabilité tendant vers 1 avec la taille de l'encodage convolutif.*

Démonstration. Supposons sans perte de généralité que le mot de code nul est envoyé par un encodage convolutif de taille N . Notons $(z^{(0)}, \dots, z^{(N-1)}, a^{(N)})$ la séquence reçue à la sortie du canal, où pour tout $t \in \llbracket 0; N-1 \rrbracket$, $z^{(t)}$ est la séquence de longueur n reçue après l'envoi de l'état de sortie nul au temps t de l'encodage convolutif. Soit

$$(y^{(0)}, \dots, y^{(e-1)}, 0_m) = \phi_e(0_m, x^{(0)}, \dots, x^{(e-1)})$$

un mot de code non nul de l'encodage convolutif d'une taille $e \geq m/k$, tel que présenté dans la démonstration précédente. Pour tout entier t compris entre 0 et $N - e$, notons E_t l'événement « la séquence $(z^{(t)}, \dots, z^{(t+e-1)})$ est plus proche

de la séquence $(y^{(0)}, \dots, y^{(e-1)})$ que de la séquence nulle au sens de la distance de Hamming $\gg$. Il est clair que si l'événement E_t est réalisé, la séquence reçue est nécessairement plus proche du mot de code

$$(0_n, \dots, 0_n, y^{(0)}, \dots, y^{(e-1)}, 0_n, \dots, 0_n, 0_m)$$

où $y^{(0)}$ apparaît au temps t que du mot de code nul, et par conséquent le décodage au maximum de vraisemblance échoue nécessairement. Si l'on note d le poids de Hamming de la séquence $(y^{(0)}, \dots, y^{(e-1)})$, la probabilité de l'événement E_t vaut :

$$p_E = \sum_{i=\lfloor d/2 \rfloor + 1}^d \binom{d}{i} p^i (1-p)^{d-i}$$

Cette probabilité est indépendante de t , et s'obtient en considérant tous les motifs d'erreurs d'un poids i strictement supérieur à $d/2$. Une condition nécessaire pour un décodage correct au maximum de vraisemblance est l'échec de tous les événements E_t . On peut alors constater que les événements E_t , où t décrit les multiples de e entre 0 et $N - e$, sont indépendants car ils portent sur des positions disjointes de la séquence reçue et car le canal considéré est sans mémoire. La probabilité de succès du décodage au maximum de vraisemblance est donc inférieure à :

$$(1 - p_E)^{\lfloor N/e \rfloor}$$

Cela prouve que le décodage au maximum de vraisemblance échoue avec une probabilité tendant vers 1 lorsque N tend vers l'infini. $\square$

Un encodage convolutif est donc incapable de protéger l'information lors d'un passage par un canal bruité. Il est en revanche possible, en le parallélisant ou en le concaténant avec un deuxième encodage convolutif, d'obtenir un encodage plus puissant. C'est le principe du turbo-encodage qui offre à la fois une propriété de distance minimale non bornée et une capacité de décodage très efficace utilisant un algorithme en temps linéaire. On présentera donc d'abord les deux schémas de turbo-encodage, suivis du résultat fondamental concernant leur distance minimale. On se focalise ensuite sur le schéma de turbo-encodage en série, pour lequel on présentera les grandes lignes de la preuve de ce résultat. On décrira finalement l'algorithme de décodage itératif et on présentera sa performance expérimentale.

1.9 Le turbo-encodage en parallèle

Passons en revue, dans cette section et la suivante, deux modèles de turbo-encodage. A l'origine des turbo-codes [10], le modèle de turbo-encodage utilisé fait intervenir des encodages convolutifs en parallèle ; ce modèle peut être nommé *turbo-encodage en parallèle*. Un autre modèle développé par la suite [7], où deux encodages convolutifs sont placés en série, est nommé *turbo-encodage en série*. Le turbo-encodage en parallèle se décompose en deux opérations successives. Il

prend en entrée un message $x = (x_1, \dots, x_{m+Nk})$ constitué de $m + Nk$ bits et produit un mot de code constitué de $Nk + d(N(n - k) + m)$ bits.

La première opération consiste à effectuer d copies du message. Puis, pour tout $i \in \llbracket 2; d \rrbracket$, une permutation $\pi_i \in S_{m+Nk}$ entrelace les positions des bits de la i -ème copie. On obtient ainsi d séquences sous la forme $x_{\pi_i} = (x_{\pi_i(1)}, \dots, x_{\pi_i(m+Nk)})$ où, pour la première séquence, π_1 est égale à l'identité.

Ensuite, un encodage convolutif systématique ϕ_N , appelé encodage convolutif interne, de taille N et de paramètres n, k et m est appliqué à chacune de ces d séquences. La première séquence obtenue est gardée intégralement, tandis que sur les $d - 1$ dernières séquences, les Nk bits correspondant à la sortie systématique sont supprimés pour ne garder que les $N(n - k) + m$ bits restants. Le mot de code produit est alors égal à la concaténation de toutes ces séquences.

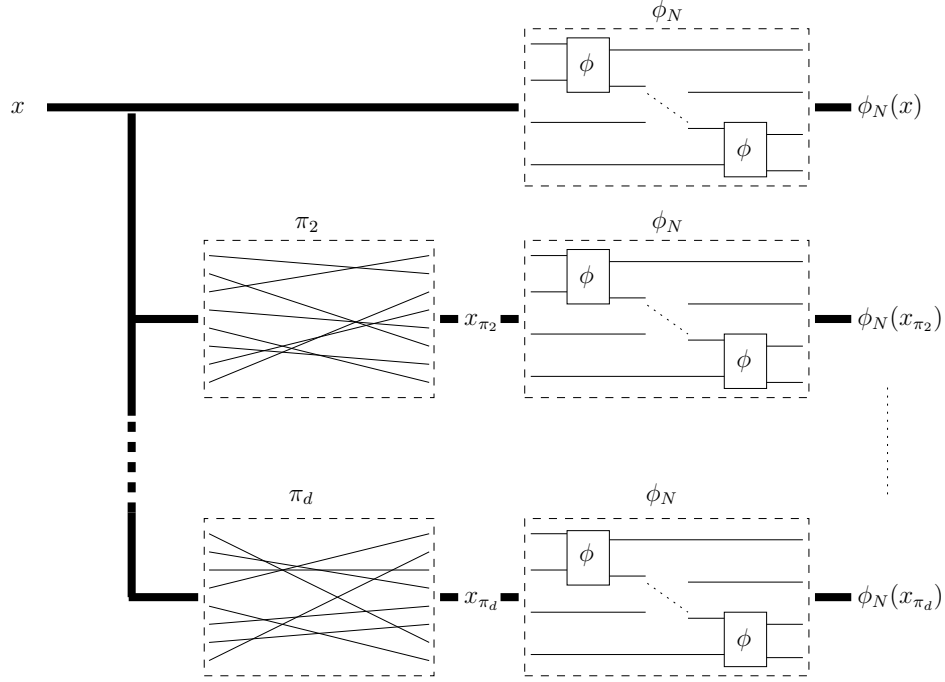


FIGURE 1.2 – Turbo-encodage en parallèle

Un turbo-encodage en parallèle est donc entièrement défini par l'encodage convolutif interne ϕ_N et par les $d - 1$ permutations π_i .

De par l'impossibilité, en mécanique quantique, de dupliquer l'état d'un bit quantique en d bits quantiques, le schéma de turbo-encodage en parallèle n'a pas d'équivalent évident dans la théorie des codes correcteurs d'erreurs quantiques. C'est pour cela que l'accent est mis davantage, dans le cadre de cette thèse, sur le schéma du turbo-encodage en série.

1.10 Le turbo-encodage en série

Présentons à présent le second schéma de turbo-encodage. Le turbo-encodage en série consiste en la concaténation de trois opérations : un encodage convolutif, un entrelacement des positions des bits, puis un second encodage convolutif. Le premier encodage convolutif est qualifié d'externe tandis que le second est qualifié d'interne. Afin de pouvoir appliquer une telle concaténation, il est nécessaire que la longueur de la séquence en sortie de l'encodage convolutif externe soit égale à la longueur de la séquence en entrée de l'encodage convolutif interne. Introduisons donc au préalable la définition suivante :

Définition 1.10.1. Lorsque les tailles respectives N_{ext} et N_{in} ainsi que les paramètres respectifs n_{ext} , k_{ext} , m_{ext} et n_{in} , k_{in} et m_{in} de deux encodages convolutifs $\phi_{ext N_{ext}}$ et $\phi_{in N_{in}}$ vérifient la relation $N_{ext}n_{ext} + m_{ext} = m_{in} + N_{in}k_{in}$, la paire d'encodages convolutifs $(\phi_{ext N_{ext}}, \phi_{in N_{in}})$ est dite compatible.

Un turbo-encodage en série se définit alors de la manière suivante à partir d'une paire compatible d'encodages convolutifs.

Définition 1.10.2. Le turbo-encodage en série de taille $N = N_{in}$ défini à partir de la paire compatible d'encodages $(\phi_{ext N_{ext}}, \phi_{in N_{in}})$ et la permutation $\pi \in S_{N_{ext}n_{ext}+m_{ext}}$ prend en entrée un message $(x_1, \dots, x_{m_{ext}+N_{ext}k_{ext}})$ de longueur $m_{ext}+N_{ext}k_{ext}$ et lui applique les opérations successives suivantes pour produire un mot de code de longueur $N_{in}n_{in} + m_{in}$:

- l'encodage convolutif externe $\phi_{ext N_{ext}}$ qui donne une séquence :

$$(y_1, \dots, y_{N_{ext}n_{ext}+m_{ext}})$$

- la permutation $\pi \in S_{N_{ext}n_{ext}+m_{ext}}$ qui entrelace la position des bits de la séquence obtenue précédemment : $(y_{\pi(1)}, \dots, y_{\pi(N_{ext}n_{ext}+m_{ext})})$
- l'encodage convolutif interne $\phi_{in N_{in}}$ qui donne le mot de code :

$$\phi_{in N_{in}}(y_{\pi(1)}, \dots, y_{\pi(N_{ext}n_{ext}+m_{ext})})$$

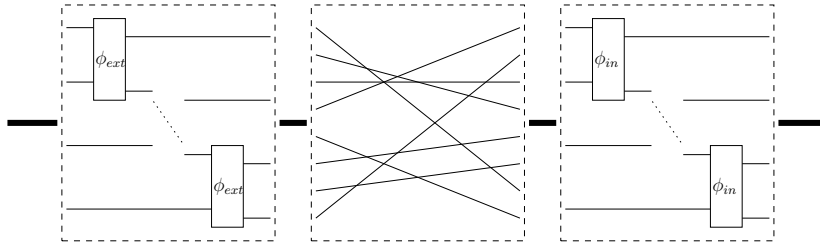


FIGURE 1.3 – Turbo-encodage en série

1.11 Distance minimale d'un turbo-encodage

Le résultat fondamental suivant est une propriété statistique portant sur la distance minimale d'un turbo-encodage. Il se décline en deux théorèmes pour les deux schémas respectifs de turbo-encodage. Les bornes polynomiales sont le résultat [22] de N.Kahale et R.Urbanke, tandis que la borne logarithmique est, dans le cas du turbo-encodage en parallèle, le résultat [12] de M. Breiling dont un raisonnement similaire s'étend au turbo-encodage en série :

Théorème 1.11.1 (Distance minimale d'un turbo-encodage en parallèle). *Soit ϕ_N un encodage convolutif récursif et systématique d'encodage de base ϕ fixé, et soient $\pi_2, \dots, \pi_d$ des permutations tirées selon une loi uniforme sur S_{m+Nk} . Alors, pour tout $t < \frac{d-2}{d}$, avec une probabilité tendant vers 1 lorsque N tend vers l'infini, la distance minimale du turbo-encodage en parallèle défini par l'encodage convolutif ϕ_N et les permutations $\pi_2, \dots, \pi_d$ est supérieure à :*

$$N^t$$

si $d \geq 3$. Si $d = 2$, alors pour toute permutation π_2 , la distance minimale obtenue est majorée par la borne :

$$O(\log N)$$

Théorème 1.11.2 (Distance minimale d'un turbo-encodage en série). *Soit $(\phi_{ext N_{ext}}, \phi_{in N_{in}})$ une paire compatible d'encodages convolutifs d'encodages de bases respectifs ϕ_{ext} et ϕ_{in} fixés. Notons d la distance libre de l'encodage convolutif externe, et supposons que l'encodage convolutif interne est récursif et systématique. Alors, pour tout $t < \frac{d-2}{d}$, avec une probabilité tendant vers 1 lorsque $N = N_{in}$ tend vers l'infini, la distance minimale du turbo-encodage en série de taille N basé sur la paire compatible d'encodages convolutifs $(\phi_{ext N_{ext}}, \phi_{in N_{in}})$ et une permutation π tirée selon une loi uniforme sur $S_{m_{in}+N_{in}k_{in}}$ est supérieure à :*

$$N^t$$

si $d = 3$. Si $d = 2$, alors pour toute permutation π , la distance minimale obtenue est majorée par la borne :

$$O(\log N)$$

Le comportement de la distance minimale de ces deux modèles de turbo-encodage est donc très proche dans les deux cas. Le nombre d de branches externes du turbo-encodage en parallèle y joue le même rôle que la distance libre du turbo-encodage en série. Ces deux théorèmes affirment alors l'existence d'un régime polynomial et un régime logarithmique pour la distance minimale, en fonction de la valeur de d . En réalité lorsque $d = 2$, dans le cas des deux modèles de turbo-encodage ci-dessus, même s'il existe une limite logarithmique pour un choix judicieux de la permutation, lorsque celle-ci est tirée aléatoirement, la distance minimale est bornée par une constante [39] [22]. Parmi les résultats démontrés dans cette thèse, figure notamment le fait que contrairement au cas des turbo-encodages classiques, le régime correspondant au cas $d = 2$ dans le cas

quantique présente une distance minimale non bornée lorsque la permutation est tirée selon une loi aléatoire uniforme. Cela démontre une des forces du modèle de turbo-encodage quantique introduit.

Soulignons par ailleurs que [22] établit que l'exposant de N avec laquelle croît la distance minimale d'un turbo-encodage en parallèle ou en série est en réalité à la fois *majoré* et *minoré* par $(d-2)/d$. Dans le cadre de cette thèse et des théorèmes ci-dessus, on s'est simplement intéressé au caractère *minorant* d'un tel exposant afin d'établir une borne inférieure sur la distance minimale d'un turbo-encodage classique ou quantique. Le caractère *majorant* n'a pas été étudié et peut faire l'objet d'une recherche complétant ce travail.

1.12 Algorithme de décodage itératif d'un turbo-encodage en série

L'intérêt du turbo-encodage réside principalement dans sa capacité à être décodé de manière efficace et permettant d'approcher les limites posées par le théorème de Shannon sur la capacité d'un canal. Cette section s'intéresse à l'algorithme de décodage itératif du turbo-encodage en série, qui sera repris sous une forme adaptée au chapitre 3 afin de simuler le fonctionnement du turbo-encodeur quantique présenté dans le cadre de cette thèse. Une introduction détaillée au fonctionnement de cet algorithme peut être trouvée dans [21].

Récapitulons dans le tableau 1.1 les notations que l'on utilise afin de désigner les différents bits et séquences en jeu, dans le processus de transmission par un canal d'un message encodé via un turbo-encodage en série. On rappelle que la lettre a est employée pour désigner les états de mémoire, tandis que les lettres x et y sont employées pour désigner respectivement des états d'entrée et des états de sortie. La lettre z sera employée pour désigner les états reçus en sortie du canal.

On rappelle les relations :

$$(y^{(t)}, a^{(t+1)}) = \phi_{ext}(a^{(t)}, x^{(t)})$$

et :

$$(y'^{(t)}, a'^{(t+1)}) = \phi_{in}(a'^{(t)}, x'^{(t)})$$

Récapitulons également dans le tableau 1.2 les notations employées pour désigner les différentes lois de probabilité intervenant au cours de l'algorithme de décodage itératif.

Bit ou séquence	Notation
convention : bit d'indice u de l'état d'indice t de la séquence x	$x_u^{(t)}$
séquence d'entrée du turbo-encodage	$x = (a^{(0)}, x^{(0)}, \dots, x^{(N_{ext}-1)})$
séquence de mémoire de l'encodage convolutif externe	$(a^{(0)}, \dots, a^{(N_{ext})})$
séquence en sortie de l'encodage convolutif externe $\phi_{ext N_{ext}}$	$y = (y^{(0)}, \dots, y^{(N_{ext}-1)}, a^{(N_{ext})})$
séquence d'entrée de l'encodage convolutif interne	$x' = (a'^{(0)}, x'^{(0)}, \dots, x'^{(N_{in}-1)})$
séquence de mémoire de l'encodage convolutif interne	$(a'^{(0)}, \dots, a'^{(N_{in})})$
séquence en sortie du turbo-encodage	$y' = (y'^{(0)}, \dots, y'^{(N_{in}-1)}, a'^{(N_{in})})$
séquence en sortie du canal	$z = (z^{(0)}, \dots, z^{(N_{in}-1)}, a_z^{(N_{in})})$

TABLE 1.1 – notations utilisées dans l'algorithme de décodage itératif

Loi de probabilité	Notation
Loi du canal	$\mathbf{P}_{canal}$
Lois temporaire, et de mise à jour avant, arrière et locale	$\mathbf{P}_{temp}, \mathbf{P}_{av}, \mathbf{P}_{ar}, \mathbf{P}_{loc}$
Loi au niveau d'une séquence x , à l'étape i du décodage	$\mathbf{P}_i^x$
Probabilité, à l'étape i , que le bit $x_u^{(t)}$ vaille α	$\mathbf{P}_i^x(x_u^{(t)} = \alpha)$
Probabilité, à l'étape i , que l'état $x^{(t)}$ vaille $\alpha = (\alpha_1, \dots, \alpha_k)$	$\mathbf{P}_i^x(x^{(t)} = \alpha)$

TABLE 1.2 – différentes probabilités en jeu lors du décodage

Les lois des séquences x , y , x' et y' sont modélisées comme des lois indépendantes calculées au niveau de chaque bit. Les lois des séquences de mémoire a et a' sont quant à elles discrétisées au niveau des états. Les lois de probabilité au lancement de l'algorithme sont les suivantes :

- loi des bits de la séquence y' donnée par la loi du canal et la séquence z reçue, c'est-à-dire

$$\begin{cases} \mathbf{P}^{y'}(y_u^{(t)} = \alpha) &= \mathbf{P}_{canal}(y_u^{(t)} = \alpha | z_u^{(t)}) \\ \mathbf{P}^{y'}(a_u^{(N_{in})} = \alpha) &= \mathbf{P}_{canal}(a_u^{(N_{in})} = \alpha | a_{zu}^{(N_{in})}) \end{cases}$$

- une distribution de départ uniforme des bits des autres séquences.

L'algorithme, donné dans le pseudo-code 1, fonctionne par itérations successives permettant d'estimer la distribution de probabilité la plus vraisemblable au niveau de la séquence y . Dans un second temps, cette distribution permet d'estimer les bits du message x . Le nombre d'itérations i_{max} est en pratique de l'ordre de la dizaine. Chaque itération consiste en une mise à jour de la loi de la séquence x' , calculée au niveau de l'encodage interne, suivie d'une mise à jour de la loi de la séquence y , calculée au niveau de l'encodage externe. Les deux fonctions principales calculées à chaque itération, *MiseAJourLogique* et *MiseAJourPhysique*, sont explicitées dans les pseudo-codes 2 et 3. Leur rôle est de calculer la loi de probabilité marginale de chaque bit, respectivement en entrée et en sortie de l'encodage convolutif passé en argument, en se basant sur les lois de probabilité en entrée et en sortie passées en argument. Les fonctions *Entrelacement* et *Entrelacement*⁻¹ consistent simplement à appliquer respectivement l'entrelacement π du turbo-encodage et l'entrelacement inverse π^{-1} à la loi de probabilité donnée en argument. La fonction *MaxVraisemblance* désigne l'estimation au maximum de vraisemblance d'une séquence pour une loi de probabilité de cette dernière donnée en argument.

Algorithme 1 Boucle Principale de l'algorithme

```

pour  $i = 0 \rightarrow i_{max} - 1$  faire                                ▷ Boucle principale
     $\mathbf{P}_{temp}^{x'} \leftarrow \text{MISEAJOURLOGIQUE}(\mathbf{P}_i^{x'}, \mathbf{P}^{y'}, \phi_{in N_{in}})$ 
     $\mathbf{P}_{temp}^y \leftarrow \text{ENTRELACEMENT}^{-1}(\mathbf{P}_{temp}^{x'})$ 
     $\mathbf{P}_{i+1}^y \leftarrow \text{MISEAJOURPHYSIQUE}(\mathbf{P}^x, \mathbf{P}_{temp}^y, \phi_{ext N_{ext}})$ 
     $\mathbf{P}_{i+1}^{x'} \leftarrow \text{ENTRELACEMENT}(\mathbf{P}_{i+1}^y)$ 
fin pour                                                        ▷ Dernière itération

     $\mathbf{P}_{temp}^{x'} \leftarrow \text{MISEAJOURLOGIQUE}(\mathbf{P}_{i_{max}}^{x'}, \mathbf{P}^{y'}, \phi_{in N_{in}})$ 
     $\mathbf{P}_{temp}^y \leftarrow \text{ENTRELACEMENT}^{-1}(\mathbf{P}_{temp}^{x'})$ 
     $\mathbf{P}_{final}^x \leftarrow \text{MISEAJOURLOGIQUE}(\mathbf{P}^x, \mathbf{P}_{temp}^y, \phi_{ext N_{ext}})$ 
     $\hat{x} \leftarrow \text{MAXVRAISEMBLANCE}(\mathbf{P}_{final}^x)$ 
retourner  $\hat{x}$ 

```

Intéressons-nous à présent aux fonctions de mise à jour. La procédure de mise à jour de chacune de ces fonctions conforme au procédé de *belief propagation* introduit par Pearl [25] [35] [36] [37] est décrite par les pseudo-codes dans les algorithmes 2 et 3. Il convient de souligner que cette procédure calcule, au niveau de chacun des encodeurs convolutifs, une distribution de probabilité de la valeur de l'erreur en chaque bit, qui correspondrait à la distribution de probabilité marginale exacte si les estimations $\mathbf{P}_i$ et $\mathbf{P}_{temp}$ passées en paramètres étaient exactes [26]. Cependant, l'information extrinsèque $\mathbf{P}_{temp}^{x'}$, calculée au niveau de l'encodeur interne et passée à l'encodeur externe, décorrèle les différents bits de x' , et il en est de même pour l'information extrinsèque $\mathbf{P}_{i+1}^y$ calculée au niveau de l'encodeur externe et passée à l'encodeur interne. Cela explique que l'algorithme de décodage itératif dans sa globalité a vocation, non pas à calculer la distribution de probabilité marginale exacte sur les bits du message x , mais à produire une estimation expérimentalement fiable de cette distribution.

Dans l'écriture des pseudo-codes 2 et 3, on considère un encodage convolutif de taille N et de paramètres $[n, k, m]$ basé sur un encodage ϕ , dont la séquence d'entrée, la séquence de sortie et la séquence des états de la mémoire sont respectivement notées x , y et a . On note $\phi = (\pi, \mu)$ de telle sorte que si

$$(y^{(t)}, a^{(t+1)}) = \phi(a^{(t)}, x^{(t)})$$

alors

$$\begin{cases} y^{(t)} &= \pi(a^{(t)}, x^{(t)}) \\ a^{(t+1)} &= \mu(a^{(t)}, x^{(t)}) \end{cases}$$

Dans les deux cas, le calcul consiste d'abord en un parcours des blocs de l'encodage convolutif vers l'avant et un parcours de ces mêmes blocs vers l'arrière. Ces parcours calculent des distributions de probabilité, notées respectivement $\mathbf{P}_{av}$ et $\mathbf{P}_{ar}$, portant sur chacun des états de la mémoire. Ces distributions de probabilité portent bien sur chacun des états $a^{(t)}$ et non sur les bits qui les constituent. Les parcours avant et arrière sont suivis d'une mise à jour locale, qui rassemble, au niveau de chaque bloc, l'information en provenance des deux parcours. Au niveau d'un bloc t , la mise à jour locale prend en compte les distributions $\mathbf{P}_{av}(a^{(t)})$ et $\mathbf{P}_{ar}(a^{(t+1)})$; dans le cas de la fonction *MiseAJourLogique*, elle calcule alors pour chaque bit $x_u^{(t)}$ de $x^{(t)}$ la distribution de probabilité marginale :

$$\mathbf{P}_{loc}(x_u^{(t)} = \alpha) = P((x_u^{(t)} = \alpha | x_{t,\bar{u}}, y)$$

où $x_{t,\bar{u}}$ est la séquence x privée du bit $x_u^{(t)}$; de façon similaire, dans le cas de la fonction *MiseAJourPhysique*, elle calcule pour chaque bit $y_u^{(t)}$ de $y^{(t)}$ la distribution de probabilité marginale :

$$\mathbf{P}_{loc}(y_u^{(t)} = \alpha) = P((y_u^{(t)} = \alpha | x, y_{t,\bar{u}})$$

où $y_{t,\bar{u}}$ est la séquence y privée du bit $y_u^{(t)}$. Les distributions de probabilité marginales calculées considèrent en tant qu'a priori les distributions de probabilité $\mathbf{P}^x$ et $\mathbf{P}^y$ passées en paramètres des fonctions, et comportent une étape de normalisation.

Algorithme 2 Fonction *MiseAJourLogique*

fonction MISEAJOURLOGIQUE($\mathbf{P}^x, \mathbf{P}^y, \phi_N$)

 $\mathbf{P}_{av}(a^{(0)}) \leftarrow \mathbf{P}^x(a^{(0)})$
 $\triangleright$ Parcoures avant

pour $t = 0 \rightarrow N - 2$ **faire**
pour tout $\alpha \in \mathbb{F}_2^m$ **faire**

$$\mathbf{P}_{av}(a^{(t+1)} = \alpha) \leftarrow \sum_{\substack{(a,x) \\ \mu(a,x)=\alpha}} \mathbf{P}_{av}(a^{(t)} = a) \mathbf{P}^x(x^{(t)} = x) \mathbf{P}^y(y^{(t)} = \pi(a, x))$$

fin pour
fin pour
 $\mathbf{P}_{ar}(a^{(N)}) \leftarrow \mathbf{P}^x(a^{(N)})$
 $\triangleright$ Parcoures arrière

pour $t = N - 1 \rightarrow 0$ **faire**
pour tout $\alpha \in \mathbb{F}_2^m$ **faire**

$$\mathbf{P}_{ar}(a^{(t)} = \alpha) \leftarrow \sum_x \mathbf{P}^x(x^{(t)} = x) \mathbf{P}^y(y^{(t)} = \pi(\alpha, x)) \mathbf{P}_{ar}(a^{(t+1)} = \mu(\alpha, x))$$

fin pour
fin pour
pour $t = 0 \rightarrow N - 1$ **faire**
 $\triangleright$ Mise à jour locale

pour $u = 1 \rightarrow k$ **faire**
pour tout $\alpha \in \mathbb{F}_2$ **faire**

$$\begin{aligned} \mathbf{P}_{loc}(x_u^{(t)} = \alpha) &\leftarrow \sum_{x: x_u = \alpha} \sum_a \mathbf{P}_{av}(a^{(t)} = a) \\ &\times \prod_{\substack{v=1 \\ v \neq u}}^k \mathbf{P}^x(x_v^{(t)} = x_v) \mathbf{P}^y(y^{(t)} = \pi(x, a)) \mathbf{P}_{ar}(a^{(t+1)} = \mu(x, a)) \end{aligned}$$

fin pour

$$somme = \sum_{\alpha \in \mathbb{F}_2} \mathbf{P}_{loc}(x_u^{(t)} = \alpha)$$

 $\triangleright$ Normalisation

$$\mathbf{P}_{loc}(x_u^{(t)}) \leftarrow \mathbf{P}_{loc}(x_u^{(t)}) / somme$$

fin pour
fin pour
retourner $\mathbf{P}_{loc}$
fin fonction

Algorithme 3 Fonction *MiseAJourPhysique*

fonction MISEAJOURPHYSIQUE($\mathbf{P}^x, \mathbf{P}^y, \phi_N$)

 $\mathbf{P}_{av}(a^{(0)}) \leftarrow \mathbf{P}^x(a^{(0)})$
 $\triangleright$ Parcours avant

pour $t = 0 \rightarrow N - 2$ **faire**

 pour tout $\alpha \in \mathbb{F}_2^m$ **faire**

$$\mathbf{P}_{av}^{t+1}(a^{(t+1)} = \alpha) \leftarrow \sum_{\substack{(a,x) \\ \mu(a,x)=\alpha}} \mathbf{P}_{av}(a^{(t)} = a) \mathbf{P}^x(x^{(t)} = x) \mathbf{P}^y(y^{(t)} = \pi(a, x))$$

fin pour
fin pour
 $\mathbf{P}_{ar}(a^{(N)}) \leftarrow \mathbf{P}^x(a^{(N)})$
 $\triangleright$ Parcours arrière

pour $t = N - 1 \rightarrow 0$ **faire**

 pour tout $\alpha \in \mathbb{F}_2^m$ **faire**

$$\mathbf{P}_{ar}(a^{(t)} = \alpha) \leftarrow \sum_x \mathbf{P}^x(x^{(t)} = x) \mathbf{P}^y(y^{(t)} = \pi(\alpha, x)) \mathbf{P}_{ar}(a^{(t+1)} = \mu(\alpha, x))$$

fin pour
fin pour
pour $t = 0 \rightarrow N - 1$ **faire**
 $\triangleright$ Mise à jour locale

 pour $u = 1 \rightarrow n$ **faire**

 pour tout $\alpha \in \mathbb{F}_2$ **faire**

$$\begin{aligned} \mathbf{P}_{loc}(y_u^{(t)} = \alpha) \leftarrow & \sum_{\substack{(x,a) \\ (\pi(x,a))_u = \alpha}} \mathbf{P}_{av}(a^{(t)} = a) \mathbf{P}^x(x^{(t)} = x) \\ & \times \prod_{\substack{v=1 \\ v \neq u}}^n \mathbf{P}^y(y_v^{(t)} = (\pi(x, a))_v) \mathbf{P}_{ar}(a^{(t+1)} = \mu(x, a)) \end{aligned}$$

fin pour

$$somme = \sum_{\alpha \in \mathbb{F}_2} \mathbf{P}_{loc}(y_u^{(t)} = \alpha)$$

 $\triangleright$ Normalisation

$$\mathbf{P}_{loc}(y_u^{(t)}) \leftarrow \mathbf{P}_{loc}(y_u^{(t)}) / somme$$

fin pour
fin pour
retourner $\mathbf{P}_{loc}$
fin fonction

Chapitre 2

Codes correcteurs d'erreurs quantiques

D'après la théorie de la mécanique quantique, l'ensemble des états dans lesquels un système physique peut se trouver est plus riche que l'ensemble des états permis par la physique classique. De tels états sont plus délicats à produire et à protéger. Il est nécessaire de les isoler de toute interaction indésirable avec l'environnement lorsque l'on souhaite les conserver ou effectuer une opération sur eux. Outre l'intérêt qu'ils procurent en soi en tant qu'objets de recherche en physique, ces états permettent d'effectuer des opérations plus vastes que celles que permet la physique classique. De telles applications s'étendent sur plusieurs domaines incluant la cryptographie, certains algorithmes, et la simulation des systèmes quantiques tels que les molécules lors d'une réaction chimique. Néanmoins, lorsqu'on effectue une mesure permettant d'observer des états permis par la mécanique quantique ou d'en tirer une information, de tels états sont modifiés de sorte qu'on ne perçoit qu'un état ou une information conforme à notre expérience classique.

Avant même que les défis technologiques liés à implémenter des systèmes manipulables dans des états quantiques, il apparaît déjà un concept théorique nouveau d'information quantique : on peut définir une information quantique comme un état pris par un système physique tel que le permettent les lois de la mécanique quantique. Cette information est manipulée et mesurée selon des lois spécifiques aux axiomes de la mécanique quantique que l'on présente ci-dessous. Après cette présentation, on redéfinira la problématique de protection de l'information quantique contre l'erreur et on aboutira sur les définitions de code et d'encodage dans le cadre quantique.

2.1 Premier axiome : états d'un système

Le premier axiome décrit la structure mathématique prise par l'ensemble des états que peut prendre un système physique. A tout système physique S cor-

respond un espace de Hilbert $\mathcal{H}$ appelé espace quantique associé à S , de sorte que l'ensemble des états que peut prendre le système S correspond à l'ensemble des droites vectorielles de $\mathcal{H}$. Autrement dit l'ensemble des états correspond à $\mathcal{H} \setminus \{0\}$ quotienté par la relation d'équivalence : $xRy \Leftrightarrow \exists \lambda \in \mathbb{C}^* : x = \lambda y$. On adopte généralement une manière équivalente et plus concrète de définir l'état d'un système, correspondant à considérer un état comme un vecteur unitaire de l'espace quantique. On énonce ceci dans le premier axiome, dans lequel on admet également une propriété permettant de déterminer l'espace quantique associé à la réunion de deux systèmes physiques en fonction des espaces quantiques propres à chacun des deux systèmes :

Axiome 1. L'état d'un système S est un vecteur unitaire d'un espace de Hilbert H associé au système et appelé **espace quantique associé à S** . L'état d'un système est donc un élément de la sphère unité $\mathcal{B}(\mathcal{H})$.

De plus, si deux systèmes physiques disjoints S_1 et S_2 ont pour espaces quantiques associés respectifs H_1 et H_2 , le système physique S constitué de la réunion de S_1 et S_2 a pour espace quantique associé le produit tensoriel $H_1 \otimes H_2$.

Il faut garder à l'esprit, à la lumière de la description précédente, que deux vecteurs unitaires de $\mathcal{H}$ représentent le même état s'ils diffèrent d'une constante multiplicative.

Par la suite, on utilise la notation $|\psi\rangle$ pour désigner l'état d'un système, et plus globalement tout vecteur (non nécessairement unitaire) d'un espace quantique. Le produit hermitien entre deux vecteurs $|\psi\rangle$ et $|\psi'\rangle$ est noté $\langle\psi|\psi'\rangle$. Le projecteur orthogonal de rang 1 sur la droite vectorielle engendrée par un vecteur non nul $|\psi\rangle$ s'écrit $|\psi\rangle\langle\psi|$.

Décrivons à présent les systèmes physiques élémentaires contenant de l'information quantique. Ce sont les systèmes qui, classiquement, peuvent contenir une quantité finie d'information. On admet le résultat suivant qui se base sur le deuxième axiome.

Propriété 2.1.1. *Un système physique permettant de représenter un symbole d'un alphabet de taille n possède un espace quantique associé dont la dimension est au moins égale à n . Il existe un tel système physique, appelé minimal, dont l'espace quantique associé est de dimension n et admet une base orthonormée d'états correspondant aux n différents symboles de l'alphabet.*

Cela permet donc d'introduire les définitions suivantes de registre quantique et de symbole quantique pour un symbole d'un alphabet fini A :

Définition 2.1.1. Un **registre à 1 symbole** est un système physique dont l'espace quantique associé est de dimension $|A|$ et admet comme base orthonormée la famille d'états $(|x\rangle)_{x \in A}$, nommée famille d'**états classiques du registre**, et où $|x\rangle$ désigne l'état du système correspondant au symbole x . Un **symbole quantique** est un état d'un registre à 1 symbole quantique.

En conséquence, un symbole quantique est un état qui s'écrit sous la forme :

$$|\psi\rangle = \sum_{x \in A} \alpha_x |x\rangle$$

où les α_x sont des coefficients complexes tels que :

$$\sum_{x \in A} |\alpha_x|^2 = 1$$

Un cas que l'on étudie en particulier dans le cadre de cette thèse est celui du **registre à 1 bit**, système minimal utilisé pour représenter un bit, et dont on note $\mathcal{H}_2$ l'espace quantique associé. $\mathcal{H}_2$ est de dimension 2 et admet pour base orthonormée la famille d'états $(|0\rangle, |1\rangle)$ désignant les états classiques du registre représentant respectivement le bit 0 et le bit 1. Un bit quantique est l'état que peut prendre un tel registre et s'écrit sous la forme :

$$|\psi\rangle = \alpha_0 |0\rangle + \alpha_1 |1\rangle$$

où $|\alpha_0|^2 + |\alpha_1|^2 = 1$.

La manière courante de représenter l'information classique est d'utiliser un ensemble de registres à 1 symbole. Dans la théorie de l'information classique, un tel ensemble de registres prend un nombre restreint d'états, qui sont ceux qui se décomposent en une séquence d'états classiques de ces registres. L'architecture similaire potentielle pour le traitement de l'information quantique consiste à utiliser également un ensemble de registres à 1 symbole, pouvant se trouver cette fois-ci dans tout état permis par la théorie de la mécanique quantique. Définissons le système physique constitué d'un tel ensemble de registres.

Définition 2.1.2. Un registre à n symboles est un système physique constitué de la réunion de n registres à 1 symbole.

Cette définition s'applique de manière directe au cas binaire pour lequel le système de n registres binaires porte le nom de **registre à n bits**. On se focalise désormais dans toute la suite sur le cas de l'alphabet binaire. Un registre à n bits étant la réunion de n registres à 1 bit, il découle du premier axiome que l'espace quantique associé à un registre à n bits est égal à $\mathcal{H}_2^{\otimes n}$. Il a pour base orthonormée canonique la famille des 2^n états classiques du registre, qui s'obtient en effectuant le produit euclidien des n familles d'états classiques respectives à chaque registre à 1 bit. Toute séquence de n bits $x = (x_1, \dots, x_n)$ est représentée par l'état de la base canonique $|x_1\rangle \otimes \dots \otimes |x_n\rangle$ que l'on notera plus simplement $|x\rangle$. On appelle séquence de n bits quantiques tout état du registre à n bits. Une séquence de n bits quantiques s'écrit donc sous la forme :

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

où les scalaires α_x vérifient $\sum_{x \in \{0,1\}^n} |\alpha_x|^2 = 1$.

Il est très important de noter le phénomène suivant dit **phénomène de non-localité**, applicable au cas d'un registre à n bits comme au cas général de toute réunion finie de systèmes physiques disjoints. A ce stade, il est utile d'introduire la définition d'un état factorisé ou intriqué.

Définition 2.1.3. L'état $|\psi\rangle$ d'une réunion de systèmes physiques disjoints $\bigcup_{i=1}^n S_i$ est dit **factorisé** s'il est décomposable en un produit d'états locaux $\bigotimes_{i=1}^n |\psi\rangle_i$ propres à chacun des systèmes S_i , ou **intriqué** s'il ne l'est pas.

L'état d'un système physique égal à une réunion de sous-systèmes n'est donc généralement pas une notion définie de manière locale. L'exemple le plus simple d'un état intriqué est l'état de Bell d'un registre à 2 bits : $|\psi\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle)$. Un tel état ne peut pas s'écrire sous la forme factorisée :

$$|\psi\rangle = (\alpha_0|0\rangle + \alpha_1|1\rangle) \otimes (\beta_0|0\rangle + \beta_1|1\rangle)$$

Contrairement à l'intuition classique, il n'est donc pas possible d'affirmer que le premier ou le deuxième registre à 1 bit se trouve dans un état particulier.

Il faut donc retenir qu'un système ne possède pas toujours un état propre à lui mais peut faire partie d'un système englobant qui se trouve dans un état intriqué.

2.2 Deuxième axiome : mesure effectuée sur un système physique

Le deuxième axiome décrit quant à lui tout processus de mesure effectué sur un système physique. On se restreint ici au cas particulier des mesures projectives. Une mesure est une opération permettant d'extraire de l'information sur l'état d'un système physique. En mécanique quantique, la mesure est une opération perturbative : elle transforme l'état du système en fonction du résultat mesuré. L'axiome permet de définir une mesure et de déterminer son action de la façon suivante.

Axiome 2. Soit S un système physique d'espace quantique associé $\mathcal{H}$. Soient J un ensemble et $(P_j)_{j \in J}$ une famille de projecteurs orthogonaux sur $\mathcal{H}$ vérifiant $\sum_{j \in J} P_j = I$. La **mesure** sur S de **famille de projecteurs associés** $(P_j)_{j \in J}$ est une opération probabiliste sur le système S agissant de la manière suivante. Si $|\psi\rangle$ est l'état du système au moment de la mesure, la mesure donne le résultat $j \in J$ et donne au système l'état projeté normalisé $\frac{P_j|\psi\rangle}{\|P_j|\psi\rangle\|}$ avec une probabilité égale à $\langle\psi|P_j|\psi\rangle = \|P_j|\psi\rangle\|^2$. J est nommé l'ensemble des résultats de la mesure.

Il est possible de montrer que les espaces images $(\text{Im } P_j)_{j \in J}$ d'une telle famille de projecteurs sont en somme directe orthogonale. Cet axiome est cohérent avec le premier axiome puisqu'il garantit la préservation du caractère unitaire d'un état au cours d'une mesure.

Parmi les mesures que permet l'axiome, on en développe quelques types qui présenteront un intérêt par la suite.

Définition 2.2.1. La mesure associée à une base $(|\psi_j\rangle)_{j \in J}$ de l'espace quantique est la mesure de famille de projecteurs associés $(P_j)_{j \in J}$, où pour tout $j \in J$, P_j est le projecteur orthogonal sur $\text{Vect}(|\psi_j\rangle)$, soit $P_j = |\psi_j\rangle\langle\psi_j|$.

L'exemple typique de mesure associée à une base est la mesure de la valeur des n bits d'un registre à n bits : c'est une mesure associée à la base $(|x\rangle)_{x \in \{0,1\}^n}$. Lorsqu'elle est appliquée à un état sous la forme $|\psi\rangle = \sum_{x \in \{0,1\}^n} \alpha_x |x\rangle$, elle donne le résultat j et met le système dans l'état $|j\rangle$ avec probabilité $|\alpha_j|^2$.

Définition 2.2.2. La mesure de la valeur du i -ème bit d'un registre à n bits est définie par la famille des deux projecteurs $(P_j)_{j \in \{0,1\}}$ où pour tout $j \in \{0,1\}$, P_j est le projecteur orthogonal sur l'espace $\mathcal{H}_2^{\otimes i-1} \otimes \text{Vect}(|j\rangle) \otimes \mathcal{H}_2^{\otimes n-i}$.

Appliquée à un état sous la forme $|\psi\rangle = \sum_{x \in \{0,1\}^n} \alpha_x |x\rangle$, cette mesure donne en résultat le bit j avec probabilité $p_j = \sum_{\substack{x \in \{0,1\}^n \\ x_i = j}} |\alpha_x|^2$, et met alors le système dans l'état :

$$\frac{1}{\sqrt{p_j}} \sum_{\substack{x \in \{0,1\}^n \\ x_i = j}} \alpha_x |x\rangle$$

Définition 2.2.3. La mesure d'un ensemble $S \subset \llbracket 1; n \rrbracket$ de s bits ($|S| = s$) d'un registre à n bits est définie par la famille des 2^s projecteurs $(P_j)_{j \in \{0,1\}^S}$, où pour tout $j \in \{0,1\}^S$, P_j est le projecteur orthogonal sur l'espace s'écrivant sous la forme du produit tensoriel $\bigotimes_{i=1}^n \mathcal{H}_i^{(j)}$ avec $\mathcal{H}_i^{(j)} = \mathcal{H}_2$ si $i \notin S$ et $\mathcal{H}_i^{(j)} = \text{Vect}(|j(i)\rangle)$ si $i \in S$.

Cette mesure correspond exactement à effectuer successivement, et dans un ordre quelconque, les s mesures de chacun des bits de l'ensemble S . Appliquée à un état sous la forme $|\psi\rangle = \sum_{x \in \{0,1\}^n} \alpha_x |x\rangle$, elle donne le résultat $j \in \{0,1\}^S$ avec probabilité $p_j = \sum_{\substack{x \in \{0,1\}^n \\ x_{/S} = j}} |\alpha_x|^2$ où $x_{/S}$ est la restriction de la séquence x à l'ensemble S , et met alors le système dans l'état :

$$\frac{1}{\sqrt{p_j}} \sum_{\substack{x \in \{0,1\}^n \\ x_{/S} = j}} \alpha_x |x\rangle$$

Il est à noter que ce dernier état est factorisé sur le registre des s bits de l'ensemble S et sur le registre des $n - s$ bits de l'ensemble $\bar{S}$ complémentaire dans $\llbracket 1; n \rrbracket$. Il s'écrit en effet :

$$\frac{1}{\sqrt{p_j}} |j\rangle \otimes \sum_{x \in \{0,1\}^{\bar{S}}} \alpha_{(j,x)} |x\rangle$$

où (j, x) désigne la séquence de $\llbracket 1; n \rrbracket$ qui coïncide avec j sur S et qui coïncide avec x sur $\bar{S}$.

2.3 Troisième axiome : évolution d'un système physique isolé

Le troisième axiome concerne l'évolution des systèmes physiques isolés. Un système physique est dit isolé s'il n'est pas en interaction avec un autre système physique, et s'il ne subit pas de mesure. Dans le cas d'un système physique S en interaction avec un autre système physique, il faut considérer le plus petit système physique isolé contenant le système S pour lui appliquer l'axiome de l'évolution. Cet axiome s'énonce comme suit.

Axiome 3. Un système S isolé, c'est-à-dire n'interagissant pas avec un autre système, évolue conformément à une transformation unitaire U agissant sur l'espace quantique $\mathcal{H}$ qui lui est associé.

En d'autres termes, si un système S est isolé entre un certain temps initial et un certain temps final, il existe une transformation unitaire U telle que si $|\psi\rangle$ est l'état du système au temps initial, son état au temps final sera $U|\psi\rangle$.

Cet axiome est lui aussi cohérent avec le premier axiome en ceci qu'il garantit la conservation du caractère unitaire de tout état d'un système au cours d'une évolution isolée.

Si un système S_1 dans un état $|\psi\rangle_1$ fait partie d'un système plus large S qui se trouve dans un état factorisé $|\psi\rangle_1 \otimes |\psi\rangle_2$, une manière de rendre son état intriqué avec celui de $S_2 = S \setminus S_1$ est d'appliquer une évolution unitaire U au système S qui ne peut pas s'écrire comme le produit $U_1 \otimes U_2$ de deux évolutions locales à chacun des systèmes S_1 et S_2 . Une telle évolution correspond à une interaction entre les deux systèmes S_1 et S_2 et peut être obtenue dès que le système S_1 n'est pas parfaitement physiquement isolé du système S_2 . On appelle ce phénomène **couplage**.

Le couplage non contrôlé de deux systèmes est un phénomène que l'on cherche à éviter si le système S_1 contient un état que l'on cherche à protéger et que le système S_2 représente l'environnement du système S_1 . En effet, l'environnement est constamment sujet à des mesures à cause de la présence d'observateurs. Or, comme on le développe dans la partie suivante traitant du canal quantique, le couplage non contrôlé du système S_1 avec son environnement, suivi d'une mesure sur l'environnement, est susceptible d'induire une erreur aléatoire sur l'état du système S_1 .

A l'inverse, le couplage est l'opération élémentaire utilisée pour encoder l'état d'un système ou pour lui appliquer un calcul. Dans le cas du registre à 2 bits, une manière simple de coupler chacun des registres à 1 bit est d'appliquer une évolution correspondant à la porte logique classique C-NOT. Cette porte prend en entrée deux bits b_1 et b_2 , et inverse l'état du bit b_2 conditionnellement à l'état du bit b_1 . Ainsi, la paire de bits obtenue en sortie d'une porte C-NOT est $(b_1, b_1 \oplus b_2)$ où $\oplus$ désigne l'opération OU-exclusif correspondant à l'évaluation de l'assertion $b_1 \neq b_2$. Cette porte possède un équivalent en mécanique quantique, qui est l'évolution unitaire U dont l'action sur la base canonique du registre à 2 bits est la suivante : pour toute paire de bits $x = (x_1, x_2)$, $U|x\rangle = |y\rangle$ où

$y = (x_1, x_1 \oplus x_2)$. Si l'on applique une porte C-NOT quantique à un registre à 2 bits dans l'état factorisé :

$$|\psi\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \otimes |0\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |10\rangle)$$

le résultat obtenu est l'état intriqué nommé *état de Bell* :

$$U|\psi\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle)$$

2.4 Erreur quantique et canal quantique

On présente ici le processus d'erreur affectant un système S_1 . Ce processus correspond au « Positive Operator-Valued Measure » (POVM), un type généralisé de mesure d'un système et qui peut se décomposer comme un couplage avec un système S_2 suivi d'une mesure projective effectuée sur le système S_2 . Un processus d'erreur est exactement un POVM dont on oublie le résultat de la mesure. Cette présentation débouchera sur la définition d'un canal quantique.

Commençons donc par présenter le principe d'un POVM. Soit S_1 un système d'espace quantique associé $\mathcal{H}^{(1)}$, formant avec un système S_2 d'espace quantique associé $\mathcal{H}^{(2)}$ un système global S . On note $|\psi\rangle_1$ l'état de départ du système S_1 , et on suppose que le système S_2 est au départ dans un état $|\psi\rangle_2$ fixé et indépendant de $|\psi\rangle_1$. En premier lieu, un couplage U s'opère entre les systèmes S_1 et S_2 , de sorte que l'état du système global S après couplage devient $U(|\psi\rangle_1 \otimes |\psi\rangle_2)$. Ce couplage est suivi d'une mesure projective effectuée sur le système S_2 et associée à une base orthonormée $(|\psi_j\rangle_2)_{j \in J}$ de $\mathcal{H}^{(2)}$. Au niveau du système S , cette mesure correspond à la mesure projective de famille de projecteurs associée $(I \otimes |\psi_j\rangle_2 \langle \psi_j|_2)_{j \in J}$, où I désigne l'opérateur identité agissant sur le système S_1 . Afin de décrire l'effet de cette mesure, notons que $(|\psi_j\rangle_2)_{j \in J}$ étant une base de $\mathcal{H}^{(2)}$, tout vecteur $|\psi\rangle \in \mathcal{H}^{(1)} \otimes \mathcal{H}^{(2)}$ se décompose sous la forme :

$$|\psi\rangle = \sum_{j \in J} |\psi_j\rangle_1 \otimes |\psi_j\rangle_2$$

où pour tout $j \in J$, $|\psi_j\rangle_1$ est un vecteur de $\mathcal{H}^{(1)}$ dépendant de manière linéaire du vecteur $|\psi\rangle$. Par conséquent, pour tout état $|\psi\rangle_1$ du système S_1 , on peut écrire l'état $U(|\psi\rangle_1 \otimes |\psi\rangle_2)$ sous la forme :

$$U(|\psi\rangle_1 \otimes |\psi\rangle_2) = \sum_{j \in J} (E_j |\psi\rangle_1) \otimes |\psi_j\rangle_2$$

où E_j est un endomorphisme de $\mathcal{H}^{(1)}$. Avec une probabilité égale à $\|E_j |\psi\rangle_1\|^2$, la mesure donne ainsi le résultat j et met le système S dans l'état factorisé

$$\frac{E_j |\psi\rangle_1}{\|E_j |\psi\rangle_1\|} \otimes |\psi_j\rangle_2$$

de sorte que le système S_1 à la fin du processus d'erreur devient

$$\frac{E_j|\psi\rangle_1}{\|E_j|\psi\rangle_1\|}$$

avec probabilité $\|E_j|\psi\rangle_1\|^2$.

Les opérateurs E_j , définis ci-dessus à partir de U , sont nommés opérateurs de Kraus. Ils vérifient une condition nécessaire et suffisante découlant du caractère unitaire de U . On obtient cette condition en requérant que, pour tout état $|\psi\rangle_1$ du système S_1 , $U(|\psi\rangle_1 \otimes |\psi\rangle_2)$ soit de norme 1. En écrivant le produit hermitien de ce dernier vecteur par lui-même, il vient :

$$\sum_{j \in J} (\langle \psi |_1 E_j^\dagger) \otimes \langle \psi |_2 \sum_{j \in J} (E_j |\psi\rangle_1) \otimes |\psi_j\rangle_2 = 1$$

où $E_j^\dagger$ est l'opérateur adjoint de E_j et où $\langle \psi |_1$ et $\langle \psi |_2$ désignent respectivement les opérations de produit hermitien à gauche par les vecteurs $|\psi\rangle_1$ et $|\psi\rangle_2$. Par orthonormalité de la famille $(|\psi_j\rangle_2)_{j \in J}$, cela équivaut à la condition :

$$\sum_{j \in J} \langle \psi |_1 E_j^\dagger E_j |\psi\rangle_1 = 1$$

valable pour tout état $|\psi\rangle_1$ du système S_1 . La condition portant sur les opérateurs E_j et que l'on nomme condition de Kraus s'écrit donc :

$$\sum_{j \in J} E_j^\dagger E_j = I$$

En conclusion, un POVM est une opération plus large qu'une mesure projective, en ceci qu'elle n'est pas associée à une famille de projecteurs orthogonaux $(P_j)_{j \in J}$ vérifiant la condition $\sum_{j \in J} P_j = I$ mais, plus largement, à une famille d'opérateurs linéaires $(E_j)_{j \in J}$ vérifiant la condition de Kraus $\sum_{j \in J} E_j^\dagger E_j = I$. Résumons cette description par la définition suivante.

Définition 2.4.1. Soit S un système d'espace quantique associé $\mathcal{H}$. Une **famille d'opérateurs de Kraus** est une famille $(E_j)_{j \in J}$ d'applications linéaires sur $\mathcal{H}$ vérifiant la **condition de Kraus** :

$$\sum_{j \in J} E_j^\dagger E_j = I$$

Un **Positive Operator Valued Measure (POVM)** sur S de famille d'opérateurs de Kraus associés $(E_j)_{j \in J}$ est une opération probabiliste agissant de la manière suivante. Si $|\psi\rangle$ est l'état du système au début de l'opération, le POVM donne le résultat $j \in J$ et donne au système l'état $\frac{E_j|\psi\rangle}{\|E_j|\psi\rangle\|}$ avec une probabilité égale à $\langle \psi | E_j^\dagger E_j | \psi \rangle = \|E_j|\psi\rangle\|^2$.

Le processus d'erreur qui affecte un système est similaire à celui qui entre en jeu dans un POVM. Le système S_1 qui subit une erreur entre en interaction avec un système S_2 représentant l'environnement du système S_1 . Le couplage qui a lieu entre les deux systèmes ainsi que la mesure projective sur l'environnement sont des opérations non contrôlées et inconnues, dont on ne connaît que les opérateurs de Kraus correspondants grâce à des expériences répétées sur le système S_1 qui révèlent des propriétés statistiques concernant le processus d'erreur. En revanche, le résultat de la mesure projective sur l'environnement est quant à lui perdu. Par exemple, on peut imaginer que le système S_1 est une boîte et que l'environnement S_2 est l'atmosphère. Dans un tel cas, la mesure projective qui s'effectue sur S_2 peut provenir de la simple interaction entre les molécules de l'air du système S_2 avec une personne présente sans que cette personne ne soit consciente de la mesure effectuée.

On retient ainsi la définition de deux notions : celle d'une erreur qui agit sur un système, ainsi que celle d'un canal quantique qui transforme l'état d'un système par l'application d'une erreur conformément à une certaine loi de probabilité.

Définition 2.4.2. Une **erreur** sur un système S d'espace quantique associé $\mathcal{H}$ est un endomorphisme de $\mathcal{H}$. **L'action d'une erreur** E sur un état $|\psi\rangle$ est définie dans le cas où $|\psi\rangle \notin \text{Ker } E$ et correspond à la transformation $|\psi\rangle \mapsto \frac{E|\psi\rangle}{\|E|\psi\rangle\|}$.

Un **canal quantique** sur S est donné par une famille d'opérateurs de Kraus $(E_j)_{j \in J}$ sur $\mathcal{H}$ appelée **famille d'erreurs associée** au canal quantique. Il transforme tout état $|\psi\rangle$ du système S en un état $\frac{E_j|\psi\rangle}{\|E_j|\psi\rangle\|}$, où E_j est tel que $E_j|\psi\rangle \neq 0$, conformément à la loi de probabilité $P(E_j|\psi) = \|E_j|\psi\rangle\|^2$.

Lorsque les erreurs quantiques agissent sur un registre à n symboles, elles ont un poids correspondant au nombre maximal de positions du registre sur lesquelles elles agissent simultanément. Pour les modèles de canaux quantiques courants que l'on présente, les familles d'erreurs associées ont une distribution de poids intéressante. Afin d'introduire formellement un tel poids, considérons une base $(E^{(0)}, \dots, E^{(k^2-1)})$ des erreurs quantiques agissant sur un registre à 1 symbole, où $k \geq 1$ désigne la taille de l'alphabet sur lequel est défini le symbole, et telle que $E^{(0)}$ correspond à l'opérateur identité. Toute erreur quantique agissant sur un registre à n symboles peut alors se décomposer sur la famille des erreurs $(E^{(0)}, \dots, E^{(k^2-1)})^{\otimes n}$, ce qui permet d'identifier son poids.

Définition 2.4.3. Toute erreur quantique E agissant sur un registre à n symboles peut s'écrire comme une combinaison linéaire :

$$E = \sum_{v \in \llbracket 0; k^2-1 \rrbracket^n} c_v E^{(v_1)} \otimes \dots \otimes E^{(v_n)}$$

où pour tout $v \in \llbracket 0; k^2-1 \rrbracket^n$, $c_v \in \mathbb{C}$. Le **poids de l'erreur quantique** E est le nombre maximal de coordonnées différentes de 0 d'une séquence $v \in \llbracket 0; k^2-1 \rrbracket^n$

telle que $c_v \neq 0$:

$$\max\{|\{i \in \llbracket 1; n \rrbracket, v_i \neq 0\}|, v \in \llbracket 0; k^2 - 1 \rrbracket^n, c_v \neq 0\}$$

Cette définition de poids est indépendante du choix de la base d'erreurs $(E^{(0)}, \dots, E^{(k^2-1)})^{\otimes n}$:

Propriété 2.4.1. Soient $(E^{(0)}, \dots, E^{(k^2-1)})$ et $(F^{(0)}, \dots, F^{(k^2-1)})$ deux bases de l'espace quantique associé à un registre à 1 symbole, telles que $E^{(0)} = F^{(0)} = I$ où I désigne l'opérateur identité. Le poids d'une erreur quantique E agissant sur un registre à n symboles est identique selon qu'on décompose E sur la première ou la deuxième base.

Démonstration. On décompose la première base sur la deuxième, en écrivant pour tout $i \in \llbracket 0; k^2 - 1 \rrbracket$:

$$E^{(i)} = \sum_{j=0}^{k^2-1} a_{i,j} F^{(j)}$$

où pour tout $j \in \llbracket 0; k^2 - 1 \rrbracket$, $a_{i,j} \in \mathbb{C}$. Soit alors E une erreur quantique agissant sur un registre à n symboles, et écrivons :

$$E = \sum_{v \in \llbracket 0; k^2 - 1 \rrbracket^n} c_v E^{(v_1)} \otimes \dots \otimes E^{(v_n)}$$

Alors chaque terme non nul $c_v E^{(v_1)} \otimes \dots \otimes E^{(v_n)}$ de la somme se décompose sur la base $(F^{(0)}, \dots, F^{(k^2-1)})^{\otimes n}$, en remplaçant chaque $E^{(v_i)} \neq I$ par $\sum_{j=0}^{k^2-1} a_{v_i,j} F^{(j)}$. On obtient de la sorte une somme partielle où tous les coefficients non nuls correspondent à des séquences v' qui contiennent au moins autant de 0 que v . Ainsi, le poids de l'erreur E telle que décomposée sur la base $(F^{(0)}, \dots, F^{(k^2-1)})$ est inférieur ou égal au poids obtenu avec la base $(E^{(0)}, \dots, E^{(k^2-1)})$.

Par le raisonnement réciproque, en écrivant d'abord E dans la deuxième base et en décomposant le résultat sur la première base, on obtient l'égalité des poids issus des deux écritures de E . $\square$

Deux exemples de canal quantique sont le canal dépolarisant et le canal à effacement. La manière la plus simple de les décrire est la suivante.

Définition 2.4.4. Le canal dépolarisant de probabilité d'erreur p est le canal quantique ayant pour famille d'erreurs associée

$$(\sqrt{1-p}\mathcal{I}, \sqrt{p/3}\mathcal{X}, \sqrt{p/3}\mathcal{Y}, \sqrt{p/3}\mathcal{Z})$$

où $\mathcal{I}, \mathcal{X}, \mathcal{Y}, \mathcal{Z}$ sont les quatre erreurs dites de Pauli, données dans la définition 2.9.1 :

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

Le canal dépolarisant de probabilité d'effacement p est un canal quantique dont l'action sur un état $|\psi\rangle \in \mathcal{H}_2$ est de le transmettre intact avec probabilité $1-p$, et de transmettre un état fixé $|e\rangle \notin \mathcal{H}_2$ avec probabilité p .

Une définition du canal dépolarisant est également possible en termes de famille d'erreurs associée, en considérant que ce canal agit non pas sur $\mathcal{H}_2$ mais sur un espace quantique plus grand, de dimension au moins 3, généré par $|0\rangle$, $|1\rangle$ et un état $|e\rangle$. Il est toutefois plus simple de le considérer de la manière décrite ci-dessus.

2.5 Encodage et code quantiques

On présente la problématique de l'encodage de l'information quantique que contient un registre à k symboles. Sans perte de généralité, on peut supposer que les symboles sont définis sur un alphabet A contenant le symbole 0. On notera $\mathcal{H}_A$ l'espace quantique de dimension $|A|$ associé à un registre à 1 symbole.

L'état d'un registre à k symboles est vulnérable aux erreurs s'il doit être transmis via un canal. Une idée empruntée de la théorie des codes classiques consiste alors à encoder au préalable cet état dans un registre à n symboles où $n \geq k$. Toute transformation d'un système physique étant décrite par une évolution unitaire, un encodage quantique se réalise de la manière suivante. Supposons que l'on veuille encoder l'état $|\psi\rangle$ d'un registre à k symboles. On annexe au registre à k symboles un registre à $n - k$ symboles initialisé à l'état $|0_{n-k}\rangle$ où 0_{n-k} désigne la séquence constituée de $n - k$ symboles égaux à 0. Le registre à n symboles obtenu est alors dans l'état factorisé $|\psi\rangle \otimes |0_{n-k}\rangle$. Ensuite, on applique une transformation unitaire U à l'état du registre à n symboles de sorte à obtenir l'état $|\psi\rangle \otimes |0_{n-k}\rangle$. Un encodage quantique est donc défini par la donnée de n , k et la transformation unitaire U .

Définition 2.5.1. Un **encodage quantique** est un triplet (n, k, U) où n et k sont des entiers vérifiant $n \geq k \geq 1$ et U est une application unitaire de $\mathcal{H}_A^{\otimes n}$. On dira également par abus de langage qu'une application unitaire U de $\mathcal{H}_A^{\otimes n}$ est un encodage quantique de paramètres n et k en désignant par cela l'encodage quantique (n, k, U) . Un **code quantique** $\mathcal{C}$ est l'ensemble des états obtenus par un encodage quantique :

$$\mathcal{C} = U(\mathcal{B}(\mathcal{H}_A^{\otimes k}) \otimes \{|0_{n-k}\rangle\})$$

Il s'agit de la sphère unité du sous-espace vectoriel $U(\mathcal{H}_A^{\otimes k} \otimes \{|0_{n-k}\rangle\})$. L'entier n est alors appelé la **longueur du code**.

Soulignons que la définition faite ci-dessus d'un code quantique est restrictive, car elle impose à l'espace vectoriel supportant le code quantique d'avoir une dimension égale à une puissance de 2. D'une manière générale, un code quantique est la sphère unité d'un sous-espace vectoriel de $\mathcal{H}_A^{\otimes n}$. La restriction imposée ci-dessus sur la dimension d'un code quantique vient du fait que l'on souhaite utiliser les codes quantiques, dans le cadre de cette thèse, pour encoder l'état d'un registre contenant un certain nombre de bits.

Soulignons également une particularité des encodages quantiques qui les rend plus riches et plus difficiles à étudier que les encodages classiques. Un

encodage quantique peut prendre un ensemble continu de valeurs données par l'ensemble des applications unitaires de l'espace quantique $\mathcal{H}_A^{\otimes n}$, alors qu'un automorphisme d'encodage classique peut prendre un nombre fini de valeurs donné par l'ensemble des bijections de A^n . Pour cette raison, les encodages stabilisateurs quantiques que l'on présente par la suite sont un outil très puissant, car ils correspondent à une catégorie d'encodages quantiques décrits par un nombre fini de paramètres.

La capacité d'un code quantique à protéger un état quantique contre certaines erreurs est examinée après avoir présenté la notion d'états distinguables dans la section suivante. En effet, on montrera que la capacité d'un code quantique à corriger un ensemble d'erreurs donné repose sur la garantie que le caractère distinguable de deux états quelconques du code quantique est préservé sous l'action d'une erreur quelconque de cet ensemble.

2.6 Transformations, états distinguables

Par la suite, le terme « transformation physique » prendra un sens précis, et désignera toute suite finie d'opérations permises par les axiomes de la mécanique quantique.

Définition 2.6.1. Une **transformation physique** est la composée d'une suite finie d'opérations sur un système physique choisies parmi les **opérations admissibles** suivantes :

- l'intégration du système considéré dans un système plus large (qui n'est pas réellement une transformation mais plutôt une adaptation de l'échelle du système étudié)
- une évolution unitaire
- une mesure projective

Lorsqu'une mesure a lieu au cours d'une transformation, son résultat peut conditionner la suite des opérations.

Contrairement au cas de la mécanique classique, deux états peuvent être différents sans que l'on puisse avec certitude les distinguer par une mesure. Outre l'intérêt en soi que présente la particularité de ce concept en mécanique quantique, il possède une propriété remarquable vis-à-vis de toute transformation physique. Il en découlera des conditions nécessaires à un code quantique afin qu'il puisse être protégé contre un ensemble d'erreurs donné. Par la suite, pour toute mesure sur un système de famille de projecteurs associés $(P_j)_{j \in J}$, et pour tout état $|\psi\rangle$ de ce système, on appellera **ensemble des résultats probables** pour $|\psi\rangle$ l'ensemble :

$$\{j \in J / P_j |\psi\rangle \neq 0\}$$

Il s'agit de l'ensemble des résultats de la mesure ayant une probabilité non nulle d'être obtenus lorsque le système est mesuré dans l'état $|\psi\rangle$.

Définition 2.6.2. On dit que deux états $|\psi\rangle$ et $|\psi'\rangle$ d'un système physique sont **distinguable**s s'il existe une mesure dont les ensembles respectifs des résultats probables pour $|\psi\rangle$ et pour $|\psi'\rangle$ sont disjoints.

Ainsi, si un système se trouve dans l'un des deux états distinguables $|\psi\rangle$ et $|\psi'\rangle$, il existe une mesure dont le résultat permet de savoir sans ambiguïté dans lequel de ces deux états il se trouve. Montrons la propriété suivante :

Propriété 2.6.1. *Deux états sont distinguables si et seulement si ils sont orthogonaux.*

Démonstration. Si $|\psi\rangle$ est orthogonal à $|\psi'\rangle$, on complète la famille de ces deux états en une base orthonormée de l'espace quantique $\mathcal{H}$. La mesure associée à cette base donne alors deux résultats différents selon que l'on mesure $|\psi\rangle$ ou $|\psi'\rangle$. Réciproquement, supposons que $|\psi\rangle$ et $|\psi'\rangle$ sont distinguables. Soit une mesure permettant de distinguer ces deux états, et notons $(P_j)_{j \in J}$ la famille de projecteurs orthogonaux qui lui est associée. Les ensembles respectifs J_1 et J_2 des résultats probables pour $|\psi\rangle$ et pour $|\psi'\rangle$ sont par hypothèse disjoints. Or, les deux identités :

$$\begin{aligned} |\psi\rangle &= \sum_{j \in J} P_j |\psi\rangle = \sum_{j \in J_1} P_j |\psi\rangle \\ |\psi'\rangle &= \sum_{j \in J} P_j |\psi'\rangle = \sum_{j \in J_2} P_j |\psi'\rangle \end{aligned}$$

montrent que $|\psi\rangle \in \bigoplus_{j \in J_1} \text{Im } P_j$ et $|\psi'\rangle \in \bigoplus_{j \in J_2} \text{Im } P_j$. Les états $|\psi\rangle$ et $|\psi'\rangle$ appartiennent donc à des espaces orthogonaux et sont bien orthogonaux. $\square$

Le caractère distinguable de deux états d'un système ne dépend pas de l'échelle du système considéré, au sens où pour un système S_1 faisant partie d'un système $S = S_1 \cup S_2$ plus large, deux états distinguables (respectivement non distinguables) correspondent à des états distinguables (respectivement non distinguables) du système S . Supposons en effet qu'un système S_1 puisse se trouver à un moment donné dans l'état $|\psi\rangle$ ou l'état $|\psi'\rangle$, et considérons un système englobant $S = S_1 \cup S_2$, où le système S_2 est dans un état noté $|\psi\rangle_2$. Le système S se trouve alors dans un état factorisé valant respectivement $|\psi\rangle \otimes |\psi\rangle_2$ ou $|\psi'\rangle \otimes |\psi\rangle_2$. Le produit hermitien entre les états $|\psi\rangle$ et $|\psi'\rangle$ a la même valeur que le produit hermitien entre les états $|\psi\rangle \otimes |\psi\rangle_2$ et $|\psi'\rangle \otimes |\psi\rangle_2$, ce qui garantit l'équivalence entre le caractère distinguable des états $|\psi\rangle$ et $|\psi'\rangle$ et celui des états $|\psi\rangle \otimes |\psi\rangle_2$ et $|\psi'\rangle \otimes |\psi\rangle_2$.

Une propriété intéressante est qu'il n'existe aucune transformation physique qui garantit le fait de rendre distinguables une paire d'états non distinguables. Dans un certain sens, cela correspond à affirmer que le caractère distinguable de deux états ne peut que se dégrader par des transformations physiques. Cette propriété vient dans le corollaire 2.6.1 et découle de la propriété suivante.

Propriété 2.6.2. *Soient $|\psi\rangle$ et $|\psi'\rangle$ deux états d'un même système. Soit T une transformation physique comportant un nombre fixe $k \geq 1$ de mesures. Alors la*

probabilité que T ne distingue pas ces deux états est minorée par $|\langle\psi'|\psi\rangle|^2$. En des termes plus précis, si l'on note $j = (j_1, \dots, j_k)$ et $j' = (j'_1, \dots, j'_k)$ les variables aléatoires correspondant à la séquence des résultats des mesures obtenues en appliquant T respectivement à $|\psi\rangle$ et à $|\psi'\rangle$, la probabilité que $j = j'$ est minorée par $|\langle\psi'|\psi\rangle|^2$.

Démonstration. Afin de simplifier le modèle et sans perte de généralité, on peut supposer que l'ensemble des résultats pour toutes les mesures est le même; notons-le J . On supposera également, sans que cela n'affecte la preuve, que la dernière étape de la transformation est une mesure. Commençons par remarquer que l'état final du système après application de T à l'état $|\psi\rangle$ est entièrement déterminé par la séquence des résultats des mesures j ; on le notera donc $|\psi_j\rangle$, et on note de même $|\psi'_{j'}\rangle$ l'état final du système après application de T à $|\psi'\rangle$.

Démontrons le résultat H_k suivant par récurrence sur k :

$$\sum_{v \in J^k} |\langle\psi'_v|\psi_v\rangle|^2 p(j = j' = v) \geq |\langle\psi'|\psi\rangle|^2$$

Comme pour tout $v \in J^k$, $|\langle\psi'_v|\psi_v\rangle|^2 \leq 1$, ce résultat suffit à démontrer la propriété.

Supposons $k = 1$. T est composée d'une suite d'opérations préservant le produit hermitien notée O , suivie d'une mesure de famille de projecteurs associée $(P_v)_{v \in J}$. Ainsi, $p(j = v) = \|P_v O|\psi\rangle\|^2$ et $p(j' = v) = \|P_v O|\psi'\rangle\|^2$, et les états finaux $|\psi_v\rangle$ et $|\psi'_v\rangle$ si le résultat de la mesure est v valent

$$|\psi_v\rangle = \frac{P_v O|\psi\rangle}{\|P_v O|\psi\rangle\|}$$

et

$$|\psi'_v\rangle = \frac{P_v O|\psi'\rangle}{\|P_v O|\psi'\rangle\|}$$

Par conséquent :

$$\begin{aligned} \sum_{v \in J} |\langle\psi'_v|\psi_v\rangle|^2 p(j = j' = v) &= \sum_{v \in J} |\langle\psi'|O^\dagger P_v P_v O|\psi\rangle|^2 \\ &= \sum_{v \in J} |\langle\psi'|P_v|\psi\rangle|^2 \\ &\geq \left| \sum_{v \in J} \langle\psi'|P_v|\psi\rangle \right|^2 \\ &= |\langle\psi'|\psi\rangle|^2 \end{aligned}$$

où l'on obtient l'avant-dernière ligne grâce à l'inégalité de Cauchy-Schwarz, et la dernière ligne grâce au fait que $\sum P_v = I$. H_1 est donc vrai.

Supposons à présent H_{k-1} vrai pour un certain $k \geq 2$ et démontrons H_k . La transformation T est composée d'une transformation T_{k-1} comportant $k-1$ mesures, suivie d'une transformation T_1 comportant une unique mesure. T_1 est

fonction des résultats des $k - 1$ premières mesures de T_{k-1} . H_k peut s'écrire sous la forme suivante :

$$\sum_{v_k \in J} \sum_{\bar{v} \in J^{k-1}} |\langle \psi'_{(\bar{v}, v_k)} | \psi_{(\bar{v}, v_k)} \rangle|^2 p(j_k = j'_k = v_k | \bar{j} = \bar{j}' = \bar{v})$$

$$p(\bar{j} = \bar{j}' = \bar{v}) \geq |\langle \psi' | \psi \rangle|^2$$

Supposons dans un premier temps $\bar{v}$ fixé. Pour tout $v_k \in J$, la probabilité $p(j_k = j'_k = v_k | \bar{j} = \bar{j}' = \bar{v})$ correspond alors, pour la transformation T_1 , à $p(j_k = j'_k = v_k)$. Appliquons l'hypothèse H_1 sur la transformation T_1 et les états $|\psi_{\bar{v}}\rangle$ et $|\psi'_{\bar{v}}\rangle$:

$$\sum_{v_k \in J} |\langle \psi'_{(\bar{v}, v_k)} | \psi_{(\bar{v}, v_k)} \rangle|^2 p(j_k = j'_k = v_k) \geq |\langle \psi'_{\bar{v}} | \psi_{\bar{v}} \rangle|^2$$

On a donc :

$$\sum_{v_k \in J} \sum_{\bar{v} \in J^{k-1}} |\langle \psi'_{(\bar{v}, v_k)} | \psi_{(\bar{v}, v_k)} \rangle|^2 p(j_k = j'_k = v_k | \bar{j} = \bar{j}' = \bar{v})$$

$$p(\bar{j} = \bar{j}' = \bar{v}) \geq \sum_{\bar{v} \in J^{k-1}} p(\bar{j} = \bar{j}' = \bar{v}) |\langle \psi'_{\bar{v}} | \psi_{\bar{v}} \rangle|^2$$

L'application de l'hypothèse H_{k-1} au membre de droite précédent montre alors que H_k est vrai. $\square$

La propriété 2.6.2 se base par commodité sur une transformation dont le nombre de mesures k est fixe. Dans le cas général, cela n'est pas vrai, car une transformation peut être conditionnée par les résultats de ses propres mesures. La propriété reste toutefois vraie dans le cas d'une transformation quelconque, puisque l'événement $j = j'$ se produit dans un ensemble d'éventualités où la transformation agit comme si elle possédait un nombre fixe de mesures.

Corollaire 2.6.1. *Si deux états $|\psi\rangle$ et $|\psi'\rangle$ d'un même système ne sont pas distinguables, alors pour toute transformation physique T , avec une probabilité minorée par $|\langle \psi' | \psi \rangle|^2$, T produit à partir de $|\psi\rangle$ et $|\psi'\rangle$ la même séquence de résultats de mesures ainsi que deux états $T(|\psi\rangle)$ et $T(|\psi'\rangle)$ non distinguables.*

Démonstration. Soit T une transformation physique appliquée à $|\psi\rangle$ et $|\psi'\rangle$. Soient j et j' les séquences des résultats des mesures respectivement pour $|\psi\rangle$ et $|\psi'\rangle$, et $|\psi_j\rangle$ et $|\psi'_{j'}\rangle$ les états obtenus respectivement après l'application de T sur $|\psi\rangle$ et $|\psi'\rangle$. Supposons alors que l'on effectue une mesure finale afin de discriminer les états $|\psi_j\rangle$ et $|\psi'_{j'}\rangle$. Si l'on note $\tilde{j}$ et $\tilde{j}'$ les résultats respectifs de cette mesure sur $|\psi_j\rangle$ et $|\psi'_{j'}\rangle$, la propriété 2.6.2 appliquée à la transformation constituée de T suivie de cette dernière mesure affirme que $j = j'$ et $\tilde{j} = \tilde{j}'$ avec une probabilité au moins égale à $|\langle \psi' | \psi \rangle|^2$. $\square$

2.7 Code correcteur d'un ensemble d'erreurs

Abordons à présent la problématique de correction d'erreurs par un code quantique. On démontre dans cette section une propriété nécessaire et suffisante que doit détenir un code quantique pour corriger un ensemble d'erreurs. Dans la thèse de Gottesman, la propriété nécessaire et suffisante qui est donnée est la suivante :

Propriété 2.7.1. *Soit $\mathcal{C}$ un code quantique, et soit $(|\psi_i\rangle)_{i \in \llbracket 1; r \rrbracket}$ une base orthonormée d'états de $\text{Vect } \mathcal{C}$. $\mathcal{C}$ corrige un ensemble d'erreurs, si et seulement si pour toute paire d'erreurs E_a et E_b de cet ensemble, il existe une constante $C_{a,b}$ telle que pour tout $(i, j) \in \llbracket 1; r \rrbracket^2$:*

$$\langle \psi_i | E_a^\dagger \circ E_b | \psi_j \rangle = C_{a,b} \delta_{i,j}$$

On choisit ici de présenter et démontrer une propriété qui lui est équivalente, et dont l'atout est d'exprimer la problématique de correction d'un ensemble d'erreurs exclusivement en termes d'états distinguables. On montre ainsi que le rôle joué par la notion d'états *distinguables* dans la théorie des codes quantiques correspond exactement au rôle joué par la notion de séquences *différentes* dans la théorie des codes classiques. On se propose donc de prouver la propriété :

Propriété 2.7.2. *Soit $\mathcal{C}$ un code quantique. $\mathcal{C}$ corrige un ensemble d'erreurs, si et seulement si pour toute paire d'erreurs E_a et E_b de cet ensemble, et pour toute paire $(|\psi\rangle, |\psi'\rangle)$ d'états distinguables de $\mathcal{C}$, les états $E_a|\psi\rangle$ et $E_b|\psi'\rangle$ sont distinguables.*

On démontre ainsi que la condition de correction d'un ensemble d'erreurs par un code quantique s'assimile à la condition correspondante de la propriété 1.3.1 pour un code classique, pour lequel il faut et il suffit que pour toute paire (y, y') d'éléments différents du code, les éléments $E_a(y)$ et $E_b(y')$ sont différents.

Commençons par définir explicitement ce que signifie la correction d'un ensemble d'erreurs pour un code quantique.

Définition 2.7.1. Soit $\{E_j, j \in J\}$ un ensemble d'erreurs opérant sur un espace quantique $\mathcal{H}_A^{\otimes n}$, et soit $\mathcal{C}$ un code quantique de longueur n . Le code quantique $\mathcal{C}$ **corrige l'ensemble d'erreurs** $\{E_j, j \in J\}$ s'il existe une transformation physique permettant, pour tout $j \in J$ et pour tout $|\psi\rangle \in \mathcal{C}$ tel que $E_j|\psi\rangle \neq 0$, de transformer $\frac{E_j|\psi\rangle}{\|E_j|\psi\rangle\|}$ en $|\psi\rangle$.

Une condition nécessaire à ce qu'un code quantique corrige un ensemble d'erreurs $\{E_j, j \in J\}$ se déduit immédiatement de la préservation du caractère non distinguable de deux états par toute transformation physique.

Propriété 2.7.3. *Si un code quantique $\mathcal{C}$ corrige un ensemble d'erreurs $\{E_j, j \in J\}$, alors pour tout $(i, j) \in J^2$ et pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ tel que $\langle \psi' | \psi \rangle = 0$, $\langle \psi' | E_i^\dagger \circ E_j | \psi \rangle = 0$.*

Démonstration. Soit $(i, j) \in J^2$ et soit $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ tel que $\langle\psi'|\psi\rangle = 0$. Si $E_j|\psi\rangle = 0$ ou $E_i|\psi'\rangle = 0$ on a clairement $\langle\psi'|E_i^\dagger \circ E_j|\psi\rangle = 0$. Sinon, il existe par hypothèse une transformation physique permettant de transformer $\frac{E_j|\psi\rangle}{\|E_j|\psi\rangle\|}$ en $|\psi\rangle$ et $\frac{E_i|\psi'\rangle}{\|E_i|\psi'\rangle\|}$ en $|\psi'\rangle$. Les deux états $|\psi\rangle$ et $|\psi'\rangle$ étant distinguables, il est donc nécessaire d'après le corollaire 2.6.1 que $\frac{E_j|\psi\rangle}{\|E_j|\psi\rangle\|}$ et $\frac{E_i|\psi'\rangle}{\|E_i|\psi'\rangle\|}$ le soient également, d'où :

$$\langle\psi'|E_i^\dagger \circ E_j|\psi\rangle = 0$$

□

Afin de prouver que la condition est également suffisante, on a besoin de démontrer l'équivalence suivante.

Propriété 2.7.4. *Pour toute erreur E , les conditions ci-dessous sont équivalentes :*

(i) *Il existe $\alpha \in \mathbb{C}$ tel que, pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$:*

$$\langle\psi'|E|\psi\rangle = \alpha\langle\psi'|\psi\rangle$$

(ii) *Pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ tel que $\langle\psi'|\psi\rangle = 0$, $\langle\psi'|E|\psi\rangle = 0$*

Démonstration. Le sens (i) $\Rightarrow$ (ii) est évident. Prouvons le sens (ii) $\Rightarrow$ (i). Supposons (ii) vérifiée. Soit $(|\psi_1\rangle, \dots, |\psi_r\rangle)$ une base orthonormée de $\text{Vect } \mathcal{C}$ où r désigne la dimension de $\text{Vect } \mathcal{C}$. Notons, pour tout $i \in \llbracket 1; r \rrbracket$, $\alpha_i = \langle\psi_i|E|\psi_i\rangle$. Montrons d'abord que cette valeur ne dépend pas de i . Soit $i \neq j$, et posons :

$$\begin{aligned} |\psi\rangle &= \frac{1}{\sqrt{2}}(|\psi_i\rangle + |\psi_j\rangle) \\ |\psi'\rangle &= \frac{1}{\sqrt{2}}(|\psi_i\rangle - |\psi_j\rangle) \end{aligned}$$

Ces deux vecteurs sont unitaires et sont donc des états du code. De plus ils sont orthogonaux. Par conséquent on a $\langle\psi'|E|\psi\rangle = 0$, soit :

$$\frac{1}{2}(\alpha_i - \alpha_j - \langle\psi_j|E|\psi_i\rangle + \langle\psi_i|E|\psi_j\rangle) = 0$$

Or $|\psi_i\rangle$ et $|\psi_j\rangle$ sont des états orthogonaux donc :

$$\langle\psi_j|E|\psi_i\rangle = \langle\psi_i|E|\psi_j\rangle = 0$$

et par conséquent $\alpha_i = \alpha_j$.

Notons $\alpha = \alpha_1 = \dots = \alpha_r$. Soit $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$, que l'on décompose sous la forme $|\psi\rangle = \sum_{i=1}^r c_i |\psi_i\rangle$ et $|\psi'\rangle = \sum_{i=1}^r c'_i |\psi_i\rangle$ où $\sum_{i=1}^r |c_i|^2 = \sum_{i=1}^r |c'_i|^2 = 1$. Alors :

$$\begin{aligned} \langle\psi'|E|\psi\rangle &= \sum_{i,j=1}^r \bar{c}'_i c_j \langle\psi_i|E|\psi_j\rangle \\ &= \sum_{i=1}^r \bar{c}'_i c_i \alpha \\ &= \alpha \langle\psi'|\psi\rangle \end{aligned}$$

où la réduction de la somme entre la deuxième à la troisième ligne se fait en remarquant que lorsque les indices i et j sont différents, les états $|\psi_i\rangle$ et $|\psi_j\rangle$ sont des états orthogonaux et $\langle\psi_i|E|\psi_j\rangle = 0$. $\square$

Montrons à présent que la condition énoncée dans la proposition 2.7.3 est également suffisante. On explicite pour cela une transformation physique permettant de reconstituer les états du code après qu'une erreur se soit produite.

Démonstration - Propriété 2.7.2. On suppose donc que pour tout $(i, j) \in J^2$, il existe $\alpha \in \mathbb{C}$ tel que, pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$:

$$\langle\psi'|E_i^\dagger E_j|\psi\rangle = \alpha \langle\psi'|\psi\rangle$$

D'après la propriété 2.7.4, pour tout $(i, j) \in J^2$, et pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$, il existe $\alpha \in \mathbb{C}$ tel que :

$$\langle\psi'|E_i^\dagger E_j|\psi\rangle = 0$$

On recherche une transformation physique permettant de reconstituer tout état du code quantique ayant été affecté par une erreur dans l'ensemble $\{E_j, j \in J\}$. Pour cela, on construit d'abord un ensemble d'erreurs particulier $(\tilde{E}_j)_{1 \leq j \leq r}$. On démontre que cet ensemble possède trois propriétés qui permettent de concevoir une transformation corrigeant toutes les erreurs de l'ensemble $\{E_j, j \in J\}$ sur le code quantique.

Soit $r = \dim \text{Vect } \{E_j, j \in J\}$. Sans perte de généralité, on peut supposer que J contient l'ensemble d'entiers $\llbracket 1; r \rrbracket$ et que la famille d'erreurs $(E_j)_{1 \leq j \leq r}$ est une base de $\text{Vect } \{E_j, j \in J\}$. Définissons alors la matrice $C \in M_n(\mathbb{C})$ dont chaque élément $C_{i,j}$ vérifie, pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$:

$$\langle\psi'|E_i^\dagger E_j|\psi\rangle = C_{i,j} \langle\psi'|\psi\rangle$$

La matrice C est hermitienne; en effet pour tout $(i, j) \in \llbracket 1; r \rrbracket^2$ et tout $|\psi\rangle \in \mathcal{C}$:

$$C_{i,j} = \langle\psi|E_i^\dagger E_j|\psi\rangle = \langle\psi|E_j^\dagger E_i|\psi\rangle = \bar{C}_{j,i}$$

Il existe donc une matrice diagonale réelle $D = \text{diag}(d_1, \dots, d_r)$ et une matrice unitaire $U = (U_{i,j})_{1 \leq i,j \leq r}$ telles que $U^\dagger C U = D$. On définit alors la famille d'erreurs $(\tilde{E}_j)_{1 \leq j \leq r}$ par les relations, pour tout $j \in \llbracket 1; r \rrbracket$:

$$\tilde{E}_j = \sum_{k=1}^r U_{k,j} E_k$$

Soit r' le rang de D . Sans perte de généralité, on peut supposer que les premiers r' éléments diagonaux de D sont tous non nuls de sorte que $(d_1, \dots, d_r) = (d_1, \dots, d_{r'}, 0, \dots, 0)$. Afin d'alléger la présentation, présentons les trois propriétés de la famille d'erreurs $(\tilde{E}_j)_{1 \leq j \leq r}$ sous la forme d'un lemme, que l'on prouve après cette démonstration.

Lemme 2.7.1. *La famille $(\tilde{E}_j)_{1 \leq j \leq r}$ possède les trois propriétés suivantes :*

- (i) Les espaces $(\tilde{E}_j(\text{Vect } \mathcal{C}))_{1 \leq j \leq r}$ sont en somme directe orthogonale
- (ii) Pour tout $j \in \llbracket r' + 1; r \rrbracket$, la restriction de $\tilde{E}_j$ à $\mathcal{C}$ est égale à 0, et pour tout $j \in \llbracket 1; r' \rrbracket$, l'application qui à tout $|\psi\rangle \in \mathcal{C}$ associe :

$$\frac{\tilde{E}_j|\psi\rangle}{\|\tilde{E}_j|\psi\rangle\|}$$

est définie et prolongeable sur $\text{Vect } \mathcal{C}$ en une application unitaire F_j

- (iii) La famille $(\tilde{E}_j)_{1 \leq j \leq r}$ est une base de $\text{Vect } \{E_j, j \in J\}$

On construit alors une transformation définie par une mesure suivie d'une transformation unitaire. Pour tout $j \in \llbracket 1; r' \rrbracket$, soit P_j le projecteur orthogonal sur l'espace $\tilde{E}_j(\text{Vect } \mathcal{C})$. Afin d'obtenir une famille de projecteurs dont la somme vaut I , complétons-la par le projecteur orthogonal P_0 sur l'espace orthogonal de $\bigoplus_{j=1}^{r'} (\tilde{E}_j(\text{Vect } \mathcal{C}))$.

La transformation est alors constituée de la mesure de famille de projecteurs associée $(P_j)_{0 \leq j \leq r'}$ suivie de l'évolution unitaire F_j^{-1} où j est le résultat, nécessairement non nul comme on le verra, de la mesure. Montrons que cette transformation permet effectivement de reconstituer tout état du code quantique après que celui-ci soit affecté d'une erreur de l'ensemble $\{E_j, j \in J\}$.

Soit un état $|\psi\rangle \in \mathcal{C}$, affecté après son passage par le canal d'une erreur E_i où $i \in J$ est tel que $E_i|\psi\rangle \neq 0$. L'état obtenu est donc :

$$|\psi_{err}\rangle = \frac{E_i|\psi\rangle}{\|E_i|\psi\rangle\|}$$

D'après la troisième propriété du lemme 2.7.1, on peut décomposer l'erreur E_i sous la forme :

$$E_i = \sum_{j=1}^r c_{i,j} \tilde{E}_j$$

où les $c_{i,j}$ sont des complexes. L'état obtenu en sortie du canal est donc :

$$|\psi_{err}\rangle = \sum_{j=1}^r \frac{c_{i,j}}{\|E_i|\psi\rangle\|} \tilde{E}_j|\psi\rangle$$

D'après la deuxième propriété du lemme 2.7.1, $\tilde{E}_j|\psi\rangle = 0$ si $j > r'$, ce qui permet d'écrire plus simplement :

$$|\psi_{err}\rangle = \sum_{j=1}^{r'} \frac{c_{i,j}}{\|E_i|\psi\rangle\|} \tilde{E}_j|\psi\rangle$$

On effectue alors la mesure de famille de projecteurs associée $(P_j)_{0 \leq j \leq r'}$. Comme $|\psi_{err}\rangle \in \bigoplus_{j=1}^{r'} (\tilde{E}_j(\text{Vect } \mathcal{C}))$, on a $P_0|\psi_{err}\rangle = 0$ et il est impossible que

le résultat de la mesure soit 0. Appelons donc j le résultat de la mesure où $j \in \llbracket 1; r' \rrbracket$. On a :

$$P_j|\psi_{err}\rangle = \frac{c_{i,j}}{\|\tilde{E}_i|\psi\rangle\|} \tilde{E}_j|\psi\rangle$$

et ce vecteur est non nul car j appartient à l'ensemble des résultats probables de la mesure. Après la mesure, le nouvel état est donc :

$$\frac{P_j|\psi_{err}\rangle}{\|P_j|\psi_{err}\rangle\|} = \frac{\tilde{E}_j|\psi\rangle}{\|\tilde{E}_j|\psi\rangle\|} = F_j|\psi\rangle$$

En appliquant ensuite l'évolution unitaire F_j^{-1} , on reconstitue donc bien l'état de départ $|\psi\rangle$. $\square$

Démonstration du lemme 2.7.1. On se base sur le calcul suivant pour démontrer les deux premières propriétés. Pour tout $(i, j) \in \llbracket 1; r \rrbracket^2$ et pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$:

$$\begin{aligned} \langle \psi' | \tilde{E}_i^\dagger \circ \tilde{E}_j | \psi \rangle &= \langle \psi' | \left(\sum_{k=1}^r U_{k,i}^- E_k^\dagger \right) \circ \left(\sum_{k'=1}^r U_{k',j} E_{k'} \right) | \psi \rangle \\ &= \sum_{k,k'=1}^r U_{k,i}^- U_{k',j} \langle \psi' | E_k^\dagger \circ E_{k'} | \psi \rangle \\ &= \sum_{k,k'=1}^r U_{k,i}^- C_{k,k'} U_{k',j} \langle \psi' | \psi \rangle \\ &= D_{i,j} \langle \psi' | \psi \rangle \\ &= \delta_{i,j} d_i \langle \psi' | \psi \rangle \end{aligned}$$

La propriété (i) se montre alors ainsi. Pour tout $(i, j) \in \llbracket 1; r \rrbracket^2$ avec $i \neq j$, tout vecteur de $\tilde{E}_i(\text{Vect } \mathcal{C})$ et tout vecteur de $\tilde{E}_j(\text{Vect } \mathcal{C})$ s'écrivent respectivement sous la forme $\tilde{E}_i|\psi'\rangle$ et $\tilde{E}_j|\psi\rangle$ où $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$, et leur produit hermitien vaut donc :

$$\langle \psi' | \tilde{E}_i^\dagger \circ \tilde{E}_j | \psi \rangle = 0$$

Cela prouve que les espaces de la famille $(\tilde{E}_j(\text{Vect } \mathcal{C}))_{j \in J}$ sont deux à deux orthogonaux donc en somme directe orthogonale.

Montrons à présent la propriété (ii). Pour tout $|\psi\rangle \in \mathcal{C}$:

$$\|\tilde{E}_j|\psi\rangle\|^2 = \langle \psi | \tilde{E}_j^\dagger \circ \tilde{E}_j | \psi \rangle = d_j \langle \psi | \psi \rangle = d_j$$

Si $j \in \llbracket r' + 1; r \rrbracket$, $d_j = 0$ donc $\tilde{E}_j|\psi\rangle = 0$ et la restriction de $\tilde{E}_j$ à $\mathcal{C}$ est égale à 0. Supposons maintenant que $j \in \llbracket 1; r' \rrbracket$. Alors $d_j > 0$ et $\|\tilde{E}_j|\psi\rangle\| = \sqrt{d_j}$. Ainsi, l'application linéaire F_j définie par la restriction à $\text{Vect } \mathcal{C}$ de $\frac{\tilde{E}_j}{\sqrt{d_j}}$ est unitaire car elle envoie tout vecteur unitaire sur un vecteur unitaire. De plus, comme

pour tout $|\psi\rangle \in \mathcal{C}$, $\|\tilde{E}_j|\psi\rangle\| = \sqrt{d_j}$, F_j coïncide sur $\mathcal{C}$ avec l'application qui à tout $|\psi\rangle \in \mathcal{C}$ associe :

$$\frac{\tilde{E}_j|\psi\rangle}{\|\tilde{E}_j|\psi\rangle\|}$$

Pour démontrer la propriété (iii), il suffit de prouver que la famille $(\tilde{E}_j)_{1 \leq j \leq r}$ engendre tous les éléments de la famille $(E_j)_{1 \leq j \leq r}$ puisque cette dernière est une base de Vect $\{E_j, j \in J\}$. Soit donc $j \in \llbracket 1; r \rrbracket$. Le calcul suivant donne :

$$\begin{aligned} \sum_{k=1}^r U_{j,k}^- \tilde{E}_k &= \sum_{k,k'=1}^r U_{j,k}^- U_{k',k} E_{k'} \\ &= \sum_{k,k'=1}^r U_{k,j}^\dagger U_{k',k} E_{k'} \\ &= \sum_{k'=1}^r \left(\sum_{k=1}^r U_{k',k} U_{k,j}^\dagger \right) E_{k'} \\ &= \sum_{k'=1}^r \delta_{k',j} E_{k'} \\ &= E_j \end{aligned}$$

ce qui prouve le résultat. $\square$

A partir de ce résultat, il vient un corollaire immédiat :

Corollaire 2.7.1. *Un code quantique qui corrige un ensemble d'erreurs corrige également l'espace vectoriel engendré par ce dernier.*

Démonstration. Si $\mathcal{C}$ corrige un ensemble d'erreurs $\{E_j, j \in J\}$, alors pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ tel que $\langle \psi' | \psi \rangle = 0$ et pour tout $(i, j) \in J^2$, $\langle \psi' | E_i^\dagger \circ E_j | \psi \rangle = 0$. Par linéarité, il vient que pour tout $(E, E') \in (\text{Vect } \{E_j, j \in J\})^2$, $\langle \psi' | E'^\dagger \circ E | \psi \rangle = 0$, ce qui montre le résultat. $\square$

2.8 Distance minimale d'un code quantique

De manière similaire au lien entre la distance minimale d d'un code classique et sa capacité à corriger toutes les erreurs d'un poids borné par $(d-1)/2$, la définition de la distance minimale d'un code quantique doit permettre d'affirmer qu'un code quantique de distance minimale d corrige toutes les erreurs de poids inférieur ou égal à $(d-1)/2$. Cette définition s'obtient en transposant de manière naturelle la définition de distance minimale dans le cadre classique et en y remplaçant la notion de différence par la notion de distinguabilité.

Définition 2.8.1. La distance minimale d'un code quantique $\mathcal{C}$ est le plus petit poids d d'une erreur quantique E telle qu'il existe $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$ et $\langle \psi' | E | \psi \rangle \neq 0$.

2.9 Erreurs de Pauli

On introduit les erreurs de Pauli car elles constituent la base de la construction des codes stabilisateurs. Il s'agit d'un ensemble d'erreurs possédant des propriétés remarquables et rendant plus simple l'étude des codes quantiques. Un code stabilisateur se définit alors comme un espace propre associé à un sous-ensemble d'erreurs de Pauli formant un groupe pour la loi de composition et appelé groupe stabilisateur. Dans le cadre de cette thèse, les codes stabilisateurs que l'on étudie sont définis sur un alphabet binaire, donc dans l'espace quantique $\mathcal{H}_2^{\otimes n}$.

Définition 2.9.1. Soient les quatre erreurs définies sur $\mathcal{H}_2$ par leurs matrices représentatives dans la base canonique ($|0\rangle, |1\rangle$) :

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

On note $\mathcal{P} = \{\mathcal{I}, \mathcal{X}, \mathcal{Y}, \mathcal{Z}\}$. Pour tout $n \geq 1$, on appelle **erreur de Pauli de taille n** toute erreur définie sur $\mathcal{H}_2^{\otimes n}$ par :

$$E = cE_1 \otimes \dots \otimes E_n$$

où $c \in \Phi = \{1, -1, i, -i\}$ et pour tout $i \in \llbracket 1; n \rrbracket$, $E_i \in \mathcal{P}$.

Propriété 2.9.1. L'ensemble $\Phi \times \mathcal{P}^{\otimes n}$ forme un groupe pour la loi de composition des endomorphismes appelé **groupe de Pauli de taille n** . On notera cette loi $\circ$ dans le cadre de ce groupe. De plus, tout élément de $\pm \mathcal{P}^{\otimes n}$ est un opérateur hermitien dont le carré est égal à $\mathcal{I}^{\otimes n}$.

Démonstration. On peut vérifier en effet que $\mathcal{I}^{\otimes n}$ est l'élément neutre pour la composition et que les opérations de composition dans $\Phi \times \mathcal{P}^{\otimes n}$ découlent des égalités $\mathcal{I}^2 = \mathcal{X}^2 = \mathcal{Y}^2 = \mathcal{Z}^2 = \mathcal{I}$, $\mathcal{X} \circ \mathcal{Y} = -\mathcal{Y} \circ \mathcal{X} = i\mathcal{Z}$, $\mathcal{Y} \circ \mathcal{Z} = -\mathcal{Z} \circ \mathcal{Y} = i\mathcal{X}$, et $\mathcal{Z} \circ \mathcal{X} = -\mathcal{X} \circ \mathcal{Z} = i\mathcal{Y}$. En particulier, tout élément de $\pm \mathcal{P}^{\otimes n}$ est son propre inverse, et tout élément E de $\pm i\mathcal{P}^{\otimes n}$ a pour inverse $-E$. Le fait que tout élément de $\pm \mathcal{P}^{\otimes n}$ est un opérateur hermitien découle du fait que tout élément de $\mathcal{P}$ est un opérateur hermitien. $\square$

Le poids d'une erreur de Pauli correspond simplement au nombre de coordonnées de l'erreur qui sont différentes de $\mathcal{I}$.

La présence d'une éventuelle phase pour toute erreur de Pauli permet de conférer à l'ensemble des erreurs de Pauli une structure de groupe pour la loi de composition. Cependant, pour plus de simplicité, il est possible de construire une structure de groupe dans laquelle il n'est pas tenu compte de l'éventuelle phase multiplicative qui apparaît lors d'une opération de composition. Formellement, cela revient à définir le groupe quotient suivant.

Définition 2.9.2. On appelle **groupe de Pauli effectif** le quotient $\Phi \times \mathcal{P}^{\otimes n} / \Phi \times \mathcal{I}^{\otimes n}$ du groupe de Pauli par le sous-groupe des phases $\Phi \times \mathcal{I}^{\otimes n}$.

Ses éléments sont les **erreurs de Pauli effectives**. Il est noté plus simplement P^n où $P = \{I, X, Y, Z\}$, en faisant correspondre pour chaque erreur de Pauli effective les symboles respectifs I, X, Y, Z aux opérateurs $\mathcal{I}, \mathcal{X}, \mathcal{Y}, \mathcal{Z}$ en chaque coordonnée. Le **poïds** d'une erreur de Pauli effective est le nombre de coordonnées différentes de I .

Comme la phase n'intervient plus dans les erreurs de Pauli effectives, il s'agit d'un groupe abélien.

Cette structure de groupe sera utile à chaque fois que l'on voudra ignorer la phase des erreurs de Pauli afin de simplifier la présentation. Cela sera par exemple le cas lorsque l'on exprimera la distance minimale d'un code stabilisateur en fonction de son groupe stabilisateur. Mais on procèdera surtout à cette simplification après l'introduction des transformations de Clifford ; ces encodages, qui laissent les erreurs de Pauli stables par conjugaison, ont pour propriété que toute erreur de Pauli affectant un état avant encodage équivaut à une erreur de Pauli correspondante qui affecte l'état après encodage. On procèdera alors à la simplification du modèle en se restreignant aux transformations de Clifford qui stabilisent $\mathcal{P}^{\otimes n}$, et ce afin de permettre d'étudier leur action par conjugaison sur le groupe de Pauli effectif.

S'il n'y a pas d'ambiguïté à propos de la valeur de n , on adaptera la notation $\mathcal{X}_i, \mathcal{Y}_i$ ou $\mathcal{Z}_i$ pour désigner respectivement l'erreur de Pauli de longueur n constituée de l'erreur $\mathcal{I}$ sur toutes les positions sauf la i -ème position et de l'erreur $\mathcal{X}, \mathcal{Y}$ ou $\mathcal{Z}$ en i -ème position. Les erreurs de Pauli effectives constituées d'un X, Y ou Z en i -ème position et du symbole I en toutes les autres positions seront notées de même respectivement X_i, Y_i ou Z_i .

Toute erreur de Pauli s'écrit comme le produit d'une phase de Φ et de la composée d'un ensemble d'erreurs parmi les $\mathcal{X}_i$ et les $\mathcal{Z}_i$. En mettant la phase de côté, cette décomposition est représentable par une séquence de bits appelée représentation simplectique.

Définition 2.9.3. La **représentation simplectique** est l'application σ :

$$\begin{aligned} \Phi \times \mathcal{P}^{\otimes n} &\rightarrow \mathbb{F}_2^{2n} \\ E = cE_1 \otimes \dots \otimes E_n &\mapsto (x_1, \dots, x_n, z_1, \dots, z_n) \end{aligned}$$

où pour tout $i \in \llbracket 1; n \rrbracket$, $x_i = 0$ si $E_i \in \{\mathcal{I}, \mathcal{Z}\}$ et $x_i = 1$ si $E_i \in \{\mathcal{X}, \mathcal{Y}\}$, tandis que $z_i = 0$ si $E_i \in \{\mathcal{I}, \mathcal{X}\}$ et $z_i = 1$ si $E_i \in \{\mathcal{Y}, \mathcal{Z}\}$.

La propriété suivante est immédiate.

Propriété 2.9.2. σ est un morphisme de groupes surjectif de $(\Phi \times \mathcal{P}^{\otimes n}, \circ)$ dans $(\mathbb{F}_2^{2n}, +)$, de noyau $\Phi \times \{\mathcal{I}^{\otimes n}\}$

Introduisons également le produit simplectique, qui permet de vérifier si deux erreurs de Pauli commutent ou anti-commutent en examinant le produit simplectique de leurs représentations simplectiques.

Définition 2.9.4. Le **produit symplectique** de deux séquences $\sigma(E) = (x_1, \dots, x_n, z_1, \dots, z_n)$ et $\sigma(E') = (x'_1, \dots, x'_n, z'_1, \dots, z'_n)$ est le bit donné par :

$$\sigma(E) \cdot \sigma(E') = \sum_{i=1}^n x_i z'_i + x'_i z_i$$

Propriété 2.9.3. Toute paire d'erreurs de Pauli $(E, E') \in (\Phi \times \mathcal{P}^{\otimes n})^2$ commute ou anti-commute, et l'on a :

$$E \circ E' = (-1)^{\sigma(E) \cdot \sigma(E')} E' \circ E$$

Démonstration. On peut vérifier en parcourant tous les cas que, pour toute paire $(E, E') \in \mathcal{P}^2$ de représentation symplectique $\sigma(E) = (x, z)$ et $\sigma(E') = (x', z')$, E et E' commutent si $xz' + x'z = 0$ et anticommulent sinon. Alors pour tout $(E, E') \in (\Phi \times \mathcal{P}^{\otimes n})^2$:

$$\begin{aligned} E \circ E' &= cc'(E_1 \circ E'_1) \otimes \dots \otimes (E_n \circ E'_n) \\ &= cc'((-1)^{\sigma(E_1) \cdot \sigma(E'_1)} E'_1 \circ E_1) \otimes \dots \otimes ((-1)^{\sigma(E_n) \cdot \sigma(E'_n)} E'_n \circ E_n) \\ &= cc'(-1)^{\sigma(E) \cdot \sigma(E')} (E'_1 \circ E_1) \otimes \dots \otimes (E'_n \circ E_n) \\ &= (-1)^{\sigma(E) \cdot \sigma(E')} E' \circ E \end{aligned}$$

□

L'opération de composition dans le groupe de Pauli effectif est commutative. On peut cependant être amené à vérifier si deux erreurs de Pauli, représentées par deux erreurs de Pauli effectives, commutent ou anticommulent. En extrapolant la définition de la représentation symplectique au groupe de Pauli effectif, on peut, grâce au produit symplectique, retrouver une telle information.

Propriété 2.9.4. La représentation symplectique σ induit un isomorphisme de P^n sur $\mathbb{F}_2^{2n}$, que l'on appellera également représentation symplectique par abus de langage.

Démonstration. Il s'agit d'un isomorphisme car P^n est obtenu en quotientant $\Phi \times \mathcal{P}^{\otimes n}$ par $\Phi \times \{\mathcal{I}^{\otimes n}\}$, qui est le noyau de la représentation symplectique du groupe de Pauli. □

La propriété qui suit concerne les espaces propres des erreurs de Pauli et sera utile pour l'analyse des codes stabilisateurs.

Propriété 2.9.5. Pour tout élément E de $\pm \mathcal{P}^{\otimes n}$ différent de $\pm \mathcal{I}^{\otimes n}$, les espaces propres $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ respectivement associés aux valeurs propres 1 et -1 sont de dimension 2^{n-1} et vérifient :

$$\mathcal{H}^{(+)} \bigotimes^{\perp} \mathcal{H}^{(-)} = \mathcal{H}_2^{\otimes n}$$

De plus, pour toute erreur de Pauli E' qui anticommute avec E , on a $E'(\mathcal{H}^{(+)}) = \mathcal{H}^{(-)}$ et $E'(\mathcal{H}^{(-)}) = \mathcal{H}^{(+)}$.

Démonstration. E est hermitien donc diagonalisable dans une base orthonormée. Comme $E^2 = \mathcal{I}$, ses valeurs propres ont pour carré 1 et ne peuvent valoir que 1 ou -1 . On sait donc que $\mathcal{H}^{(+)} \otimes^\perp \mathcal{H}^{(-)} = \mathcal{H}_2^{\otimes n}$.

Montrons d'abord que l'ensemble des erreurs de Pauli qui anticommulent avec E n'est pas vide. Comme $E \notin \Phi \times \mathcal{I}^{\otimes n}$, il existe une coordonnée E_i de E différente de $\mathcal{I}$. Soit alors $E'_i \in \mathcal{P}$ tel que E_i anticommute avec E'_i , et soit $E' \in \mathcal{P}^{\otimes n}$ dont toutes les coordonnées valent $\mathcal{I}$ sauf la i -ème coordonnée qui vaut E'_i . E'_i anticommute donc avec E .

Soit alors E' une erreur de Pauli qui anticommute avec E . Pour tout $|\psi\rangle \in \mathcal{H}^{(+)}$:

$$E \circ E' |\psi\rangle = -E' \circ E |\psi\rangle = -E' |\psi\rangle$$

Cela montre que $E' |\psi\rangle \in \mathcal{H}^{(-)}$, donc que $E'(\mathcal{H}^{(+)}) \subset \mathcal{H}^{(-)}$. De même, pour tout $|\psi\rangle \in \mathcal{H}^{(-)}$:

$$E \circ E' |\psi\rangle = -E' \circ E |\psi\rangle = E' |\psi\rangle$$

donc $E'(\mathcal{H}^{(-)}) \subset \mathcal{H}^{(+)}$. Comme E' est un isomorphisme, on en déduit que $E'(\mathcal{H}^{(+)}) = \mathcal{H}^{(-)}$ et $E'(\mathcal{H}^{(-)}) = \mathcal{H}^{(+)}$. Cela montre également que $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ sont de même dimension 2^{n-1} . $\square$

Le poids d'une erreur de Pauli est le nombre de coordonnées E_i différentes de l'identité. La propriété suivante est immédiate et résulte du fait que les quatre erreurs de $\mathcal{P}$ forment une base de $\mathcal{L}(\mathcal{H}_2^{\otimes n})$. Il en découle un corollaire intéressant qui facilite le travail de vérification de la distance minimale d'un code quantique.

Propriété 2.9.6. $\mathcal{P}^{\otimes n}$ est une base de $\mathcal{L}(\mathcal{H}_2^{\otimes n})$. De même, pour tout entier $d \leq n$, l'ensemble des erreurs de $\mathcal{P}^{\otimes n}$ de poids inférieur ou égal à d est une base de l'espace des erreurs quantiques de poids inférieur ou égal à d sur $\mathcal{L}(\mathcal{H}_2^{\otimes n})$.

Corollaire 2.9.1. La distance minimale d'un code quantique $\mathcal{C}$ est le plus petit poids d d'une erreur de Pauli $E \in \mathcal{P}^{\otimes n}$ telle qu'il existe $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$ et $\langle \psi' | E | \psi \rangle \neq 0$.

Démonstration. Soit d la distance minimale de $\mathcal{C}$. Alors d'une part, pour toute erreur quantique E de poids strictement inférieur à d , et pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$, on a $\langle \psi' | E | \psi \rangle = 0$. Cela est donc vrai en l'occurrence pour toute erreur de $\mathcal{P}^{\otimes n}$ de poids strictement inférieur à d .

Par l'absurde, supposons à présent que pour toute erreur $E \in \mathcal{P}^{\otimes n}$ de poids d et pour tout $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$, on a $\langle \psi' | E | \psi \rangle = 0$. Cette propriété est alors vérifiée pour toute erreur $E \in \mathcal{P}^{\otimes n}$ de poids inférieur ou égal à d , donc par linéarité, pour toute erreur quantique de poids d . Cela contredit le fait que d est la distance minimale de $\mathcal{C}$. Il existe donc bien une erreur $E \in \mathcal{P}^{\otimes n}$ de poids d et une paire $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$ et $\langle \psi' | E | \psi \rangle \neq 0$. $\square$

2.10 Codes stabilisateurs

Les codes stabilisateurs introduits par Gottesman [19] présentent des propriétés particulières vis-à-vis des erreurs de Pauli. Ils sont définis comme des espaces

propres associés à un sous-ensemble commutatif d'erreurs de Pauli appelé groupe stabilisateur. Le formalisme sous-jacent aux codes stabilisateurs est très puissant car il permet de concevoir toute une classe de codes quantiques qui, par construction, ont une distance minimale déductible directement du groupe stabilisateur, et pour lesquels il existe une procédure de décodage par syndrome très proche du décodage par syndrome classique.

Cette section expose des résultats essentiels sur les codes stabilisateurs, et est organisée de la manière suivante. On définit d'abord un groupe stabilisateur et un code stabilisateur. On introduit également la définition du commutateur d'un groupe stabilisateur, qui permet de déduire des propriétés sur la taille du groupe stabilisateur, puis d'exprimer la dimension du code stabilisateur. On montre finalement comment s'exprime la distance minimale d'un code stabilisateur en fonction du groupe stabilisateur et de son commutateur.

Définition 2.10.1. Un **groupe stabilisateur $\mathcal{S}$ de longueur n** est un sous-groupe abélien du groupe de Pauli de taille n qui ne contient pas $-\mathcal{I}^{\otimes n}$. Un **code stabilisateur** de groupe stabilisateur $\mathcal{S}$ est la sphère unité $\mathcal{C} = \mathcal{B}(\mathcal{H})$ du sous-espace vectoriel $\mathcal{H}$ de $\mathcal{H}_2^{\otimes n}$ invariant par tous les éléments de $\mathcal{S}$. En d'autres termes :

$$\mathcal{C} = \{|\psi\rangle \in \mathcal{B}(\mathcal{H}_2^{\otimes n}) : \forall E \in \mathcal{S}, E|\psi\rangle = |\psi\rangle\}$$

Afin qu'un code stabilisateur corresponde à la définition 2.5.1 d'un code quantique, c'est-à-dire l'image d'un encodage quantique, il faut encore montrer que la dimension du sous-espace vectoriel $\mathcal{H}$ dont le code stabilisateur est la sphère unité est égale à une puissance de 2. On prouve ce résultat par la suite.

Propriété 2.10.1. *Tout groupe stabilisateur $\mathcal{S}$ est généré par une famille minimale appelée **base**, et $\mathcal{S}$ est de cardinal 2^t où t est le cardinal de toute base de $\mathcal{S}$.*

Démonstration. La restriction de σ à $\mathcal{S}$ est un isomorphisme. En effet, $\mathcal{S} \cap \text{Ker } \sigma = \{\mathcal{I}^{\otimes n}\}$, car d'une part, par définition $-\mathcal{I}^{\otimes n} \notin \mathcal{S}$, et d'autre part, si $\pm i\mathcal{I}^{\otimes n} \in \mathcal{S}$ alors par composition on aurait $\mathcal{I}^{\otimes n} \in \mathcal{S}$. De plus, $\mathcal{S}$ étant un groupe, $\sigma(\mathcal{S})$ est un sous-groupe de $\mathbb{F}_2^{2n}$ donc un sous-espace vectoriel de $\mathbb{F}_2^{2n}$. $\sigma(\mathcal{S})$ possède donc une base d'un certain cardinal t , et $\sigma(\mathcal{S})$ est alors de cardinal 2^t . Par isomorphisme, $\mathcal{S}$ est donc également engendré par une famille minimale de cardinal t , et a pour cardinal 2^t . $\square$

Définition 2.10.2. Soit $\mathcal{S}$ un groupe stabilisateur. Le **commutateur** de $\mathcal{S}$ noté $N(\mathcal{S})$ est le sous-groupe des erreurs de Pauli qui commutent avec tous les éléments de $\mathcal{S}$.

Propriété 2.10.2. $\mathcal{S} \subset N(\mathcal{S})$. De plus, le cardinal d'une base de $\mathcal{S}$ s'écrit sous la forme $n - k$ où $k \geq 0$, et $N(\mathcal{S})$ est engendré par le sous-groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$ et une famille minimale de cardinal $n + k$.

Démonstration. Comme tout élément $\mathcal{S}$ commute avec tous les éléments de $\mathcal{S}$, on a $\mathcal{S} \subset N(\mathcal{S})$. Soit $(\bar{Z}_1, \dots, \bar{Z}_t)$ une base de $\mathcal{S}$. Par définition, $N(\mathcal{S})$ est l'ensemble

des erreurs de Pauli qui commutent avec cette base, soit l'image réciproque par σ du noyau de l'application linéaire :

$$\begin{aligned} \mathbb{F}_2^{2n} & \rightarrow \mathbb{F}_2^t \\ (x_1, \dots, x_n, z_1, \dots, z_n) & \mapsto ((x_1, \dots, x_n, z_1, \dots, z_n) \cdot \sigma(\bar{Z}_i))_{1 \leq i \leq t} \end{aligned}$$

Comme $(\sigma(\bar{Z}_1), \dots, \sigma(\bar{Z}_t))$ forme une famille libre dans $\mathbb{F}_2^{2n}$, le rang de cette application est égal à t , et d'après le théorème du rang, son noyau $\sigma(N(\mathcal{S}))$ est donc un espace vectoriel de dimension $2n - t$. Comme $\sigma(\mathcal{S}) \subset \sigma(N(\mathcal{S}))$, on a par conséquent $t \leq 2n - t$, et donc $t = n - k$ où $k \geq 0$. De plus, $N(\mathcal{S})$ est l'image réciproque par σ de cet espace, donc il est généré par $\text{Ker } \sigma = \Phi \times \{\mathcal{I}^{\otimes n}\}$ et un antécédent par σ d'une base de $\sigma(N(\mathcal{S}))$, qui est une famille minimale de cardinal $2n - t = n + k$. $\square$

Propriété 2.10.3. *Soit $\mathcal{S}$ un groupe stabilisateur de base $(\bar{Z}_1, \dots, \bar{Z}_{n-k})$. L'espace laissé invariant par $\mathcal{S}$:*

$$\mathcal{H}^{(n-k)} = \{|\psi\rangle \in \mathcal{H}_2^{\otimes n} : \forall E \in \mathcal{S}, E|\psi\rangle = |\psi\rangle\}$$

est de dimension 2^k .

Démonstration. On a :

$$\mathcal{H}^{(n-k)} = \{|\psi\rangle \in \mathcal{H}_2^{\otimes n} : \forall i \in \llbracket 1; n-k \rrbracket, \bar{Z}_i|\psi\rangle = |\psi\rangle\}$$

Pour tout $t \in \llbracket 0; n-k \rrbracket$, soit $\mathcal{H}^t$ l'espace :

$$\mathcal{H}^{(t)} = \{|\psi\rangle \in \mathcal{H}_2^{\otimes n} : \forall i \in \llbracket 1; t \rrbracket, \bar{Z}_i|\psi\rangle = |\psi\rangle\}$$

Montrons par récurrence sur t que $\dim \mathcal{H}^{(t)} = 2^{n-t}$. Pour $t = 0$, la propriété revient à affirmer que $\dim \mathcal{H}_2^{\otimes n} = 2^n$ et est donc vraie. Supposons à présent que la propriété est vraie pour un certain $t \in \llbracket 0; n-k-1 \rrbracket$. Par définition :

$$\mathcal{H}^{(t+1)} = \mathcal{H}^{(t)} \cap \mathcal{H}^{(t+1,+)}$$

où $\mathcal{H}^{(t+1,+)}$ est l'espace propre de $\bar{Z}_{t+1}$ associé à la valeur propre 1. Montrons l'existence d'une erreur de Pauli E qui anticommute avec $\bar{Z}_{t+1}$ et qui commute avec tous les opérateurs de la famille $(\bar{Z}_1, \dots, \bar{Z}_t)$. La condition à satisfaire par E correspond à :

$$(\sigma(E) \cdot \sigma(\bar{Z}_i))_{1 \leq i \leq t+1} = (0, \dots, 0, 1)$$

Or, la famille $(\sigma(\bar{Z}_i))_{1 \leq i \leq t+1}$ étant libre, l'application :

$$(x_1, \dots, x_n, z_1, \dots, z_n) \mapsto ((x_1, \dots, x_n, z_1, \dots, z_n) \cdot \sigma(\bar{Z}_i))_{1 \leq i \leq t+1}$$

est surjective, ce qui montre qu'une telle erreur de Pauli existe. Comme E anticommute avec $\bar{Z}_{t+1}$, on a $E(\mathcal{H}^{(t+1,+)}) = \mathcal{H}^{(t+1,-)}$. De plus, comme E commute avec les opérateurs de la famille $(\bar{Z}_1, \dots, \bar{Z}_t)$, on a $E(\mathcal{H}^{(t)}) = \mathcal{H}^{(t)}$. Par conséquent, $E(\mathcal{H}^{(t)} \cap \mathcal{H}^{(t+1,+)}) = \mathcal{H}^{(t)} \cap \mathcal{H}^{(t+1,-)}$, ce qui montre que ces deux

espaces sont de même dimension. Etant donné que $\mathcal{H}^{(t)} \cap \mathcal{H}^{(t+1,+)} \oplus \mathcal{H}^{(t)} \cap \mathcal{H}^{(t+1,-)} = \mathcal{H}^{(t)}$, on en déduit que :

$$\dim(\mathcal{H}^{(t+1)}) = \frac{1}{2} \dim(\mathcal{H}^{(t)}) = 2^{n-t-1}$$

ce qui prouve la propriété au rang $t + 1$, et donc en l'occurrence au rang $t = n - k$. $\square$

Corollaire 2.10.1. *Un code stabilisateur $\mathcal{C}$ de groupe stabilisateur $\mathcal{S}$, où $\mathcal{S}$ a une base de taille $n - k$, est un code quantique qui permet d'encoder l'état d'un registre à k bits.*

On introduit désormais les images respectives S et $N(S)$ du groupe stabilisateur $\mathcal{S}$ et de son commutateur $N(\mathcal{S})$ par la surjection canonique du groupe de Pauli sur le groupe de Pauli effectif. Cette définition sera réutilisée lors de la présentation du décodage d'un code stabilisateur et dans la rubrique présentant la notion d'*encodeur*.

Définition 2.10.3. On note S et $N(S)$ les représentants des ensembles $\mathcal{S}$ et $N(\mathcal{S})$ dans P^n .

La distance minimale d'un code stabilisateur se formule de manière simple en fonction de ces deux ensembles d'erreurs de Pauli effectives.

Propriété 2.10.4. *La distance minimale d'un code stabilisateur $\mathcal{C}$ est le plus petit poids d d'un élément de $N(S) \setminus S$.*

Démonstration. Identifions S et $N(S)$ aux ensembles respectifs de leurs représentants dans le groupe de Pauli $\mathcal{P}^{\otimes n}$. On cherche le plus petit poids d d'une erreur $E \in \mathcal{P}^{\otimes n}$ telle qu'il existe $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$ et $\langle \psi' | E | \psi \rangle \neq 0$. On va d'abord montrer qu'aucune erreur de Pauli dans $S \cup \mathcal{P}^{\otimes n} \setminus N(S)$ ne vérifie cette condition, et qu'à l'opposé, toutes les erreurs de $N(S) \setminus S$ la vérifient. Cela prouvera que d est le plus petit poids d'une erreur de $N(S) \setminus S$.

Soit $E \in S$. Il existe donc $c \in \Phi$ tel que $cE \in \mathcal{S}$. Soit $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$. On a alors $E|\psi\rangle = c^{-1}|\psi\rangle$ donc $\langle \psi' | E | \psi \rangle = 0$.

De même, soit $E \in \mathcal{P}^{\otimes n} \setminus N(S)$. Soit $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$ vérifiant $\langle \psi' | \psi \rangle = 0$. Il existe $E' \in S$ tel que E anticommute avec E' . Par conséquent, si $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ sont les espaces propres orthogonaux de E' associés respectivement aux valeurs propres 1 et -1 , on a $E(\mathcal{H}^{(+)}) = \mathcal{H}^{(-)}$ de sorte que $E|\psi\rangle \in \mathcal{H}^{(-)}$, et comme par ailleurs $|\psi'\rangle \in \mathcal{H}^{(+)}$, cela prouve que $\langle \psi' | E | \psi \rangle = 0$.

Soit à présent $E \in N(S) \setminus S$. Comme E commute avec tous les éléments de $\mathcal{S}$, il stabilise leurs espaces propres associés à la valeur propre 1, et donc E stabilise $\text{Vect } \mathcal{C}$. Notons $I_{\mathcal{C}}$ l'opérateur identité sur $\text{Vect } \mathcal{C}$, et $E_{\mathcal{C}}$ la restriction de E à $\text{Vect } \mathcal{C}$. $E_{\mathcal{C}}^2 = I_{\mathcal{C}}$ et $E_{\mathcal{C}}$ est hermitien donc $E_{\mathcal{C}}$ possède deux espaces propres orthogonaux $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ associés aux valeurs propres respectives 1 et -1 . Comme $E \notin \pm \mathcal{S}$, $E_{\mathcal{C}}$ n'agit ni comme l'identité $I_{\mathcal{C}}$ ni comme son opposé $-I_{\mathcal{C}}$. Par conséquent, aucun des sous-espaces $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ n'est réduit à $\{0\}$, et $\mathcal{H}^{(+)}$ et $\mathcal{H}^{(-)}$ contiennent respectivement un vecteur unitaire $|\psi^{(+)}\rangle$ et un vecteur

unitaire $|\psi^{(-)}\rangle$. Posons alors $|\psi\rangle = \frac{1}{\sqrt{2}}(|\psi^{(+)}\rangle + |\psi^{(-)}\rangle)$ et $|\psi'\rangle = \frac{1}{\sqrt{2}}(|\psi^{(+)}\rangle - |\psi^{(-)}\rangle)$. $|\psi\rangle$ et $|\psi'\rangle$ étant des vecteurs unitaires, on a $(|\psi\rangle, |\psi'\rangle) \in \mathcal{C}^2$. De plus, $\langle\psi'|\psi\rangle = 0$. Cependant, $E|\psi\rangle = |\psi'\rangle$ ce qui prouve que $\langle\psi'|E|\psi\rangle = 1$. $\square$

Terminons cette section de présentation des codes stabilisateurs par deux remarques importantes. Par définition, un groupe stabilisateur $\mathcal{S}$ agit trivialement sur son code stabilisateur. En particulier, toute erreur de $\mathcal{S} \setminus \{\mathcal{I}^{\otimes n}\}$ est une erreur non triviale qui agit pourtant trivialement sur le code. Ce phénomène n'a pas d'équivalent en théorie des codes classiques, où une erreur non triviale affecte nécessairement le code. La prise en compte de cet ensemble d'erreurs ayant un effet trivial sur le code stabilisateur introduit une flexibilité au niveau du décodage : si une erreur E de Pauli affecte le code, la véritable question du décodage consiste non pas à trouver l'erreur $\hat{E}$ la plus probable mais à trouver l'erreur $\hat{E}$ dont la classe d'équivalence modulo $\mathcal{S}$ est la plus probable. Cette remarque sera reprise à la définition 2.12.2 du décodage au maximum de classe de vraisemblance.

La seconde remarque est que pour certains codes stabilisateurs, il existe un élément de $\mathcal{S} \setminus \{\mathcal{I}^{\otimes n}\}$ dont le poids est strictement inférieur à la distance minimale du code. Un tel code est dit *dégénéré*.

Définition 2.10.4. Si le poids minimal d'un élément de $\mathcal{S} \setminus \{I^n\}$ est strictement inférieur à la distance minimale du code stabilisateur, ce dernier est dit *dégénéré*.

Dans cette thèse, on montre que la distance minimale d'un turbo-code quantique n'est pas fonction de la distance minimale du code externe comme dans le résultat classique, mais plutôt du plus petit poids d'un élément de $N(\mathcal{S}) \setminus \{I^n\}$. Si le code externe est non dégénéré, cette grandeur correspond effectivement à sa distance minimale, mais s'il est dégénéré, elle correspond au plus petit poids d'un élément de $\mathcal{S} \setminus \{I^n\}$. On nommera cette grandeur *distance dégénérée* au chapitre suivant.

2.11 Transformations de Clifford

Dans cette section, on présente un type de transformations permettant d'encoder les codes stabilisateurs et descriptibles par des paramètres discrets. Il s'agit des transformations de Clifford, des évolutions unitaires laissant le groupe de Pauli stable par conjugaison. Sous ce modèle, l'étude des erreurs qui affectent un code stabilisateur se trouve grandement simplifiée. En effet, l'effet d'une erreur de Pauli sur l'état d'un code stabilisateur équivaut à celui d'une erreur de Pauli correspondante qui affecte l'état avant encodage. La correspondance entre les erreurs de Pauli équivalentes avant et après encodage est donnée par l'action par conjugaison de la transformation de Clifford sur le groupe de Pauli. Cette équivalence permet de faire abstraction de l'encodage lui-même, qui est un opérateur sur un espace continu, pour se concentrer sur l'effet discret de l'encodage sur les erreurs de Pauli. Il apparaît ainsi un encodage virtuel qui se

rapporte uniquement aux erreurs de Pauli et que l'on introduira ensuite. Comme on le présentera dans la section suivante, il existe de plus une procédure de décodage par syndrome, où le syndrome permet d'acquérir de l'information sur l'erreur de Pauli ayant affecté l'état du code stabilisateur. Cela permet ainsi de se représenter le problème par un modèle très proche de la théorie des codes classiques : un encodage défini sur les séquences d'un alphabet fini, et pour lequel existe une procédure de décodage par syndrome qui livre de l'information sur cette séquence. Une telle représentation définie purement sur la séquence des erreurs de Pauli, et faisant abstraction de l'espace continu des états sous-jacent, fera l'objet d'une présentation spécifique qui permettra d'appréhender directement, par un lecteur familier des notions de théorie des codes classiques, la problématique du décodage des codes stabilisateurs et de leur distance minimale.

Définition 2.11.1. Soit U une application unitaire sur $\mathcal{H}_2^{\otimes n}$. U est une **transformation de Clifford** si, pour tout $E \in \Phi \times \mathcal{P}^{\otimes n}$, $UEU^\dagger \in \Phi \times \mathcal{P}^{\otimes n}$.

Il est possible de représenter une transformation de Clifford via son action par conjugaison sur le groupe de Pauli. Une telle représentation permet de faire abstraction de l'espace quantique et de se concentrer sur le groupe des erreurs de Pauli. On montre dans la propriété suivante les conditions que doit vérifier toute action sur le groupe de Pauli afin qu'elle corresponde à l'action par conjugaison d'une transformation de Clifford. La propriété affirme ainsi qu'une telle action est un isomorphisme du groupe de Pauli, qui conserve la commutativité, et qui laisse le groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$ invariant.

Propriété 2.11.1. Soit $\mathcal{E}$ une application du groupe de Pauli $\Phi \times \mathcal{P}^{\otimes n}$ dans lui-même. Les deux conditions suivantes sont équivalentes :

- (i) Il existe une transformation de Clifford U , unique à une phase près, telle que pour tout $E \in \Phi \times \mathcal{P}^{\otimes n}$:

$$UEU^\dagger = \mathcal{E}(E)$$

- (ii) $\mathcal{E}$ est un isomorphisme qui conserve la commutativité, c'est-à-dire que pour tout $(E, E') \in (\Phi \times \mathcal{P}^{\otimes n})^2$:

$$\sigma(\mathcal{E}(E)).\sigma(\mathcal{E}(E')) = \sigma(E).\sigma(E')$$

et qui laisse le groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$ invariant.

Démonstration. Supposons qu'il existe une transformation de Clifford U vérifiant la condition en (i). Soit alors $(E, E') \in (\Phi \times \mathcal{P}^{\otimes n})^2$. On a :

$$\mathcal{E}(E) \circ \mathcal{E}(E') = (UEU^\dagger) \circ (UE'U^\dagger) = U(E \circ E')U^\dagger = \mathcal{E}(E \circ E')$$

Donc $\mathcal{E}$ est un morphisme. De plus, l'action $E \mapsto UEU^\dagger$ est inversible d'inverse $E \mapsto U^\dagger EU$, ce qui montre que $\mathcal{E}$ est un isomorphisme. Par ailleurs, si $E \circ E' = E' \circ E$:

$$\mathcal{E}(E) \circ \mathcal{E}(E') = U(E \circ E')U^\dagger = U(E' \circ E)U^\dagger = \mathcal{E}(E') \circ \mathcal{E}(E)$$

Et si à l'inverse, $E \circ E' = -E' \circ E$, un calcul similaire montre que $\mathcal{E}(E) \circ \mathcal{E}(E') = -\mathcal{E}(E') \circ \mathcal{E}(E)$. Ainsi $\mathcal{E}$ conserve la commutativité. Finalement, comme l'action $E \mapsto UEU^\dagger$ agit comme l'identité sur le groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$, $\mathcal{E}$ laisse également le groupe des phases invariant.

Montrons à présent la réciproque. Supposons que $\mathcal{E}$ respecte la condition (ii). Montrons d'abord l'unicité à une phase près de U , puis son existence. Si U et U' vérifient $UEU^\dagger = \mathcal{E}(E)$ et $U'EU'^\dagger = \mathcal{E}(E)$ pour toute erreur de Pauli E , on pose $V = U'^\dagger U$. Alors V est unitaire, et pour toute erreur de Pauli E :

$$VEV^{-1} = VEV^\dagger = U'^\dagger UEU^\dagger U' = U'^\dagger \mathcal{E}(E) U' = E$$

Par linéarité, on en conclut que $VEV^{-1} = E$ et donc $VE = EV$ pour tout $E \in \mathcal{L}(\mathcal{H}_2^{\otimes n})$. Ainsi V commute avec tous les endomorphismes de $\mathcal{L}(\mathcal{H}_2^{\otimes n})$, et V est donc un multiple de l'identité. Comme V est unitaire, on peut écrire $V = c\mathcal{I}^{\otimes n}$ où $|c| = 1$, et donc $U' = cU$.

Montrons maintenant l'existence de U . Pour tout $i \in \llbracket 1; n \rrbracket$, soient $\bar{\mathcal{X}}_i = \mathcal{E}(\mathcal{X}_i)$ et $\bar{\mathcal{Z}}_i = \mathcal{E}(\mathcal{Z}_i)$. On va construire une transformation de Clifford U vérifiant la condition (i) à partir de ces erreurs de Pauli. Le schéma général de la construction est le suivant. On montre d'abord que pour toute séquence de n valeurs propres dans $\{-1, 1\}$, les espaces propres respectifs de la famille d'opérateurs $(\mathcal{Z}_1, \dots, \mathcal{Z}_n)$ et $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n)$ associés à cette séquence de valeurs propres sont des droites vectorielles. Ensuite, on définit U comme l'application linéaire envoyant, pour toute séquence de valeurs propres, un état de la première droite sur un état convenablement choisi de la seconde droite.

Remarquons au préalable que pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$, comme $\mathcal{E}$ respecte la commutativité, $\bar{\mathcal{X}}_i$ et $\bar{\mathcal{Z}}_j$ anticommulent si $i = j$ et commutent si $i \neq j$.

Soit $b = (b_1, \dots, b_n) \in \{0, 1\}^n$. La famille $((-1)^{b_1} \mathcal{Z}_1, \dots, (-1)^{b_n} \mathcal{Z}_n)$ est libre, commutative, et n'engendre pas l'erreur $-\mathcal{I}^{\otimes n}$. $\mathcal{E}$ étant un isomorphisme qui conserve la commutativité et qui laisse le groupe des phases invariant, ces mêmes propriétés se retrouvent chez la famille $((-1)^{b_1} \bar{\mathcal{Z}}_1, \dots, (-1)^{b_n} \bar{\mathcal{Z}}_n)$. Par conséquent, ces deux familles génèrent chacune un groupe stabilisateur. D'après la propriété 2.10.3 appliquée au cas $k = 0$, l'espace $\mathcal{H}^{(b)}$ laissé invariant par la famille $((-1)^{b_1} \mathcal{Z}_1, \dots, (-1)^{b_n} \mathcal{Z}_n)$ et l'espace $\mathcal{H}^{(\bar{b})}$ laissé invariant par la famille $((-1)^{b_1} \bar{\mathcal{Z}}_1, \dots, (-1)^{b_n} \bar{\mathcal{Z}}_n)$ sont donc de dimension 1. $\mathcal{H}^{(b)}$ est engendré par l'état $|b\rangle$ tandis que $\mathcal{H}^{(\bar{b})}$ est engendré par un état $|\bar{b}\rangle$ que l'on construit grâce au lemme suivant.

Lemme 2.11.1. *Pour tout $b \in \{0, 1\}^n$, on pose :*

$$|\bar{b}\rangle = \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} |\bar{0}\rangle$$

où $\bar{\mathcal{X}}_i^{b_i}$ vaut $\bar{\mathcal{X}}_i$ si $b_i = 1$ et l'identité si $b_i = 0$. Alors pour tout $b \in \{0, 1\}^n$, $|\bar{b}\rangle$ est un état qui engendre l'espace $\mathcal{H}^{(\bar{b})}$. De plus, la famille $(|\bar{b}\rangle)_{b \in \{0, 1\}^n}$ est une base orthonormée de $\mathcal{H}_2^{\otimes n}$.

La preuve de ce lemme fait suite à cette preuve. Soit alors U l'application linéaire définie par $U|b\rangle = |\bar{b}\rangle$ pour tout $b \in \{0, 1\}^n$. Comme $(|\bar{b}\rangle)_{b \in \{0, 1\}^n}$ est une base orthonormée de $\mathcal{H}_2^{\otimes n}$, U est unitaire.

Montrons à présent que toute erreur de Pauli $E \in \Phi \times \mathcal{P}^{\otimes n}$ vérifie l'égalité :

$$UEU^\dagger = \mathcal{E}(E)$$

ce qui garantira également que U est une transformation de Clifford. Il suffit de prouver cette propriété pour tout $E \in \{\pm \mathcal{I}, \pm i\mathcal{I}, \mathcal{X}_1, \dots, \mathcal{X}_n, \mathcal{Z}_1, \dots, \mathcal{Z}_n\}$. En effet, toute erreur de Pauli E se décompose comme la composée d'éléments de cet ensemble, et la propriété $UEU^\dagger = \mathcal{E}(E)$ est stable par composition puisque pour tout $(E, E') \in (\Phi \times \mathcal{P}^{\otimes n})^2$ tel que $UEU^\dagger = \mathcal{E}(E)$ et $UE'U^\dagger = \mathcal{E}(E')$, on a :

$$U(E \circ E')U^\dagger = (UEU^\dagger) \circ (UE'U^\dagger) = \mathcal{E}(E) \circ \mathcal{E}(E') = \mathcal{E}(E \circ E')$$

L'égalité $UEU^\dagger = \mathcal{E}(E)$ est vérifiée par hypothèse pour tout $E \in \{\pm \mathcal{I}, \pm i\mathcal{I}\}$. Soit à présent $i \in \llbracket 1; n \rrbracket$. Pour tout $b \in \{0, 1\}^n$, on a :

$$U\mathcal{Z}_i U^\dagger |\bar{b}\rangle = U\mathcal{Z}_i |b\rangle = U(-1)^{b_i} |b\rangle = (-1)^{b_i} |\bar{b}\rangle = \bar{\mathcal{Z}}_i |\bar{b}\rangle$$

Ainsi, $U\mathcal{Z}_i U^\dagger = \bar{\mathcal{Z}}_i = \mathcal{E}(\mathcal{Z}_i)$, et l'égalité $UEU^\dagger = \mathcal{E}(E)$ est vérifiée pour tout $E \in \{\mathcal{Z}_1, \dots, \mathcal{Z}_n\}$. De même, si l'on note $b^{(i)}$ la séquence binaire qui concorde avec b en toutes les positions sauf la i -ème position, on a $|b^{(i)}\rangle = \mathcal{X}_i |b\rangle$, et :

$$U\mathcal{X}_i U^\dagger |\bar{b}\rangle = U\mathcal{X}_i |b\rangle = U|b^{(i)}\rangle = |\bar{b}\rangle = \bar{\mathcal{X}}_i |\bar{b}\rangle$$

la dernière égalité provenant du fait que :

$$|\bar{b}\rangle = \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} |\bar{0}\rangle$$

et :

$$|b^{(i)}\rangle = \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ (\bar{\mathcal{X}}_i^{b_i} \circ \bar{\mathcal{X}}_i) \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} |\bar{0}\rangle$$

Par conséquent l'égalité $UEU^\dagger = \mathcal{E}(E)$ est également vérifiée pour tout $E \in \{\mathcal{X}_1, \dots, \mathcal{X}_n\}$, et donc pour tout $E \in \Phi \times \mathcal{P}^{\otimes n}$. $\square$

Preuve du lemme 2.11.1. Montrons d'abord que le vecteur $|\bar{b}\rangle$ ainsi défini est un état qui engendre l'espace $\mathcal{H}^{(\bar{b})}$. D'abord, $|\bar{b}\rangle$ est bien un vecteur unitaire donc un état, puisqu'il est l'image de l'état $|b\rangle$ par une composée d'erreurs de Pauli qui sont des applications unitaires. De plus, pour tout $i \in \llbracket 1; n \rrbracket$, comme $\bar{\mathcal{Z}}_i$ anticommute avec $\bar{\mathcal{X}}_i$ et commute avec tous les autres $\bar{\mathcal{X}}_j$ où $j \neq i$:

$$\begin{aligned} (-1)^{b_i} \bar{\mathcal{Z}}_i |\bar{b}\rangle &= (-1)^{b_i} \bar{\mathcal{Z}}_i \circ \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} |\bar{0}\rangle \\ &= \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} \circ \bar{\mathcal{Z}}_i |\bar{0}\rangle \\ &= \bar{\mathcal{X}}_1^{b_1} \circ \dots \circ \bar{\mathcal{X}}_n^{b_n} \circ |\bar{0}\rangle \\ &= |\bar{b}\rangle \end{aligned}$$

Cela montre bien que $|\bar{b}\rangle \in \mathcal{H}^{(\bar{b})}$ et que $|\bar{b}\rangle$ est un état qui engendre $\mathcal{H}^{(\bar{b})}$. De plus, soit $(b, b') \in \{0, 1\}^n$ avec $b \neq b'$. Soit $i \in \llbracket 1; n \rrbracket$ tel que $b_i \neq b'_i$. Notons $\mathcal{H}^{(0)}$ et $\mathcal{H}^{(1)}$ les espaces propres de l'opérateur $\bar{\mathcal{Z}}_i$ associés respectivement aux valeurs propres 1 et -1 . Ces deux espaces sont orthogonaux, et comme $|\bar{b}\rangle \in \mathcal{H}^{(b_i)}$ tandis que $|\bar{b}'\rangle \in \mathcal{H}^{(b'_i)}$, il vient que $|\bar{b}\rangle$ et $|\bar{b}'\rangle$ sont orthogonaux. Cela prouve que la famille $(|\bar{b}\rangle)_{b \in \{0, 1\}^n}$ est une base orthonormée de $\mathcal{H}_2^{\otimes n}$. $\square$

Remarque 1. La conservation de la commutativité par $\mathcal{E}$ n'est pas une condition suffisante pour que $\mathcal{E}$ corresponde à l'action par conjugaison d'une transformation de Clifford. Il est nécessaire de mentionner l'invariance du groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$ par $\mathcal{E}$. En témoigne l'exemple suivant, pour lequel la commutativité est conservée mais le groupe des phases n'est pas laissé invariant. Soit $\mathcal{E}$ l'isomorphisme du groupe de Pauli qui, à toute erreur $E = cE_1 \otimes \dots \otimes E_n$, associe $\mathcal{E}(E) = \bar{c}E_1 \otimes \dots \otimes E_n$. $\mathcal{E}$ n'agit que sur la phase des erreurs de Pauli et conserve donc la commutativité; cependant, il n'existe aucune application unitaire U telle que, par exemple, $Ui\mathcal{I}^{\otimes n}U^\dagger = \mathcal{E}(i\mathcal{I}^{\otimes n})$ puisque $Ui\mathcal{I}^{\otimes n}U^\dagger = i\mathcal{I}^{\otimes n}$ et $\mathcal{E}(i\mathcal{I}^{\otimes n}) = -i\mathcal{I}^{\otimes n}$.

On aura besoin, pour construire un turbo-encodage quantique, du corollaire suivant, vrai car l'action par conjugaison requise conserve la commutativité.

Corollaire 2.11.1. *Pour toute permutation s de l'ensemble $\{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}$, il existe une transformation de Clifford U_s telle que pour tout $E \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}$, $U_s E U_s^\dagger = s(E)$.*

On va se baser sur cette propriété afin de construire une transformation de Clifford permettant d'encoder un code stabilisateur. L'indexation des erreurs de la base du groupe stabilisateur est dorénavant modifiée par commodité d'écriture : ces erreurs correspondent en effet à l'image par conjugaison des erreurs $(\mathcal{Z}_{k+1}, \dots, \mathcal{Z}_n)$ qui laissent invariant l'espace quantique avant encodage $\mathcal{H}_2^{\otimes k} \otimes \{|0_{n-k}\rangle\}$.

Propriété 2.11.2. *Soit $\mathcal{C}$ le code stabilisateur d'un groupe stabilisateur $\mathcal{S}$ généré par une base $(\bar{\mathcal{Z}}_{k+1}, \dots, \bar{\mathcal{Z}}_n)$. Il existe une transformation de Clifford U vérifiant $\bar{\mathcal{Z}}_i = U \mathcal{Z}_i U^{-1}$ pour tout $i \in \llbracket k+1; n \rrbracket$. L'encodage (n, k, U) engendre alors $\mathcal{C}$.*

Démonstration. Construisons une famille libre

$$(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_n)$$

complétant la famille

$$(\bar{\mathcal{Z}}_{k+1}, \dots, \bar{\mathcal{Z}}_n)$$

et qui vérifie les relations de commutation $\sigma(\bar{\mathcal{X}}_i) \cdot \sigma(\bar{\mathcal{Z}}_j) = \delta_{i,j}$, $\sigma(\bar{\mathcal{Z}}_i) \cdot \sigma(\bar{\mathcal{Z}}_j) = 0$ et $\sigma(\bar{\mathcal{X}}_i) \cdot \sigma(\bar{\mathcal{X}}_j) = 0$.

On complète d'abord itérativement la famille

$$(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n)$$

Si $k = 0$, il n'y a rien à faire. Sinon, supposons que pour un certain $t \in \llbracket 2; k+1 \rrbracket$ la famille $(\bar{\mathcal{Z}}_t, \dots, \bar{\mathcal{Z}}_n)$ est construite. La famille $(\sigma(\bar{\mathcal{Z}}_t), \dots, \sigma(\bar{\mathcal{Z}}_n))$ étant libre dans $\mathbb{F}_2^{2n}$, l'application de $\mathbb{F}_2^{2n}$ dans $\mathbb{F}_2^{n-t+1}$:

$$(x_1, \dots, x_n, z_1, \dots, z_n) \mapsto ((x_1, \dots, x_n, z_1, \dots, z_n) \cdot \sigma(\bar{\mathcal{Z}}_i))_{t \leq i \leq n}$$

est de rang $n - t + 1$. Son noyau est de dimension $n + t - 1$, et contient donc une séquence n'appartenant pas à Vect $\{\sigma(\bar{\mathcal{Z}}_t), \dots, \sigma(\bar{\mathcal{Z}}_n)\}$. Soit alors $\bar{\mathcal{Z}}_{t-1}$ un

antécédent par σ d'une telle séquence. Ainsi, la famille $(\bar{\mathcal{Z}}_{t-1}, \dots, \bar{\mathcal{Z}}_n)$ est toujours libre et commutative, et il suffit d'itérer ce processus jusqu'à construire toute la famille $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n)$.

On complète également itérativement la famille

$$(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_n)$$

Supposons que la famille $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_t)$, libre et respectant les relations de commutation, est construite pour un certain $t \in \llbracket 0; n-1 \rrbracket$. La famille $(\sigma(\bar{\mathcal{Z}}_1), \dots, \sigma(\bar{\mathcal{Z}}_n), \sigma(\bar{\mathcal{X}}_1), \dots, \sigma(\bar{\mathcal{X}}_t))$ est libre dans $\mathbb{F}_2^{2n}$. Ainsi, l'application de $\mathbb{F}_2^{2n}$ dans $\mathbb{F}_2^{n+t}$ qui à toute séquence $(x_1, \dots, x_n, z_1, \dots, z_n)$ associe la séquence des $n+t$ produits symplectiques respectifs de $(x_1, \dots, x_n, z_1, \dots, z_n)$ avec chaque séquence de la famille $(\sigma(\bar{\mathcal{Z}}_1), \dots, \sigma(\bar{\mathcal{Z}}_n), \sigma(\bar{\mathcal{X}}_1), \dots, \sigma(\bar{\mathcal{X}}_t))$ est de rang $n+t$ donc surjective. Il existe donc une séquence $(x_1, \dots, x_n, z_1, \dots, z_n)$ dont l'image vaut $(0, \dots, 0, 1, 0, \dots, 0)$ où le 1 est placé en $t+1$ -ième position. Soit alors $\bar{\mathcal{X}}_{t+1}$ une erreur de Pauli ayant $(x_1, \dots, x_n, z_1, \dots, z_n)$ pour représentation symplectique. Alors $\bar{\mathcal{X}}_{t+1}$ anticommute avec $\bar{\mathcal{Z}}_{t+1}$ et commute avec toutes les autres erreurs de la séquence $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_t)$. Cela garantit que la famille $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_{t+1})$ respecte les relations de commutation exigées. De plus, comme $\bar{\mathcal{X}}_{t+1}$ anticommute avec $\bar{\mathcal{Z}}_{t+1}$, il n'appartient pas au sous-groupe engendré par la famille $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_t)$ puisque toutes les erreurs de cette famille commutent avec $\bar{\mathcal{Z}}_{t+1}$. Ainsi, la famille $(\bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_{t+1})$ est également libre.

A présent, soit $\mathcal{E}$ le morphisme du groupe de Pauli défini par $\mathcal{E}(\mathcal{X}_i) = \bar{\mathcal{X}}_i$ et $\mathcal{E}(\mathcal{Z}_i) = \bar{\mathcal{Z}}_i$ pour tout $i \in \llbracket 1; n \rrbracket$, et par le fait que $\mathcal{E}$ laisse le groupe des phases $\Phi \times \{\mathcal{I}^{\otimes n}\}$ invariant. $\mathcal{E}$ est alors un isomorphisme qui conserve la commutativité. Par conséquent, il existe une transformation de Clifford U telle que, pour tout $E \in \Phi \times \mathcal{P}^{\otimes n}$:

$$UEU^\dagger = \mathcal{E}(E)$$

Considérons alors le code quantique $\mathcal{C}$ engendré par l'encodage (n, k, U) , c'est-à-dire :

$$\mathcal{C} = U(\mathcal{B}(\mathcal{H}_2^{\otimes k}) \otimes \{|0_{n-k}\rangle\})$$

Pour tout état $|\psi\rangle$ de l'espace quantique $\mathcal{H}_2^{\otimes n}$, on a alors l'équivalence :

$$\begin{aligned} |\psi\rangle \in \mathcal{C} &\Leftrightarrow U^{-1}|\psi\rangle \in \mathcal{H}_2^{\otimes k} \otimes \{|0_{n-k}\rangle\} \\ &\Leftrightarrow \forall i \in \llbracket k+1; n \rrbracket, \mathcal{Z}_i U^{-1}|\psi\rangle = U^{-1}|\psi\rangle \\ &\Leftrightarrow \forall i \in \llbracket k+1; n \rrbracket, U \mathcal{Z}_i U^{-1}|\psi\rangle = |\psi\rangle \\ &\Leftrightarrow \forall i \in \llbracket k+1; n \rrbracket, \bar{\mathcal{Z}}_i |\psi\rangle = |\psi\rangle \\ &\Leftrightarrow \forall E \in \mathcal{S}, E|\psi\rangle = |\psi\rangle \end{aligned}$$

Cela montre que le code quantique $\mathcal{C}$ est le code stabilisateur de groupe stabilisateur $\mathcal{S}$. $\square$

La propriété suivante permet, pour un code stabilisateur défini par une transformation de Clifford, de caractériser son groupe stabilisateur ainsi que son commutateur directement à partir de la transformation de Clifford.

Propriété 2.11.3. Soit $\mathcal{C}$ un code stabilisateur engendré par un encodage (n, k, U) où U est une transformation de Clifford. Le groupe stabilisateur $\mathcal{S}$ de $\mathcal{C}$ est généré par la famille $(\bar{\mathcal{Z}}_{k+1}, \dots, \bar{\mathcal{Z}}_n)$ et son commutateur $N(\mathcal{S})$ est généré par la famille $(i\mathcal{I}^{\otimes n}, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_k, \bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n)$ où $\bar{\mathcal{Z}}_i = U\mathcal{Z}_iU^{-1}$ et $\bar{\mathcal{X}}_i = U\mathcal{X}_iU^{-1}$ pour tout $i \in \llbracket 1; n \rrbracket$.

Démonstration. Notons $\mathcal{E}$ l'isomorphisme du groupe de Pauli donné par $U \mapsto UEU^\dagger$. Soit E une erreur de Pauli. Alors :

$$\begin{aligned} E \in \mathcal{S} &\Leftrightarrow \forall |\psi\rangle \in \mathcal{C}, E|\psi\rangle = |\psi\rangle \\ &\Leftrightarrow \forall |\psi\rangle \in \mathcal{B}(\mathcal{H}_2^{\otimes k}) \otimes \{|0_{n-k}\rangle\}, EU|\psi\rangle = U|\psi\rangle \\ &\Leftrightarrow \forall |\psi\rangle \in \mathcal{B}(\mathcal{H}_2^{\otimes k}) \otimes \{|0_{n-k}\rangle\}, \mathcal{E}^{-1}(E)|\psi\rangle = |\psi\rangle \end{aligned}$$

La dernière condition équivaut à ce que $\mathcal{E}^{-1}(E)$ est une erreur de Pauli dont la restriction à l'espace $\mathcal{H}_2^{\otimes k} \otimes \{|0_{n-k}\rangle\}$ vaut l'identité, ce qui est le cas si et seulement si $\mathcal{E}^{-1}(E)$ s'exprime comme une composée d'erreurs de la famille $(\mathcal{Z}_{k+1}, \dots, \mathcal{Z}_n)$. Cela montre que $\mathcal{S}$ est généré par la famille $(\bar{\mathcal{Z}}_{k+1}, \dots, \bar{\mathcal{Z}}_n)$.

Par ailleurs, une erreur de Pauli E appartient à $N(\mathcal{S})$ si et seulement si elle commute avec tous les éléments de $\mathcal{S}$. Comme $\mathcal{E}$ conserve la commutativité, cette condition équivaut à ce que $\mathcal{E}^{-1}(E)$ commute avec tous les éléments de la famille $(\mathcal{Z}_{k+1}, \dots, \mathcal{Z}_n)$. Cela équivaut à ce que $\mathcal{E}^{-1}(E)$ est la composée d'erreurs de la famille $(i\mathcal{I}^{\otimes n}, \mathcal{X}_1, \dots, \mathcal{X}_k, \mathcal{Z}_1, \dots, \mathcal{Z}_n)$, où la présence de l'erreur $i\mathcal{I}^{\otimes n}$ dans la famille précédente permet d'introduire un déphasage quelconque. Cela montre que $N(\mathcal{S})$ est généré par la famille $(i\mathcal{I}^{\otimes n}, \bar{\mathcal{X}}_1, \dots, \bar{\mathcal{X}}_k, \bar{\mathcal{Z}}_1, \dots, \bar{\mathcal{Z}}_n)$. $\square$

L'isomorphisme $\mathcal{E}$ du groupe de Pauli donné par $E \mapsto UEU^\dagger$ permet de suivre l'évolution de l'effet d'une erreur de Pauli au cours d'un encodage, comme en témoigne cette propriété.

Propriété 2.11.4. Soit U une transformation de Clifford sur $\mathcal{H}_2^{\otimes n}$ et $\mathcal{E}$ l'isomorphisme du groupe de Pauli associé à U , c'est-à-dire l'opération sur le groupe de Pauli $E \mapsto UEU^\dagger$. Si l'image par U d'un état $|\psi\rangle$ est l'état $|\psi'\rangle = U|\psi\rangle$, alors pour toute erreur de Pauli E , l'image par U de l'état $E|\psi\rangle$ est l'état $\mathcal{E}(E)|\psi'\rangle$.

Démonstration. On a en effet $UE|\psi\rangle = \mathcal{E}U|\psi\rangle = \mathcal{E}|\psi'\rangle$. $\square$

Finalement, introduisons l'isomorphisme du groupe de Pauli effectif induit par une transformation de Clifford. Cela sera utile pour simplifier la description du procédé de décodage dans la section suivante, et pour introduire ensuite la notion d'encodeur.

Propriété 2.11.5. Soit U une transformation de Clifford. L'application $\mathcal{E}$ de P^n dans P^n qui, à la classe d'équivalence de toute erreur $E \in P^n$, associe la classe d'équivalence de l'erreur $UEU^\dagger$, est bien définie et correspond à un isomorphisme appelé **isomorphisme du groupe de Pauli effectif induit** par U .

Démonstration. Afin de vérifier que l'application $\mathcal{E}$ est bien définie, il suffit de remarquer que si deux erreurs de Pauli E et E' ne diffèrent que d'une phase, il en est de même pour les erreurs de Pauli $UEU^\dagger$ et $UE'U^\dagger$. Ainsi, l'image d'une erreur de Pauli effective est bien définie et indépendante de l'erreur de Pauli choisie pour l'application de $UEU^\dagger$.

C'est de toute évidence un morphisme, puisque la classe de toute erreur $UE \circ E'U^\dagger$ est le produit des classes des erreurs $UEU^\dagger$ et $UE'U^\dagger$. Comme de plus, la classe de l'identité n'est atteinte que lorsque E appartient à la classe de l'identité, ce qui montre que $\mathcal{E}$ est un isomorphisme. $\square$

Il convient de noter le fait immédiat mais très important, qui sera repris lors de la définition 2.13.1 de la notion d'encodeur, qu'un tel isomorphisme de Pauli effectif doit nécessairement conserver la commutativité.

2.12 Décodage d'un encodage stabilisateur

On décrit ici un procédé de décodage au maximum de vraisemblance applicable aux encodages stabilisateurs. Ce procédé est optimal lorsque le modèle du canal quantique est descriptible par des erreurs de Pauli, comme c'est le cas pour le canal quantique dépolarisant et le canal quantique à effacement. Il reste applicable pour tout autre modèle de canal quantique sans être nécessairement optimal dans ce cas.

Il sera plus simple de travailler directement sur le groupe de Pauli effectif lors de la présentation du procédé de décodage. En effet, deux erreurs de Pauli qui ne diffèrent que par leur phase ont le même effet sur l'espace quantique, puisqu'une telle différence n'a pour effet que d'introduire un déphasage sur les états, et que des états qui ne diffèrent que par une phase sont en réalité identiques. On fera donc de nouveau usage de la définition 2.10.3 des ensembles S et $N(S)$ des erreurs de Pauli effectives correspondant au groupe stabilisateur $\mathcal{S}$ et à son commutateur $N(\mathcal{S})$. Le modèle d'erreur du canal quant à lui sera à l'échelle des erreurs de Pauli effectives, c'est-à-dire qu'on comptabilisera, pour toute erreur de Pauli effective, la somme des probabilités des erreurs de Pauli correspondantes. Le terme « erreur » dans cette section désignera exclusivement une erreur de Pauli effective.

Il faut commencer par souligner une particularité du décodage qui n'existe pas dans la théorie des codes classiques. Plusieurs erreurs peuvent avoir la même action sur un code stabilisateur. Cela vient du fait que par définition, l'ensemble des erreurs de S laissent le code stabilisateur invariant. Ainsi, le fait de composer une erreur par n'importe quelle erreur de S ne modifie en rien l'action d'une telle erreur sur le code stabilisateur. Lors du décodage au maximum de vraisemblance, la tâche ne consistera donc pas à estimer l'erreur la plus probable, mais la classe d'erreurs la plus probable modulo S . Afin de souligner cette subtilité, appelons le procédé **décodage au maximum de classe de vraisemblance**.

Le procédé de décodage consiste en substance à appliquer la transformation de Clifford inverse sur l'état reçu à la sortie du canal, puis à effectuer une mesure sur les $n - k$ bits de l'état obtenu. Cette mesure constitue le syndrome qui livre une information sur la classe modulo S à laquelle appartient l'erreur E qui a été introduite par le canal. Ensuite, il s'agit de retrouver la classe d'erreurs la plus probable parmi celles qui correspondent à ce syndrome. Définissons donc le syndrome d'une erreur :

Définition 2.12.1. Soit (n, k, U) un encodage quantique où U est une transformation de Clifford, et soit $\mathcal{E}$ l'isomorphisme du groupe de Pauli effectif induit par U . Le **syndrome** d'une erreur de Pauli effective $E \in P^n$ est la séquence :

$$(\sigma(E) \cdot \sigma(\mathcal{E}(Z_{k+1})), \dots, \sigma(E) \cdot \sigma(\mathcal{E}(Z_n)))$$

Le syndrome d'une erreur est donc le produit symplectique de l'erreur par la famille canonique des générateurs de S . On verra grâce à la description ci-dessous que le syndrome d'une erreur peut être obtenu grâce la mesure des $n - k$ derniers bits après avoir appliqué U^{-1} sur l'état reçu.

Considérons donc un code stabilisateur $\mathcal{C}$ engendré par un encodage (n, k, U) où U est une transformation de Clifford, et notons $\mathcal{E}$ l'isomorphisme du groupe de Pauli effectif induit par U . On suppose que le canal quantique introduit une erreur $E \in P^n$ suivant une distribution de probabilité $P(E)$ connue. Cette distribution de probabilité peut être connue a priori comme dans le cas du canal dépolarisant, ou a posteriori comme dans le cas du canal à effacement. Comme $\mathcal{E}$ est un isomorphisme, on peut écrire toute erreur E sous la forme :

$$E = \mathcal{E}(L, S)$$

où $L \in P^k$ et $S \in P^{n-k}$. La lettre L est employée pour désigner l'erreur située sur les k premières positions porteuses de l'information quantique et appelées *logiques*, tandis que la lettre S est employée pour désigner l'erreur située sur les $n - k$ dernières positions sur lesquelles on mesurera le *syndrome*.

Supposons que l'on envoie via ce canal un état $|\psi'\rangle = U(|\psi\rangle \otimes |0_{n-k}\rangle)$ où $|\psi\rangle \in \mathcal{H}_2^{\otimes k}$ représente l'état d'un registre à k bits que l'on cherche à protéger. Avec probabilité $P(E)$, une erreur E est introduite par le canal. Lorsque l'on applique U^{-1} à l'état reçu, cela correspond donc à l'introduction de l'erreur (L, S) sur l'état $|\psi\rangle \otimes |0_{n-k}\rangle$.

Soit $(x_1, \dots, x_{n-k}, z_1, \dots, z_{n-k})$ la représentation symplectique de S . S s'écrit comme le produit $S^{(X)}S^{(Z)}$ où $S^{(X)}$ et $S^{(Z)}$ sont les erreurs dont la i -ème coordonnée vaut, respectivement, X si $x_i = 1$ et 0 sinon, et Z si $z_i = 1$ et I sinon. Comme $S^{(Z)}$ agit trivialement sur l'état $|0_{n-k}\rangle$ des $n - k$ derniers bits du registre, S envoie l'état $|0_{n-k}\rangle$ sur l'état $|x\rangle$ où $x = (x_1, \dots, x_{n-k})$.

On constate que la séquence x correspond au syndrome de l'erreur $E = \mathcal{E}(L, S)$. En effet, pour tout $i \in \llbracket k + 1; n \rrbracket$, on a :

$$\sigma(\mathcal{E}(L, S)) \cdot \sigma(\mathcal{E}(Z_i)) = \sigma(L, S) \cdot \sigma(Z_i)$$

et (L, S) anticommute avec Z_i si et seulement si $x_i = 1$.

Ainsi en fonction de l'erreur E introduite par le canal, l'état obtenu après application de U^{-1} à l'état reçu par le canal appartient à l'un des espaces $\mathcal{H}_2^{\otimes k} \otimes \{|x\rangle\}$ où $x \in \{0, 1\}^{n-k}$. La famille des projecteurs orthogonaux sur ces espaces est la famille associée à la mesure des $n - k$ derniers bits du registre, comme cela a été présenté dans la définition 2.2.3. En effectuant la mesure des $n - k$ derniers bits du registre, le résultat x de la mesure révèle ainsi le syndrome de l'erreur E .

Une fois le syndrome x mesuré sur les $n - k$ dernières positions, on est alors intéressé par connaître l'erreur L la plus vraisemblable sur les k premières positions du registre, afin d'appliquer ensuite la correction L^{-1} sur ces positions du registre. On cherche donc une erreur $\hat{L}$ maximisant la probabilité que l'erreur E du canal s'écrive sous la forme $\mathcal{E}(\hat{L}, S)$ où la partie en X de la représentation symplectique de S correspond à la séquence x :

$$\hat{L} \in \operatorname{argmax}_{L \in P^k} \sum_{S^{(Z)} \in \{I, Z\}^{n-k}} P(\mathcal{E}(L, S^{(X)} S^{(Z)}))$$

où $S^{(X)}$ est l'erreur dont la i -ème coordonnée vaut X si $x_i = 1$ et I si $x_i = 0$. Résumons donc le procédé :

Définition 2.12.2. Soit un canal quantique de loi de probabilité d'erreur $P(E)$ où $E \in P^n$. On rappelle que $P(E)$ désigne ainsi la somme des probabilités d'occurrence d'une erreur de Pauli appartenant à l'erreur de Pauli effective E . Soit (n, k, U) un encodage quantique où U est une transformation de Clifford, et soit $\mathcal{C}$ le code stabilisateur qu'il engendre. Le **décodage au maximum de classe de vraisemblance** d'un état $U(|\psi\rangle \otimes \{0_{n-k}\}) \in \mathcal{C}$, où $|\psi\rangle \in \mathcal{H}_2^{\otimes k}$, envoyé par le canal consiste à effectuer successivement ces trois étapes :

- L'application de U^{-1} à l'état reçu en sortie du canal
- La mesure de la séquence x des $n - k$ derniers bits de cet état, conformément à la définition 2.2.3
- L'estimation de l'erreur la plus probable sur les k premiers bits :

$$\hat{L} \in \operatorname{argmax}_{L \in P^k} \sum_{S^{(Z)} \in \{I, Z\}^{n-k}} P(\mathcal{E}(L, S^{(X)} S^{(Z)}))$$

où $S^{(X)}$ est l'erreur dont la i -ème coordonnée vaut X si $x_i = 1$ et I si $x_i = 0$.

Ce procédé permet de rétablir, avec probabilité maximale, l'état $|\psi\rangle$ porteur de l'information quantique en appliquant $\hat{L}^{-1} = \hat{L}$ sur les k premières positions du registre.

Un tel procédé de décodage est applicable également aux modèles de canaux quantiques dont la loi de probabilité n'est pas définie uniquement sur les erreurs de Pauli. En effet, si l'on suppose qu'une erreur quelconque E est introduite par

le canal quantique, on peut toujours séparer les erreurs de Pauli formant $U^\dagger EU$ en fonction de leur syndrome x . Si l'on note $S^{(x)}$ l'erreur de $\mathcal{P}^{\otimes(n-k)}$ formée de $\mathcal{X}$ aux positions où $x_i = 1$ et de $\mathcal{I}$ aux positions où $x_i = 0$, on a ainsi :

$$U^\dagger EU = \sum_{x \in \{0,1\}^{n-k}} \sum_{L \in \mathcal{P}^{\otimes k}} \sum_{S^{(z)} \in \{\mathcal{I}, \mathcal{Z}\}^{\otimes(n-k)}} c_{L \otimes (S^{(x)} \circ S^{(z)})} L \otimes (S^{(x)} \circ S^{(z)})$$

Après avoir appliqué U^{-1} à l'état reçu en sortie du canal, l'état obtenu n'appartient pas nécessairement à l'un des espaces $\mathcal{H}_2^{\otimes k} \otimes \{|x\rangle\}$ où $x \in \{0,1\}^{n-k}$, mais à une combinaison linéaire d'états de ces espaces. En effectuant toutefois la mesure des $n-k$ derniers bits et en notant x le syndrome mesuré, on obtient alors un état factorisé sur le registre des k premiers bits et le registre des $n-k$ derniers bits, dont l'état sur les k premiers bits correspond à l'état $|\psi\rangle$ affecté de l'erreur :

$$\sum_{L \in \mathcal{P}^{\otimes k}} \sum_{S^{(z)} \in \{\mathcal{I}, \mathcal{Z}\}^{\otimes(n-k)}} c_{L \otimes (S^{(x)} \circ S^{(z)})} L$$

Alors, en connaissant la loi de probabilité de l'erreur E introduite par le canal quantique, et par conséquent la loi de probabilité conditionnelle de l'erreur ci-dessus sachant le résultat x de la mesure, on recherche l'erreur la plus probable affectant le registre des k premiers bits dans cette situation. Cela ne garantit cependant pas que l'état $|\psi\rangle$ puisse être rétabli, quand bien même la bonne erreur sur le registre des k premiers bits a été trouvée, car dans le cas général, il reste toujours la possibilité que cette erreur soit non inversible.

2.13 Encodeurs quantiques

Les définitions de cette section sont présentées de manière à être lisibles à partir des seules bases de théorie des codes classiques, sans nécessiter de connaissance en mécanique quantique. Le premier paragraphe permet néanmoins de justifier l'approche ainsi faite et d'en expliquer l'intérêt scientifique pour les codes stabilisateurs.

On a montré que tout code stabilisateur peut être engendré par un encodage qui correspond à une transformation de Clifford, c'est-à-dire une évolution unitaire U qui laisse le groupe de Pauli invariant par l'opération de conjugaison $E \mapsto UEU^\dagger$. L'intérêt principal d'un encodage induit par une transformation de Clifford réside dans le fait qu'il permet de suivre l'effet d'une erreur de Pauli sur un état de l'espace quantique avant et après encodage, cet effet étant décrit par l'isomorphisme $E \mapsto UEU^\dagger$ du groupe de Pauli comme présenté dans la propriété 2.11.4. De plus, comme on l'a vu aussi bien pour la propriété 2.10.4 de distance minimale que pour le procédé de décodage, il n'est pas utile de conserver l'information portant sur la phase des erreurs de Pauli. Il est donc plus simple de s'intéresser à l'isomorphisme du groupe de Pauli effectif induit par la transformation de Clifford U , défini dans la propriété 2.11.5, un tel objet permettant de décrire toutes les propriétés essentielles de l'encodage sans

s'encombrer du signe des erreurs de Pauli. En faisant abstraction de la structure de l'encodage quantique sous-jacente, on va donc définir simplement la notion d'**encodeur quantique**, isomorphisme sur le groupe de Pauli effectif dont toutes les propriétés essentielles peuvent être extraites.

Au-delà de sa capacité à décrire le décodage ainsi que la distance minimale de l'encodage et du code stabilisateur sous-jacents, l'encodeur quantique est surtout un modèle reposant sur des paramètres discrets, puisqu'il s'agit d'une action sur P^n où $P = \{I, X, Y, Z\}$ est un alphabet fini, et dont la présentation est proche des encodages dans la théorie des codes classiques. Ainsi, toute la suite de l'exposé peut être suivie pour un lecteur uniquement familiarisé aux codes correcteurs d'erreurs classiques. C'est également le point fort ayant permis d'effectuer des rapprochements entre le modèle du turbo-encodage classique et le modèle du turbo-encodeur quantique que l'on introduit dans le cadre de la thèse, et ayant permis d'analyser les propriétés qu'un tel turbo-encodeur doit posséder afin de lui garantir de bonnes propriétés de distance minimale et de décodage.

Afin que toute la suite soit auto-suffisante, on rappelle les faits suivants.

Rappel 2.13.1. L'ensemble $P = \{I, X, Y, Z\}$ est un groupe abélien dont la loi de groupe est donnée par $I^2 = X^2 = Y^2 = Z^2 = I$, $XY = YX = Z$, $YZ = ZY = X$, $ZX = XZ = Y$ et $IE = EI = E$ pour tout $E \in P$. Pour tout $n \in \mathbb{N}^*$, le groupe P^n obtenu par produit cartésien est nommé le **groupe de Pauli effectif** de taille n , et ses éléments sont les **erreurs de Pauli effectives** de taille n , que l'on pourra nommer plus brièvement **erreurs** de taille n .

Le **poids** d'une erreur est le nombre de coordonnées de cette erreur qui sont différentes de I .

La **représentation symplectique** σ est l'application de P dans $\mathbb{F}_2^2$ donnée par $\sigma(I) = (0, 0)$, $\sigma(X) = (1, 0)$, $\sigma(Y) = (1, 1)$ et $\sigma(Z) = (0, 1)$. Il s'agit d'un isomorphisme. Par extension, pour tout $n \in \mathbb{N}^*$, la **représentation symplectique** d'une erreur de taille n $E = (E_1, \dots, E_n)$ est la séquence de $2n$ bits $\sigma(E) = (x_1, \dots, x_n, z_1, \dots, z_n)$ où, pour tout $i \in \llbracket 1; n \rrbracket$, $(x_i, z_i) = \sigma(E_i)$.

Le **produit symplectique** de deux séquences de $2n$ bits, $(x_1, \dots, x_n, z_1, \dots, z_n)$ et $(x'_1, \dots, x'_n, z'_1, \dots, z'_n)$, est le bit :

$$(x_1, \dots, x_n, z_1, \dots, z_n) \cdot (x'_1, \dots, x'_n, z'_1, \dots, z'_n) = \sum_{i=1}^n x_i z'_i + x'_i z_i$$

On rappelle également que l'on écrit X_i , Y_i et Z_i pour désigner une erreur de taille n connue valant respectivement X , Y et Z en i -ème position et I en toutes les autres positions.

On peut à présent définir un encodeur quantique.

Définition 2.13.1. Un **encodeur quantique** est un triplet $(n, k, \mathcal{E})$ où n et k sont des entiers vérifiant $n \geq k \geq 1$ et $\mathcal{E}$ est un isomorphisme du groupe de Pauli effectif P^n qui conserve le produit symplectique, c'est-à-dire tel que pour tout $(E, E') \in (P^n)^2$:

$$\sigma(\mathcal{E}(E)) \cdot \sigma(\mathcal{E}(E')) = \sigma(E) \cdot \sigma(E')$$

Il est important de noter qu'un encodeur quantique doit posséder la propriété de conservation du produit symplectique au sens ci-dessus, puisqu'il correspond à un isomorphisme de Pauli effectif induit par une transformation de Clifford. Cette condition dans la définition ne jouera toutefois aucun rôle par la suite. La valeur de k dans le triplet $(n, k, \mathcal{E})$ est quant à elle nécessaire à mentionner car elle permet de différencier les k premières positions porteuses de l'information quantique des $n - k$ dernières positions sur lesquelles est mesuré le syndrome lors du décodage.

Un encodeur quantique étant un modèle virtuel construit à partir d'un encodage stabilisateur, rattachons-lui les notions de groupe stabilisateur, groupe commutateur, et de distance minimale qui correspondent à l'encodage et au code stabilisateur sous-jacents. Les groupes stabilisateur et commutateur d'un encodeur quantique sont définis en conformité avec la propriété 2.11.3 et la définition 2.10.3, tandis que la distance minimale d'un encodeur quantique est définie en se servant de la propriété 2.10.4.

Définition 2.13.2. Le groupe stabilisateur S et le groupe commutateur $N(S)$ d'un encodeur quantique $(n, k, \mathcal{E})$ sont les sous-groupes de P^n engendrés par les familles respectives $(\mathcal{E}(Z_{k+1}), \dots, \mathcal{E}(Z_n))$ et $(\mathcal{E}(X_1), \dots, \mathcal{E}(X_k), \mathcal{E}(Z_1), \dots, \mathcal{E}(Z_n))$. En d'autres termes :

$$\begin{aligned} S &= \mathcal{E}(\{I\}^k \times \{I, Z\}^{n-k}) \\ N(S) &= \mathcal{E}(P^k \times \{I, Z\}^{n-k}) \end{aligned}$$

La distance minimale d'un encodeur quantique est le plus petit poids d'un élément de $N(S) \setminus S$ c'est-à-dire de l'ensemble :

$$\mathcal{E}((P^k \setminus \{I\}^k) \times \{I, Z\}^{n-k})$$

Dans le modèle d'encodage sous-jacent à l'encodeur quantique, toute erreur de l'ensemble $\{I, Z\}$ a un effet trivial sur les $n - k$ dernières positions de l'entrée de l'encodage. Ainsi, $N(S)$ est l'ensemble des erreurs stabilisant le code tandis que S est l'ensemble des erreurs laissant le code invariant. On dira ainsi que toute erreur de $N(S)$ est **non détectée** et toute erreur de S est **bénigne**.

Il est nécessaire de compléter le modèle en mentionnant le canal quantique ainsi que le décodage. Avant l'utilisation du canal quantique, l'encodeur n'est censé véhiculer aucune erreur, puisqu'on suppose être dans une situation où l'encodage est parfait c'est-à-dire ne comportant pas d'erreurs en soi. Le canal quantique introduit alors une erreur, dont l'estimation revient au procédé de décodage. On se restreint ici aux canaux quantiques modélisables par une loi de probabilité définie sur les erreurs de Pauli effectives, même s'il est possible de lever cette restriction comme cela a été présenté dans la section précédente traitant du décodage d'un code stabilisateur.

Définition 2.13.3. Un **canal quantique** introduit une erreur de Pauli effective $E \in P^n$ conformément à une certaine loi de probabilité. Il est dit **sans mémoire**

si sa loi de probabilité est le produit cartésien d'une loi de probabilité identique en chacune de ses n positions.

L'information qui peut être acquise sur une erreur E introduite par le canal quantique s'appelle le syndrome de l'erreur. Le syndrome est une séquence de $n - k$ bits qui permet de savoir si chacune des $n - k$ dernières coordonnées respectives de $\mathcal{E}^{-1}(E)$ appartient à $\{I, Z\}$ ou à $\{X, Y\}$. Contrairement au cas des encodeurs classiques, on ne connaît donc qu'une information partielle sur ces coordonnées, qui consiste à savoir si elles ont un effet bénin ou non. On réfère à la section 2.12 traitant du décodage des codes stabilisateurs pour plus de détails sur la signification physique du syndrome.

Définition 2.13.4. Soit $(n, k, \mathcal{E})$ un encodeur. Pour toute erreur $E \in P^n$, le syndrome $\text{syn}(E)$ de E est la séquence $x = (x_{k+1}, \dots, x_n)$ extraite de la représentation symplectique $(x_1, \dots, x_n, z_1, \dots, z_n)$ de $\mathcal{E}^{-1}(E)$. En d'autres termes, il s'agit de la séquence des $n - k$ bits dont le i -ème élément vaut 0 si la $k + i$ -ème coordonnée de $\mathcal{E}^{-1}(E)$ appartient à $\{I, Z\}$ et 1 si la $k + i$ -ème coordonnée de $\mathcal{E}^{-1}(E)$ appartient à $\{X, Y\}$.

Remarque 2. La définition précédente du syndrome est conforme à celle donnée dans la définition 2.12.1 en étendant le produit symplectique aux erreurs de Pauli effectives. Le syndrome est donc un morphisme de P^n dans $\mathbb{F}_2^n$, et le syndrome d'une erreur E est nul si et seulement si $E \in N(S)$, cette dernière propriété étant due à la conservation de la commutativité par l'encodeur quantique.

Si le syndrome de E est connu, chacune des $n - k$ dernières coordonnées de $\mathcal{E}^{-1}(E)$ est déterminée modulo $\{I, Z\}$: elle vaut I modulo $\{I, Z\}$ si $x_i = 0$, et X modulo $\{I, Z\}$ si $x_i = 1$. On peut donc écrire E sous la forme $\mathcal{E}(L, S^{(X)}S^{(Z)})$, où L est une erreur de taille k , $S^{(Z)}$ est une erreur de taille $n - k$ dont toutes les coordonnées appartiennent à $\{I, Z\}$, et $S^{(X)}$ est l'erreur de taille $n - k$ dont la i -ème coordonnée vaut X si $x_i = 1$ et I si $x_i = 0$.

Comme toute classe d'erreurs modulo S a la même action sur le code sous-jacent au modèle de l'encodeur, le problème du décodage est de déterminer la classe des erreurs modulo S la plus probable connaissant le syndrome x de l'erreur. À la lumière de la description ci-dessus, toute classe d'erreurs est déterminée de manière unique grâce à L et grâce au syndrome qui fixe $S^{(Z)}$. Le problème du décodage équivaut donc simplement à rechercher les k premières coordonnées L de $\mathcal{E}^{-1}(E)$ connaissant le syndrome x de l'erreur E introduite par le canal quantique.

Définition 2.13.5. Le décodage d'une erreur E , introduite par un canal quantique de loi de probabilité P et dont on connaît le syndrome $x = \text{syn}(E)$, consiste à produire une estimation des k premières coordonnées L de $\mathcal{E}^{-1}(E)$. Le **décodage au maximum de classe de vraisemblance** donne une erreur $\hat{L}$ qui maximise la probabilité de la classe des erreurs E de syndrome x et dont

l'antécédent $\mathcal{E}^{-1}(E)$ commence par $\hat{L}$:

$$\hat{L} \in \operatorname{argmax}_{L \in P^k} \sum_{S^{(Z)} \in \{I, Z\}^{n-k}} P(\mathcal{E}(L, S^{(X)} S^{(Z)}))$$

où $S^{(X)}$ est l'erreur dont la i -ème coordonnée vaut X si $x_i = 1$ et I si $x_i = 0$.

Chapitre 3

Le turbo-encodeur formel, construction dans un formalisme commun aux cadres classique et quantique

3.1 Un formalisme d'erreur et d'encodeur commun aux cadres classique et quantique

Dans la section 1.5 et la section 2.13, on a présenté l'*encodeur* qui permet de modéliser l'évolution d'une erreur affectant l'entrée d'un encodage lorsque l'encodage lui est appliqué. Dans le premier cas, il s'agissait d'un encodage classique linéaire tandis que dans le second cas, l'encodeur était relié à un encodage quantique correspondant à une transformation de Clifford. Des points communs existent entre ces deux modèles, à commencer par le caractère fini de l'ensemble d'erreurs. Il est possible de rapprocher ces deux modèles d'encodeurs sous un seul cadre abstrait, qui permet de traiter aussi bien le premier cas que le second. Les objets considérés sous ce cadre, tels que l'ensemble d'erreurs ou l'encodeur, seront qualifiés de *formels*. On retient le cadre commun le plus simple qui suffit à définir ensuite un turbo-encodeur *formel* avec de bonnes propriétés de distance minimale. Le turbo-encodeur formel que l'on entreprend de définir dans le formalisme commun, correspond, lorsqu'il est décliné sous le modèle classique, au turbo-encodage classique déjà introduit par [10], tandis que sous le modèle quantique, il permet de définir une nouvelle classe de turbo-encodage quantique dont la distance minimale tend vers l'infini avec la taille de l'encodage et pour lequel il existe un décodage de complexité linéaire

livrant de bonnes performances. Le but de la démarche de rapprochement des modèles des encodeurs classique et quantique est de bien mettre en évidence les points de similitude qui existent entre le turbo-encodage classique présenté par [10] et le turbo-encodage quantique que l'on introduit, ainsi que les subtilités théoriques particulières rencontrées dans le modèle quantique qui rendent le turbo-encodage quantique différent en certains points de son homologue classique. Toutes les propriétés de l'encodeur quantique ne sont notamment pas reprises, en particulier la propriété de conservation de la commutativité, car elle n'est pas nécessaire pour établir les caractéristiques de l'encodeur formel.

Le formalisme commun aux modèles d'encodeurs classique et quantique repose sur un groupe fini G nommé ensemble des erreurs, qui se décline dans le modèle classique en $\mathbb{F}_2 = \{0, 1\}$ et dans le modèle quantique en $P = \{I, X, Y, Z\}$. Un encodeur est alors un triplet $(n, k, \mathcal{E})$ où $\mathcal{E}$ est un automorphisme de G^n . Rappelons que le rôle de k est de spécifier le nombre de positions porteuses de l'information à encoder en entrée.

On a de surcroît constaté, dans le modèle d'encodage dont dérive l'encodeur quantique, que le sous-groupe d'erreurs $\{I, Z\}$ a une action triviale sur les $n - k$ dernières positions de l'entrée. L'ensemble des erreurs qui joue le même rôle dans le modèle classique est réduit à 0. Une manière d'absorber cette différence dans le formalisme commun est de définir un sous-groupe strict H de G que l'on nomme ensemble des erreurs non détectées, qui correspond à $\{I, Z\}$ dans le modèle quantique, et est réduit à $\{0\}$ dans le modèle classique. Cela permet, dans le formalisme commun, de redéfinir les ensembles S et $N(S)$ des erreurs bénignes et non détectées, et d'exprimer la distance minimale d'un encodeur.

Dans le résultat de distance minimale que l'on démontre sur le turbo-encodeur formel, il sera nécessaire de faire intervenir, en plus de la distance minimale, une autre distance liée à l'encodeur formel que l'on nomme distance dégénérée.

Définition 3.1.1. L'ensemble formel d'erreurs est un groupe fini G d'élément neutre I possédant un sous-groupe strict H nommé **groupe des erreurs non détectées**. Le **poids** d'une erreur $E = (E_1, \dots, E_n) \in G^n$ est le nombre de coordonnées E_i différentes de I ; il est noté $|E|$.

Un **encodeur formel** est un triplet $(n, k, \mathcal{E})$ où $n \geq k \geq 1$ et $\mathcal{E}$ est un automorphisme de G^n .

L'ensemble des erreurs non détectées $N(S)$ d'un encodeur formel $(n, k, \mathcal{E})$ est l'ensemble :

$$N(S) = \mathcal{E}(G^k \times H^{n-k})$$

Cet ensemble se divise en deux catégories, **l'ensemble des erreurs bénignes**

$$S = \mathcal{E}(\{I\}^k \times H^{n-k})$$

et **l'ensemble des erreurs nocives**

$$N(S) \setminus S = \mathcal{E}((G^k \setminus \{I\}^k) \times H^{n-k})$$

La **distance minimale** d'un encodeur formel $(n, k, \mathcal{E})$ est le plus petit poids d'une erreur nocive, soit une erreur de $N(S) \setminus S$. La **distance dégénérée** d'un

encodeur formel $(n, k, \mathcal{E})$ est le plus petit poids d'une erreur non détectée non triviale, soit une erreur de $N(S) \setminus \{I\}^n$.

D'un point de vue pratique, il est utile de séparer les différentes positions d'entrée et de sortie d'un encodeur formel selon leur nature. A l'entrée, les k premières positions dites *logiques* donnent l'erreur sur l'information (classique ou quantique) que l'on cherche à transmettre via l'encodage modélisé par l'encodeur formel, tandis que les $n - k$ dernières positions dites *de syndrome* sont celles sur lesquelles le syndrome est mesuré. Comme le syndrome sur ces positions révèle si chaque coordonnée de l'erreur est détectée, c'est-à-dire appartient à $G \setminus H$, il sera également utile de donner une définition particulière du poids *de syndrome* d'une erreur en ces $n - k$ positions. L'erreur de sortie d'un encodeur formel est quant à elle intégralement dite *physique*, car cette erreur est effectivement celle qui se produit dans le canal.

Définition 3.1.2. Ecrivons les erreurs d'entrée E et de sortie $\mathcal{E}(E)$ d'un encodeur formel $(n, k, \mathcal{E})$ sous les formes respectives $E = (L, S)$ et $\mathcal{E}(E) = P$, où $L \in G^k$ est une **erreur logique**, $S \in G^{n-k}$ est une **erreur de syndrome** et $P \in G^n$ est nommée **erreur physique**. On appelle alors :

- **poids logique** de E , noté $|E|_L$, le poids de L
- **poids de syndrome** de E , noté $|E|_{\bar{H}}$, le nombre de coordonnées de S appartenant à $\bar{H} = G \setminus H$
- **poids physique** de P , simplement, le poids de P

Lorsque l'on considèrera un encodeur *convolutif* formel, cas particulier d'un encodeur formel, les erreurs et poids logiques et physiques n'auront pas la même définition.

Dans le cas de l'encodeur convolutif formel, qui correspond à un encodeur formel divisible en plusieurs étapes, les erreurs intermédiaires seront séparées en un maximum de quatre groupes de positions de natures différentes, incluant des positions jouant le rôle spécifique *de mémoire*. L'étiquette de mémoire est également appliquée par commodité aux premières positions de l'entrée ainsi qu'aux dernières positions de la sortie de l'encodeur convolutif formel, ce qui modifiera la définition d'une erreur logique et d'une erreur physique. Cela sera repris dans la section suivante.

Avant de passer à la suite, on définit un objet dont la pertinence sera justifiée lors de la démonstration de la distance minimale d'un turbo-encodeur formel, notamment aux sections 4.3 et 4.4, lorsqu'il s'agira d'étudier les *encodeurs convolutifs formels anti-récursifs*.

Définition 3.1.3. Un **encodeur formel inversé** est un triplet $(n', k', \mathcal{E}')$ où $n' \leq k'$ et où $\mathcal{E}'$ est un morphisme de $G^{k'}$ dans $G^{n'}$. Son entrée E' et sa sortie $\mathcal{E}'(E')$ sont respectivement des erreurs **logique** et **physique**. L'erreur de syndrome de son entrée est réduite à la séquence vide. Les mêmes définitions et notations de poids s'appliquent dans ce cadre.

L'encodeur formel inversé réalise en quelque sorte une inversion de l'entrée et de la sortie par rapport à un encodeur formel. Les termes *logique* et *physique* restent cependant respectivement associés à l'entrée et à la sortie de l'encodeur formel inversé, et ce afin de simplifier la présentation des résultats de la section 4.3.

3.2 Turbo-encodeur formel

La construction d'un turbo-encodeur formel va ressembler à celle d'un turbo-encodage classique à quelques exceptions près.

D'abord, comme il s'agit d'un encodeur formel et non d'un encodage, il sera nécessaire d'introduire des positions auxiliaires en entrée afin d'égaliser le nombre de positions en entrée et en sortie ; ce sont ces positions sur lesquelles s'effectue la mesure du syndrome aussi bien dans le cas classique que dans le cas quantique. De plus, l'opération intermédiaire d'entrelacement des positions est légèrement plus complexe. En effet, le rôle de cette opération dans un turbo-encodage classique est d'obtenir en sortie de l'entrelacement une erreur dont la loi de probabilité, moyennée sur l'ensemble des entrelacements possibles, ne dépend que de son poids. Dans le cadre de l'encodeur formel, on souhaite définir un entrelaceur qui permute aléatoirement les positions des erreurs, mais qui modifie aussi de manière aléatoire tous les symboles de l'erreur différents de I . Cela permet d'obtenir en sortie de l'entrelaceur une erreur dont la loi de probabilité dépend exclusivement de son poids. L'entrelaceur consiste donc en une permutation aléatoire des positions de l'erreur, suivie d'une séquence de permutations de l'ensemble G , laissant I invariant, appliquées en chaque position. La possibilité d'implémenter physiquement de telles permutations provient du corollaire 2.11.1.

On procède également pour une raison pratique à un amendement du schéma de turbo-encodage, en plaçant à l'endroit de l'encodeur externe, plutôt qu'un encodeur convolutif, une séquence d'encodeurs identiques de taille constante placés en parallèle. En effet, il est montré dans l'article [41] qu'un encodage convolutif quantique ne peut être récursif et non catastrophique. Or, dans la construction de turbo-encodage standard dans laquelle tous les bits quantiques en sortie de l'encodage externe sont envoyés en entrée de l'encodage interne, le caractère non catastrophique de l'encodage interne est une condition nécessaire pour amorcer le décodage itératif. Cela incitera à proposer dans la section 3.5 une modification du schéma du turbo-encodeur, en envoyant une fraction des positions de l'encodeur externe directement vers le canal afin de permettre l'amorce du décodage. Afin de permettre à cette solution de fonctionner, une condition particulière doit être remplie par l'encodeur externe, qu'il est plus commode de rechercher s'il consiste en une séquence d'encodeurs identiques placés en parallèle.

Commençons donc d'abord par définir un encodeur convolutif formel ainsi qu'un entrelaceur formel, avant de définir le turbo-encodeur formel. A toutes les étapes de l'encodeur convolutif formel, l'erreur se décompose en une séquence d'erreurs jouant un rôle particulier tel qu'on le présente dans la définition 3.2.1.

Ainsi, une erreur obtenue en sortie d'une étape donnée et injectée en entrée de l'étape suivante constitue une erreur *de mémoire* ; en d'autres termes, les erreurs de mémoire sont celles qui entrent en jeu dans deux étapes successives de l'opération de convolution. Par extension, certaines erreurs en entrée et en sortie de l'encodeur convolutif formel seront également appelées erreurs de mémoire. Une erreur *logique* est une erreur située en des positions porteuses de l'information dans l'encodage convolutif associé, et une erreur *de syndrome* est une erreur située en des positions sur lesquelles se mesure le syndrome. Finalement, une erreur *physique* est une erreur envoyée sur la sortie de l'encodeur convolutif formel.

On est également amené à définir, en parallèle de l'encodeur convolutif formel, l'encodeur convolutif formel inversé, qui se rapporte à l'encodeur formel inversé présenté dans la définition 3.1.3. On utilisera l'encodeur convolutif formel inversé, en exploitant sa propriété de symétrie par rapport à l'encodeur convolutif formel, dans la section 4.4 qui établit des bornes portant sur l'encodeur convolutif formel interne du turbo-encodeur formel. L'encodeur convolutif formel inversé rend compte de l'opération qui consiste à inverser un encodeur convolutif formel, en supprimant les coordonnées de syndrome dans l'erreur de sortie. Cette symétrie entre les structures convolutives directe et inversée est exploitée lorsqu'on fera usage, dans la preuve de la distance minimale d'un turbo-encodeur formel, des caractères *récuratif* et *anti-récuratif* de l'encodeur convolutif formel interne. En effet, une symétrie similaire existe entre le caractère *récuratif* d'un encodeur convolutif formel et le caractère *anti-récuratif* que l'on introduit dans cette thèse. On exploitera cette symétrie afin de déduire, dans la section 4.4, une borne portant sur les encodeurs convolutifs formels anti-récuratifs à partir d'une borne portant sur les encodeurs convolutifs formels récuratifs.

Définition 3.2.1. L'**encodeur convolutif formel** de taille N et de paramètres $[n, k, m]$ où $n \geq k \geq 1$, basé sur un encodeur formel $(n + m, k + m, \mathcal{E})$ appelé **encodeur formel de base**, est l'encodeur formel $(m + Nn, m + Nk, \mathcal{E}_N)$ agissant sur une erreur d'entrée selon le procédé ci-dessous.

Soit une erreur d'entrée $E \in G^{m+Nn}$ écrite sous la forme suivante :

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)})$$

où $M^{(0)}$ est une **erreur de mémoire** de taille m , et pour tout $t \in \llbracket 0; N-1 \rrbracket$ dit **étape** t ou **bloc** t , $L^{(t)}$ et $S^{(t)}$ sont respectivement une **erreur logique** de taille k et une **erreur de syndrome** de taille $n - k$.

Définissons alors la séquence d'**erreurs de mémoire intermédiaires** de taille m $(M^{(t)})_{0 \leq t \leq N}$ et la séquence d'**erreurs physiques** de taille n $(P^{(t)})_{0 \leq t \leq N-1}$ qui vérifient les relations, pour tout $t \in \llbracket 0; N-1 \rrbracket$:

$$(P^{(t)}, M^{(t+1)}) = \mathcal{E}(M^{(t)}, L^{(t)}, S^{(t)})$$

L'erreur de sortie est alors donnée par :

$$\mathcal{E}_N(E) = (P^{(0)}, \dots, P^{(N-1)}, M^{(N)})$$

Pour tout $t \in \llbracket 0; N \rrbracket$, on appelle **erreur intermédiaire à l'étape t** l'erreur obtenue après avoir appliqué t fois l'encodeur formel $\mathcal{E}$, c'est-à-dire :

$$(P^{(0)}, \dots, P^{(t-1)}, M^{(t)}, L^{(t)}, \dots, L^{(N-1)}, S^{(t)}, \dots, S^{(N-1)})$$

En l'occurrence, les erreurs intermédiaires à l'étape 0 et N correspondent respectivement à l'erreur d'entrée et l'erreur de sortie de l'encodeur convolutif formel.

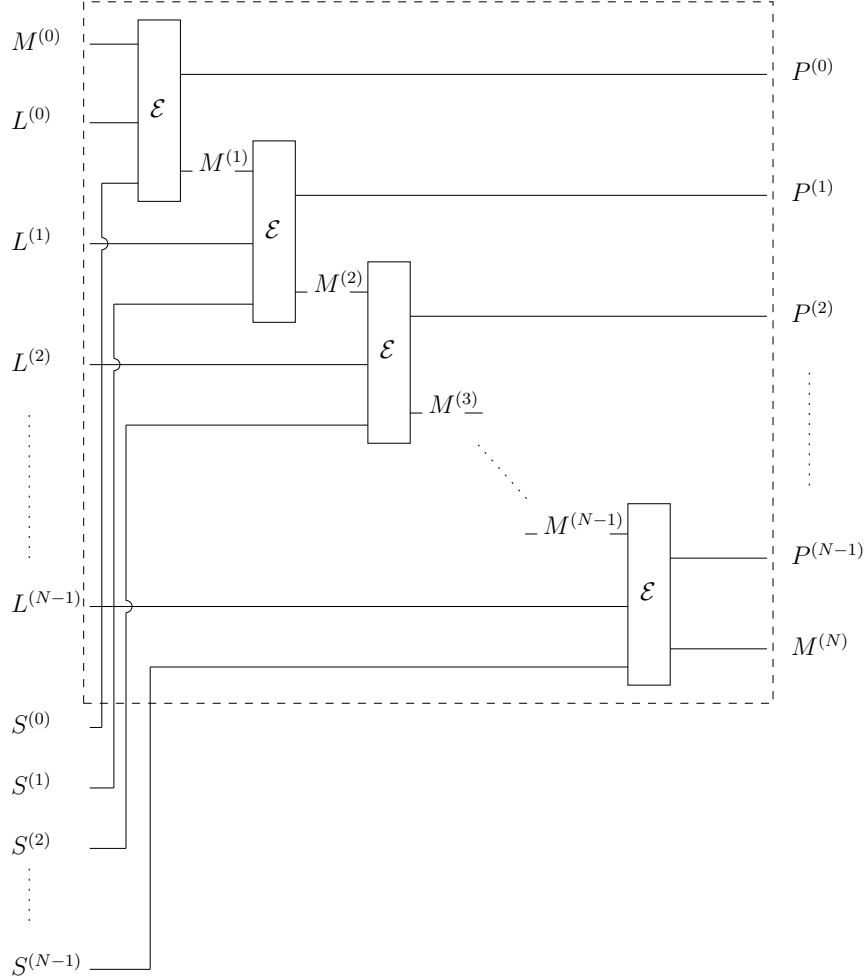


FIGURE 3.1 – Encodeur convolutif formel $\mathcal{E}_N$ de taille N basé sur l'encodeur formel $\mathcal{E}$

D'après la définition précédente, il apparaît que les erreurs intermédiaires sont divisibles en des erreurs de quatre natures différentes. Cela amène à la définition suivante.

Définition 3.2.2. Le **poids physique** $|E|_P$, le **poids logique** $|E|_L$, le **poids de mémoire** $|E|_M$, et le **poids de syndrome** $|E|_{\bar{H}}$ d'une erreur intermédiaire E , ou généralement de toute erreur E , sont respectivement le poids de la partie physique, logique, et de mémoire de E ainsi que le nombre de positions de la partie de syndrome de E appartenant à $\bar{H}$.

Une telle séparation de toute erreur intermédiaire en erreurs de natures différentes s'applique aussi en particulier à l'encodeur formel de base. Ainsi, pour un tel encodeur formel, l'entrée s'écrit sous la forme $E = (M, L, S)$ et la sortie sous la forme $\mathcal{E} = (P, M')$, où M et M' sont des erreurs de mémoire de taille m , L est une erreur logique de taille k , S une erreur de syndrome de poids $n - k$, et P une erreur physique de poids n .

On peut également généraliser à une taille infinie un encodeur convolutif formel. Afin de pouvoir écrire convenablement l'erreur d'entrée, le format de cette dernière est légèrement différent et juxtapose chaque paire d'erreurs logique et de syndrome d'indice associé.

Définition 3.2.3. L'**encodeur convolutif formel infini** de paramètres $[n, k, m]$, basé sur un encodeur formel $(n + m, k + m, \mathcal{E})$, est l'endomorphisme $\mathcal{E}_\infty$ de l'ensemble des suites de G , qui à une erreur d'entrée de taille infinie E écrite sous la forme de la suite :

$$E = (M^{(0)}, L^{(0)}, S^{(0)}, L^{(1)}, S^{(1)}, \dots)$$

où $M^{(0)} \in G^m$ et pour tout $t \in \mathbb{N}$, $L^{(t)} \in G^k$ et $S^{(t)} \in G^{n-k}$, associe l'erreur de sortie :

$$\mathcal{E}_\infty = (P^{(0)}, P^{(1)}, \dots)$$

définie par la relation de récurrence :

$$(P^{(t)}, M^{(t+1)}) = \mathcal{E}(M^{(t)}, L^{(t)}, S^{(t)})$$

L'encodeur formel inversé que l'on a introduit dans la définition 3.1.3 permet lui aussi de construire une structure convolutive, selon la même construction que l'on a définie ci-dessus qui s'applique aux encodeurs formels. Cette structure convolutive inversée permettra aux sections 4.3 et 4.4 d'exploiter l'effet du caractère *anti-récursif* de l'encodeur convolutif formel interne dans la démonstration de la distance minimale d'un turbo-encodeur formel. L'encodeur formel inversé a un intérêt limité à l'étude de l'encodeur convolutif formel interne, et il ne sera pas évoqué à l'échelle du turbo-encodeur formel.

Définition 3.2.4. L'**encodeur convolutif formel inversé** de taille N et de paramètres $[n', k', m']$ basé sur un encodeur formel inversé $(n' + m', k' + m', \mathcal{E}')$ est l'encodeur formel inversé $(m' + Nn', m' + Nk', \mathcal{E}'_N)$ agissant sur une erreur d'entrée selon le procédé ci-dessous.

Soit une erreur d'entrée $E' \in G^{m' + Nn'}$ écrite sous la forme suivante :

$$E' = (M'^{(0)}, L'^{(0)}, \dots, L'^{(N-1)})$$

où $M'^{(0)}$ est une **erreur de mémoire** de taille m' , et pour tout $t \in \llbracket 0; N-1 \rrbracket$, $L'^{(t)}$ est une **erreur logique** de taille k' .

Définissons alors la séquence d'erreurs de mémoire intermédiaires $(M'^{(t)})_{0 \leq t \leq N}$ et la séquence d'**erreurs physiques** de taille n' $(P'^{(t)})_{0 \leq t \leq N}$ qui vérifient les relations, pour tout $t \in \llbracket 0; N-1 \rrbracket$:

$$(P'^{(t)}, M'^{(t+1)}) = \mathcal{E}'(M'^{(t)}, L'^{(t)})$$

L'erreur de sortie est alors donnée par :

$$\mathcal{E}'_N(E') = (P'^{(0)}, \dots, P'^{(N-1)}, M'^{(N)})$$

Définition 3.2.5. L'**encodeur convolutif formel inversé infini** de paramètres $[n', k', m']$, basé sur un encodeur formel inversé $(n' + m', k' + m', \mathcal{E}')$, est l'endomorphisme $\mathcal{E}'_\infty$ de l'ensemble des suites de G , qui à une erreur d'entrée de taille infinie E' écrite sous la forme de la suite :

$$E' = (M'^{(0)}, L'^{(0)}, L'^{(1)} \dots)$$

où $M'^{(0)} \in G^{m'}$ et pour tout $t \in \mathbb{N}$, $L'^{(i)} \in G^{k'}$, associe l'erreur de sortie :

$$\mathcal{E}'_\infty = (P'^{(0)}, P'^{(1)}, \dots)$$

définie par la relation de récurrence :

$$(P'^{(t)}, M'^{(t+1)}) = \mathcal{E}'(M'^{(t)}, L'^{(t)})$$

La définition suivante et le lemme qui s'ensuit trivialement rendent compte du lien entre les structures convolutives directe et inversée.

Définition 3.2.6. Soit un encodeur formel $(n + m, k + m, \mathcal{E})$ où m est implicitement connu. L'**encodeur formel inversé réciproque** de $(n + m, k + m, \mathcal{E})$ est l'encodeur formel inversé $(k + m, n + m, \mathcal{E}')$, où $\mathcal{E}'$ est le morphisme de G^{n+m} dans G^{k+m} défini de la façon suivante. Pour toute erreur d'entrée de l'encodeur formel

$$E = (M, L, S)$$

où $M \in G^m$, $L \in G^k$ et $S \in G^{n-k}$, dont l'image s'écrit

$$\mathcal{E}(E) = (P, M')$$

où $P \in G^n$ et $M' \in G^m$, l'encodeur formel inversé associe alors à l'erreur d'entrée

$$E' = (M', P)$$

l'image

$$\mathcal{E}'(E') = (L, M)$$

Les erreurs M' et $L' = P$ sont respectivement la partie **de mémoire** et la partie **logique** de E' , tandis que les erreurs $P' = L$ et M sont respectivement la partie **physique** et la partie **de mémoire** de $\mathcal{E}'(E')$.

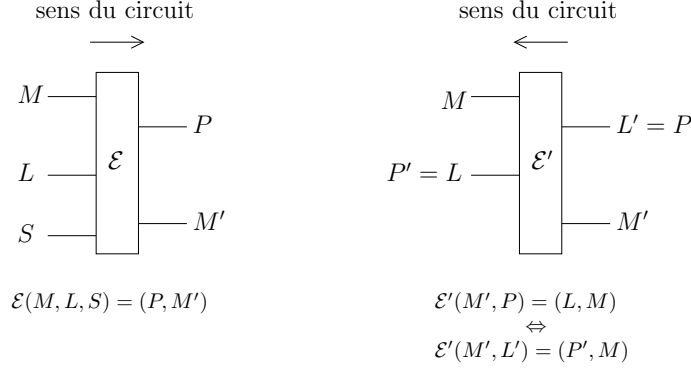


FIGURE 3.2 – Encodeur formel $\mathcal{E}$ et son encodeur formel inversé réciproque $\mathcal{E}'$

Ainsi l'encodeur formel inversé réciproque fait correspondre à l'erreur de sortie de l'encodeur formel son erreur d'entrée privée de sa partie de syndrome, en prenant soin d'inverser les positions des erreurs de mémoire en entrée et en sortie, afin de permettre d'écrire l'opération de convolution qui s'en suit d'une manière proche de l'écriture employée dans le cadre de l'encodeur convolutif formel. Le lemme suivant est alors une conséquence directe des définitions précédentes. Il relie, pour un encodeur convolutif formel basé sur un encodeur formel donné, l'erreur d'entrée privée de la partie de syndrome et l'erreur de sortie aux erreurs respectives de sortie et d'entrée de l'encodeur convolutif formel inversé basé sur l'encodeur formel inversé réciproque.

Lemme 3.2.1. *Soit $(m + Nn, m + Nk, \mathcal{E}_N)$ un encodeur convolutif formel de taille N et de paramètres $[n, k, m]$ basé sur un encodeur formel $(n + m, k + m, \mathcal{E})$. Soit $(k + m, n + m, \mathcal{E}')$ l'encodeur formel inversé réciproque de $(n + m, k + m, \mathcal{E})$, et soit $(m + Nk, m + Nn, \mathcal{E}'_N)$ l'encodeur convolutif formel inversé de taille N et de paramètres $[k, n, m]$ basé sur $(k + m, n + m, \mathcal{E}')$. Pour toute erreur d'entrée de $(m + Nn, m + Nk, \mathcal{E}_N)$, notée sous la forme*

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)})$$

et dont l'image est notée sous la forme

$$\mathcal{E}_N(E) = (P^{(0)}, \dots, P^{(N-1)}, M^{(N)})$$

l'encodeur convolutif formel inversé $(m + Nk, m + Nn, \mathcal{E}'_N)$ associe à l'erreur

$$E' = (M^{(N)}, P^{(N-1)}, \dots, P^{(0)})$$

l'image suivante :

$$\mathcal{E}'_N(E') = (L^{(N-1)}, \dots, L^{(0)}, M^{(0)})$$

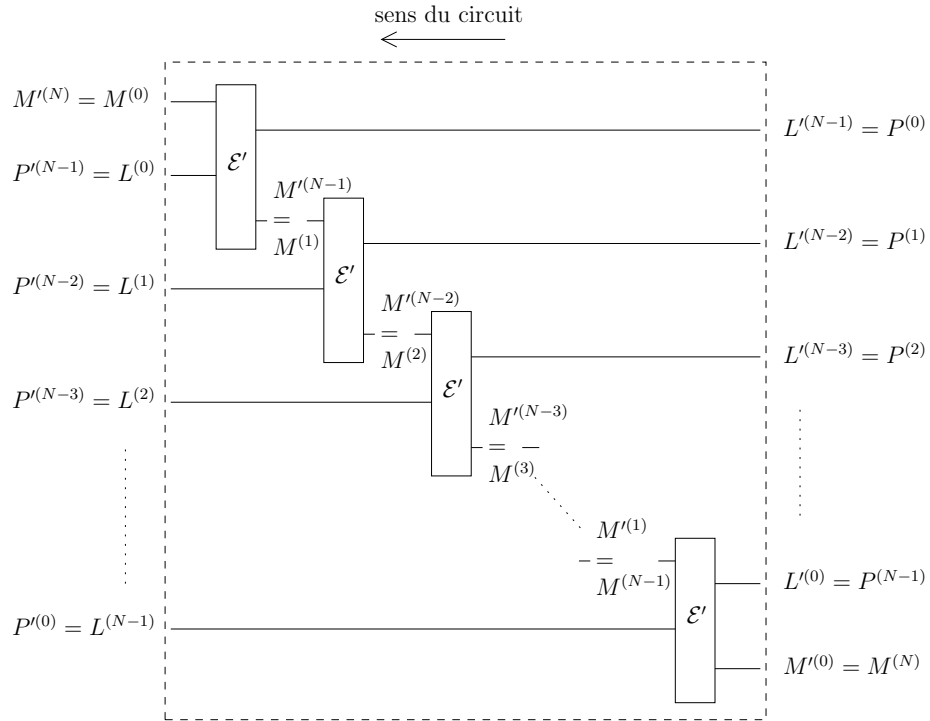


FIGURE 3.3 – Encodeur convolutif formel inversé, réciproque de l'encodeur convolutif formel donné en figure 3.1

Notons que les erreurs notées $L^{(i)}$ et $P^{(i)}$, qui jouent un rôle respectivement logique et physique vis-à-vis de l'encodeur convolutif formel, correspondent aux erreurs respectives $P^{(N-1-i)}$ et $L^{(N-1-i)}$ dans la définition 3.2.4, et ces dernières jouent un rôle respectivement physique et logique vis-à-vis de l'encodeur convolutif formel inversé. L'utilité de ce lemme apparaîtra notamment à la section 4.4 comme précisé en début de section. On introduit à présent l'entrelaceur formel afin de définir ensuite le turbo-encodeur formel.

Définition 3.2.7. On dit qu'un encodeur formel (n, n, Π) est un **entrelaceur formel** de taille n s'il existe une permutation π de l'ensemble $\llbracket 1; n \rrbracket$ et une séquence de permutations $(\pi_i)_{1 \leq i \leq n}$ de l'ensemble G laissant I invariant, tels que pour tout $E = (E_1, \dots, E_n) \in P^n$:

$$\Pi(E) = (\pi_1(E_{\pi(1)}), \dots, \pi_n(E_{\pi(n)}))$$

Un entrelaceur formel est dit **aléatoire** si toutes les permutations π et $(\pi_i)_{1 \leq i \leq n}$ qui le définissent sont choisies selon une loi de probabilité uniforme.

On peut à présent définir un turbo-encodeur formel. Ce dernier se base sur un encodeur formel externe de taille n_{ext} employé en parallèle N_{ext} fois pour produire une erreur de taille $n_{ext}N_{ext}$, suivi d'un entrelaceur formel de taille $n_{ext}N_{ext}$, puis d'un encodeur convolutif formel de taille N_{in} et de paramètres $[n_{in}, k_{in}, m_{in}]$ tel que $m_{in} + N_{in}n_{in} = n_{ext}N_{ext}$.

Définition 3.2.8. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ et $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ deux encodeurs formels nommés respectivement **encodeur formel externe** et **encodeur formel interne**, où l'on suppose la valeur de m_{in} implicitement connue. Soit également (N_{ext}, N_{in}) une paire d'entiers vérifiant :

$$n_{ext}N_{ext} = m_{in} + k_{in}N_{in}$$

et soit $(n_{ext}N_{ext}, n_{ext}N_{ext}, \Pi)$ un entrelaceur formel. Le **turbo-encodeur formel de taille N_{in}** , basé sur les encodeurs formels externe et interne et l'entrelaceur formel ci-dessus, est l'encodeur formel $(m_{in} + n_{in}N_{in}, k_{ext}N_{ext}, T_{N_{in}})$ dont l'action sur une erreur d'entrée est décrite par le procédé ci-dessous.

Soit une erreur d'entrée $E \in G^{m_{in} + n_{in}N_{in}}$ écrite sous la forme suivante :

$$E = (L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)}, S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

où pour tout $i \in \llbracket 0, N_{ext} - 1 \rrbracket$, $L_{ext}^{(i)}$ et $S_{ext}^{(i)}$ sont respectivement une erreur logique de taille k_{ext} et une erreur de syndrome de taille $n_{ext} - k_{ext}$, et pour tout $i \in \llbracket 0, N_{in} - 1 \rrbracket$, $S_{in}^{(i)}$ est une erreur de syndrome de taille $n_{in} - k_{in}$.

- On applique d'abord N_{ext} fois l'encodeur formel externe sur chaque couple $(L_{ext}^{(i)}, S_{ext}^{(i)})$. A la suite de cette première opération on obtient l'erreur nommée **erreur intermédiaire externe** :

$$E_{ext} = (P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

où pour tout $i \in \llbracket 0, N_{ext} - 1 \rrbracket$, $P_{ext}^{(i)} = \mathcal{E}_{ext}(L_{ext}^{(i)}, S_{ext}^{(i)})$.

- On applique ensuite l'entrelaceur formel sur l'erreur physique :

$$(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$$

Cette erreur étant de taille $n_{ext}N_{ext} = m_{in} + k_{in}N_{in}$, on peut écrire la sortie de l'entrelaceur sous la forme :

$$\Pi(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)}) = (M_{in}^{(0)}, L_{in}^{(0)}, \dots, L_{in}^{(N_{in}-1)})$$

où $M_{in}^{(0)}$ est une erreur de mémoire de taille m_{in} , et pour tout $i \in \llbracket 0; N_{in} - 1 \rrbracket$, $L_{in}^{(i)}$ est une erreur logique de taille k_{in} . Ainsi après application de l'entrelaceur formel, on obtient l'erreur nommée **erreur intermédiaire interne** :

$$E_{in} = (M_{in}^{(0)}, L_{in}^{(0)}, \dots, L_{in}^{(N_{in}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

- Finalement, on applique l'encodeur convolutif formel interne

$$(m_{in} + n_{in}N_{in}, m_{in} + k_{in}N_{in}, \mathcal{E}_{inN_{in}})$$

de taille N_{in} à l'erreur intermédiaire interne :

$$T_{N_{in}}(E) = \mathcal{E}_{inN_{in}}(E_{in}) = (P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$$

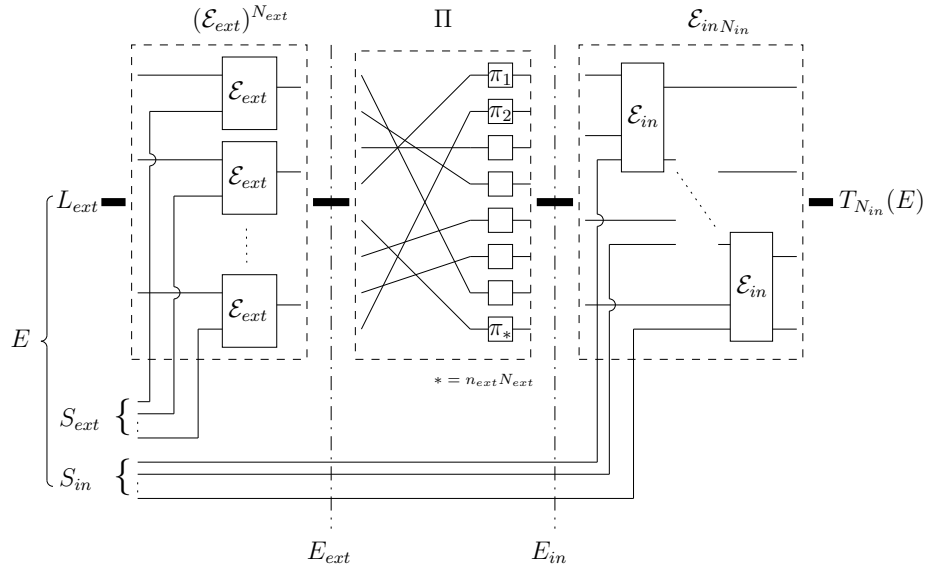


FIGURE 3.4 – Turbo-encodeur formel

Dans cette définition, les étapes d'application de l'encodeur formel externe et de l'encodeur formel interne requièrent chacune d'utiliser une erreur de syndrome propre à cette étape. On a ainsi indexé par *ext* et *in* les erreurs utilisées respectivement en entrée de l'encodeur formel externe et de l'encodeur convolutif formel interne.

Il faut remarquer que, comme l'entrelaceur formel conserve le poids, l'erreur intermédiaire externe a un poids physique identique à la somme des poids logique et de mémoire de l'erreur intermédiaire interne. Soulignons également le fait que, étant donnés un encodeur formel externe $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ et un encodeur formel interne $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ où le paramètre m_{in} est implicitement connu, l'équation $n_{ext}N_{ext} = m_{in} + k_{in}N_{in}$ d'inconnues N_{ext} et N_{in} admet un nombre nul ou infini de solutions. Ainsi, s'il est possible de construire un turbo-encodeur formel à partir de ces deux encodeurs formels, il est possible d'en construire une infinité. L'étude asymptotique de la distance minimale d'un turbo-encodeur formel se fera sur une telle suite indexée par le paramètre de taille N_{in} .

3.3 Propriétés recherchées pour un encodeur convolutif formel

Dans un turbo-encodage classique, les propriétés que satisfait l'encodage convolutif interne sont le caractère récursif et le caractère systématique afin d'obtenir une distance minimale conforme au théorème 1.11.2. Etant donné que le caractère systématique implique le caractère non catastrophique, le bon fonctionnement de l'algorithme de décodage itératif est alors garanti de fait. Dans le cadre du formalisme commun, on va redéfinir ces trois propriétés, et introduire une quatrième propriété nommée *anti-récursivité*.

Les définitions présentées établissent un lien entre le poids en entrée et le poids en sortie de l'encodeur convolutif formel, et se restreignent à l'ensemble des erreurs en entrée dont le poids de syndrome est nul. Cela permet ainsi dans le modèle classique de se restreindre aux erreurs qui s'écrivent comme un élément du code. Dans le modèle de l'encodage quantique, les définitions se déclinent en revanche de manière plus complexe dans la mesure où on autorise à l'erreur en entrée d'avoir des coordonnées égales à I ou Z sur les positions du syndrome, puisque ces erreurs ont un effet trivial. Par conséquent, le problème de rechercher un encodeur convolutif formel possédant de telles propriétés devient plus difficile dans le modèle quantique.

Exposons les motivations ayant conduit à l'introduction de la notion d'*anti-récursivité*. Il a été prouvé dans [41] que dans le cadre quantique, un encodeur convolutif récursif est nécessairement catastrophique, et donc a fortiori non systématique. Or, dans le cadre classique, le caractère systématique de l'encodeur convolutif interne garantit que toute erreur de sortie de poids d du turbo-encodeur est nécessairement l'image par l'encodeur convolutif interne d'une erreur de poids inférieur à d ; cette propriété permet ainsi de contrôler en partie

le nombre d'erreurs de poids d donné en sortie, et représente avec le caractère récursif un élément essentiel de la preuve de la distance minimale du turbo-encodeur. Sans le caractère systématique dans le cadre quantique, il devient ainsi nécessaire de rechercher une contrainte alternative si l'on veut mettre en place un modèle de turbo-encodeur dont la distance minimale correspond à celle du modèle classique. Cette thèse introduit donc le concept d'encodeur convolutif formel *anti-récursif* qui permet, en se combinant avec le caractère récursif, de pallier à l'absence du caractère systématique. Cette notion ressemble au premier abord à une notion d'encodeur convolutif formel récursif dans laquelle les rôles de l'entrée logique et de la sortie sont inversés. L'idée ayant permis d'introduire cette notion est née en recherchant un encodeur convolutif quantique récursif dont l'encodeur de base est le plus simple possible. Cette recherche a abouti à l'encodeur de base que l'on donne dans la section 3.5, est de taux $k/n = 1$ et ne possède donc pas de positions de syndrome, et a la propriété élégante d'être égal à son propre inverse à une interversion près sur ses positions d'entrée logique. L'introduction de la notion d'anti-récursivité est ainsi née en cherchant à exploiter cette propriété de symétrie. Cependant, à la différence du caractère récursif, le caractère anti-récursif laisse toute liberté à la valeur de l'erreur en entrée aux positions du syndrome. La nécessité de relâcher la contrainte en ces positions est un point technique que l'on mettra en évidence dans la preuve de la distance minimale du turbo-encodeur formel.

L'impossibilité de réunir à la fois les caractères récursif et non catastrophique dans le modèle quantique pose quant à lui problème pour l'algorithme de décodage, et on ne peut pallier au caractère catastrophique de l'encodeur convolutif qu'en proposant un changement du schéma général du turbo-encodeur comme on le présente dans la section 3.5. L'incompatibilité de ces deux caractères n'a donc pas de répercussion dans cette section, ni dans la distance minimale d'un turbo-encodeur formel.

Définition 3.3.1. Un encodeur convolutif formel de taille N et paramètres $[n, k, m]$, basé sur un encodeur formel $(n + m, k + m, \mathcal{E})$, est dit :

- **systématique** si, pour toute erreur d'entrée $E \in G^{m+Nn}$ de poids de syndrome nul, $|\mathcal{E}_N(E)| \geq |E|_L$.
- **récursif** si, pour toute erreur $E \in G^{\mathbb{N}}$ d'entrée de l'encodeur infini $\mathcal{E}_{\infty}$ de poids de syndrome nul, si $|E|_M = 0$ et $|E|_L = 1$ alors $|\mathcal{E}_{\infty}(E)| = \infty$.
- **catastrophique** s'il existe une erreur $E \in G^{\mathbb{N}}$ d'entrée de l'encodeur infini $\mathcal{E}_{\infty}$ de poids de syndrome nul, vérifiant $|\mathcal{E}_{\infty}(E)| < \infty$ et $|E|_L = \infty$.
- **anti-récursif** si, pour toute erreur $P \in G^n$ de poids 1, le poids de $\mathcal{E}_N^{-1}(I_n, \dots, I_n, P, I_m)$ tend vers l'infini lorsque $N \rightarrow \infty$.

Contrairement à la définition 1.7.3, on n'a conservé ici dans la définition du caractère systématique que la propriété reliant le poids de l'entrée logique au poids de la sortie. Le fait que la partie logique de l'entrée se trouve recopiée dans la sortie n'est en effet pas nécessaire afin d'obtenir le résultat donné au

théorème 3.4.1. Observons également que lorsqu'on qualifie un encodeur formel $(n + m, k + m, \mathcal{E})$ d'un des adjectifs précédents, on suppose implicitement la valeur du triplet $[n, k, m]$ connue, ou plus simplement la valeur de m connue. Cela sera le cas à chaque fois qu'un encodeur formel sera écrit sous la forme $(n + m, k + m, \mathcal{E})$.

Etendons la définition du caractère récursif aux encodeurs convolutifs formels inversés.

Définition 3.3.2. Un encodeur convolutif formel inversé de paramètres $[n, k, m]$, basé sur un encodeur formel inversé $(n' + m', k' + m', \mathcal{E}')$, est dit **récursif** si, pour toute erreur $E' \in G^{\mathbb{N}}$ d'entrée de l'encodeur inversé infini $\mathcal{E}_{\infty}$, si $|E'|_M = 0$ et $|E'|_L = 1$ alors $|\mathcal{E}'_{\infty}(E')| = \infty$.

Le lemme suivant exprime l'équivalence entre le caractère anti-récursif d'un encodeur convolutif formel et le caractère récursif de son encodeur convolutif formel inversé réciproque. Il découle immédiatement des définitions précédentes et du lemme 3.2.1, qui relie la structure convolutive directe et la structure convolutive inversée.

Lemme 3.3.1. *Un encodeur formel $(n + m, k + m, \mathcal{E})$, où m est implicitement connu, est anti-récursif si et seulement si l'encodeur formel inversé réciproque de $(n + m, k + m, \mathcal{E})$ est récursif.*

Ce lemme sera l'argument de base lors de la démonstration de la seconde borne de la section 4.4 portant sur les encodeurs convolutifs internes anti-récursifs.

3.4 Distance minimale d'un turbo-encodeur formel

On énonce dans cette section un résultat principal de cette thèse, concernant la distance minimale que vérifie un turbo-encodeur formel s'il est choisi selon certaines conditions. La preuve de ce résultat est longue ; on y consacre le chapitre 4 dans lequel apparaissent à la fois la preuve de la distance minimale dans le cas classique tel que démontré par Kahale et Urbanke [22] et la preuve dans le modèle quantique introduit dans cette thèse. Le résultat se décline en deux théorèmes qui se déclinent de manière naturelle respectivement sur le modèle de turbo-encodage classique et le modèle de turbo-encodage quantique. Afin d'alléger les notations, on notera N la valeur de N_{int} dans la construction du turbo-encodeur formel.

Théorème 3.4.1. *Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ un encodeur formel externe de distance dégénérée d_{deg} , $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ un encodeur formel interne récursif et systématique. Si $d_{deg} > 2$, alors pour tout $\alpha < \frac{d_{deg}-2}{d_{deg}}$, et avec probabilité tendant vers 1 lorsque N tend vers l'infini, la distance minimale d'un*

turbo-encodeur formel de taille $N_{in} = N$, basé sur ces encodeurs formels externe et interne et un entrelaceur formel aléatoire de taille N , est plus grande que

$$N^\alpha$$

Théorème 3.4.2. *On considère le cas où $|G| \geq 4$. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ un encodeur formel externe de distance minimale d_{min} et de distance dégénérée d_{deg} , $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ un encodeur formel interne récursif et anti-récursif. Si $d_{deg} > 2$, alors pour tout $\alpha < \frac{d_{deg}-2}{d_{deg}}$, et avec probabilité tendant vers 1 lorsque N tend vers l'infini, la distance minimale d'un turbo-encodeur formel de taille $N_{in} = N$, basé sur ces encodeurs formels externe et interne et un entrelaceur formel aléatoire de taille N , est plus grande que*

$$N^\alpha$$

Si en revanche, $d_{deg} = 2$ et $d_{min} > d_{deg}$, alors pour tout $\alpha < d_{min} - 2$, et avec probabilité tendant vers 1 lorsque N tend vers l'infini, ce même turbo-encodeur formel a une distance minimale plus grande que

$$\alpha \frac{\log N}{\log \log N}$$

3.5 Constructions expérimentales d'un turbo-encodeur quantique

On décrit dans cette section diverses constructions se basant sur un turbo-encodeur quantique imaginés dans le cadre de cette thèse. Les décodages de ces différents montages ont été simulés sur un canal quantique dépolarisant, d'intensité de dépolarisation p modulable selon les tests. On rappelle qu'un tel canal introduit en chaque position une erreur X , Y ou Z avec une probabilité de $p/3$ pour chacune de ces erreurs, et l'erreur triviale I avec une probabilité de $1 - p$. Les décodages ont été effectués grâce à un algorithme de décodage itératif inspiré de celui présenté au chapitre 1 et présenté dans la section suivante.

On présente également le résultat des simulations de ces montages. Soulignons que l'algorithme de décodage itératif employé évalue, au niveau de l'encodeur externe, l'erreur physique la plus vraisemblable et non la classe d'erreurs la plus vraisemblable modulo le groupe stabilisateur S . Ainsi il ne permet pas de tirer pleinement profit de la caractéristique quantique du décodage au maximum de classe de vraisemblance 2.12.2. Des améliorations dans ce sens n'ont pas été étudiées dans le cadre de cette thèse mais représentent une piste intéressante d'investigation. Un article de Pelchat et Poulin [38] présente un algorithme de complexité linéaire permettant de calculer la classe d'erreurs la plus vraisemblable pour un encodeur convolutif quantique ; cependant, il n'est pas évident d'incorporer cette information dans un algorithme de décodage itératif, dans lequel l'information échangée entre les encodeurs externe et interne est une distribution de probabilité indépendante en chaque position de l'erreur intermédiaire.

Par conséquent, la performance du décodage des différents montages est à confronter à la limite théorique du niveau de bruit permettant de connaître l'erreur la plus vraisemblable, désignée par les termes *capacité cohérente* ou *hashing bound*. Cette limite se formule par l'équivalent quantique du théorème de codage en présence de bruit de Shannon [42], donné par le théorème Lloyd-Shor-Devetak [51] démontré séparément par Lloyd [29], Shor [47] et Devetak [15]. Elle affirme dans le cas du canal dépolarisant que la capacité cohérente vaut :

$$1 - H(p)$$

où $H(p)$ désigne l'entropie du symbole d'erreur de Pauli I , X , Y ou Z en sortie du canal :

$$H(p) = \sum_{E \in \{I, X, Y, Z\}} -p(E) \log_2(p(E)) = -p \log_2(p) - (1-p) \log_2(1-p) + p \log_2(3)$$

L'intensité de dépolarisation p maximale du canal pour laquelle cette capacité est positive vaut environ 0,1893. Comme cette limite correspond à la borne supérieure pouvant être atteinte lorsque le décodage est effectué au maximum de vraisemblance et non au maximum de classe de vraisemblance, elle peut théoriquement être dépassée, même si cela s'avère difficile en pratique. En témoigne par exemple [48], dans lequel un protocole d'encodage et de décodage permet de communiquer jusqu'à $p = 0,1903$. Le tableau 3.1 donne quelques valeurs de la probabilité d'erreur p maximale du canal pouvant être tolérée en fonction du taux d'encodage. Les variations intermédiaires peuvent être approchées par une fonction affine.

Taux d'encodage	p_{max} donné par la <i>hashing bound</i> , arrondi au millième
0	0,189
1/8	0,157
1/7	0,152
1/6	0,147
1/5	0,139

TABLE 3.1 – intensité maximale de dépolarisation du canal tolérée par la *hashing bound* en fonction du taux d'encodage

3.5.1 Schéma de turbo-encodeur modifié

Dans le schéma de concaténation en série conforme à la définition 3.2.8, le décodage itératif ne peut être amorcé à cause du phénomène suivant. Au niveau de l'encodeur convolutif interne, le caractère catastrophique assure l'existence d'un motif d'erreur de poids borné en sortie, qui ne laisse aucune trace sur les positions du syndrome et qui affecte, lorsque la taille de cet encodeur convolutif

tend vers l'infini, une infinité de positions de l'entrée logique. Ainsi, ce motif d'erreur ainsi que l'ensemble de ses translatés pénalisent la première itération du décodage itératif : la distribution de probabilité de l'erreur, calculée en chaque position logique de l'entrée de l'encodeur convolutif interne, tend en moyenne vers une distribution de probabilité uniforme lorsque la taille de l'encodeur convolutif interne tend vers l'infini. Cette distribution est envoyée au niveau des positions de l'erreur de sortie de l'encodeur externe. La distribution de probabilité marginale calculée à son tour par l'encodeur externe en chaque position de sa sortie est alors elle aussi uniforme. Ainsi, à la fin de la première itération du décodage itératif, aucune information n'a été acquise sur les différentes erreurs en jeu dans le turbo-encodeur, et ce même phénomène se répète à toutes les itérations suivantes.

Afin de pallier au caractère catastrophique de l'encodeur interne, le schéma du turbo-encodeur est légèrement modifié. En chacun des N_{ext} blocs de l'encodeur externe, une position de l'erreur de sortie est directement envoyée vers le canal, tandis que les $n_{ext} - 1$ positions restantes entrent dans l'entrelaceur avant d'être envoyées vers l'encodeur interne. Le turbo-encodeur ainsi construit correspond à la définition suivante, dans laquelle $n_c = 1$ désigne le nombre de positions en sortie de chaque bloc de l'encodeur externe envoyées directement vers le canal. On l'appelle *turbo-encodeur modifié* et il est représenté dans la figure 3.5. Les erreurs qui interviennent aux différentes étapes sont schématisées dans la figure 3.6.

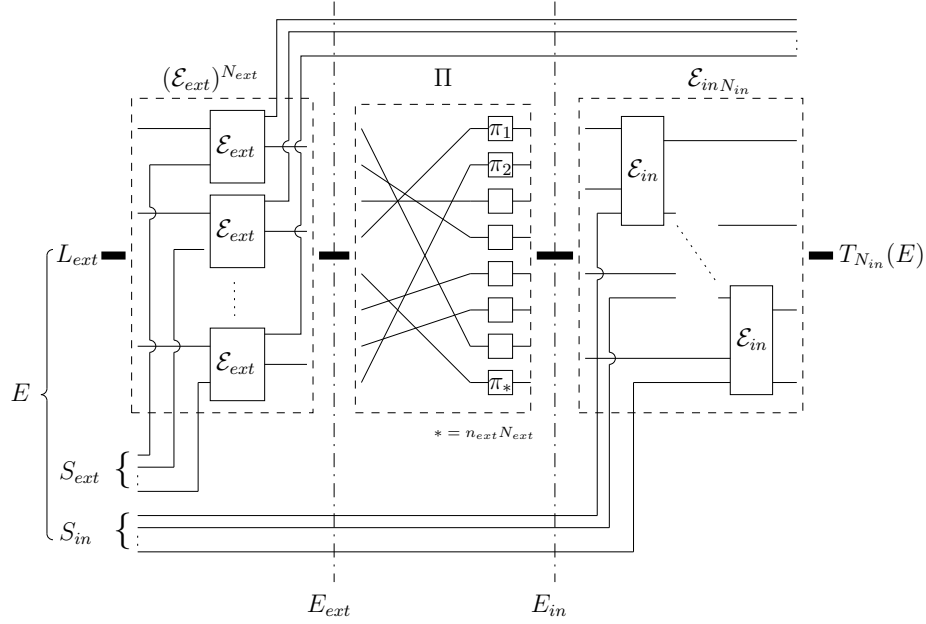


FIGURE 3.5 – Turbo-encodeur modifié

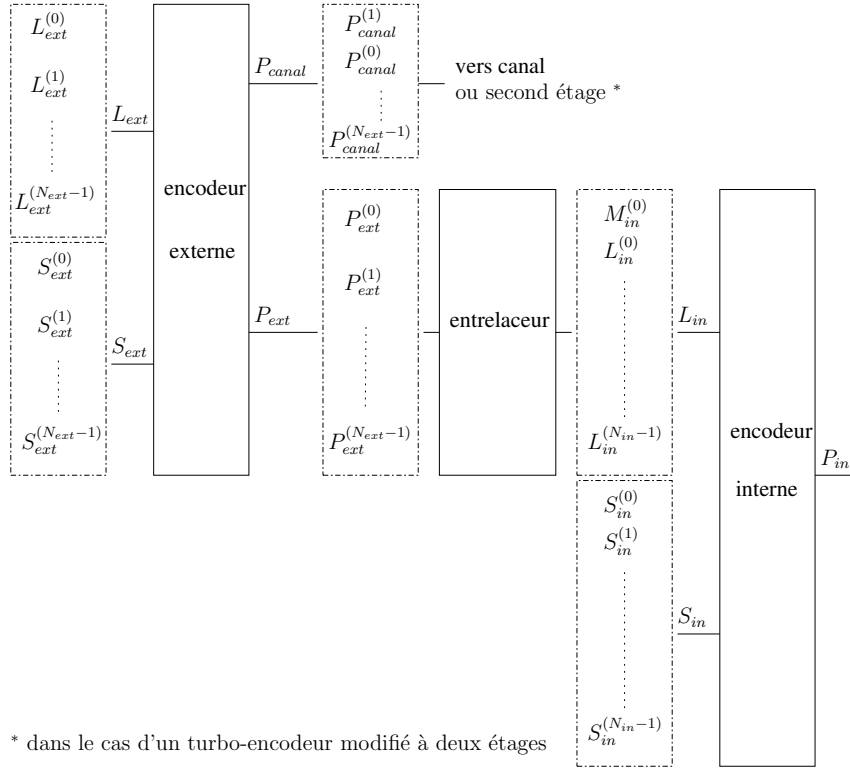


FIGURE 3.6 – Erreurs en jeu dans le turbo-encodeur modifié

Définition 3.5.1. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ et $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ deux encodeurs nommés respectivement **encodeur externe** et **encodeur interne**, où l'on suppose la valeur de m_{in} implicitement connue. Soit $n_c \in \llbracket 1; n_{ext} - 1 \rrbracket$. Soit également (N_{ext}, N_{in}) une paire d'entiers vérifiant :

$$(n_{ext} - n_c)N_{ext} = m_{in} + k_{in}N_{in}$$

et soit $((n_{ext} - n_c)N_{ext}, (n_{ext} - n_c)N_{ext}, \Phi)$ un entrelaceur.

Le **turbo-encodeur modifié de taille N_{in} et de paramètre de canal n_c** , basé sur les encodeurs externe et interne et l'entrelaceur ci-dessus, est l'encodeur $(m_{in} + n_{in}N_{in}, k_{ext}N_{ext}, T_{N_{in}})$ dont l'action sur une erreur d'entrée est décrite par le procédé ci-dessous.

Soit une erreur d'entrée $E \in G^{m_{in} + n_{in}N_{in}}$ écrite sous la forme suivante :

$$E = (L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)}, S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

où pour tout $i \in \llbracket 0, N_{ext} - 1 \rrbracket$, $L_{ext}^{(i)}$ et $S_{ext}^{(i)}$ sont respectivement une erreur logique de taille k_{ext} et une erreur de syndrome de taille $n_{ext} - k_{ext}$, et pour tout $i \in \llbracket 0, N_{in} - 1 \rrbracket$, $S_{in}^{(i)}$ est une erreur de syndrome de taille $n_{in} - k_{in}$.

- On applique d'abord N_{ext} fois l'encodeur externe sur chaque couple $(L_{ext}^{(i)}, S_{ext}^{(i)})$.
A la suite de cette première opération on obtient l'erreur nommée **erreur intermédiaire externe** :

$$E_{ext} = (P_{canal}^{(0)}, P_{ext}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)}, P_{ext}^{(N_{ext}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

où pour tout $i \in \llbracket 0, N_{ext} - 1 \rrbracket$, $P_{canal}^{(i)}$ et $P_{ext}^{(i)}$ sont les erreurs de tailles respectives n_c et $n_{ext} - n_c$ telles que $(P_{canal}^{(i)}, P_{ext}^{(i)}) = \mathcal{E}_{ext}(L_{ext}^{(i)}, S_{ext}^{(i)})$.

- On applique ensuite l'entrelaceur sur l'erreur physique suivante :

$$(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$$

Cette erreur étant de taille $(n_{ext} - n_c)N_{ext} = m_{in} + k_{in}N_{in}$, on peut écrire la sortie de l'entrelaceur sous la forme :

$$\Phi(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)}) = (M_{in}^{(0)}, L_{in}^{(0)}, \dots, L_{in}^{(N_{in}-1)})$$

où $M_{in}^{(0)}$ est une erreur de mémoire de taille m_{in} , et pour tout $i \in \llbracket 0; N_{in} - 1 \rrbracket$, $L_{in}^{(i)}$ est une erreur logique de taille k_{in} . Ainsi après application de l'entrelaceur, on obtient l'erreur nommée **erreur intermédiaire interne** :

$$E_{in} = (M_{in}^{(0)}, L_{in}^{(0)}, \dots, L_{in}^{(N_{in}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N_{in}-1)})$$

- On applique l'encodeur convolutif interne

$$(m_{in} + n_{in}N_{in}, m_{in} + k_{in}N_{in}, \mathcal{E}_{inN_{in}})$$

de taille N_{in} à l'erreur intermédiaire interne :

$$= \mathcal{E}_{inN_{in}}(E_{in}) = (P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$$

L'image $T_{N_{in}}(E)$ de E par le turbo-encodeur est alors l'erreur :

$$T_{N_{in}}(E) = (P_{canal}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)}, P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$$

On a intérêt, dans un tel schéma, à définir les *distances épurées* minimale et dégénérée, ce qui sera notamment utile dans la discussion qui suit, afin d'éclairer la construction à deux étages qui sera proposée ainsi que le choix de l'encodeur externe. Ces distances s'appliquent à la fois au turbo-encodeur modifié et à l'encodeur externe. Elles nécessitent au préalable de manipuler le groupe stabilisateur S et le groupe commutateur $N(S)$ de chacun de ces encodeurs.

Définition 3.5.2. Soit un turbo-encodeur modifié, de longueur N_{in} et paramètre de canal n_c . Si S et $N(S)$ désignent les groupes stabilisateur et commutateur du turbo-encodeur modifié, notons $\tilde{S}$ et $\tilde{N}(S)$ les sous-groupes respectifs de S et $N(S)$ constitués des erreurs

$$(P_{canal}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)}, P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$$

telles que $(P_{canal}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)}) = I_{n_c(N_{ext}-1)}$.

La **distance épurée minimale** du turbo-encodeur modifié est le plus petit poids d'un élément de $N(\bar{S}) \setminus \bar{S}$; sa **distance épurée dégénérée** est le plus petit poids d'un élément de $N(\bar{S}) \setminus \{I_{n_{in}N_{in}+m_{in}}\}$.

De même, si S' et $N(S')$ désignent les groupes stabilisateur et commutateur du turbo-encodeur modifié, notons $\bar{S}'$ et $N(\bar{S}')$ les sous-groupes respectifs de S' et $N(S')$ constitués des erreurs

$$(P_{canal}, P_{in})$$

telles que $P_{canal} = I_{n_c}$.

La **distance épurée minimale** de l'encodeur externe est le plus petit poids d'un élément de $N(\bar{S}') \setminus \bar{S}'$; sa **distance épurée dégénérée** est le plus petit poids d'un élément de $N(\bar{S}') \setminus \{I_{n_{in}N_{in}+m_{in}}\}$.

On considérera toujours dans la suite que le paramètre de canal vaut $n_c = 1$.

3.5.2 Propriétés et choix de l'encodeur externe

La modification du schéma du turbo-encodeur a pour objectif d'introduire de l'information de la part du canal auprès de l'encodeur externe afin d'amorcer la première itération du décodage itératif. Pour satisfaire cet objectif, il faut concevoir un encodeur $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ possédant une propriété particulière. Celui-ci doit posséder un stabilisateur tel que, lors de la mesure du syndrome, l'information sur l'erreur connue sur la première position en sortie de $\mathcal{E}_{ext}$ permet de déduire de l'information sur l'une des autres positions en sortie de $\mathcal{E}_{ext}$ qui seront envoyées vers l'encodeur interne. Il faut pour cela que le groupe stabilisateur lié à l'encodeur $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ contienne une erreur $\tilde{E}$ de poids 2, faisant intervenir la première position de sortie et une autre position.

Cependant, si cette première condition est réalisée, la partie de l'erreur $\tilde{E}$ envoyée vers l'encodeur interne est de poids 1. L'encodeur externe agit alors vis-à-vis de la structure de turbo-encodage comme un encodeur de distance dégénérée égale à 1. Le turbo-encodeur a alors une distance minimale bornée, ce qui nuit aux performances asymptotiques de décodage du turbo-encodeur. Ce problème peut être surmonté grâce à une seconde idée : pour chaque bloc de l'encodeur externe, l'idée consiste à envoyer la première position de sortie précédemment notée $P_{canal}^{(i)}$ non pas directement vers le canal, mais en entrée d'un second encodeur, dans le but de réduire fortement le taux d'erreur en cette position. Dite autrement, l'idée consiste à concaténer au turbo-encodeur un étage supplémentaire encodant les positions

$$(P_{canal}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)})$$

et laissant intactes les positions

$$(P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$$

On représente une telle construction à deux étages dans la figure 3.7. Afin d'alléger cette figure, les positions de syndrome de l'encodeur au second étage ne sont pas représentées.

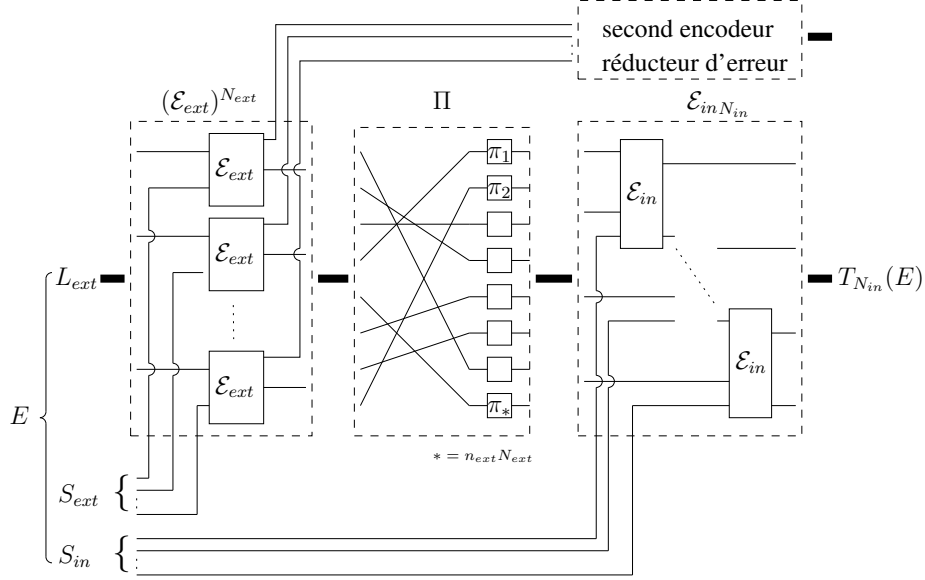


FIGURE 3.7 – Turbo-encodeur modifié, assisté par un encodeur réducteur d'erreur

Afin de guider le raisonnement, supposons dans un premier temps qu'une telle concaténation permet de connaître précisément l'erreur en chaque position envoyée vers le second encodeur. Une fois effectué le décodage au niveau du second encodeur, il reste, pour connaître l'erreur en sortie du turbo-encodeur modifié, une incertitude modulo le sous-groupe $N(S)$ défini en 3.5.2. Cela équivaut, pour chaque bloc de l'encodeur externe, à une incertitude modulo $N(S')$, sur les $n_{ext} - 1$ positions $P_{ext}^{(i)}$. Cette équivalence se traduit également par le lien suivant entre les distances épurées de l'encodeur externe et la distance épurée minimale du turbo-encodeur modifié. Ce lien est une conséquence du théorème 3.4.2.

Lemme 3.5.1. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ un encodeur externe de distance épurée minimale d_{min} et de distance épurée dégénérée d_{deg} , $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ un encodeur interne récursif et anti-récursif. Si $d_{deg} > 2$, alors pour tout $\alpha < \frac{d_{deg}-2}{d_{deg}}$, et avec probabilité tendant vers 1 lorsque N tend vers l'infini, la distance épurée minimale d'un turbo-encodeur modifié de taille $N_{in} = N$, basé sur ces encodeurs externe et interne et un entrelaceur aléatoire de taille N , est plus grande que

$$N^\alpha$$

Si en revanche, $d_{deg} = 2$ et $d_{min} > d_{deg}$, alors pour tout $\alpha < d_{min} - 2$, et avec

probabilité tendant vers 1 lorsque N tend vers l'infini, ce même turbo-encodeur modifié a une distance épurée minimale plus grande que

$$\propto \frac{\log N}{\log \log N}$$

Les conditions recherchés chez l'encodeur externe se formalisent donc ainsi.

Problème 1. On pose $n_c = 1$ pour le paramètre de canal. Rechercher un encodeur $(n, k, \mathcal{E})$, de groupe stabilisateur S et de distances épurées minimale d_{min} et dégénérée d_{deg} , qui possède les propriétés suivantes :

- (i) Il existe $\tilde{E} \in S$ de poids 2 et de première coordonnée différente de I .
- (ii) $d_{deg} \geq 3$, ou $d_{deg} = 2$ et $d_{min} \geq 3$.

La condition (i) permet l'amorce du décodage itératif de la manière suivante. Supposons sans perte de généralité que $\tilde{E}$ est constituée d'un X sur les deux premières positions et de $n - 2$ I sur les positions restantes, et notons E l'erreur qui correspond à la sortie d'un des blocs de l'encodeur externe. En mesurant le syndrome de E , on peut connaître le produit symplectique des représentations symplectiques de E et de $\tilde{E}$; en d'autres termes, on sait si les erreurs de Pauli réelles que représentent E et $\tilde{E}$ commutent ou anti-commutent. Si ce produit symplectique est égal à 0, cela signifie que l'on a soit $E_1 \in \{I, X\}$ et $E_2 \in \{I, X\}$, soit $E_1 \in \{Y, Z\}$ et $E_2 \in \{Y, Z\}$. S'il est en revanche égal à 1, cela signifie que l'on a soit $E_1 \in \{I, X\}$ et $E_2 \in \{Y, Z\}$, soit $E_1 \in \{Y, Z\}$ et $E_2 \in \{I, X\}$. Sachant que l'on connaît la distribution de probabilité de la coordonnée E_1 de E , qui correspond à la distribution de probabilité donnée par le canal, la première étape du décodage itératif permet ainsi de déduire une distribution de probabilité non uniforme sur la coordonnée E_2 , qui servira d'information extrinsèque pour l'encodeur interne lors de la deuxième étape du décodage itératif.

D'autre part, la condition (ii) garantit que le turbo-encodeur modifié possède une *distance épurée minimale* non bornée. Sous caution que le second étage appliqué aux positions $P_{canal}^{(i)}$ permette une forte réduction du taux d'erreur par position, cette distance épurée minimale non bornée se reflète alors en termes de performance de décodage.

En recherchant des encodeurs qui satisfont le problème 1, deux encodeurs ont été trouvés. Les résultats des simulations que l'on étaye ci-dessous montrent que les deux turbo-encodeurs de paramètre de canal 1 basés sur ces encodeurs externes respectifs sont de très bons *réducteurs* d'erreurs. Ainsi, dans une construction de turbo-encodeur modifié à deux étages telle que dans la figure 3.7, ils peuvent être assemblés pour faire à la fois office d'encodeur réducteur d'erreur au second étage et de turbo-encodeur modifié au premier étage.

Le premier encodeur trouvé, avec $k = 1$ et $n = 8$, possède les propriétés suivantes :

- il existe deux erreurs vérifiant la condition (i) ;

- toutes les erreurs de $N(\bar{S})$ sont de poids supérieur ou égal à 3, sauf une erreur de poids 2.

Les résultats des simulations que l'on reporte plus bas montrent qu'un turbo-encodeur modifié basé sur cet encodeur externe réalise une très puissante réduction du taux d'erreur par bit après décodage. Le groupe stabilisateur de l'encodeur trouvé est généré par les erreurs suivantes :

$$\begin{aligned}
\bar{Z}_2 &= X & X & I & I & I & I & I \\
\bar{Z}_3 &= X & I & X & I & I & I & I \\
\bar{Z}_4 &= Z & Z & Z & Z & I & Z & I \\
\bar{Z}_5 &= X & I & I & X & Z & Z & X \\
\bar{Z}_6 &= I & I & I & I & X & Z & Z \\
\bar{Z}_7 &= I & I & I & X & I & X & Z \\
\bar{Z}_8 &= I & I & I & Z & X & I & X
\end{aligned}$$

En respectant les relations de commutation exposées au chapitre 2, on peut compléter ce groupe stabilisateur afin de définir un encodeur dont l'action sur la base canonique de G^8 est la suivante.

$$\begin{aligned}
\bar{Z}_1 &= X & I & I & X & I & I & X & Y & \bar{X}_1 &= I & I & I & X & X & X & X & X \\
\bar{Z}_2 &= X & X & I & I & I & I & I & I & \bar{X}_2 &= X & Y & I & I & I & I & I & I \\
\bar{Z}_3 &= X & I & X & I & I & I & I & I & \bar{X}_3 &= X & I & Y & I & I & I & I & I \\
\bar{Z}_4 &= Z & Z & Z & Z & I & Z & I & I & \bar{X}_4 &= X & I & I & I & I & I & I & I \\
\bar{Z}_5 &= X & I & I & X & Z & Z & X & I & \bar{X}_5 &= I & I & I & I & X & I & I & I \\
\bar{Z}_6 &= I & I & I & I & X & Z & Z & X & \bar{X}_6 &= X & I & I & I & X & X & I & I \\
\bar{Z}_7 &= I & I & I & X & I & X & Z & Z & \bar{X}_7 &= X & I & I & I & X & X & X & I \\
\bar{Z}_8 &= I & I & I & Z & X & I & X & Z & \bar{X}_8 &= X & I & I & X & I & I & I & I
\end{aligned}$$

Ce groupe stabilisateur possède deux erreurs indépendantes $\bar{Z}_2$ et $\bar{Z}_3$ de poids 2 et faisant intervenir la première position. L'existence des erreurs $\bar{Z}_2$ et $\bar{Z}_3$ peut être un atout lors de l'amorce du décodage, en ce sens qu'elles permettent d'acquérir de l'information sur les deuxième et troisième positions de la sortie de chaque bloc de l'encodeur externe. Toutefois, l'erreur $\bar{Z}_2 \circ \bar{Z}_3 = (IXXIIIII)$ est de poids 2 et appartient à $N(\bar{S})$; ainsi, la distance épurée minimale du turbo-encodeur modifié n'est pas polynomiale mais sous-logarithmique.

Le second groupe stabilisateur trouvé remplit les conditions posées par le problème 1 et est de longueur 7. Il possède de plus deux erreurs indépendantes $\bar{Z}_2$ et $\bar{Z}_3$ intervenant toutes deux sur les deux premières positions. Par conséquent, en mesurant le syndrome d'une erreur E au niveau d'un bloc de l'encodeur externe, on peut connaître totalement la deuxième position de l'erreur à partir de la première position.

$$\begin{aligned}
\bar{Z}_2 &= X & X & I & I & I & I \\
\bar{Z}_3 &= Z & Z & I & I & I & I \\
\bar{Z}_4 &= I & I & X & Z & Z & X \\
\bar{Z}_5 &= I & I & I & X & Z & Z \\
\bar{Z}_6 &= I & I & X & I & X & Z \\
\bar{Z}_7 &= I & I & Z & X & I & X
\end{aligned}$$

Le groupe stabilisateur est celui d'un encodeur de taille $n = 7$ et encodant $k = 1$ qubit, factorisable en un encodeur de paramètres $n = 2$ et $k = 0$ agissant sur les deux premières positions, et un encodeur de paramètres $n = 5$ et $k = 1$ qui réalise un encodage du code à cinq qubits de Gottesman [19]. Voici dans l'ordre deux tels encodeurs réalisant cette factorisation, donnés par leur action sur les bases canoniques respectives de G^2 et G^5 .

$$\begin{aligned}
\bar{Z}_1 &= X & X & \bar{X}_1 &= Z & I \\
\bar{Z}_2 &= Z & Z & \bar{X}_2 &= I & X
\end{aligned}$$

$$\begin{aligned}
\bar{Z}_1 &= X & I & I & X & Y & \bar{X}_1 &= X & X & X & X & X \\
\bar{Z}_2 &= X & Z & Z & X & I & \bar{X}_2 &= I & X & I & I & I \\
\bar{Z}_3 &= I & X & Z & Z & X & \bar{X}_3 &= I & X & X & I & I \\
\bar{Z}_4 &= X & I & X & Z & Z & \bar{X}_4 &= I & X & X & X & I \\
\bar{Z}_5 &= Z & X & I & X & Z & \bar{X}_5 &= X & I & I & I & I
\end{aligned}$$

Ces deux encodeurs peuvent être vus comme deux modules effectuant des tâches dédiées. Le premier encodeur ne permet d'encoder aucun qubit mais acquiert, grâce à ses deux positions de syndrome, de l'information permettant de connaître exactement la deuxième position de l'erreur de sortie en fonction de la première position en provenance du canal. Ainsi, ce module fournit à l'algorithme de décodage une position en entrée de l'encodeur interne dont la loi de probabilité correspond à la loi du canal. Le second encodeur permet quant à lui d'encoder un qubit et de fournir le résultat à l'entrée de l'encodeur interne. Les tâches d'amorce du décodage itératif et d'encodage de l'information sont ainsi clairement séparées, et cette solution présente par rapport à celle de l'encodeur de taille 8 précédent une plus grande facilité de mise en oeuvre. Soulignons également le fait que le premier encodeur peut être implémenté grâce à un circuit composé d'une porte C-NOT et d'une porte de Hadamard, cette dernière étant la porte effectuant sur un qubit l'opération linéaire suivante :

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

3.5.3 L'encodeur interne

Les deux encodeurs externes donnés par les encodeurs précédents ont été testés avec le même encodeur interne. Celui-ci est un encodeur convolutif de paramètres $n = 2$, $k = 2$ et $m = 1$ basé sur l'encodeur $\mathcal{E}_{in}$ dont l'action sur

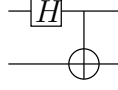


FIGURE 3.8 – Circuit de l'encodeur à 2 qubits

la base canonique de Pauli est la suivante :

$$\begin{aligned} \mathcal{E}_{in}(X_1) &= (XXX) & \mathcal{E}_{in}(Z_1) &= (ZZZ) \\ \mathcal{E}_{in}(X_2) &= (IXX) & \mathcal{E}_{in}(Z_2) &= (ZIZ) \\ \mathcal{E}_{in}(X_3) &= (XIX) & \mathcal{E}_{in}(Z_3) &= (IZZ) \end{aligned}$$

Cet encodeur possède un circuit d'encodage particulièrement simple donné par des portes $C - NOT$ représenté dans la figure 3.9.

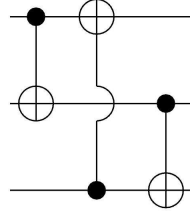


FIGURE 3.9 – Circuit implémentant l'encodeur interne

Il s'agit de l'encodeur le plus élémentaire que l'on peut réaliser de sorte à être récursif et anti-récursif. La mémoire est constituée d'un seul bit quantique, qui, conformément au schéma de l'encodeur convolutif, est le premier bit quantique de l'entrée et le dernier bit quantique de la sortie. Cet encodeur a un taux d'encodage k/n égal à 1, et a de plus la propriété d'être égal à son propre inverse modulo une interversion des deux positions logiques de son entrée ou des deux positions physiques de sa sortie. Grâce à cette propriété, il suffit de démontrer que l'encodeur est récursif afin qu'il soit également anti-récursif.

Lemme 3.5.2. *L'encodeur de base $\mathcal{E}_{in}$ de paramètres $[2, 2, 1]$ est récursif et anti-récursif.*

Démonstration. Conformément à la remarque précédente, démontrons simplement le caractère récursif de cet encodeur. Dans le cas présent, la partie de syndrome de l'encodeur est de taille nulle. L'encodeur est donc récursif si, pour toute erreur d'entrée de taille infinie sous la forme

$$E = (I, II, \dots, II, L, II, \dots)$$

où L est une erreur logique de taille 2 et de poids 1, l'erreur de sortie $\mathcal{E}_{in\infty}(E)$ est de poids infini. Les étapes précédant L ne jouant aucun rôle, on peut supposer que L intervient à la première étape de l'encodeur convolutif.

Or, on constate que si $L = (XI)$ ou $L = (IX)$, l'erreur de mémoire intermédiaire $\mu(I, L)$ vaut alors X , et comme $\mathcal{E}_{in}(X_1) = (XXX)$, toutes les erreurs de mémoire intermédiaires suivantes valent également X , et la sortie de l'encodeur convolutif donne à chaque étape l'erreur physique XX . L'erreur de sortie est donc de poids infini. De même, si $L = (ZI)$ ou $L = (IZ)$, et si $L = (YI)$ ou $L = (IY)$ le même raisonnement s'applique en remplaçant X respectivement par Z et par Y . $\square$

3.5.4 Performances de réduction d'erreur des turbo-encodeurs modifiés

Dans un premier temps, on s'est intéressé aux performances de décodage du turbo-encodeur modifié. On a testé deux configurations pour l'encodeur externe correspondant à l'encodeur de taille 8 et à la paire d'encodeurs de tailles 2 et 5.

On a d'abord utilisé, en tant qu'encodeur externe, des blocs correspondant à l'encodeur de taille 8, dont la première position de sortie est envoyée vers le canal et les 7 autres positions vers l'encodeur interne. Trois tailles N_{in} de turbo-encodeur modifié ont été testés : 500, 4000 et 60000 ; ces tailles permettent d'encoder un nombre K de bits quantiques valant respectivement 143, 1143 et 17143. Pour chacune de ces valeurs, on a complété la définition du turbo-encodeur modifié avec un entrelaceur fixé par un tirage aléatoire. Les taux d'erreurs de décodage et d'erreurs par bit sont représentés dans les figures 3.10 et 3.11. La taille $N_{in} = 200000$ correspondant à $K = 57143$ a également été testée sur un nombre plus réduit de points, afin d'observer la proximité de la courbe de taux d'erreur par bit pour $K = 17143$ par rapport à la courbe asymptotique.

Les taux d'erreurs par bit montrent une caractéristique attendue : les performances de décodage sont meilleures avec la taille de l'encodeur, et il existe un seuil d'efficacité du décodage situé autour de $p = 0,142$, soit 0,015 point en dessous de la capacité cohérente pour un taux d'encodage de $1/8$. En ce point, le taux d'erreur par bit est d'environ $1,1 * 10^{-3}$. Ces performances atteignent une courbe limite lorsque la taille de l'encodeur tend vers l'infini, due à la présence d'un bruit de canal irréductible en première position de sortie de chaque bloc de l'encodeur externe. Les taux d'erreurs de décodage sont quant à eux inhabituels, car ils augmentent avec la taille de l'encodeur contrairement aux taux d'erreurs par bit. Il s'agit ici de l'effet du bruit irréductible en sortie de l'encodeur externe, qui introduit une probabilité d'échec du décodage d'autant plus grande que le turbo-encodeur est long. Le turbo-encodeur modifié ainsi réalisé est en conclusion est mauvais code correcteur d'erreurs, mais un excellent réducteur du taux d'erreur par bit.

On a ensuite utilisé un encodeur externe basé sur la paire d'encodeurs de tailles 2 et 5. Comme décrit précédemment, ces deux encodeurs sont dédiés à des tâches séparées. L'encodeur de taille 2 n'est composé que de positions de syndrome, n'encode aucune information mais aide à l'amorce du décodage itératif ; une de ses positions de sortie est envoyée vers le canal, tandis que la seconde est envoyée vers l'encodeur interne. L'encodeur de taille 5 encode 1

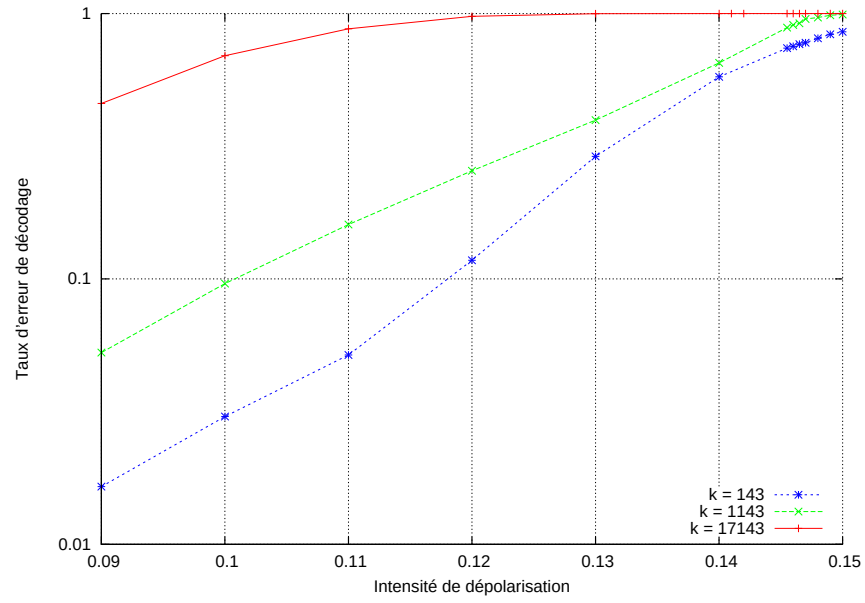


FIGURE 3.10 – Taux d'erreur de décodage

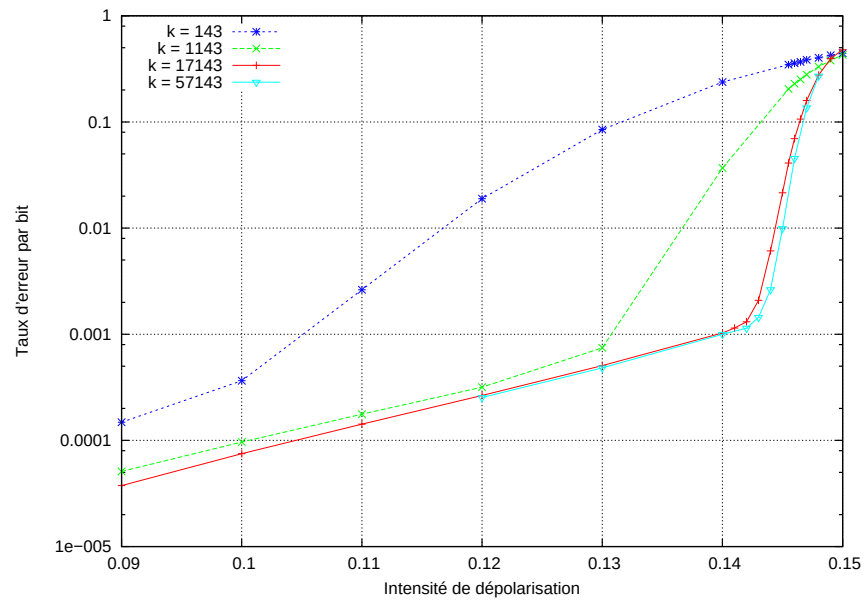


FIGURE 3.11 – Taux d'erreur par bit

position d'information, sa première position de sortie est envoyée vers le canal, et les 4 autres sont envoyées vers l'encodeur interne. On a réalisé une série de tests à nombre de bits quantiques encodés fixé à $K = 17143$. Il y a donc $K = 17143$ encodeurs de taille 5. Le nombre d'encodeurs de taille 2 peut quant à lui être modulé librement. Augmenter leur nombre facilite l'amorce du décodage itératif, mais cela se fait au détriment d'une baisse du rendement du code et d'une augmentation du bruit irréductible qui subsiste à la fin du décodage. On a ainsi testé un encodeur externe contenant les répartitions suivantes :

- $K = 17143$ encodeurs de taille 5;
- $17136 - 2448 * i$ encodeurs de taille 2;

où i est un paramètre entier conçu pour varier de 0, pour une répartition quasi-égale des deux encodeurs, à 7, pour ne garder aucun encodeur de taille 2. La valeur $i = 7$ n'a pas été testée au vu des mauvaises performances du cas $i = 6$. Les taux d'erreurs de décodage et d'erreurs par bit sont représentés dans les figures 3.12 et 3.13.

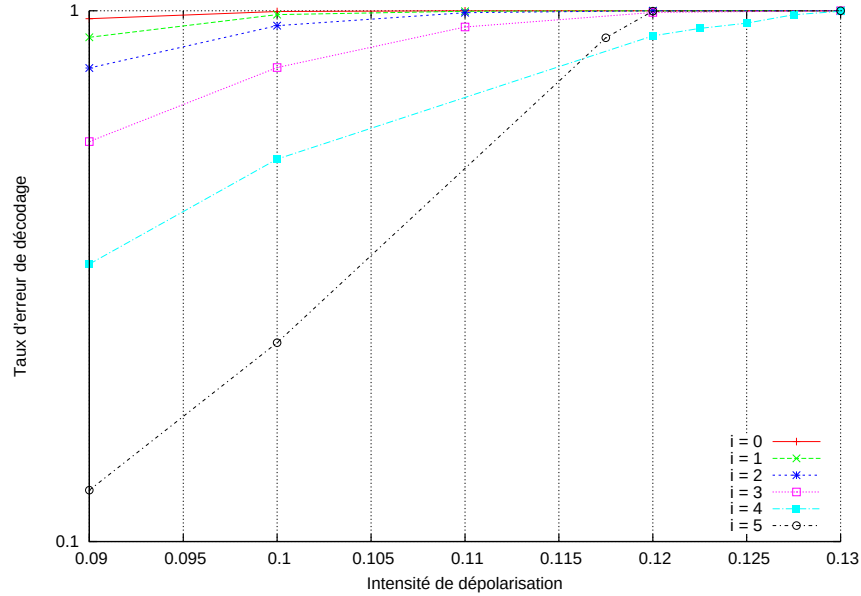


FIGURE 3.12 – Taux d'erreur de décodage

On constate que les performances de décodage par bit présentent un seuil d'efficacité qui baisse lorsque le taux d'incorporation d'encodeurs de taille 2 diminue, tandis que les performances de décodage en deçà de ce seuil deviennent meilleures. Le taux d'incorporation correspondant à $i = 3$ est celui qui présente un écart minimal entre le seuil de performance de décodage, proche de 0,13, et

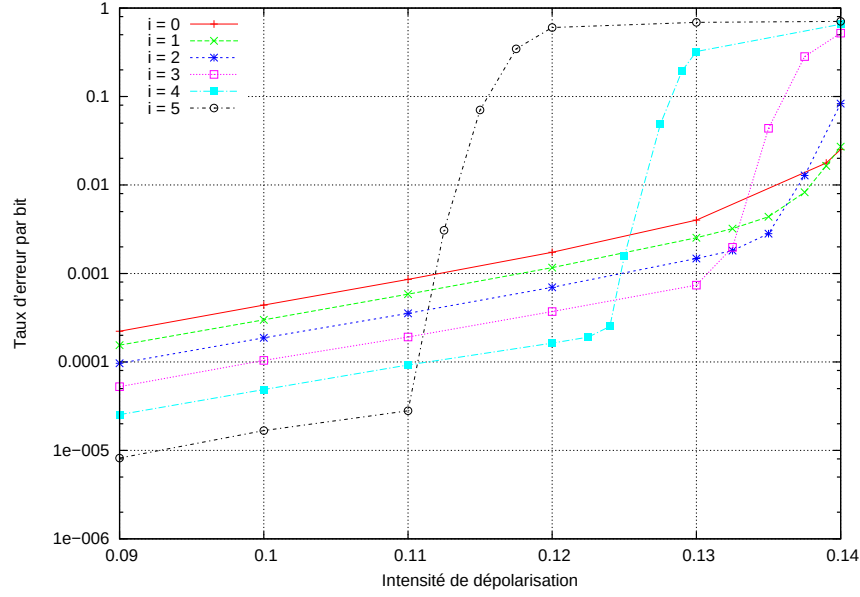


FIGURE 3.13 – Taux d'erreur par bit

la capacité cohérente, de 0,147. Cet écart est du même ordre que celui de 0,015 point constaté dans la construction précédente, avec un taux d'erreur par bit au seuil légèrement inférieur à 10^{-3} . Dans les deux constructions, les pentes des courbes en deçà du seuil sont équivalentes, et se traduisent par une diminution d'un facteur 10 tous les 0,035 point.

Dans les deux situations, on observe que le turbo-encodeur modifié est inadapté pour faire de la *correction* d'erreur, avec des taux d'échec de décodage très élevés. Cependant, les différentes constructions réalisent une remarquable *réduction* du taux d'erreur par bit. Elles peuvent par conséquent être utilisées en tant que *second étage* réducteur d'erreur, sur lesquelles se base, au premier étage, un turbo-encodeur modifié de distance épurée non bornée, conformément à la description faite à la sous-section précédente.

On peut mettre en perspective ces performances de réduction d'erreur en les comparant avec ce que l'on obtiendrait si, au lieu d'utiliser un turbo-encodeur modifié, on utilise simplement l'encodeur de taille 8 ou l'encodeur de taille 5. Cette comparaison permet de montrer le saut de performance gagné grâce à l'entrelaceur aléatoire et à l'encodeur convolutif interne de rendement 1, alors même que ces opérations n'ajoutent aucune redondance. Les figures 3.14 et 3.15 montrent ainsi les taux d'erreurs par bit respectifs obtenus pour le simple encodeur de taille 8 et pour le simple encodeur de taille 5.

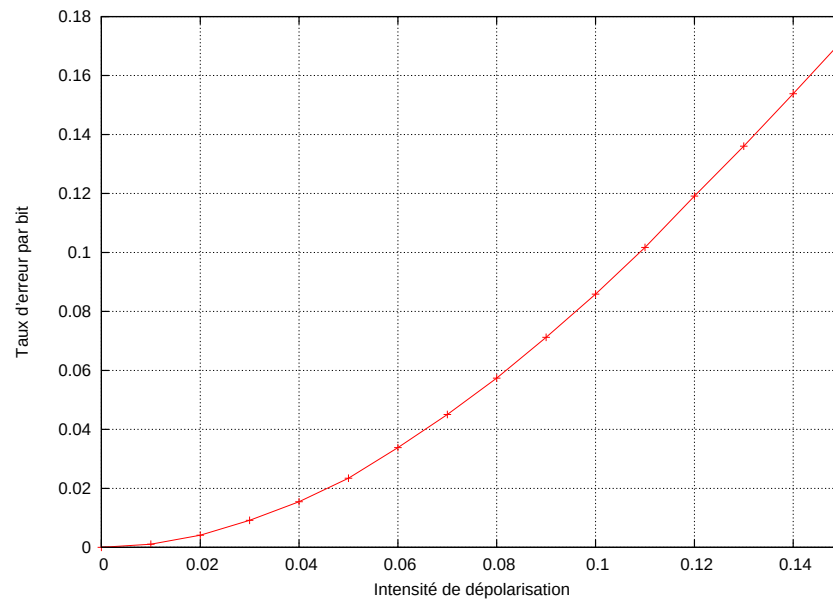


FIGURE 3.14 – Taux d'erreur par bit de l'encodeur de taille 8

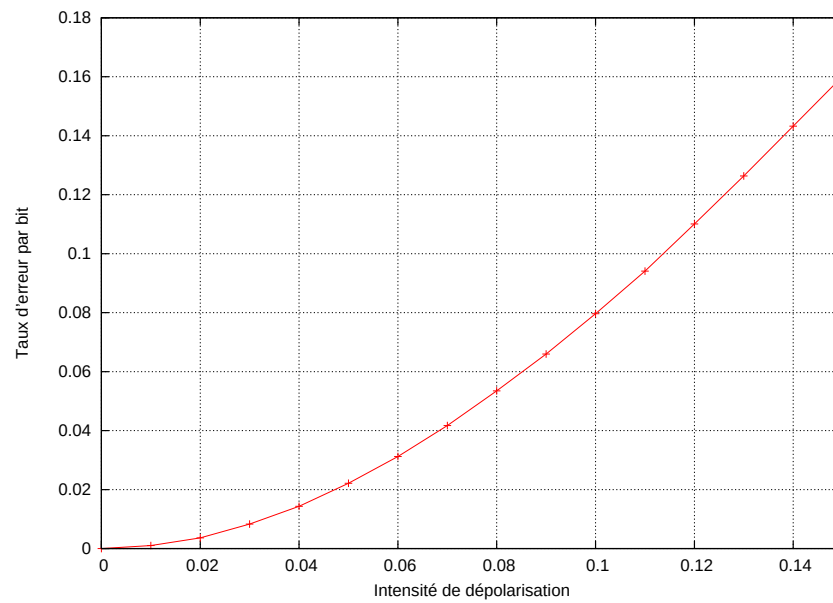


FIGURE 3.15 – Taux d'erreur par bit de l'encodeur de taille 5

3.5.5 Turbo-encodeurs modifiés à deux étages

Les tests réalisés dans un second temps ont porté sur les deux turbo-encodeurs modifiés, assistés par un second étage tel qu'explicité dans la figure 3.7. Au vu des performances brutes des turbo-encodeurs modifiés observées précédemment, on a également incorporé au *second étage* un turbo-encodeur modifié. Ainsi, la construction consiste en un turbo-encodeur modifié de distance épurée non bornée, assisté au second étage par un turbo-encodeur modifié.

Au premier étage, on a successivement testé le turbo-encodeur modifié basé sur l'encodeur de taille 8, puis celui basé sur un mélange d'encodeurs de tailles 2 et 5 pour différentes valeurs de i .

Au second étage, le turbo-encodeur modifié choisi est basé sur l'encodeur externe de taille 8. En effet, celui-ci présente, au voisinage des intensités de dépolarisation les plus élevées, une performance absolue supérieure à celle du turbo-encodeur modifié basé sur un mélange d'encodeurs de tailles 2 et 5, indépendamment de la valeur de i pour ce dernier. De plus, la différence entre les rendements de ces deux encodeurs, de $1/8$ pour le premier et d'environ $1/6$ pour le second ($1/6$, 14 pour une approximation au centième), a un faible impact sur le rendement global de la construction à deux étages. Pour s'en convaincre, on peut effectuer un calcul approximatif rapide pour quelques choix de turbo-encodeur modifié au premier étage. Le tableau 3.2 donne les différents rendements obtenus, où l'on a supposé que les encodeurs comparés au second étage ont pour rendements $1/8$ et $1/6$, et qu'au premier étage, une variation de i de 0 à 7 fait varier le taux d'incorporation de l'encodeur de taille 2 de 0,5 à 0.

	Encodeur de taille 8	Encodeurs de taille 2 et 5, $i = 0$	Encodeurs de taille 2 et 5, $i = 4$
$r = 1/8$	1/15	1/12	$7/62 \approx 1/8, 9$
$r = 1/6$	1/13	1/10	1/8

TABLE 3.2 – Comparaison des rendements de la construction pour des rendements respectifs de $r = 1/8$ et $r = 1/6$ de l'encodeur au second étage. En chaque colonne est renseigné le choix du turbo-encodeur modifié au premier étage, en indiquant l'encodeur externe sur lequel il est basé.

Les figures 3.16 et 3.17 représentent les performances relevées lorsque le turbo-encodeur modifié au premier étage est basé sur l'encodeur externe de taille 8. Elles montrent des taux d'erreurs par bit extrêmement plus bas que ceux de la même construction sans second étage. Les taux d'erreurs de décodage suivent cette fois-ci une courbe similaire aux taux d'erreurs par bit, et révèlent que l'encodeur est apte à corriger des erreurs avec un taux de succès supérieur à 99% jusqu'à une intensité de dépolarisation de 0,143. Dans la figure 3.16, deux courbes supplémentaires sont tracées, et représentent la part des décodages pour lesquels il subsiste respectivement au moins 2 et au moins 3 positions en entrée mal décodées. Cela montre qu'en ajoutant à la construction une couche

supplémentaire, de rendement proche de 1 et capable de corriger 1 voire 2 erreurs, par exemple un code BCH quantique [20], les performances de décodage peuvent être substantiellement améliorées

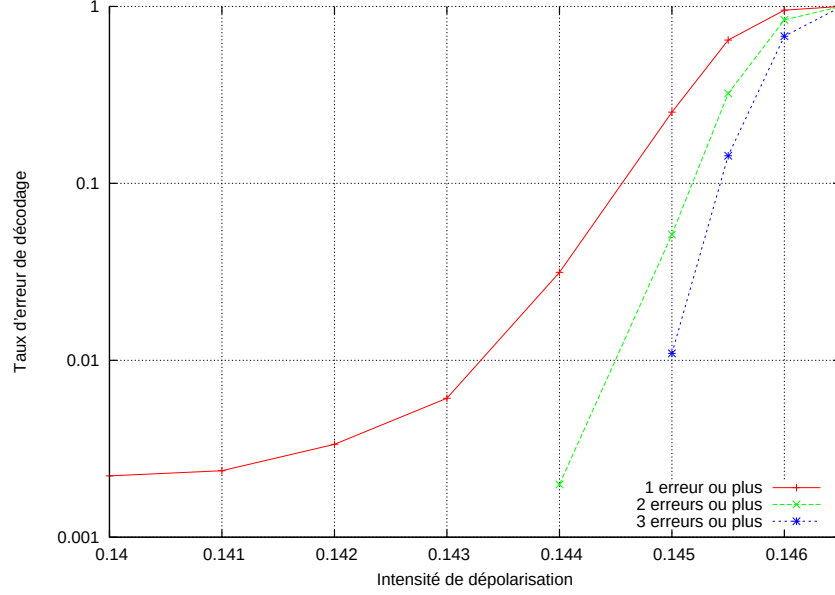


FIGURE 3.16 – Taux de décodages présentant au moins 1, 2 ou 3 positions erronées

Les figures 3.18 et 3.19 représentent les performances relevées lorsque le turbo-encodeur modifié au premier étage est basé sur un encodeur externe constitué d'un mélange des encodeurs de tailles 2 et 5, selon diverses valeurs de i . L'observation des résultats montre, pour les fortes intensités de dépolarisation, que la valeur d'incorporation i optimale de l'encodeur de taille 2 vaut 3. On observe en effet que $i = 3$ optimise le décodage pour les fortes intensités de dépolarisation, jusqu'à une intensité comprise entre 0,135 et 0,14, en-dessous de laquelle la valeur $i = 4$ donne des performances légèrement meilleures. On observe également clairement que, par rapport à $i = 3$, augmenter ou diminuer la valeur de i a pour effet de reculer le seuil de performance de décodage.

On peut à présent comparer les performances de décodage associées aux constructions à deux étages dont le turbo-encodeur modifié au premier étage a un encodeur externe basé, respectivement, sur l'encodeur de taille 8, et sur le mélange des encodeurs de tailles 2 et 5 défini par $i = 3$. La figure 3.20 effectue une telle comparaison.

Les rendements exacts de ces deux constructions sont respectivement de $1/15$ pour l'encodeur de taille 8 et de $7/71 \approx 1/10,14$ pour le mélange d'encodeurs de tailles 2 et 5. Les capacités cohérentes respectives pour de tels rendements

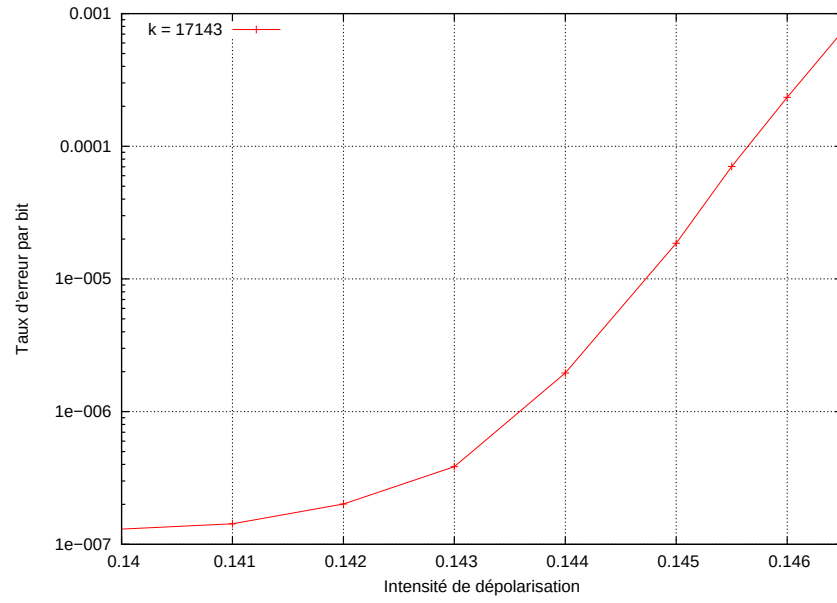


FIGURE 3.17 – Taux d'erreur par bit

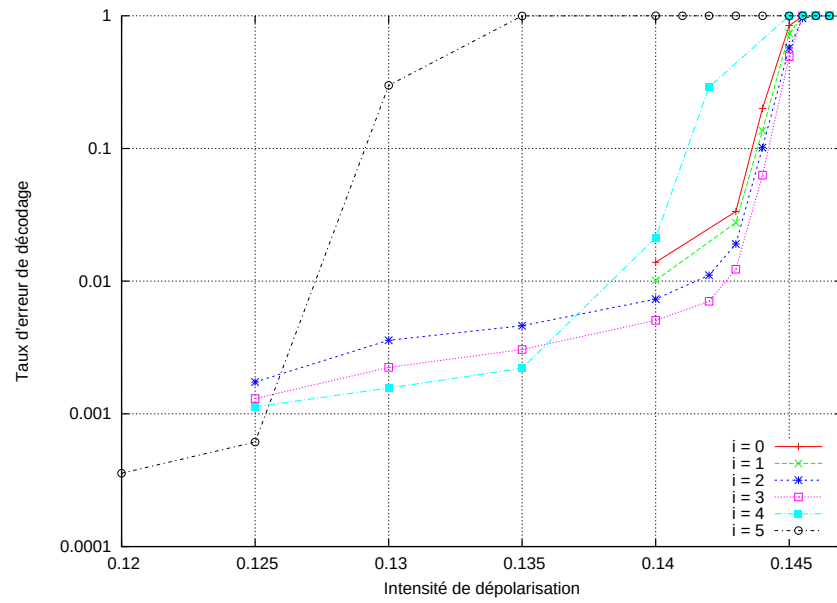


FIGURE 3.18 – Taux d'erreur de décodage

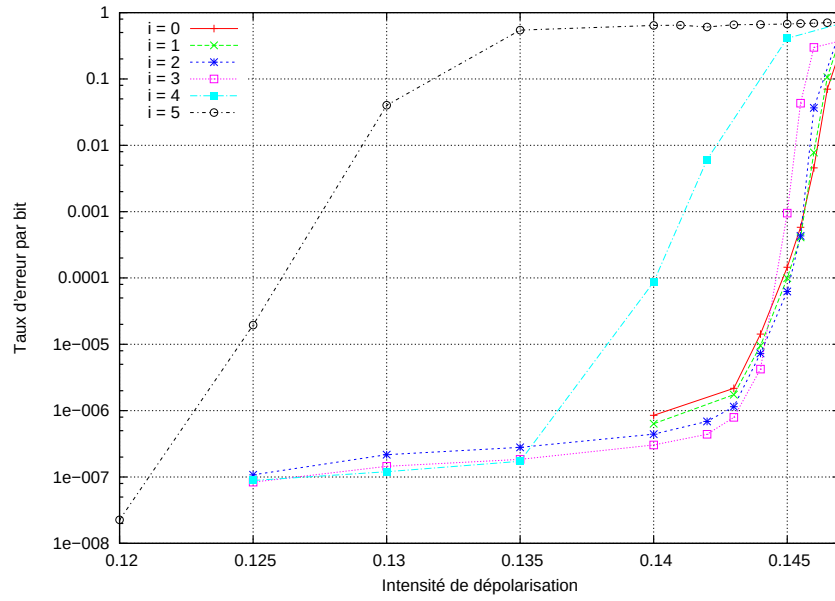


FIGURE 3.19 – Taux d'erreur par bit

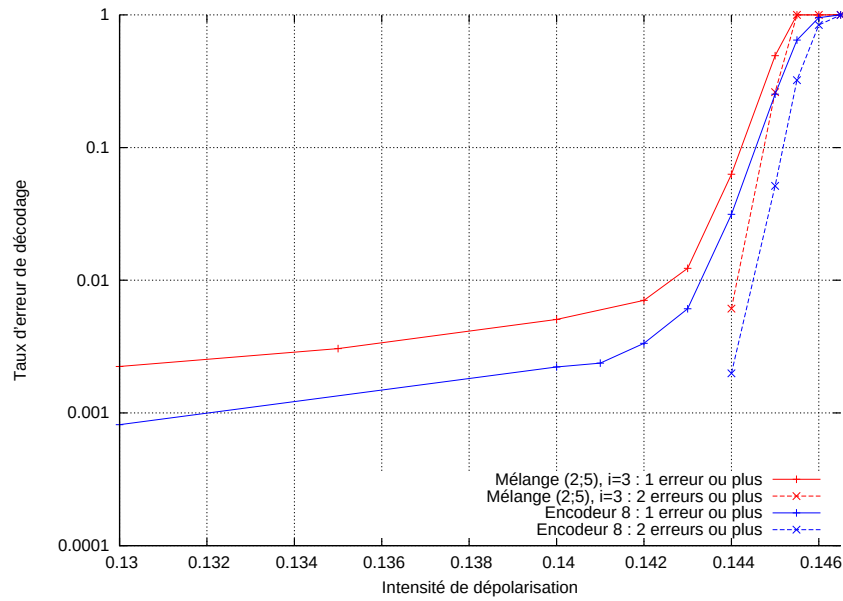


FIGURE 3.20 – Taux de décodages présentant au moins 1 ou 2 positions erronées : comparaison des deux constructions

sont d'environ 0.1634 et 0.1716. En prenant en compte l'écart entre ces deux capacités, on constate que pour les fortes intensités de dépolarisation, le mélange d'encodeurs de tailles 2 et 5 permet un décodage légèrement meilleur que l'encodeur de taille 8, tandis que cette tendance s'inverse pour les faibles intensités de dépolarisation. Les deux constructions présentent de très bonnes performances, notamment si elles sont assistées d'un encodeur supplémentaire de rendement tendant vers 1 permettant de corriger 1 erreur, avec un taux d'erreur inférieur à 1% pour une intensité de dépolarisation de 0,144.

3.6 Algorithme de décodage itératif du turbo-encodeur quantique

L'algorithme de décodage quantique que l'on emploie est décrit en partie dans [41] et correspond à un algorithme de *propagation de croyance* ou *belief propagation* quantique [23] [40]. L'algorithme décrit en [41] s'applique à une construction plus simple que celle que l'on propose dans le cadre de la thèse.

L'algorithme ci-présent diffère en quelques points de l'algorithme classique présenté au chapitre 1. Dans le montage quantique, une fraction des bits en provenance de la sortie de l'encodeur externe n'entre pas dans l'encodeur interne. Lorsque l'on simule le turbo-encodeur modifié, cette fraction de bits est envoyée directement vers le canal, tandis que lorsque l'on simule le turbo-encodeur modifié à deux étages, on suppose qu'elle est envoyée vers un canal simulant la réduction de bruit effectuée par l'encodeur au second étage. L'information extrinsèque permettant d'effectuer la mise à jour au niveau de l'encodeur externe provient donc en partie du canal simulé et en partie de l'encodeur interne. Par ailleurs, l'encodeur externe n'est pas convolutif mais est constitué de N_{ext} blocs appliqués en parallèle. Ainsi, la mise à jour effectuée au niveau de l'encodeur externe est parallélisée sur les N_{ext} blocs qui le constituent. Il y aura donc trois fonctions de mise à jour : la fonction *MiseAJourLogiqueInterne* pour l'encodeur interne –la loi au niveau de l'erreur physique en provenance du canal n'étant jamais mise à jour –, ainsi que les fonctions *MiseAJourLogiqueExterne* et *MiseAJourPhysiqueExterne* pour l'encodeur externe. La première fonction est constituée, comme au chapitre 1, d'un parcours avant et un parcours arrière suivis d'une mise à jour locale, alors que les deux dernières fonctions consistent simplement en une mise à jour locale en chacun des blocs. Finalement, les calculs de distributions de probabilités marginales se font par sommation sur l'ensemble des syndromes dont les coordonnées appartiennent à $\{I, Z\}$.

Le tableau 3.3 reprend les notations de la définition 3.5.1 et de la figure 3.6, et présente les types d'erreurs intervenant lors du décodage. Ces notations prennent le soin de distinguer les erreurs selon leur rôle dans le schéma général ; notamment, les erreurs de syndrome sont séparées, car elles ne sont pas l'objet du décodage mais permettent de révéler une information à ce dernier grâce à la mesure du syndrome binaire. Le tableau 3.4 reprend quant à lui les notations

employées afin de désigner les différentes lois de probabilité en jeu.

Erreur	Notation
convention : coordonnée u de l'erreur E	E_u
syndrome binaire correspondant à l'erreur de syndrome S	$S_z \in \{0, 1\}^{n-k}$ tel que $S_{z_u} = 0$ si et seulement si $S_u = I$ ou Z
erreur de syndrome d'entrée de l'encodeur externe	$(S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)})$
erreur de syndrome d'entrée de l'encodeur interne	$(S_{in}^{(0)}, \dots, S_{in}^{(N_{ext}-1)})$
erreur logique d'entrée du turbo-encodeur modifié	$L_{ext} = (L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)})$
erreur physique envoyée vers l'entrelaceur	$P_{ext \rightarrow in} = (P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$
erreur physique envoyée vers le canal ou vers l'encodeur au second étage	$P_{canal} = (P_{canal}^{(0)}, \dots, P_{canal}^{(N_{ext}-1)})$
erreur de mémoire et logique d'entrée de l'encodeur convolutif interne	$L_{in} = (M_{in}^{(0)}, L_{in}^{(0)}, \dots, L_{in}^{(N_{in}-1)})$
t -ème erreur de mémoire intermédiaire de l'encodeur convolutif interne	$M_{in}^{(t)}$
erreur physique en sortie de l'encodeur convolutif interne	$P_{in} = (P_{in}^{(0)}, \dots, P_{in}^{(N_{in}-1)}, M_{in}^{(N_{in})})$

TABLE 3.3 – notations utilisées dans l'algorithme de décodage itératif

Loi de probabilité	Notation
Loi du canal	$\mathbf{P}_{canal}$
Lois temporaire, et de mise à jour avant, arrière et locale	$\mathbf{P}_{temp}, \mathbf{P}_{av}, \mathbf{P}_{ar}, \mathbf{P}_{loc}$
Loi au niveau de l'erreur L_{in} , à l'étape i du décodage	$\mathbf{P}_i^{L_{in}}$, ou $\mathbf{P}_i$ si le contexte porte sans équivoque sur L_{in}
Probabilité, à l'étape i , que l'erreur $L_{in\ u}^{(t)}$ vaille E	$\mathbf{P}_i(L_{in\ u}^{(t)} = E)$
Probabilité, à l'étape i , que l'erreur $L_{in}^{(t)}$ vaille E	$\mathbf{P}_i(L_{in}^{(t)} = E)$

TABLE 3.4 – différentes probabilités en jeu lors du décodage

Les hypothèses au lancement de l'algorithme sont similaires à celles du chapitre 1, à l'exception des trois différences suivantes. Premièrement, il faut garder à l'esprit que les symboles dont les lois de probabilité sont estimées sont des symboles d'erreur et non des symboles d'information. Ainsi, au lancement de l'algorithme, on connaît dans le cas classique présenté à la section 1.12 la séquence reçue z . Par conséquent, la loi de probabilité donnée en hypothèse sur les symboles en provenance du canal s'exprime ainsi :

$$\mathbf{P}^{y'}(y_u'^{(t)} = \alpha) = \mathbf{P}_{canal}(y_u'^{(t)} = \alpha | z_u^{(t)})$$

Dans le cas présent, en revanche, on ne connaît pas l'erreur en provenance du canal ; la loi de probabilité donnée en hypothèse est donc la suivante :

$$\begin{aligned} \mathbf{P}(P_{in\ u}^{(t)} = E) &= \mathbf{P}_{canal}(P_{in\ u}^{(t)} = E) \\ \mathbf{P}(P_{canal\ u}^{(t)} = E) &= \mathbf{P}_{canal}(P_{canal\ u}^{(t)} = E) \end{aligned}$$

Il faut toutefois mitiger cette remarque en rappelant que l'on dispose, dans le cas présent, des syndromes mesurés d'une part au niveau de l'encodeur interne, et d'autre part au niveau de l'encodeur externe. Cela se formule en supposant que l'on connaît, au lancement de l'algorithme, les deux séquences de syndrome binaires :

$$\begin{cases} s_{ext} = (s_{ext}^{(0)}, \dots, s_{ext}^{(N_{ext}-1)}) & \in \{0, 1\}^{N_{ext}(n_{ext}-k_{ext})} \\ s_{in} = (s_{in}^{(0)}, \dots, s_{in}^{(N_{in}-1)}) & \in \{0, 1\}^{N_{in}(n_{in}-k_{in})} \end{cases}$$

qui correspondent respectivement aux erreurs de syndrome S_{ext} et S_{in} . La prise en compte de la séquence s_{in} afin de modéliser l'algorithme de décodage relève surtout d'un intérêt intellectuel, car comme cela a été exposé dans la section précédente, l'encodeur interne conçu afin de faire les simulations est de rendement 1. La troisième remarque concerne également la loi de canal, et est

une conséquence de la construction à deux étages proposée. Lorsque, dans un premier temps, on lance une première série de simulations du turbo-encodeur modifié, la loi de canal est simplement celle d'un canal dépolarisant d'intensité p . Lorsqu'on simule par la suite le turbo-encodeur modifié assisté par un turbo-encodeur modifié au second étage, les résultats obtenus lors de la première série de simulations permettent de modéliser le second étage : d'une part, la loi de canal portant sur les positions de l'erreur P_{in} est celle d'un canal dépolarisant d'intensité p , et d'autre part, la loi de canal portant sur les positions de l'erreur P_{canal} est celle d'un canal dépolarisant d'intensité donnée par la courbe de taux d'erreur par bit 3.11.

La boucle principale de l'algorithme est donnée par le pseudo-code 4. Les pseudo-codes 5, 5 et 5 montrent le fonctionnement des différentes mises à jour. Les indices *ext* et *in* y sont omis pour alléger les notations. De plus, on note $\mathcal{E} = (\pi, \mu)$ de telle sorte que si

$$(P^{(t)}, M^{(t+1)}) = \mathcal{E}(M^{(t)}, L^{(t)})$$

alors

$$\begin{cases} P^{(t)} &= \pi(M^{(t)}, L^{(t)}) \\ M^{(t+1)} &= \mu(M^{(t)}, L^{(t)}) \end{cases}$$

Dans la fonction *MiseAJourPhysiqueExterne*, seules les positions de l'erreur $P_{ext \rightarrow in}$, qui sont envoyées vers l'encodeur interne, et non vers le canal ou le second étage, sont concernées par la mise à jour. Les positions de l'erreur P_{canal} ont en effet une loi de probabilité invariante tout le long du décodage. Les positions concernées en chaque bloc de l'encodeur externe varient donc de la $n_c + 1$ -ième à la n_{ext} -ième position, où $n_c = 1$ est le paramètre de canal.

D'autre part, les fonctions de mise à jour au niveau de l'encodeur externe, *MiseAJourLogiqueExterne* et *MiseAJourPhysiqueExterne*, ne nécessitent pas de prendre en entrée la loi de probabilité $\mathbf{P}^{(L_{ext})}$ qui porte sur les positions de l'erreur logique d'entrée, celle-ci étant une loi uniforme. En revanche, elles prennent toutes deux en entrées les lois $\mathbf{P}_{canal}$ et $\mathbf{P}_{temp}^{P_{ext \rightarrow in}}$ portant sur les positions physiques de l'encodeur externe, qui proviennent respectivement d'une part du canal ou du second étage, et d'autre part de l'encodeur interne.

Algorithme 4 Boucle Principale de l'algorithme

pour $i = 0 \rightarrow i_{max} - 1$ **faire** ▷ Boucle Principale

$\mathbf{P}_{temp}^{L_{in}} \leftarrow \text{MISEAJOURLOGIQUEINTERNE}(\mathbf{P}_i^{L_{in}}, \mathbf{P}^{P_{in}}, \mathcal{E}_{inN_{in}})$

$\mathbf{P}_{temp}^{P_{ext \rightarrow in}} \leftarrow \text{ENTRELACEMENT}^{-1}(\mathbf{P}_{temp}^{L_{in}})$

$\mathbf{P}_{i+1}^{P_{ext \rightarrow in}} \leftarrow \text{MISEAJOURPHYSIQUEEXTERNE}(\mathbf{P}_{canal}, \mathbf{P}_{temp}^{P_{ext \rightarrow in}}, \mathcal{E}_{extN_{ext}})$

$\mathbf{P}_{i+1}^{L_{ext}} \leftarrow \text{ENTRELACEMENT}(\mathbf{P}_{i+1}^{P_{ext \rightarrow in}})$

fin pour

▷ Dernière itération

$\mathbf{P}_{temp}^{L_{in}} \leftarrow \text{MISEAJOURLOGIQUEINTERNE}(\mathbf{P}_{i_{max}}^{L_{in}}, \mathbf{P}^{P_{in}}, \mathcal{E}_{inN_{in}})$

$\mathbf{P}_{temp}^{P_{ext \rightarrow in}} \leftarrow \text{ENTRELACEMENT}^{-1}(\mathbf{P}_{temp}^{L_{in}})$

$\mathbf{P}_{final}^x \leftarrow \text{MISEAJOURLOGIQUEEXTERNE}(\mathbf{P}_{canal}, \mathbf{P}_{i_{max}}^{P_{ext \rightarrow in}}, \mathcal{E}_{extN_{ext}})$

$\hat{L}_{ext} \leftarrow \text{MAXVRAISEMBLANCE}(\mathbf{P}_{final}^{L_{ext}})$

retourner $\hat{L}_{ext}$

Algorithme 5 Fonction *MiseAJourLogiqueInterne*

fonction MISEAJOURLOGIQUEINTERNE($\mathbf{P}^L, \mathbf{P}^P, \mathcal{E}_N$)
 $\mathbf{P}_{av}(M^{(0)}) \leftarrow \mathbf{P}^L(M^{(0)})$ ▷ Parcours avant
pour $t = 0 \rightarrow N - 2$ **faire**
 pour tout $M' \in G^m$ **faire**

$$\begin{aligned}
 \mathbf{P}_{av}(M^{(t+1)} = M') \leftarrow & \sum_{S: S_z = s^{(t)}} \sum_{\substack{(M, L) \\ \mu(M, L, S) = M'}} \mathbf{P}_{av}(M^{(t)} = M) \\
 & \times \mathbf{P}^L(L^{(t)} = L) \mathbf{P}^P(P^{(t)} = \pi(M, L, S))
 \end{aligned}$$

fin pour
fin pour
 $\mathbf{P}_{ar}(M^{(N)}) \leftarrow \mathbf{P}^L(M^{(N)})$ ▷ Parcours arrière
pour $t = N - 1 \rightarrow 0$ **faire**
 pour tout $M \in G^m$ **faire**

$$\begin{aligned}
 \mathbf{P}_{ar}(M^{(t)} = M) \leftarrow & \sum_{S: S_z = s^{(t)}} \sum_L \mathbf{P}^L(L^{(t)} = L) \\
 & \times \mathbf{P}^P(P^{(t)} = \pi(M, L, S)) \mathbf{P}_{ar}(M^{(t+1)} = \mu(M, L, S))
 \end{aligned}$$

fin pour
fin pour
pour $t = 0 \rightarrow N - 1$ **faire** ▷ Mise à jour locale
 pour $u = 1 \rightarrow k$ **faire**
 pour tout $E \in G$ **faire**

$$\begin{aligned}
 \mathbf{P}_{loc}(L_u^{(t)} = E) \leftarrow & \sum_{S: S_z = s^{(t)}} \sum_{L: L_u = E} \sum_M \mathbf{P}_{av}(M^{(t)} = M) \\
 & \times \prod_{\substack{v=1 \\ v \neq u}}^k \mathbf{P}^L(L_v^{(t)} = L_v) \mathbf{P}^P(P^{(t)} = \pi(L, M, S)) \mathbf{P}_{ar}((M^{(t+1)} = \mu(L, M, S)))
 \end{aligned}$$

fin pour
 $somme = \sum_{E \in G} \mathbf{P}_{loc}(L_u^{(t)} = E)$ ▷ Normalisation
 $\mathbf{P}_{loc}(L_u^{(t)}) \leftarrow \mathbf{P}_{loc}(L_u^{(t)}) / somme$
 fin pour
fin pour
retourner $\mathbf{P}_{loc}$
fin fonction

Algorithme 6 Fonction *MiseAJourPhysiqueExterne*

fonction MISEAJOURPHYSIQUEEXTERNE($\mathbf{P}_{canal}$, $\mathbf{P}^P$, $\mathcal{E}_N$)
 pour $t = 0 \rightarrow N - 1$ **faire**
 pour $u = n_c + 1 \rightarrow n$ **faire**
 pour tout $E \in G$ **faire**

$$\mathbf{P}_{temp}(P_u^{(t)} = E) \leftarrow \sum_{S: S_z = s^{(t)}} \sum_{L: (\mathcal{E}(L, S))_u = E} \prod_{v=1}^{n_c} \mathbf{P}_{canal}(P_v^{(t)} = (\mathcal{E}(L, S))_v) \\ \times \prod_{\substack{v=n_c+1 \\ v \neq u}}^n \mathbf{P}^P(P_v^{(t)} = (\mathcal{E}(L, S))_v)$$

fin pour
 $somme = \sum_{E \in G} \mathbf{P}_{temp}(P_u^{(t)} = E)$ ▷ Normalisation
 $\mathbf{P}_{temp}(P_u^{(t)} = E) \leftarrow \mathbf{P}_{temp}(P_u^{(t)} = E) / somme$
 fin pour
 fin pour
 retourner $\mathbf{P}_{temp}$
fin fonction

Algorithme 7 Fonction *MiseAJourPhysiqueInterne*

fonction MISEAJOURPHYSIQUEINTERNE($\mathbf{P}_{canal}$, $\mathbf{P}^P$, $\mathcal{E}_N$)
 pour $t = 0 \rightarrow N - 1$ **faire**
 pour $u = 1 \rightarrow k$ **faire**
 pour tout $E \in G$ **faire**

$$\mathbf{L}_{temp}(L_u^{(t)} = E) \leftarrow \sum_{S: S_z = s^{(t)}} \sum_{L: L_u = E} \prod_{v=1}^{n_c} \mathbf{P}_{canal}(P_v^{(t)} = (\mathcal{E}(L, S))_v) \\ \times \prod_{v=n_c+1}^n \mathbf{P}^P(P_v^{(t)} = (\mathcal{E}(L, S))_v)$$

fin pour
 $somme = \sum_{E \in G} \mathbf{P}_{temp}(L_u^{(t)} = E)$ ▷ Normalisation
 $\mathbf{P}_{temp}(L_u^{(t)}) \leftarrow \mathbf{P}_{temp}(L_u^{(t)}) / somme$
 fin pour
 fin pour
 retourner $\mathbf{L}_{temp}$
fin fonction

Chapitre 4

Preuve de la distance minimale d'un turbo-encodeur formel

4.1 Schéma global de la preuve

Il est à noter que dans le schéma de turbo-encodage classique, l'encodage externe est un encodage convolutif, contrairement au schéma de turbo-encodeur que l'on adopte où l'encodeur externe est une séquence d'encodeurs identiques placés en parallèle, pour les raisons que l'on a évoquées à la section 3.2. Cependant, cela n'affecte pas le schéma global de la preuve. L'idée de la preuve de [22] consiste à majorer pour un turbo-encodage T_N donné de taille N la probabilité $p(\leq d)$ d'existence d'une séquence d'entrée x dont l'image $T_N(x)$ est de poids inférieur à $d > 0$. Pour cela, on note w le poids de la séquence intermédiaire x' , c'est-à-dire la séquence obtenue après application de l'encodage externe à x . L'approche suivie dans [22] pour obtenir cette borne est de majorer dans un premier temps les quantités suivantes :

- d'une part, le nombre $a_{ext}(w)$ de séquences d'entrée x dont l'image intermédiaire x' est de poids w ,
- d'autre part, le nombre $a_{in}(w, \leq d)$ de séquences intermédiaires x' de poids w et dont l'image finale $T_N(x)$ est de poids inférieur à d .

N.Kahale et R.Urbanke montrent ainsi que :

$$a_{ext}(w) \leq O(1)^w \binom{N}{\lfloor w/d^* \rfloor}$$

où d^* représente la distance libre de l'encodage externe, et :

$$a_{in}(w, \leq d) \leq O(1)^w \left(\frac{Nd}{w^2} \right)^{w/2}$$

Par un argument probabiliste lié au caractère aléatoire de l'entrelaceur, la probabilité $p(w, \leq d)$ qu'il existe une séquence x telle que x' est de poids w et une sortie $T_N(x)$ est de poids inférieur à d est alors majorée par :

$$p(w, \leq d) \leq \frac{a_{ext}(w)a_{in}(w, \leq d)}{n(w)} \quad (4.1.1)$$

où $n(w)$ est le nombre de séquences x' de poids w . On somme alors la borne obtenue sur toutes les valeurs de w :

$$p(\leq d) \leq \sum_w p(w, \leq d)$$

pour la valeur $d = N^\alpha$ correspondant à la borne sur la distance minimale à démontrer, et le résultat obtenu, négligeable devant 1 lorsque N tend vers l'infini, permet ainsi de démontrer cette borne.

Dans le cas présent, on applique essentiellement la même idée aux erreurs formelles d'entrée E , intermédiaire E' et de sortie $T_N(E)$, avec certaines particularités que l'on détaillera. On montre que $a_{ext}(w)$ est majoré par deux bornes (1E) et (2E), et $a_{in}(w, \leq d)$ est majorée par deux bornes (1I) et (2I). Les bornes indexées par le chiffre 1 sont essentiellement du même type que les bornes obtenues par Kahale et Urbanke, et servent à démontrer le théorème 3.4.1 qui correspond au cadre classique et qui consiste en une reformulation du théorème 1.11.2. Ces bornes restent valables sous les hypothèses du théorème 3.4.2, et doivent être combinées avec les bornes indexées par le chiffre 2 afin de démontrer le théorème 3.4.2 qui correspond au cadre quantique. Comme on le verra plus précisément, les bornes (1E) et (2E) permettent d'établir respectivement les relations suivantes :

$$a_{ext}(w) \leq O(1)^w \left(\frac{N}{w} \right)^{w/d_{deg}-1/2}$$

où d_{deg} est la *distance dégénérée* de l'encodeur externe, et :

$$a_{ext}(w) \leq c^d n(w)$$

où $n(w)$ est défini ci-dessus et où $c < 1$. Les bornes (1I) et (2I) ont quant à elles respectivement les formes suivantes :

$$a_{in}(w, \leq d) \leq O(1)^w \frac{N^{\frac{w}{2}} (w+d)^{\frac{w}{2}}}{w^w}$$

et :

$$a_{in}(w, d) \leq O(1)^d \frac{N^{\frac{d}{2}} (w+d)^{\frac{d}{2}}}{d^d}$$

où l'absence du symbole $\leq$ devant d correspond au décompte de l'ensemble des erreurs E dont le poids de l'image $T_N(E)$ est *égal* à d . Le fait que les deux derniers membres de droite se déduisent l'un de l'autre par l'intervention de w

et d est dû au rôle principalement symétrique que jouent les poids de l'entrée et de la sortie entre les définitions des caractères récursif et anti-récursif d'un encodeur convolutif formel. La table 4.1 dresse un récapitulatif des hypothèses sous lesquelles s'applique chacune de ces quatre bornes, ainsi que des théorèmes pour lesquels elles sont utilisées. On peut s'assurer que les hypothèses mentionnées dans cette table se retrouvent bien dans les énoncés des théorèmes correspondants. En effet, le théorème 3.4.1 requiert bien que $d_{deg} \geq 2$ et que l'encodeur convolutif formel interne est récursif, et le théorème 3.4.2 requiert bien que $|G| \geq 4$, $d_{deg} \geq 2$ et que l'encodeur convolutif formel interne est anti-récursif.

Borne	hypothèses de validité	théorèmes concernés
$1E$	$d_{deg} \geq 2$	3.4.1, 3.4.2
$2E$	$ G \geq 4, d_{deg} \geq 2$	3.4.2
$1I$	encodeur formel interne récursif	3.4.1, 3.4.2
$2I$	encodeur formel interne anti-récursif	3.4.2

TABLE 4.1 – hypothèses et théorèmes correspondant aux différentes bornes

Une différence majeure par rapport au schéma de Urbanke et Kahale apparaît au moment de la sommation des bornes obtenues sur w et d . Dans le cas classique, qui correspond également à celui du théorème 3.4.1, la somme s'effectue sur l'ensemble des couples (w, d) vérifiant $w \leq d$, les cas où $w > d$ étant éliminés grâce au caractère systématique de l'encodeur convolutif formel interne. Cependant, dans le cas du théorème 3.4.2 où cet encodeur formel n'est pas systématique, un nouveau régime apparaît dans lequel $w > d$ et que les bornes classiques ne suffisent pas à dominer. C'est ce qui nécessite l'introduction de la nouvelle condition d'encodeur convolutif formel *anti-récursif*, qui permet de dominer la somme obtenue dans ce régime et d'obtenir un résultat négligeable devant 1 lorsque N tend vers l'infini. Les figures suivantes montrent l'utilisation des différentes bornes dans les différents domaines d'existence du couple (w, d) .

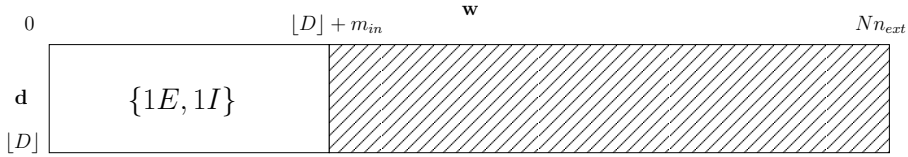


FIGURE 4.1 – Domaines d'utilisation des bornes dans le cas du théorème 3.4.1

D désigne la minoration sur la distance minimale que l'on cherche à prouver. D désigne ainsi tour à tour N^α où $\alpha < \frac{d_{deg}-2}{d_{deg}}$ lorsque l'on montre la borne polynomiale des théorèmes 3.4.1 et 3.4.2, et $\alpha \log N / (\log \log N)$ où $\alpha < d_{min} - 2$ lorsque l'on montre la borne sous-logarithmique du théorème 3.4.2. La limite $[D] + m_{in}$ comprend le terme m_{in} , qui correspond à la taille de la mémoire de

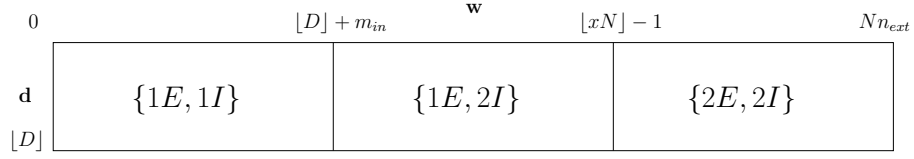


FIGURE 4.2 – Domaines d'utilisation des bornes dans le cas du théorème 3.4.2

l'encodeur convolutif formel interne. Ce terme se justifie par le fait que le schéma convolutif que l'on adopte suppose que l'état initial de la mémoire figure dans l'entrée de l'encodeur convolutif formel interne. La limite $Nk_{in} + m_{in}$ correspond quant à elle à la taille de la partie porteuses de l'information en entrée de l'encodeur convolutif formel interne, soit la partie de mémoire et logique de cette entrée.

Dans ces schémas, il apparaît que le domaine $w > [D] + m_{in}$ est lui-même composé de deux parties : un régime où w est sub-linéaire en N , et un régime où w est linéaire en N . x est une constante qui doit être choisie judicieusement afin de permettre de dominer les sommes correspondant à ces deux parties. Cela sera exposé lors de la preuve des théorèmes, dans la section 4.5.

À présent, définissons précisément les quantités présentées ci-dessus, afin d'établir une inégalité correspondant à l'inégalité (4.1.1). Notons dès à présent que, comme la distance minimale d'un encodeur formel est le plus petit poids d'une de ses erreurs nocives, et qu'une erreur nocive est l'image d'une erreur de poids logique non nul et de poids de syndrome nul, on s'intéressera toujours à de telles erreurs lors du décompte.

Définition 4.1.1. Soit un turbo-encodeur formel T_N de taille N basé sur deux encodeurs formels externe et interne fixés et un entrelaceur formel aléatoire de taille N . Pour tout couple d'entiers w et d , on considère les probabilités des événements suivants :

- $p_N(d)$ et $p_N(\leq d)$: il existe une erreur d'entrée E dont l'image $T_N(E)$ est une erreur nocive de poids respectif d et inférieur ou égal à d .
- $p_N(w, d)$ et $p_N(w, \leq d)$: il existe une erreur d'entrée E dont l'image intermédiaire externe est de poids physique w , et dont l'image $T_N(E)$ est une erreur nocive de poids respectif d et inférieur ou égal à d .

Comme on l'a précisé, ces définitions correspondent à la probabilité d'existence d'une erreur d'entrée E , de poids de syndrome nul et de poids logique non nul, dont l'image $T_N(E)$ et l'image intermédiaire vérifient les contraintes de poids énoncées ci-dessus. De plus, comme l'entrelaceur formel est aléatoire, ces probabilités ne sont fonction que des encodeurs formels externe et interne et de N . Elles sont définies pour tout couple d'entiers w et d , mais il sera commode de constater que $p_N(\leq d)$ et $p_N(w, \leq d)$ sont également définies pour les valeurs réelles de d .

Définissons également les répartitions de poids externe et interne que l'on cherchera à majorer. Ces définitions dépendent encore une fois de l'encodeur formel considéré, même si ce dernier n'apparaît pas dans les notations afin de les alléger.

Définition 4.1.2. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ et $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ deux encodeurs formels. Pour tout entier N et tout couple d'entiers (w, d) on considère :

- $a_{ext}^N(d)$ le nombre d'erreurs

$$E = (L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{Nk_{ext}} \times H^{N(n_{ext}-k_{ext})}$$

de poids de syndrome nul et de poids logique non nul, et dont l'image $(P^{(0)}, \dots, P^{(N-1)})$ donnée par $P_i = \mathcal{E}_{ext}(L_i, S_i)$ est de poids d .

- $a_{in}^N(w, d)$ et $a_{in}^N(w, \leq d)$ le nombre d'erreurs

$$E_{ML} = (M, L^{(0)}, \dots, L^{(N-1)}) \in G^{m_{in}+Nk_{in}}$$

de poids w , qui peuvent être complétées en une erreur d'entrée de l'encodeur convolutif formel $\mathcal{E}_{inN}$

$$E = (M, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{m_{in}+Nk_{in}} \times H^{N(n_{in}-k_{in})}$$

de poids de syndrome nul et dont l'image par $\mathcal{E}_{inN}$ est de poids respectif d et inférieur ou égal à d .

L'indice ML dans E_{ML} réfère simplement au fait que l'erreur est constituée d'une partie de mémoire et une partie logique. La définition de $a_{in}^N(w, \leq d)$ s'étend également aux valeurs réelles de d . On remarquera que dans les définitions de $a_{in}^N(w, d)$ et $a_{in}^N(w, \leq d)$, seule la partie de mémoire et logique est comptabilisée. L'argument de dénombrement principal décrit en début de section correspond alors à ce lemme.

Lemme 4.1.1. Soient $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ et $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ deux encodeurs formels. Un turbo-encodeur formel T_N de taille $N = N_{in}$ basé sur ces encodeurs formels externe et interne et un entrelaceur formel aléatoire de taille N vérifie :

$$\forall (w, d) \in \mathbb{N} \times \mathbb{N}, p_N(w, d) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext}n_{ext}}{w}} \quad (4.1.2)$$

$$\forall (w, d) \in \mathbb{N} \times \mathbb{R}, p_N(w, \leq d) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, \leq d)}{(|G| - 1)^w \binom{N_{ext}n_{ext}}{w}} \quad (4.1.3)$$

où

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

désigne le coefficient binomial des entiers n et k .

Démonstration. Démontrons l'inégalité 4.1.2, sachant qu'un raisonnement similaire permet de démontrer l'inégalité 4.1.3. $p_N(w, d)$ est la probabilité de l'événement : « il existe une erreur E de poids de syndrome nul et de poids logique non nul, dont l'image intermédiaire externe a pour poids physique w et l'image par T_N a pour poids d ». Pour qu'il existe une erreur

$$E = (L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)}, S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)}, S_{in}^{(0)}, \dots, S_{in}^{(N-1)})$$

vérifiant cette propriété, il est nécessaire qu'il existe une erreur

$$(L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)}, S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)})$$

qui :

- d'une part, figure parmi les $a_{ext}^{N_{ext}}(w)$ erreurs de poids de syndrome nul et de poids logique non nul dont l'image $(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$ donnée par $P_{ext}^{(i)} = \mathcal{E}_{ext}(L_{ext}^{(i)}, S_{ext}^{(i)})$ est de poids w .
- d'autre part, est telle que l'erreur obtenue en appliquant l'entrelaceur aléatoire à $(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$ puisse être complétée en une erreur d'entrée de l'encodeur convolutif formel $\mathcal{E}_{inN}$ dont le poids de syndrome est nul et dont l'image par $\mathcal{E}_{inN}$ est de poids d .

Il existe $a_{ext}^{N_{ext}}(w)$ erreurs vérifiant la première condition. Pour chacune de ces erreurs, comme l'entrelaceur formel est aléatoire, l'erreur obtenue en appliquant l'entrelaceur formel à $(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$ vaut de manière équiprobable l'une des $(|G| - 1)^w \binom{N_{ext}n_{ext}}{w}$ erreurs de taille $N_{ext}n_{ext}$ et de poids w . Ainsi, la probabilité qu'elle puisse être complétée en une erreur d'entrée de $\mathcal{E}_{inN}$ de poids de syndrome nul et dont l'image par $\mathcal{E}_{inN}$ est de poids d vaut :

$$\frac{a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext}n_{ext}}{w}}$$

D'après l'inégalité de Boole, la probabilité $p_N(w, d)$ de l'événement est donc inférieure à :

$$a_{ext}^{N_{ext}}(w) \frac{a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext}n_{ext}}{w}}$$

□

Avant de passer à la suite de la présentation, revenons à ce stade sur l'encodeur formel inversé et la structure convolutive associée que l'on a introduits dans les définitions 3.1.3, 3.2.4, et 3.2.6. La symétrie évoquée entre les structures convolutives directe et inversée va permettre, dans la section 4.4, de déduire la seconde borne sur $a_{in}^N(w, d)$ à partir de la première borne. Une telle manipulation requiert d'extrapoler la définition de $a_{in}^N(w, d)$ et $a_{in}^N(w, \leq d)$ au cas des encodeurs convolutifs formels inversés.

Définition 4.1.3. Soit $(n'_{in} + m'_{in}, k'_{in} + m'_{in}, \mathcal{E}'_{in})$ un encodeur formel inversé. Pour tout entier N et tout couple d'entiers (w, d) on considère $a'^N_{in}(w, d)$ et $a'^N_{in}(w, \leq d)$ le nombre d'erreurs

$$E' = (M', L'^{(0)}, \dots, L'^{(N-1)}) \in G^{m'_{in} + Nk'_{in}}$$

de poids w , dont l'image par $\mathcal{E}'_{inN}$ est de poids respectif d et inférieur ou égal à d .

La symétrie évoquée précédemment entre les deux structures convolutives se traduit sur les distributions de poids correspondantes de la manière suivante.

Lemme 4.1.2. Soit $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ un encodeur formel où m_{in} est implicitement connu, dont on note $(n'_{in} + m'_{in}, k'_{in} + m'_{in}, \mathcal{E}'_{in})$ l'encodeur formel inversé réciproque. Notons a^N_{in} et a'^N_{in} les distributions de poids définies précédemment associées respectivement à chacun d'eux. Alors pour tout couple d'entiers (w, d) :

$$a^N_{in}(w, d) \leq a'^N_{in}(d, w)$$

Démonstration. Notons ρ et ρ' les opérations suivantes, agissant respectivement sur l'entrée de mémoire et logique et sur la sortie de l'encodeur convolutif formel de taille N :

$$\begin{cases} \rho : E_{ML} = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}) & \mapsto (L^{(N-1)}, \dots, L^{(0)}, M^{(0)}) \\ \rho' : (P^{(0)}, \dots, P^{(N-1)}, M^{(N)}) & \mapsto (M^{(N)}, P^{(N-1)}, \dots, P^{(0)}) \end{cases}$$

Soit (w, d) un couple d'entiers. Notons $A_{w,d}$ l'ensemble des erreurs

$$E_{ML} = (M, L^{(0)}, \dots, L^{(N-1)}) \in G^{m_{in} + Nk_{in}}$$

de poids w , qui peuvent être complétées en une erreur d'entrée de l'encodeur convolutif formel

$$E = (M, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)})$$

de poids de syndrome nul et dont l'image $\mathcal{E}_{inN}(E)$ est de poids d . Soit f une application de $A_{w,d}$ dans $G^{m_{in} + Nn_{in}}$, associant à tout $E_{ML} \in A_{w,d}$ une telle erreur $\mathcal{E}_{inN}(E)$ de poids d où E est une extension de E_{ML} de poids de syndrome nul. $\mathcal{E}_{inN}$ étant une bijection, f est injective.

Notons par ailleurs $A'_{d,w}$ l'ensemble des erreurs d'entrée de l'encodeur convolutif formel inversé

$$E' = (M', L'^{(0)}, \dots, L'^{(N-1)}) \in G^{m'_{in} + Nn'_{in}}$$

de poids d , et dont l'image par $\mathcal{E}'_{inN}$ est de poids w . On observe alors, d'après le lemme 3.2.1, que $\text{Im}(\rho' \circ f) \subset A'_{d,w}$. En effet, considérons une erreur $E' \in \text{Im}(\rho' \circ f)$. Elle s'écrit sous la forme $\rho' \circ \mathcal{E}_{inN}(E)$, où E est une erreur de poids de syndrome nul et dont la partie de mémoire et logique E_{ML} est de poids w ,

et où $\mathcal{E}_{inN}(E)$ est de poids d . Par conséquent E' est de poids d , et d'après le lemme 3.2.1, E' est une erreur d'entrée de l'encodeur convolutif formel inversé, dont l'image $\mathcal{E}'_N(E')$ vérifie $\rho \circ \mathcal{E}'_N(E') = E_{ML}$ et est donc de poids w .

Cela démontre que $|A_{d,w}| \leq |A'_{d,w}|$, ou en d'autres termes $a'_{in}(w, d) \leq a'_{in}(d, w)$. $\square$

La suite du raisonnement que l'on développe dans les sections suivantes s'articule comme suit. Dans un premier temps, on établit à la section 4.2 les deux majorations sur la valeur de $a_{ext}^{N_{ext}}(w)$, d'abord en supposant que l'encodeur formel externe a pour distance dégénérée $d_{deg} > 2$, et ensuite en supposant que $|G| > 2$ et que l'encodeur formel externe a une distance minimale d_{min} et une distance dégénérée d_{deg} vérifiant $d_{min} > d_{deg} = 2$.

Ensuite, on s'intéresse à l'encodeur convolutif formel interne. La section 4.3 a pour but de mettre en place les notions de *période* d'un encodeur convolutif formel et de *trace* d'une erreur d'entrée d'un encodeur convolutif formel, qui permettront d'étudier les erreurs dont l'image par l'encodeur convolutif formel ont un poids d ou inférieur ou égal à d donné. La notion de *période* se retrouve sans démonstration dans la preuve de [22], mais dans le cas présent, à cause de l'existence d'erreurs de syndrome non triviales qui n'existent pas dans le cas classique, cet outil ainsi que celui de *trace* que l'on introduit dans cette thèse revêtent un caractère non trivial qui met en jeu une séparation des erreurs de mémoire en deux catégories. Cela permet alors, dans la section 4.4, d'établir deux majorations, l'une portant sur $a_{in}^N(w, \leq d)$ dans le cas où l'encodeur formel convolutif interne est récursif, et l'autre portant sur $a_{in}^N(w, d)$ dans le cas où il est anti-récursif.

Finalement, on rassemble les différentes majorations obtenues dans la section 4.5 en se basant sur le lemme 4.1.1 afin de démontrer les théorèmes 3.4.1 et 3.4.2.

On notera I_n l'erreur de taille n constituée de I sur toutes ses coordonnées.

4.2 Dénombrement d'erreurs au niveau de l'encodeur formel externe

On rappelle que l'encodeur formel externe consiste en une séquence d'encodeurs formels $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$ identiques placés en parallèle, et dont l'action, telle que présentée dans la définition 3.2.8, consiste à agir sur une erreur d'entrée sous la forme

$$(L_{ext}^{(0)}, \dots, L_{ext}^{(N_{ext}-1)}, S_{ext}^{(0)}, \dots, S_{ext}^{(N_{ext}-1)})$$

en lui associant l'erreur de sortie

$$(P_{ext}^{(0)}, \dots, P_{ext}^{(N_{ext}-1)})$$

définie par les relations $P_{ext}^{(i)} = \mathcal{E}_{ext}(L_{ext}^{(i)}, S_{ext}^{(i)})$. Afin de simplifier la présentation, on néglige dans cette section l'indice *ext*. Plaçons-nous d'abord dans les hy-

pothèses liées au théorème 3.4.1 qui correspond au cadre classique. On a la borne suivante.

Théorème 4.2.1. *Soit $(n, k, \mathcal{E})$ un encodeur formel de distance minimale d_{min} et de distance dégénérée $d_{deg} \geq 2$. Pour tous entiers d et N :*

$$a_{ext}^N(d) \leq \begin{cases} O(1)^d \left(\frac{N}{d}\right)^{\frac{d-d_{min}}{d_{deg}}+1} & \text{si } d \geq d_{min} \\ 0 & \text{si } d < d_{min} \end{cases} \quad (\text{Borne 1E})$$

Démonstration. Il s'agit de majorer le cardinal de l'ensemble $A_{N,d}$ d'erreurs

$$E = (L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{Nk} \times H^{N(n-k)}$$

où $|E|_L > 0$ et telle que l'image

$$P = (P^{(0)}, \dots, P^{(N-1)})$$

donnée par $P^{(i)} = \mathcal{E}(L^{(i)}, S^{(i)})$ est de poids d .

Si $d > N$, la borne est clairement vraie car on a alors $a_{ext}^N(d) = 0$. On suppose donc dorénavant que $d \leq N$. Pour toute erreur $E \in A_{N,d}$, soit $j(E)$ le nombre d'indices i tels que $(L^{(i)}, S^{(i)}) \neq I_n$, ce nombre étant non nul puisque $|E|_L > 0$. On va d'abord montrer que nécessairement, $d \geq d_{min}$ ce qui prouve la borne dans le cas où $d < d_{min}$, et que :

$$j(E) \leq \lfloor \frac{d - d_{min}}{d_{deg}} \rfloor + 1$$

Soit $E \in A_{N,d}$, et soit $j = j(E)$. Pour chaque indice i tel que $(L^{(i)}, S^{(i)}) \neq I_n$, $P^{(i)}$ est une erreur non triviale de $\mathcal{E}(G^k \times H^{n-k})$, donc $|P^{(i)}| \geq d_{deg}$ tel que présenté dans la définition 3.1.1 du formalisme commun. De plus, comme $|E|_L > 0$, il existe au moins un de ces j indices tel que $L^{(i)} \neq I_k$, de sorte que $P^{(i)}$ est une erreur nocive et $|P^{(i)}| \geq d_{min}$. Cela implique d'une part que $d \geq d_{min}$, ce qui prouve la borne $a_{ext}^N(d) = 0$ dans le cas où $d < d_{min}$. D'autre part, comme $d_{min} \geq d_{deg}$, cela implique que :

$$d \geq (j - 1)d_{deg} + d_{min}$$

puis :

$$j \leq \lfloor \frac{d - d_{min}}{d_{deg}} \rfloor + 1$$

Cela permet de majorer $a_{ext}^N(d)$. En effet, majorons d'abord, pour chaque valeur possible de j , le nombre d'erreurs $E \in A_{N,d}$ telles que $j(E) = j$. Il existe au plus $\binom{N}{j}$ manières de choisir les positions des j indices i . Ensuite, en chacune de ces j positions, l'erreur $(L^{(i)}, S^{(i)})$ ne peut prendre plus de $|G|^n$ valeurs, tandis qu'en toutes les autres positions, $(L^{(i)}, S^{(i)})$ est fixée à I_n . Il y a donc au plus

$$\binom{N}{j} |G|^{nj}$$

erreurs $E \in A_{N,d}$ telles que $j(E) = j$. Par conséquent :

$$a_{ext}^N(d) \leq \sum_{j=1}^{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} \binom{N}{j} |G|^{nj}$$

Comme $d_{min} \geq d_{deg} \geq 2$, on a pour tout indice j de la somme :

$$j \leq \lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1 \leq \lfloor \frac{d}{2} \rfloor$$

donc :

$$|G|^{nj} \leq |G|^{nd/2}$$

et comme d'autre part $d \leq N$, on a :

$$j \leq \lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1 \leq \frac{N}{2}$$

ce qui permet d'écrire :

$$\binom{N}{j} \leq \binom{N}{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1}$$

puis

$$\binom{N}{j} |G|^{nj} \leq \binom{N}{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} |G|^{nd/2}$$

et donc :

$$\begin{aligned} a_{ext}^N(d) &\leq \frac{d}{2} \binom{N}{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} |G|^{nd/2} \\ &\leq O(1)^d \binom{N}{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} \end{aligned}$$

Utilisons à présent la propriété suivante pour tout couple d'entiers $u \geq v$:

$$\binom{u}{v} \leq \left(\frac{ue}{v} \right)^v$$

Appliquée au cas présent, cela donne :

$$\begin{aligned} a_{ext}^N(d) &\leq O(1)^d \left(\frac{Ne}{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} \right)^{\lfloor \frac{d-d_{min}}{d_{deg}} \rfloor + 1} \\ &\leq O(1)^d \left(\frac{N}{d} \right)^{\frac{d-d_{min}}{d_{deg}} + 1} \end{aligned}$$

□

On va à présent démontrer la seconde borne, qui s'applique dans les conditions du théorème 3.4.2 correspondant au cadre quantique. Comme on le présente ci-dessous, cette borne comporte l'expression :

$$(|G| - 1)^d \binom{Nn}{d}$$

qui désigne le nombre d'erreurs de taille N et de poids d . Cette borne revient ainsi à affirmer que le nombre $a_{ext}^N(d)$ d'erreurs nocives de poids d est exponentiellement négligeable en d devant le nombre d'erreurs de taille N et de poids d .

Théorème 4.2.2. *On suppose que $|G| \geq 4$. Soit $(n, k, \mathcal{E})$ un encodeur formel de distance dégénérée $d_{deg} \geq 2$. Il existe une constante $c \in]0; 1[$ telle que pour tous entiers d et N :*

$$a_{ext}^N(d) \leq c^d (|G| - 1)^d \binom{Nn}{d} \quad (\text{Borne } 2E)$$

Démonstration. Il s'agit de nouveau de majorer le cardinal de l'ensemble $A_{N,d}$ d'erreurs

$$E = (L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{Nk} \times H^{N(n-k)}$$

où $|E|_L > 0$ et telle que l'image

$$P = (S^{(0)}, \dots, S^{(N-1)})$$

donnée par $P^{(i)} = \mathcal{E}(L^{(i)}, S^{(i)})$ est de poids d .

Le groupe G possède la propriété suivante utile pour la suite : il existe trois éléments distincts X, Y et Z de $G \setminus \{I\}$ vérifiant $X \circ Y = Z$. Il suffit pour cela de choisir un élément $X \in G \setminus \{I\}$, puis comme $|G| \geq 4$, un élément $Y \in G \setminus \{I\}$ différent de X et de l'inverse de X . Alors l'élément Z défini par $Z = X \circ Y$ ne peut être égal à X ou à Y , sinon on aurait respectivement $Y = I$ et $X = I$, ni à I sinon Y serait l'inverse de X , ce qui montre bien que X, Y et Z sont distincts et différents de I .

Définissons pour tout $p \in \llbracket 0; n \rrbracket$ l'ensemble

$$C_p = \{E \in \mathcal{E}(G^k \times H^{n-k}), |E| = i\}$$

et le réel

$$c_p = \frac{|C_p|}{(|G| - 1)^p \binom{n}{p}}$$

Montrons alors que $c_0 = 1$ et pour tout $p \in \llbracket 1; n \rrbracket$, $c_p < 1$.

Le fait que $c_0 = 1$ est immédiat car $C_0 = \{I_n\}$. Soit à présent $p \in \llbracket 1; n \rrbracket$. Le fait que $c_p < 1$ équivaut à ce que $\mathcal{E}(G^k \times H^{n-k})$ ne contient pas toutes les erreurs de poids p . Pour montrer cela, supposons que $\mathcal{E}(G^k \times H^{n-k})$ contient l'erreur de poids p

$$E^{(X)} = (X, X, \dots, X, I, \dots, I)$$

dont les p premières coordonnées valent X et les $n - p$ restantes valent I , et l'erreur de poids p

$$E^{(Z)} = (Z, X, \dots, X, I, \dots, I)$$

dont la première coordonnée vaut Z , les $p - 1$ suivantes valent X et les $n - p$ restantes valent I . Comme $\mathcal{E}(G^k \times H^{n-k})$ est un sous-groupe de G^n , il contient alors l'erreur $E^{(X)^{-1}} \circ E^{(Z)}$ dont la première coordonnée vaut Y et les $n - 1$ restantes valent I . Cette erreur est de poids 1, ce qui contredit le fait que $d_{deg} \geq 2$ est le plus petit poids d'une erreur non triviale de $\mathcal{E}(G^k \times H^{n-k})$. Ainsi $\mathcal{E}(G^k \times H^{n-k})$ ne contient pas à la fois $E^{(X)}$ et $E^{(Z)}$, et $c_p < 1$.

Posons alors $c = \max\{c_p^{1/p}, 1 \leq p \leq n\}$, de sorte que $c \in]0; 1[$ et pour tout $p \in \llbracket 0; n \rrbracket$, $c_p \leq c^p$. Pour toute erreur $E \in A_{N,d}$, si pour tout $i \in \llbracket 0; N - 1 \rrbracket$ on note $p_i = |P^{(i)}|$, alors $P^{(i)} \in C_{p_i}$ car $P^{(i)} \in \mathcal{E}(G^k \times H^{n-k})$; ainsi :

$$P \in C_{p_0} \times \dots \times C_{p_{N-1}}$$

où

$$\sum_{i=0}^{N-1} p_i = d$$

Par conséquent :

$$\begin{aligned} a_{ext}^N(d) &\leq \sum_{\substack{(p_0, \dots, p_{N-1}) \\ 0 \leq p_i \leq n \\ \sum p_i = d}} \prod_{i=0}^{N-1} |C_{p_i}| \\ &\leq \sum_{\substack{(p_0, \dots, p_{N-1}) \\ 0 \leq p_i \leq n \\ \sum p_i = d}} \prod_{i=0}^{N-1} c^{p_i} (|G| - 1)^{p_i} \binom{n}{p_i} \\ &\leq c^d (|G| - 1)^d \sum_{\substack{(p_0, \dots, p_{N-1}) \\ 0 \leq p_i \leq n \\ \sum p_i = d}} \prod_{i=0}^{N-1} \binom{n}{p_i} \\ &\leq c^d (|G| - 1)^d \binom{Nn}{d} \end{aligned}$$

où le passage à la dernière ligne s'obtient en constatant que choisir d éléments parmi Nn revient, en considérant les Nn éléments comme N blocs successifs de n éléments, à choisir p_i éléments dans chaque bloc i de sorte que

$$\sum_{i=0}^{N-1} p_i = d$$

et que pour tout i , $p_i \in \llbracket 0; n \rrbracket$.

□

4.3 L'encodeur convolutif interne : période de l'encodeur, et trace d'une erreur d'entrée

Tous les résultats de cette section s'appliquent aussi bien aux *encodeurs formels* qu'aux *encodeurs formels inversés* introduits dans la définition 3.1.3. Afin de rendre plus simple la présentation, le terme *encodeur formel* désigne dans cette section aussi bien un encodeur formel qu'un encodeur formel inversé, et de même, le terme *encodeur convolutif formel* (de taille finie ou infinie) désigne aussi bien un encodeur convolutif formel qu'un encodeur convolutif formel inversé. Dans le cas d'un encodeur formel inversé, toutes les séquences d'erreurs de syndrome correspondent à la séquence vide, ce qui revient simplement à reprendre les mêmes énoncés ci-dessous sans intégrer les erreurs de syndrome. Cela entraîne entre autres que les résultats que l'on obtient sur le nombre d'erreurs $a_{in}^N(w, d)$ lié à l'encodeur convolutif formel se transposent sur le nombre d'erreurs $a'_{in}^N(w, d)$ lié à l'encodeur convolutif formel inversé.

De la même manière que précédemment, on rend implicite l'indice in dans cette section et la suivante qui traitent de l'encodeur convolutif formel interne.

Tentons une description approchée du raisonnement que l'on va suivre. On va établir une partition, qui dépend de l'encodeur convolutif formel, de l'ensemble G^m en deux ensembles $\mathbb{M}_0$ et $\mathbb{M}_1$. Schématiquement, on montrera, grâce à la notion de *période* de l'encodeur convolutif formel, que le poids de l'erreur de sortie croît de manière proportionnelle avec le nombre d'étapes d'indice i où l'erreur de mémoire intermédiaire donnée par la définition 3.2.1 appartient à $\mathbb{M}_1$. De plus, on montrera dans le cas *récuratif* que dès qu'une erreur logique d'entrée L_i de poids 1 est rencontrée, si l'erreur de mémoire intermédiaire à l'étape i appartient à $\mathbb{M}_0$, elle passe nécessairement à $\mathbb{M}_1$ à l'étape $i + 1$, et y reste au moins aussi longtemps qu'une nouvelle erreur logique d'entrée de poids 1 n'est pas rencontrée. Si le poids de l'erreur de sortie doit rester dominé par une certaine valeur d , cela permet d'établir une contrainte sur la répartition des positions en entrée où la coordonnée de l'erreur logique est non triviale. Cette contrainte permet ensuite d'aboutir à une majoration de $a_{in}^N(w, d)$ et $a_{in}^N(w, \leq d)$. On attire l'attention sur le fait que le raisonnement décrit ci-dessus ne correspond pas au raisonnement exact, qui manipule en réalité des erreurs de mémoire différentes des erreurs de mémoire intermédiaires. En effet, le raisonnement exact tente de décrire plus finement l'effet de chaque coordonnée logique non triviale de l'erreur d'entrée sur la croissance du poids de l'erreur de sortie, ce qui nécessite d'introduire la notion de *troncature* d'une erreur d'entrée comme on le précise par la suite.

Intuitivement, l'appartenance d'une erreur de mémoire à $\mathbb{M}_0$ ou $\mathbb{M}_1$ ressemble à l'état d'un interrupteur. Une erreur de mémoire d'entrée dans $\mathbb{M}_1$ engendre une erreur de sortie dont le poids est infini, et comme on l'énoncera plus précisément, dont la restriction à une taille donnée en sortie est de poids essentiellement linéaire avec cette taille. Cela correspond à l'intuition d'un interrupteur fermé qui produit un poids proportionnel au temps d'exécution. Dans le cas d'une erreur de mémoire d'entrée dans $\mathbb{M}_0$, l'erreur de sortie engendrée est

de poids nul à partir d'une certaine coordonnée, ce qui correspond à l'intuition d'un interrupteur ouvert.

Voici par conséquent comment s'organise cette section. On commence par introduire la partition de l'ensemble des erreurs de mémoire en $\mathbb{M}_0$ et $\mathbb{M}_1$, puis on démontre l'existence de la *période* η de l'encodeur convolutif formel. Ensuite, on introduit les *troncatures* successives d'une erreur d'entrée, opération qui consiste à couper l'erreur d'entrée en une position où la coordonnée logique diffère de I , et à la compléter par des coordonnées triviales pour obtenir une erreur de format adéquat pour l'encodeur convolutif formel. On définit alors des erreurs de mémoire particulières *liées à ces troncatures*, dont l'esprit se rapproche des erreurs de mémoire intermédiaires et sur lesquelles portera précisément l'étude. Afin d'étudier le comportement de ces erreurs de mémoire et leur effet sur le poids de l'erreur de sortie, on définit la *trace* d'une erreur d'entrée, séquence binaire permettant de savoir à quel ensemble parmi $\mathbb{M}_0$ et $\mathbb{M}_1$ appartient chacune de ces erreurs de mémoire. On montre alors que dans le cas récursif, les coordonnées logiques non triviales de l'erreur d'entrée ont pour effet de découper la trace d'une erreur d'entrée en *détours*, dont on présente les conséquences sur le poids de l'erreur de sortie. L'application des résultats de cette étude à la section suivante permet de majorer les cardinaux $a_{in}^N(w, d)$ et $a_{in}^N(w, \leq d)$.

Avant de commencer, établissons un lemme qui servira pour de nombreuses manipulations. Introduisons d'abord la notation suivante permettant de séparer la sortie physique et la sortie de mémoire d'un encodeur convolutif formel.

Définition 4.3.1. Considérons un encodeur convolutif formel $(m + Nn, m + Nk, \mathcal{E}_N)$ de taille N et paramètres $[n, k, m]$ basé sur un encodeur formel $(n + m, k + m, \mathcal{E})$. On note π_N et μ_N les morphismes associant respectivement la sortie physique et la sortie de mémoire à une erreur d'entrée. Ainsi pour tout $E \in G^{m+Nn}$:

$$\mathcal{E}_N(E) = (\pi_N(E), \mu_N(E))$$

où $\pi_N(E) \in G^{Nn}$ et $\mu_N(E) \in G^m$.

Le lemme suivant découle immédiatement de la structure convolutive.

Lemme 4.3.1. Soit $(n + m, k + m, \mathcal{E})$ un encodeur formel et soit $(N_1, N_2) \in \mathbb{N}^2$ où $N_1 < N_2$. Pour toute erreur d'entrée de $\mathcal{E}_{N_2}$ écrite sous la forme

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N_2-1)}, S^{(0)}, \dots, S^{(N_2-1)})$$

si l'on note

$$\bar{E} = (M^{(0)}, L^{(0)}, \dots, L^{(N_1-1)}, S^{(0)}, \dots, S^{(N_1-1)})$$

on a :

$$\mathcal{E}_{N_2}(E) = (\pi_{N_1}(\bar{E}), \mathcal{E}_{N_2-N_1}(\mu_{N_1}(\bar{E}), L^{(N_1)}, \dots, L^{(N_2-1)}, S^{(N_1)}, \dots, S^{(N_2-1)}))$$

De même, pour toute erreur d'entrée de $\mathcal{E}_\infty$ écrite sous la forme

$$E = (M^{(0)}, L^{(0)}, S^{(0)}, L^{(1)}, S^{(1)}, \dots)$$

si l'on note

$$\bar{E} = (M^{(0)}, L^{(0)}, \dots, L^{(N_1-1)}, S^{(0)}, \dots, S^{(N_1-1)})$$

on a :

$$\mathcal{E}_\infty(E) = (\pi_{N_1}(\bar{E}), \mathcal{E}_\infty(\mu_{N_1}(\bar{E}), L^{(N_1)}, S^{(N_1)}, L^{(N_1+1)}, S^{(N_1+1)}, \dots))$$

A présent, établissons la partition de l'ensemble G^m des erreurs de mémoire. On rappelle que l'encodeur convolutif formel infini introduit à la définition 3.2.3 prend des erreurs d'entrée sous un format particulier à cause de son caractère infini, alternant erreurs logiques et erreurs de syndrome.

Définition 4.3.2. Soit $(n + m, k + m, \mathcal{E})$ un encodeur formel. On définit alors les ensembles :

$$\mathbb{M}_0 = \{M \in G^m, \exists (S^{(0)}, S^{(1)}, \dots) \in (H^{n-k})^{\mathbb{N}} / |\mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots)| < \infty\}$$

$$\mathbb{M}_1 = \{M \in G^m, \forall (S^{(0)}, S^{(1)}, \dots) \in (H^{n-k})^{\mathbb{N}}, |\mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots)| = \infty\}$$

Notons que $\mathbb{M}_0$ est un sous-groupe de G^m . On fera usage par la suite de la propriété suivante de stabilité de l'ensemble $\mathbb{M}_1$.

Lemme 4.3.2. Pour tout $i > 0$ et pour tout $(S^{(0)}, \dots, S^{(i-1)}) \in (H^{n-k})^i$, l'application

$$\begin{aligned} G^m &\rightarrow G^m \\ M &\mapsto \mu_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)}) \end{aligned}$$

stabilise l'ensemble $\mathbb{M}_1$.

Démonstration. Soient $i > 0$, $(S^{(0)}, \dots, S^{(i-1)}) \in (H^{n-k})^i$ et $M \in \mathbb{M}_1$. On veut montrer que l'erreur de mémoire

$$M' = \mu_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)})$$

appartient à $\mathbb{M}_1$. Pour toute suite $(S^{(i)}, S^{(i+1)}, \dots) \in (H^{n-k})^{\mathbb{N}}$ d'erreurs de syndrome, on a d'après le lemme 4.3.1 :

$$\begin{aligned} \mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots) &= (\pi_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)}), \\ &\mathcal{E}_\infty(M', L^{(i)}, S^{(i)}, L^{(i+1)}, S^{(i+1)}, \dots)) \end{aligned}$$

ce qui montre que :

$$|\mathcal{E}_\infty(M', L^{(i)}, S^{(i)}, L^{(i+1)}, S^{(i+1)}, \dots)| = \infty$$

et que $M' \in \mathbb{M}_1$. □

Présentons à présent la notion de *période* d'un encodeur convolutif formel. Ce terme est choisi car dans le cas d'un encodage classique, une sortie de poids infini générée à partir d'une entrée de poids fini est effectivement périodique à partir

d'une certaine coordonnée. Cette régularité n'est pas garantie dans le contexte présent à cause de la liberté de choix existant sur les erreurs de syndrome de l'entrée, mais l'idée selon laquelle l'erreur de sortie contient toujours un poids non nul sur toute fenêtre de taille fixée reste valable. L'existence de la période est prouvée par le théorème suivant.

Théorème 4.3.1. *Soit $(n+m, k+m, \mathcal{E})$ un encodeur formel. Il existe un entier positif η tel que :*

$$\forall M \in \mathbb{M}_1, \forall (S^{(0)}, \dots, S^{(\eta-1)}) \in (H^{n-k})^\eta, |\pi_\eta(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(\eta-1)})| \geq 1$$

*Pour tout encodeur convolutif formel basé sur l'encodeur formel $(n+m, k+m, \mathcal{E})$, le plus petit entier η vérifiant cette propriété est nommé **période de l'encodeur convolutif formel**.*

Démonstration. On raisonne par l'absurde en supposant que pour tout $\eta \in \mathbb{N}^*$:

$$\begin{aligned} \exists M^\eta \in \mathbb{M}_1, \exists (S^{(0),\eta}, \dots, S^{(\eta-1),\eta}) \in (H^{n-k})^\eta / \\ |\pi_\eta(M, I_k, \dots, I_k, S^{(0),\eta}, \dots, S^{(\eta-1),\eta})| = 0 \end{aligned}$$

Construisons alors une erreur de mémoire $M \in \mathbb{M}_1$ et une suite d'erreurs de syndrome $(S^{(i)})_{i \in \mathbb{N}} \in (H^{n-k})^\mathbb{N}$ telles que

$$|\mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots)| = 0$$

ce qui contredira le fait que $M \in \mathbb{M}_1$.

Comme la suite $(M^\eta)_{\eta \in \mathbb{N}^*}$ est à valeurs dans l'ensemble fini $\mathbb{M}_1$, il existe $M \in \mathbb{M}_1$ tel que l'ensemble :

$$A = \{\eta \in \mathbb{N}^* / M^\eta = M\}$$

est de cardinal infini.

A présent, pour tout $i \in \mathbb{N}$ et tout $(S^{(0)}, \dots, S^{(i)}) \in (H^{n-k})^{i+1}$, définissons l'ensemble :

$$A_{(S^{(0)}, \dots, S^{(i)})} = \{\eta \in A, \eta > i, (S^{(0),\eta}, \dots, S^{(i),\eta}) = (S^{(0)}, \dots, S^{(i)})\}$$

Il est clair que ces ensembles vérifient les relations suivantes :

$$A = \bigcup_{S \in H^{n-k}} A_{(S)} \quad (4.3.1)$$

et pour tout $i \in \mathbb{N}$ et pour tout $(S^{(0)}, \dots, S^{(i)}) \in (H^{n-k})^{i+1}$:

$$A_{(S^{(0)}, \dots, S^{(i)})} \setminus \{i+1\} = \bigcup_{S \in H^{n-k}} A_{(S^{(0)}, \dots, S^{(i)}, S)} \quad (4.3.2)$$

Ce constat permet de construire une suite $(S^{(i)})_{i \in \mathbb{N}} \in (H^{n-k})^\mathbb{N}$ telle que pour tout $i \in \mathbb{N}$, $A_{(S^{(0)}, \dots, S^{(i)})}$ est de cardinal infini. En effet l'égalité (4.3.1) garantit

l'existence d'une erreur de syndrome $S^{(0)}$ telle que $A_{(S^{(0)})}$ est de cardinal infini, puis à chaque étape i , $A_{(S^{(0)}, \dots, S^{(i)})}$ étant de cardinal infini, l'égalité (4.3.2) garantit l'existence d'une erreur de syndrome $S^{(i+1)}$ telle que $A_{(S^{(0)}, \dots, S^{(i+1)})}$ est de cardinal infini.

L'erreur M et la suite d'erreurs $(S^{(i)})_{i \in \mathbb{N}}$ produisent alors la contradiction recherchée. En effet, pour tout $i \in \mathbb{N}$, on peut choisir un entier $\eta \in A_{(S^{(0)}, \dots, S^{(i)})}$ ce qui permet d'écrire :

$$\begin{aligned} |\pi_{i+1}(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i)})| &= |\pi_{i+1}(M^\eta, I_k, \dots, I_k, S^{(0), \eta}, \dots, S^{(i), \eta})| \\ &= 0 \end{aligned}$$

Comme pour tout $i \in \mathbb{N}$,

$$\pi_{i+1}(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i)})$$

représente les $n(i+1)$ premières coordonnées de

$$\mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots)$$

cela prouve que

$$|\mathcal{E}_\infty(M, I_k, S^{(0)}, I_k, S^{(1)}, \dots)| = 0$$

et contredit le fait que $M \in \mathbb{M}_1$. □

A partir de ce théorème et de la stabilité de l'ensemble $\mathbb{M}_1$ présentée au lemme 4.3.2, on obtient la relation suivante.

Lemme 4.3.3. *Pour tout $i > 0$, $M \in \mathbb{M}_1$ et $(S^{(0)}, \dots, S^{(i-1)}) \in (H^{n-k})^i$:*

$$|\pi_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)})| \geq \lfloor i/\eta \rfloor$$

Démonstration. Prouvons le résultat par récurrence sur $p = \lfloor i/\eta \rfloor$.

Le résultat est évident lorsque $p = 0$. Supposons à présent que le résultat est vrai pour un certain rang $p - 1$ où $p \geq 1$. Soit i tel que $\lfloor i/\eta \rfloor = p$ et soient $M \in \mathbb{M}_1$ et $(S^{(0)}, \dots, S^{(i-1)}) \in (H^{n-k})^i$. Notons

$$M' = \mu_\eta(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(\eta-1)})$$

D'après le lemme 4.3.2, $M' \in \mathbb{M}_1$. Et d'après le lemme 4.3.1 :

$$\begin{aligned} \pi_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)}) &= (\pi_\eta(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(\eta-1)}), \\ &\quad \pi_{i-\eta}(M', I_k, \dots, I_k, S^{(\eta)}, \dots, S^{(i-1)})) \end{aligned}$$

Or d'après le théorème 4.3.1 :

$$|\pi_\eta(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(\eta-1)})| \geq 1$$

et comme $M' \in \mathbb{M}_1$, l'hypothèse de récurrence au rang $p - 1$ donne :

$$|\pi_{i-\eta}(M', I_k, \dots, I_k, S^{(\eta)}, \dots, S^{(i-1)})| \geq p - 1$$

Par conséquent :

$$|\pi_i(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(i-1)})| \geq p$$

ce qui prouve le résultat au rang p . □

On définit à présent les troncatures de l'erreur d'entrée ainsi que les erreurs de mémoire liées à ces troncatures. L'idée poussant à définir la notion de troncature est la suivante. On souhaite appliquer l'encodeur convolutif formel sur E , étape par étape, jusqu'à l'étape N_i où l'on rencontre la i -ème coordonnée logique non triviale de E . On souhaite alors obtenir une sorte d'« arrêt sur image » permettant de décrire l'état de la sortie, et notamment l'état de la mémoire en sortie, à l'instant précis où cette coordonnée en entrée est rencontrée. On s'intéresse en d'autres termes à affiner le modèle des erreurs de mémoire intermédiaires, en examinant l'erreur de mémoire à l'étape virtuelle où la i -ème coordonnée logique non triviale de E est rencontrée. Or, si l'on note p_i la position de cette coordonnée, p_i n'est pas nécessairement un multiple de k , et l'opération de convolution ne peut s'arrêter à cette position. C'est pourquoi on complète les p_i premières coordonnées logiques de E par un nombre minimal de coordonnées égales à I pour former une erreur dont le nombre de coordonnées logiques est un multiple de k . La partie de syndrome de E est quant à elle reprise jusqu'à la position correspondant à l'application de l'étape N_i de l'encodeur convolutif formel. Par conséquent, si E s'écrit sous la forme suivante :

$$E = (M, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)})$$

la i -ème troncature de E aura la forme ci-dessous :

$$E^{\rightarrow i} = (M, L^{(0)}, \dots, L^{(N_i-1)}, L^{\rightarrow i}, S^{(0)}, \dots, S^{(N_i)})$$

où l'erreur $L^{\rightarrow i}$ peut différer de l'erreur $L^{(N_i-1)}$ sur ses dernières positions à cause de la coupure décrite précédemment. Ainsi, $L^{\rightarrow i}$ est constituée des $p_i - N_i k$ premières coordonnées de $L^{(N_i-1)}$, complétées de $N_i(k+1) - p_i$ coordonnées égales à I . L'erreur de mémoire liée à la i -ème troncature de E est alors simplement l'erreur de mémoire obtenue en appliquant l'encodeur convolutif formel, jusqu'à l'étape N_i , sur la i -ème troncature de E . Illustrons la notion de troncature par l'exemple suivant.

Exemple 1. On suppose que $G = \{I, X, Y, Z\}$, et on considère un encodeur convolutif formel $(26, 16, \mathcal{E}_5)$ de taille $N = 5$ et de paramètres $[n, k, m] = [5, 3, 1]$ basé sur un encodeur formel $(6, 4, \mathcal{E})$. Soit E l'erreur d'entrée suivante où les virgules sont indiquées par commodité :

$$E = (X, IZZ, III, XII, YZI, III, YY, XZ, II, ZZ, ZI)$$

E est de poids logique 5 et possède donc 5 troncatures. Les positions des coordonnées logiques non triviales de E sont :

$$p_1 = 2, \quad p_2 = 3, \quad p_3 = 7, \quad p_4 = 10, \quad p_5 = 11$$

Notons bien qu'on recense ces positions en considérant uniquement la partie logique de E , de sorte que ces positions sont comprises entre 1 et Nk soit 15 dans le cas présent. Les étapes de l'encodeur convolutif formel où ces coordonnées sont rencontrées sont :

$$N_1 = 0, \quad N_2 = 0, \quad N_3 = 2, \quad N_4 = 3, \quad N_5 = 3$$

Les troncatures de E sont les erreurs suivantes :

$$\begin{aligned}
E^{\rightarrow 1} &= (X, IZI, YY) \\
E^{\rightarrow 2} &= (X, IZZ, YY) \\
E^{\rightarrow 3} &= (X, IZZ, III, XII, YY, XZ, II) \\
E^{\rightarrow 4} &= (X, IZZ, III, XII, YII, YY, XZ, II, ZZ) \\
E^{\rightarrow 5} &= (X, IZZ, III, XII, YZI, YY, XZ, II, ZZ)
\end{aligned}$$

et les erreurs logiques notées $L^{\rightarrow i}$ sont les suivantes :

$$L^{\rightarrow 1} = IZI, \quad L^{\rightarrow 2} = IZZ, \quad L^{\rightarrow 3} = XII, \quad L^{\rightarrow 4} = YII, \quad L^{\rightarrow 5} = YZI$$

Voici la définition des troncatures de E et des erreurs de mémoire liées à ces troncatures.

Définition 4.3.3. Soit

$$E = (M, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)})$$

une erreur d'entrée d'un encodeur convolutif formel de taille N et de paramètres $[n, k, m]$, et soit $w_L = |E|_L$. Les w_L coordonnées différentes de I de l'erreur $(L^{(0)}, \dots, L^{(N-1)})$ se situent à des positions que l'on note

$$1 \leq p_1 < \dots < p_{w_L} \leq Nk$$

et font partie des erreurs respectives $L^{(N_1)}, \dots, L^{(N_{w_L})}$, de sorte que pour tout $i \in \llbracket 1; w_L \rrbracket$:

$$N_i k + 1 \leq p_i \leq N_i k + k$$

Pour tout $i \in \llbracket 1; w_L \rrbracket$, notons $L^{\rightarrow i} \in G^k$ l'erreur qui coïncide avec $L^{(N_i)}$ sur les $p_i - N_i k$ premières coordonnées et dont le reste des coordonnées vaut I .

Pour tout $i \in \llbracket 1; w_L \rrbracket$, on définit alors **la i -ème troncature de E** :

$$E^{\rightarrow i} = (M, L^{(0)}, \dots, L^{(N_i-1)}, L^{\rightarrow i}, S^{(0)}, \dots, S^{(N_i)})$$

ainsi que **l'erreur de mémoire liée à la i -ème troncature** :

$$M^{\rightarrow i} = \mu_{N_i+1}(E^{\rightarrow i})$$

Par convention pour $i = 0$, on note également :

$$E^{\rightarrow 0} = M^{\rightarrow 0} = M$$

ainsi que :

$$p_0 = N_0 = 0$$

Pour une erreur d'entrée E donnée, la *trace* de E renseigne alors sur l'appartenance à $\mathbb{M}_0$ ou $\mathbb{M}_1$ des erreurs de mémoire des troncatures de E .

Définition 4.3.4. La **trace** d'une erreur E d'entrée d'un encodeur convolutif formel, de poids logique w_L , est le $w_L + 1$ -uplet $(b_0, \dots, b_{w_L}) \in \{0; 1\}^{w_L+1}$ tel que, pour tout $i \in \llbracket 0; w_L \rrbracket$:

$$M^{\rightarrow i} \in \mathbb{M}_{b_i}$$

L'effet du caractère récursif de l'encodeur convolutif formel apparaît alors ainsi.

Lemme 4.3.4. Soit E une erreur d'entrée d'un encodeur convolutif formel récursif, et soit $(b_0, \dots, b_{w_L})$ la trace de E . Alors, pour tout $i \in \llbracket 0; w_L - 1 \rrbracket$, si $b_i = 0$ alors $b_{i+1} = 1$.

Démonstration. Reprenons les notations de la définition 4.3.3. Soit $i \in \llbracket 0; w_L - 1 \rrbracket$ tel que $b_i = 0$. Cela signifie que $M^{\rightarrow i} \in \mathbb{M}_0$, donc il existe une suite d'erreurs de syndrome

$$(S'^{(0)}, S'^{(1)}, \dots) \in (H^{n-k})^{\mathbb{N}}$$

telle que :

$$|\mathcal{E}_\infty(M^{\rightarrow i}, I_k, S'^{(0)}, I_k, S'^{(1)}, \dots)| < \infty$$

Ainsi si l'on définit l'erreur d'entrée de taille infinie qui vaut

$$E_\infty^{\rightarrow i} = (M, L^{(0)}, S^{(0)}, \dots, L^{(N_i-1)}, S^{(N_i-1)}, L^{\rightarrow i}, S^{(N_i)}, \\ I_k, S'^{(0)}, I_k, S'^{(1)}, \dots)$$

si $i \geq 1$, et

$$E_\infty^{\rightarrow i} = (M, I_k, S'^{(0)}, I_k, S'^{(1)}, \dots)$$

si $i = 0$, le lemme 4.3.1 permet d'écrire :

$$|\mathbb{E}_\infty(E_\infty^{\rightarrow i})| < \infty$$

On veut montrer que $b_i = 1$, c'est-à-dire que $M^{\rightarrow i+1} \in \mathbb{M}_1$, ou encore que pour toute suite d'erreurs de syndrome

$$(S''^{(0)}, S''^{(1)}, \dots) \in (H^{n-k})^{\mathbb{N}}$$

on a :

$$|\mathcal{E}_\infty(M^{\rightarrow i+1}, I_k, S''^{(0)}, I_k, S''^{(1)}, \dots)| = \infty$$

Considérons une telle suite d'erreurs de syndrome et notons également

$$E_\infty^{\rightarrow i+1} = (M, L^{(0)}, S^{(0)}, \dots, L^{(N_{i+1}-1)}, S^{(N_{i+1}-1)}, L^{\rightarrow i+1}, S^{(N_{i+1})}, \\ I_k, S''^{(0)}, I_k, S''^{(1)}, \dots)$$

Encore une fois d'après le lemme 4.3.1, cela revient à montrer que :

$$|\mathcal{E}_\infty(E_\infty^{\rightarrow i+1})| = \infty$$

Or, les erreurs de taille infinie $E_\infty^{\rightarrow i}$ et $E_\infty^{\rightarrow i+1}$ ont la même erreur de mémoire, des erreurs de syndrome dans H , et d'après les définitions des troncatures $E^{\rightarrow i}$

et $E^{\rightarrow i+1}$, diffèrent sur leur partie logique en une position unique correspondant à la $i + 1$ -ème coordonnée logique non triviale de E . On peut donc écrire le produit suivant dans $G^{\mathbb{N}}$:

$$E_{\infty}^{\rightarrow i+1} = E_{\infty}^{\rightarrow i} E_{\Delta}$$

où l'erreur de taille infinie E_{Δ} vérifie $|E_{\Delta}|_{\bar{H}} = 0$, $|E_{\Delta}|_M = 0$ et $|E_{\Delta}|_L = 1$. D'après la définition 3.3.1 de la récursivité on a donc :

$$|\mathcal{E}_{\infty}(E_{\Delta})| = \infty$$

et par conséquent :

$$|\mathcal{E}_{\infty}(E_{\infty}^{\rightarrow i+1})| = \infty$$

□

Dans le cas d'un encodeur convolutif formel récursif, la trace d'une erreur d'entrée se découpe donc en *détours* dont la définition est la suivante.

Définition 4.3.5. Un **détour** est une séquence binaire finie non vide dont les éléments valent 1, éventuellement complétée par un élément égal à 0. Un détour est dit **régulier** s'il est complété par un 0 et **terminal** sinon.

Le découpage en détours de la trace d'une erreur d'entrée s'énonce ainsi.

Lemme 4.3.5. Soit E une erreur d'entrée d'un encodeur convolutif formel récursif. On note $w_L = |E|_L$ en supposant que $w_L \geq 1$, et $w = |E|_M + |E|_L$. Soit $(b_0, \dots, b_{w_L})$ la trace de E . Alors la séquence

$$\begin{cases} (b_0, \dots, b_{w_L}) & \text{si } b_0 = 1 \\ (b_1, \dots, b_{w_L}) & \text{si } b_0 = 0 \end{cases}$$

consiste en une concaténation de $c - 1$ détours réguliers, où $c \geq 1$, suivis d'un détour. De plus :

$$\begin{cases} c \leq \lfloor \frac{w+1}{2} \rfloor & \text{si le dernier détour est terminal} \\ c \leq \lfloor \frac{w}{2} \rfloor & \text{si le dernier détour est régulier} \end{cases}$$

Démonstration. Une telle partition en détours est une conséquence immédiate du fait que pour tout $i \in \llbracket 0; w_L - 1 \rrbracket$, si $b_i = 0$ alors $b_{i+1} = 1$. On remarque ensuite que la concaténation de ces détours, qui vaut $(b_1, \dots, b_{w_L})$ si $b_0 = 1$ et $(b_0, \dots, b_{w_L})$ si $b_0 = 0$, contient au plus w éléments car d'une part, si $b_0 = 1$ alors $w_L \leq w$ et d'autre part, si $b_0 = 0$ alors $|E|_M \geq 1$ et donc $w_L + 1 \leq w$. De plus, un détour régulier contient au moins deux éléments et un détour terminal au moins un élément. On compte ainsi dans la séquence précédente au moins $2c - 1$ éléments si le dernier détour est terminal et au moins $2c$ éléments s'il est régulier ce qui prouve la majoration de c . □

Pour une erreur d'entrée E dont les positions des coordonnées logiques non triviales sont données par $p_1 < \dots < p_{w_L}$, on va montrer que le poids de l'erreur de sortie est minoré par une fonction affine de la somme, sur tous les détours de la trace de E , des écarts entre les positions p_i et $p_{i'}$ qui correspondent aux indices de début i et de fin i' de chaque détour.

Théorème 4.3.2. *Soit E une erreur d'entrée d'un encodeur convolutif formel récursif de taille N , de paramètres $[n, k, m]$ et de période η . On note*

$$1 \leq p_1 < \dots < p_{w_L} \leq Nk$$

les positions des $w_L = |E|_L$ coordonnées logiques non triviales de E ,

$$0 \leq v_1 < \dots < v_c \leq w_L$$

les indices de départ des c détours de la trace de E , et

$$v_{c+1} = w_L + 1$$

Le poids d de l'erreur de sortie obéit alors à l'inégalité :

$$d \geq \frac{1}{\eta k} \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i}) - w_L$$

Si le dernier détour de la trace de E est terminal, on a par ailleurs :

$$d \geq \frac{Nk - p_{v_c}}{\eta k} - (w_L + 1)$$

Démonstration. Reprenons les notations introduites dans la définition 4.3.3, de sorte que la j -ème troncature de E s'écrit :

$$E^{\rightarrow j} = (M, L^{(0)}, \dots, L^{(N_j-1)}, L^{\rightarrow j}, S^{(0)}, \dots, S^{(N_j)})$$

et l'erreur de mémoire liée à la j -ème troncature est :

$$M^{\rightarrow j} = \mu_{N_j+1}(E^{\rightarrow j})$$

Reprenons également la convention $p_0 = N_0 = 0$ et notons $N_{w_L+1} = N$. Ecrivons :

$$(P^{(0)}, \dots, P^{(N-1)}) = \pi_N(E)$$

Soit $(b_0, \dots, b_{w_L})$ la trace de E . Pour tout $j \in \llbracket 1; w_L \rrbracket$, soit

$$d_j = |(P^{(0)}, \dots, P^{(N_j)})|$$

et soient $d_0 = 0$ et $d_{w_L+1} = d$. Montrons alors que pour tout $j \in \llbracket 0; w_L \rrbracket$ tel que $b_j = 1$:

$$d_{j+1} - d_j \geq \frac{N_{j+1} - N_j}{\eta} - 1 \quad (4.3.3)$$

Traisons d'abord le cas où $j = 0$. On a alors $M \in \mathbb{M}_1$. Comme la première coordonnée logique non triviale de E fait partie de l'erreur $L^{(N_1)}$, toutes les erreurs logiques précédentes de E sont triviales. Ainsi d'après le lemme 4.3.3 :

$$\begin{aligned} d_1 &= |(P^{(0)}, \dots, P^{(N_1)})| \\ &\geq |(P^{(0)}, \dots, P^{(N_1-1)})| \\ &= |\pi_{N_1}(M, I_k, \dots, I_k, S^{(0)}, \dots, S^{(N_1-1)})| \\ &\geq \lfloor \frac{N_1}{\eta} \rfloor \end{aligned}$$

ce qui prouve l'inégalité (4.3.3) dans ce cas.

A présent supposons que $j \geq 1$; on distingue deux cas. Si $N_{j+1} = N_j$, le résultat est juste car le terme de gauche est nul et le terme de droite est négatif. Si en revanche $N_{j+1} > N_j$, la j -ème coordonnée logique non triviale de E fait partie de l'erreur $L^{(N_j)}$, tandis que la $j+1$ -ème coordonnée logique non triviale de E fait partie de l'erreur $L^{(N_{j+1})}$ si $j \leq w-1$ et n'existe pas si $j = w$. Cela signifie que d'une part, $L^{\rightarrow j} = L^{(N_j)}$ c'est-à-dire :

$$E^{\rightarrow j} = (M, L^{(0)}, \dots, L^{(N_j)}, S^{(0)}, \dots, S^{(N_j)})$$

et d'autre part :

$$(L^{(0)}, \dots, L^{(N_{j+1}-1)}) = (L^{(0)}, \dots, L^{(N_j)}, I_k, \dots, I_k)$$

Donc d'après le lemme 4.3.1 :

$$\begin{aligned} (P^{(0)}, \dots, P^{(N_{j+1}-1)}) &= \pi_{N_{j+1}}(M, L^{(0)}, \dots, L^{(N_j)}, I_k, \dots, I_k, S^{(0)}, \dots, S^{(N_{j+1}-1)}) \\ &= (\pi_{N_{j+1}}(E^{\rightarrow j}), \\ &\quad \pi_{N_{j+1}-N_j-1}(M^{\rightarrow j}, I_k, \dots, I_k, S^{(N_j+1)}, \dots, S^{(N_{j+1}-1)})) \\ &= (P^{(0)}, \dots, P^{(N_j)}, \\ &\quad \pi_{N_{j+1}-N_j-1}(M^{\rightarrow j}, I_k, \dots, I_k, S^{(N_j+1)}, \dots, S^{(N_{j+1}-1)})) \end{aligned}$$

Comme $M^{\rightarrow j} \in \mathbb{M}_1$, d'après le lemme 4.3.3 :

$$\begin{aligned} d_{j+1} - d_j &\geq |(P^{(0)}, \dots, P^{(N_{j+1}-1)})| - |(P^{(0)}, \dots, P^{(N_j)})| \\ &= |\pi_{N_{j+1}-N_j-1}(M^{\rightarrow j}, I_k, \dots, I_k, S^{(N_j+1)}, \dots, S^{(N_{j+1}-1)})| \\ &\geq \lfloor \frac{N_{j+1} - N_j - 1}{\eta} \rfloor \end{aligned}$$

Par définition de la partie entière, le terme

$$\frac{N_{j+1} - N_j - 1}{\eta} - \lfloor \frac{N_{j+1} - N_j - 1}{\eta} \rfloor$$

est strictement inférieur à 1. De plus il appartient à $\frac{1}{\eta}\mathbb{Z}$. Donc :

$$\frac{N_{j+1} - N_j - 1}{\eta} - \lfloor \frac{N_{j+1} - N_j - 1}{\eta} \rfloor \leq \frac{\eta - 1}{\eta}$$

ce qui implique :

$$\frac{N_{j+1} - N_j}{\eta} - 1 \leq \lfloor \frac{N_{j+1} - N_j - 1}{\eta} \rfloor$$

et démontre l'inégalité (4.3.3).

On somme cette inégalité sur les éléments valant 1 dans chaque détour. Pour tout $i \in \llbracket 1; c \rrbracket$:

$$\begin{aligned} d_{v_{i+1}-1} - d_{v_i} &= \sum_{j=v_i}^{v_{i+1}-2} d_{j+1} - d_j \\ &\geq \sum_{j=v_i}^{v_{i+1}-2} \left(\frac{N_{j+1} - N_j}{\eta} - 1 \right) \\ &= \frac{N_{v_{i+1}-1} - N_{v_i}}{\eta} - (v_{i+1} - v_i - 1) \end{aligned}$$

et si le dernier détour est terminal on a également :

$$\begin{aligned} d - d_{v_c} &= \sum_{j=v_c}^{w_L} d_{j+1} - d_j \\ &\geq \sum_{j=v_c}^{w_L} \left(\frac{N_{j+1} - N_j}{\eta} - 1 \right) \\ &= \frac{N - N_{v_c}}{\eta} - (w_L - v_c + 1) \end{aligned}$$

Or pour tout $j \in \llbracket 1; w_L \rrbracket$, on a par définition de N_j et p_j :

$$N_j k + 1 \leq p_j \leq N_j k + k$$

ce qui donne la minoration :

$$N_j \geq \frac{p_j}{k} - 1$$

et la majoration :

$$N_j \leq \frac{p_j}{k}$$

cette dernière majoration restant valable si $j = 0$. Donc pour tout $i \in \llbracket 1; c \rrbracket$:

$$\begin{aligned} d_{v_{i+1}-1} - d_{v_i} &\geq \frac{p_{v_{i+1}-1} - p_{v_i}}{\eta k} - \frac{1}{\eta} - (v_{i+1} - 1 - v_i) \\ &\geq \frac{p_{v_{i+1}-1} - p_{v_i}}{\eta k} - (v_{i+1} - v_i) \end{aligned}$$

Si le dernier détour est terminal, on a également :

$$d - d_{v_c} \geq \frac{Nk - p_{v_c}}{\eta k} - (w_L - v_c + 1)$$

d'où on tire la deuxième inégalité du lemme :

$$d \geq \frac{Nk - p_{v_c}}{\eta k} - (w_L + 1)$$

Dans le cas général, en constatant que

$$d \geq \sum_{i=1}^c d_{v_{i+1}-1} - d_{v_i}$$

et que

$$w_L \geq \sum_{i=1}^c v_{i+1} - v_i$$

on obtient :

$$d \geq \frac{1}{\eta k} \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i}) - w$$

□

4.4 Dénombrement d'erreurs au niveau de l'encodeur formel interne

Les deux bornes que l'on démontre reposent sur le théorème suivant.

Théorème 4.4.1. *Soit un encodeur convolutif formel récursif de taille N , de paramètres $[n, k, m]$ et de période η .*

$$a_{in}^N(w, \leq d) \leq O(1)^w \binom{m + Nk}{\lfloor w/2 \rfloor} \binom{\lfloor \eta k(w + d) \rfloor}{\lceil w/2 \rceil}$$

Ce théorème se prouve en majorant d'abord le nombre d'erreurs possibles pour un nombre de détours c imposé. Pour cela on sépare deux cas, selon que le dernier détour est régulier ou terminal. La suite de la section est organisée comme suit. Les lemmes correspondant à ces deux cas sont d'abord énoncés et démontrés, et vient ensuite la preuve du théorème 4.4.1. Finalement, on présente les deux bornes qui correspondent aux hypothèses de chacun des deux théorèmes 3.4.1 et 3.4.2.

Lemme 4.4.1 (Lemme 4.4.1 : c détours dont le dernier est régulier). *Soit un encodeur convolutif formel récursif de taille N , de paramètres $[n, k, m]$ et de période η . Soient $w \in \llbracket 1; kN \rrbracket$ et $c \in \llbracket 1; \lfloor w/2 \rfloor \rrbracket$, et soit d un réel supérieur à 1. Considérons l'ensemble $R_{w,c,\leq d}$ d'erreurs*

$$(M^{(0)}, L^{(0)}, \dots, L^{(N-1)}) \in G^{m+Nk}$$

de poids w , qui peuvent être complétées en une erreur

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{m+Nk} \times H^{N(n-k)}$$

telle que $|\mathcal{E}_N(E)| \leq d$ et dont la trace contient c détours dont le dernier est régulier. Le cardinal de $R_{w,c,\leq d}$ est majoré par :

$$O(1)^w \binom{m+Nk}{c} \binom{\min(m+Nk, \lfloor \eta k(w+d) \rfloor)}{w-c}$$

où $O(1)$ est majorable par une constante indépendante de w , d et N .

Démonstration. Considérons une erreur

$$E_{ML} = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}) \in R_{w,c,d}$$

de poids w , et notons w_L son poids logique. Considérons une extension E de E_{ML} telle que $|\mathcal{E}_N(E)| \leq d$ et dont la trace $(b_0, \dots, b_{w_L})$ contient c détours. Notons

$$1 \leq p_1 < \dots < p_{w_L} \leq Nk$$

les positions des w_L coordonnées logiques non triviales de E , et

$$0 \leq v_1 < \dots < v_c \leq w_L$$

les indices de départ des c détours de la trace de E . On reprend la convention $p_0 = 0$ et $v_{c+1} = w_L + 1$. On suppose d'abord w_L fixé, et on s'intéresse dans un premier temps à majorer le nombre de valeurs possibles de la séquence

$$(p_1, \dots, p_{w_L})$$

D'après le lemme 4.3.5, l'indice v_1 vaut 0 ou 1. Pour chacun de ces choix, les $c-1$ indices $v_2 < \dots < v_c$ sont compris dans l'intervalle $\llbracket 1; w_L \rrbracket$, et le nombre de choix possibles de ces indices est donc majoré par

$$\binom{w}{c-1}$$

Pour chacun de ces choix, la séquence des c positions $(p_{v_1}, \dots, p_{v_c})$ prend au plus

$$\binom{Nk}{c} \leq \binom{m+Nk}{c}$$

valeurs si $v_1 = 1$ tandis qu'elle prend au plus

$$\binom{Nk}{c-1} \leq \binom{m+Nk}{c-1}$$

valeurs si $v_1 = 0$ car $p_0 = 0$. Pour chacun de ces choix, on va majorer le nombre de choix possibles de la séquence des positions restantes :

$$p_{reste} = (p_{v_1+1}, \dots, p_{v_2-1}, p_{v_2+1}, \dots, p_{v_3-1}, [\dots], p_{v_c+1}, \dots, p_{w_L})$$

Cette séquence contient $w_L - c$ éléments lorsque $v_1 = 1$ et $w_L - c + 1$ éléments lorsque $v_1 = 0$. On peut alors lui associer de manière inversible la séquence

$$p'_{reste} = (p'_{v_1+1}, \dots, p'_{v_2-1}, p'_{v_2+1}, \dots, p'_{v_3-1}, [\dots], p'_{v_c+1}, \dots, p'_{w_L})$$

construite de la manière suivante :

$$\begin{cases} \text{pour tout } j \in \llbracket v_1 + 1; v_2 - 1 \rrbracket, & p'_j = p_j - p_{v_1} \\ \text{pour tout } i \in \llbracket 2; c \rrbracket \text{ et } j \in \llbracket v_i + 1; v_{i+1} - 1 \rrbracket, & p'_j = p_j - p_{v_i} + p'_{v_i-1} \end{cases}$$

L'opération inverse est simplement la suivante :

$$\begin{cases} \text{pour tout } j \in \llbracket v_1 + 1; v_2 - 1 \rrbracket, & p_j = p'_j + p_{v_1} \\ \text{pour tout } i \in \llbracket 2; c \rrbracket \text{ et } j \in \llbracket v_i + 1; v_{i+1} - 1 \rrbracket, & p_j = p'_j + p_{v_i} - p'_{v_i-1} \end{cases}$$

La séquence p'_{reste} obtenue est strictement croissante. En effet l'opération effectuée translate chaque séquence d'ordre i

$$(p_{v_i+1}, \dots, p_{v_{i+1}-1})$$

et conserve l'ordre entre les séquences d'ordres successifs puisque pour tout $i \in \llbracket 2; c \rrbracket$, la relation

$$p'_{v_i+1} = p_{v_i+1} - p_{v_i} + p'_{v_i-1}$$

implique l'inégalité :

$$p'_{v_i+1} - p'_{v_i-1} = p_{v_i+1} - p_{v_i} > 0$$

De plus les éléments de la séquence p'_{reste} sont compris entre 1 et

$$\min(Nk, \lfloor \eta k(w_L + d) \rfloor)$$

En effet, d'une part, le plus petit élément de la séquence p'_{reste} vérifie :

$$p'_{v_1+1} = p_{v_1+1} - p_{v_1} \geq 1$$

D'autre part, la séquence d'ordre i

$$(p'_{v_i+1}, \dots, p'_{v_{i+1}-1})$$

est contenue dans un intervalle de taille

$$p'_{v_{i+1}-1} - p'_{v_i+1} = p_{v_{i+1}-1} - p_{v_i+1}$$

et deux séquences successives d'ordres $i-1$ et i sont distantes de

$$p'_{v_i+1} - p'_{v_i-1} = p_{v_i+1} - p_{v_i}$$

ce qui permet, par sommation, de noter que le plus grand élément de la séquence p'_{reste} vérifie :

$$\begin{aligned} p'_{w_L} &= \sum_{i=1}^c (p'_{v_{i+1}-1} - p'_{v_i+1}) + \sum_{i=2}^c (p'_{v_i+1} - p'_{v_{i-1}}) + p'_{v_1+1} \\ &= \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i+1}) + \sum_{i=2}^c (p_{v_i+1} - p_{v_i}) + p_{v_1+1} - p_{v_1} \\ &= \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i+1}) + \sum_{i=1}^c (p_{v_{i+1}} - p_{v_i}) \\ &= \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i}) \end{aligned}$$

Comme d'après le théorème 4.3.2,

$$d \geq \frac{1}{\eta k} \sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i}) - w_L$$

on a :

$$\sum_{i=1}^c (p_{v_{i+1}-1} - p_{v_i}) \leq \lfloor \eta k(w_L + d) \rfloor$$

Comme par ailleurs cette même somme est inférieure à p_{w_L} donc à Nk , cela donne bien la majoration :

$$p'_{w_L} \leq \min(Nk, \lfloor \eta k(w_L + d) \rfloor)$$

Par conséquent, le nombre de valeurs possibles pour la séquence p'_{reste} , et donc la séquence p_{reste} est majoré par

$$\begin{cases} \binom{\min(Nk, \lfloor \eta k(w_L + d) \rfloor)}{w_L - c} & \text{si } v_1 = 1 \\ \binom{\min(Nk, \lfloor \eta k(w_L + d) \rfloor)}{w_L - c + 1} & \text{si } v_1 = 0 \end{cases}$$

En utilisant le fait que, pour tout quadruplet d'entiers positifs $u \geq v$ et $\delta_u \geq \delta_v$,

$$\binom{u}{v} \leq \binom{u + \delta_u}{v + \delta_v}$$

et en notant que le poids de l'erreur de mémoire $w - w_L$ est inférieur à m , on peut majorer de nouveau le nombre de valeurs possibles de la séquence p_{reste} par :

$$\begin{cases} \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} & \text{si } v_1 = 1 \\ \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c + 1} & \text{si } v_1 = 0 \end{cases}$$

Cela montre que le nombre de valeurs que peut prendre la séquence de positions

$$(p_1, \dots, p_{w_L})$$

est majoré par :

$$\begin{aligned} & \binom{w}{c-1} \left[\binom{m + Nk}{c} \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} \right. \\ & \quad \left. + \binom{m + Nk}{c-1} \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c + 1} \right] \\ & \leq 2 \binom{w}{c-1} \binom{m + Nk}{c} \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} \end{aligned}$$

où l'on obtient la deuxième ligne grâce à la relation pour tout $u > v$:

$$\binom{u}{v+1} / \binom{u}{v} = \frac{u-v}{v+1}$$

qui donne, en notant $\min([\dots]) = \min(m + Nk, \lfloor \eta k(w + d) \rfloor)$ et en remarquant que $c \leq w - c + 1$:

$$\begin{aligned} \frac{\binom{m+Nk}{c} \binom{\min([\dots])}{w-c}}{\binom{m+Nk}{c-1} \binom{\min([\dots])}{w-c+1}} &= \frac{m + Nk - (c - 1)}{c} \frac{w - c + 1}{\min([\dots]) - (w - c)} \\ &= \frac{m + Nk - (c - 1)}{\min([\dots]) - (w - c)} \frac{w - c + 1}{c} \geq 1 \end{aligned} \quad (4.4.1)$$

Finalement, une fois les positions des coordonnées logiques non triviales de E_{ML} fixées, il reste au plus

$$(|G| - 1)^{w_L} \leq |G|^{w_L}$$

choix possibles des valeurs de ces coordonnées, tandis que les coordonnées logiques restantes de E_{ML} sont fixées à I . L'erreur de mémoire de E_{ML} , de poids $w_M = w - w_L$, ne peut prendre quant à elle plus de

$$\binom{m}{w_M} |G|^{w_M}$$

valeurs possibles. Le nombre d'erreurs E_{ML} possibles pour w_L fixé est donc majoré par :

$$2 \binom{m}{w_M} |G|^w \binom{w}{c-1} \binom{m+Nk}{c} \binom{\min(m+Nk, \lfloor \eta k(w+d) \rfloor)}{w-c}$$

La somme de cette expression sur les valeurs possibles de w_M démontre que :

$$\begin{aligned} |R_{w,c,\leq d}| &\leq 2^{m+1} |G|^w \binom{w}{c-1} \binom{m+Nk}{c} \binom{\min(m+Nk, \lfloor \eta k(w+d) \rfloor)}{w-c} \\ &\leq O(1)^w \binom{m+Nk}{c} \binom{\min(m+Nk, \lfloor \eta k(w+d) \rfloor)}{w-c} \end{aligned}$$

□

Lemme 4.4.2 (Lemme 4.4.2 : c détours dont le dernier est terminal). *Soit un encodeur convolutif formel récursif de taille N , de paramètres $[n, k, m]$ et de période η . Soient $w \in \llbracket 1; kN \rrbracket$ et $c \in \llbracket 1; \lfloor \frac{w+1}{2} \rfloor \rrbracket$, et soit d un réel supérieur à 1. Considérons l'ensemble $T_{w,c,\leq d}$ d'erreurs*

$$(M^{(0)}, L^{(0)}, \dots, L^{(N-1)}) \in G^{m+Nk}$$

de poids w , qui peuvent être complétées en une erreur

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{m+Nk} \times H^{N(n-k)}$$

telle que $|\mathcal{E}_N(E)| \leq d$ et dont la trace contient c détours dont le dernier est terminal. Le cardinal de $T_{w,c,\leq d}$ est majoré par :

$$O(1)^w d \binom{m+Nk}{c-1} \binom{\min(m+Nk, \lfloor \eta k(w+d) \rfloor)}{w-c}$$

où $O(1)$ est majorable par une constante indépendante de w , d et N .

Démonstration. Les mêmes majorations effectuées afin de prouver le lemme 4.4.1 s'appliquent ici. Cependant, le nombre de choix possibles de la séquence des c positions

$$(p_{v_1}, \dots, p_{v_c})$$

peut être majoré de manière plus fine. En effet, d'après le théorème 4.3.2 :

$$d \geq \frac{Nk - p_{v_c}}{\eta k} - (w + 1)$$

ce qui implique :

$$Nk - p_{v_c} \leq \lfloor \eta k(w + d + 1) \rfloor$$

p_{v_c} prend donc au plus $\lfloor \eta k(w + d + 1) \rfloor + 1$ valeurs possibles. La séquence des $c - 1$ éléments restants, si elle n'est pas vide, peut prendre au plus

$$\binom{Nk}{c - 1}$$

valeurs si $v_1 = 1$, et

$$\binom{Nk}{c - 2}$$

valeurs si $v_1 = 0$ car $p_0 = 0$. Le reste du raisonnement conserve sa validité, en observant néanmoins un détail : si le dernier détour ne contient qu'un élément, les séquences p_{reste} et p'_{reste} se terminent à la position d'indice $v_c - 1$, et la majoration de la valeur de $p'_{v_c - 1}$ se fait alors de la même manière que précédemment en arrêtant l'indice de sommation i à la valeur $c - 1$. Le nombre de valeurs que peut prendre la séquence de positions

$$(p_1, \dots, p_{w_L})$$

est ainsi majoré dans le cas présent par :

$$\begin{aligned} & \binom{w}{c - 1} (\lfloor \eta k(w + d + 1) \rfloor + 1) \left[\binom{Nk}{c - 1} \binom{\min(Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} \right. \\ & \quad \left. + \binom{Nk}{c - 2} \binom{\min(Nk, \lfloor \eta k(w + d) \rfloor)}{w - c + 1} \right] \\ & \leq 2 \binom{w}{c - 1} (\lfloor \eta k(w + d + 1) \rfloor + 1) \binom{Nk}{c - 1} \binom{\min(Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} \end{aligned}$$

où la deuxième ligne s'obtient par une comparaison similaire à l'inégalité (4.4.1) dans la preuve du lemme précédent. En majorant, de la même manière que précédemment, le nombre de valeurs possibles pour l'erreur E_{ML} une fois la séquence $(p_1, \dots, p_{w_L})$ fixée, puis en sommant le résultat sur les $m + 1$ valeurs possibles de w_L , cela donne la majoration suivante du cardinal de $T_{w,c,\leq d}$:

$$\begin{aligned} & 2^{m+1} |G|^w \binom{w}{c - 1} (\lfloor \eta k(w + d + 1) \rfloor + 1) \binom{m + Nk}{c - 1} \binom{\min(\dots)}{w - c} \\ & = O(1)^w d \binom{m + Nk}{c - 1} \binom{\min(m + Nk, \lfloor \eta k(w + d) \rfloor)}{w - c} \end{aligned}$$

où l'on a utilisé le fait que, comme $d \geq 1$:

$$\lfloor \eta k(w + d + 1) \rfloor + 1 \leq \eta k(w + d + 2) \leq \eta k(w + 3)d \leq O(1)^w d$$

□

Preuve du théorème 4.4.1. Il s'agit de majorer le cardinal $a_{in}^N(w, \leq d)$ de l'ensemble $A_{w, \leq d}$ d'erreurs

$$E_{ML} = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}) \in G^{m+Nk}$$

de poids w , qui peuvent être complétées en une erreur

$$E = (M^{(0)}, L^{(0)}, \dots, L^{(N-1)}, S^{(0)}, \dots, S^{(N-1)}) \in G^{m+Nk} \times H^{N(n-k)}$$

telle que $|\mathcal{E}_N(E)| \leq d$. En fonction du nombre c de détours de la trace de E et du caractère régulier ou terminal du dernier détour, E_{ML} appartient à un certain ensemble $R_{w,c,\leq d}$ où $c \in \llbracket 1; \lfloor w/2 \rfloor \rrbracket$, ou $T_{w,c,\leq d}$ où $c \in \llbracket 1; \lfloor (w+1)/2 \rfloor \rrbracket$, les bornes sur c étant obtenues grâce au lemme 4.3.5. Par mesure de simplicité notons $\min(\dots) = \min(m + Nk, \lfloor \eta k(w + d) \rfloor)$. D'après les lemmes 4.4.1 et 4.4.2, on peut donc majorer $a_{in}^N(w, \leq d)$ par :

$$\begin{aligned} & \sum_{c=1}^{\lfloor w/2 \rfloor} |R_{w,c,\leq d}| + \sum_{c=1}^{\lfloor (w+1)/2 \rfloor} |T_{w,c,\leq d}| \\ &= O(1)^w \left[\sum_{c=1}^{\lfloor w/2 \rfloor} \binom{m+Nk}{c} \binom{\min(\dots)}{w-c} + \sum_{c=1}^{\lfloor (w+1)/2 \rfloor} d \binom{m+Nk}{c-1} \binom{\min(\dots)}{w-c} \right] \\ &= O(1)^w \left[\max_{c \in \llbracket 1; \lfloor w/2 \rfloor \rrbracket} \binom{m+Nk}{c} \binom{\min(\dots)}{w-c} \right. \\ & \quad \left. + d \max_{c \in \llbracket 1; \lfloor (w+1)/2 \rfloor \rrbracket} \binom{m+Nk}{c-1} \binom{\min(\dots)}{w-c} \right] \\ &= O(1)^w \left[\max_{c \in \llbracket 1; \lfloor w/2 \rfloor \rrbracket} \binom{m+Nk}{c} \binom{\min(\dots)}{w-c} \right. \\ & \quad \left. + \eta d \max_{c \in \llbracket 1; \lfloor w/2 \rfloor + 1 \rrbracket} \binom{m+Nk}{c-1} \binom{\min(\dots)}{w-c} \right] \end{aligned}$$

D'après l'inégalité (4.4.1) dans la preuve du lemme 4.4.1, pour tout c tel que $c \leq w - c + 1$:

$$\binom{m+Nk}{c} \binom{\min(\dots)}{w-c} \geq \binom{m+Nk}{c-1} \binom{\min(\dots)}{w-c+1}$$

Ainsi le premier maximum est atteint en $c = \lfloor w/2 \rfloor$. De manière similaire, si

$c \in \llbracket 2; \lfloor w/2 \rfloor + 1 \rrbracket$, on a $c - 1 \leq w - c + 1$ et :

$$\begin{aligned} \frac{\binom{m+Nk}{c-1} \binom{\min(\lfloor \dots \rfloor)}{w-c}}{\binom{m+Nk}{c-2} \binom{\min(\lfloor \dots \rfloor)}{w-c+1}} &= \frac{m+Nk-(c-2)}{c-1} \frac{w-c+1}{\min(\lfloor \dots \rfloor) - (w-c)} \\ &= \frac{m+Nk+1-(c-1)}{\min(\lfloor \dots \rfloor) + 1 - (w-c+1)} \frac{w-c+1}{c-1} \geq 1 \end{aligned}$$

Le second maximum est donc atteint en $c = \lfloor w/2 \rfloor + 1$, et $a_{in}^N(w, \leq d)$ est majoré par :

$$\begin{aligned} &O(1)^w \left[\binom{m+Nk}{\lfloor w/2 \rfloor} \binom{\min(\lfloor \dots \rfloor)}{\lceil w/2 \rceil} + d \binom{m+Nk}{\lfloor w/2 \rfloor} \binom{\min(\lfloor \dots \rfloor)}{\lceil w/2 \rceil - 1} \right] \\ &= O(1)^w \left[\binom{m+Nk}{\lfloor w/2 \rfloor} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil} + d \binom{m+Nk}{\lfloor w/2 \rfloor} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil - 1} \right] \end{aligned}$$

La relation suivante, valable pour tout couple d'entiers $u \geq v$:

$$\binom{u}{v-1} = \frac{v}{u-v+1} \binom{u}{v}$$

entraîne :

$$\begin{aligned} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil - 1} &= \frac{\lceil w/2 \rceil}{\lfloor \eta k(w+d) \rfloor - \lceil w/2 \rceil + 1} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil} \\ &\leq \frac{\lceil w/2 \rceil}{1/2 \lfloor \eta k(w+d) \rfloor} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil} \\ &= \frac{O(1)^w}{d} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil} \end{aligned}$$

Cela permet de conclure :

$$a_{in}^N(w, \leq d) = O(1)^w \binom{m+Nk}{\lfloor w/2 \rfloor} \binom{\lfloor \eta k(w+d) \rfloor}{\lceil w/2 \rceil}$$

□

Le théorème 4.4.1 implique alors les deux bornes suivantes.

Corollaire 4.4.1. *Soit $(n+m, k+m, \mathcal{E})$ un encodeur formel récursif.*

$$a_{in}^N(w, \leq d) = O(1)^w \frac{N^{\frac{w}{2}} (w+d)^{\frac{w}{2}}}{w^w} \quad (\text{Borne 1I})$$

Démonstration. On se sert de la borne du théorème 4.4.1 et on utilise la propriété pour tout couple d'entiers $u \geq v$:

$$\binom{u}{v} \leq \left(\frac{ue}{v} \right)^v$$

Cela donne :

$$\begin{aligned}
a_{in}^N(w, \leq d) &= O(1)^w \left(\frac{(m + Nk)e}{\lfloor w/2 \rfloor} \right)^{\lfloor w/2 \rfloor} \left(\frac{(\lfloor \eta k(w + d) \rfloor)e}{\lceil w/2 \rceil} \right)^{\lceil w/2 \rceil} \\
&= O(1)^w \frac{(m + Nk)^{\lfloor w/2 \rfloor}}{(\lfloor w/2 \rfloor)^{\lfloor w/2 \rfloor}} \frac{(\lfloor \eta k(w + d) \rfloor)^{\lceil w/2 \rceil}}{(\lceil w/2 \rceil)^{\lceil w/2 \rceil}} \\
&= O(1)^w \frac{(m + Nk)^{\lfloor w/2 \rfloor} (\eta k(w + d))^{\lceil w/2 \rceil}}{(\lfloor w/2 \rfloor)^w} \\
&= O(1)^w \frac{N^{\lfloor w/2 \rfloor} (w + d)^{\lceil w/2 \rceil}}{w^w} \\
&= O(1)^w \frac{N^{w/2} (w + d)^{w/2}}{w^w}
\end{aligned}$$

Justifions le passage à la dernière ligne. La dernière expression s'obtient en multipliant l'expression qui la précède par le facteur :

$$\left(\frac{N}{w + d} \right)^{w/2 - \lfloor w/2 \rfloor}$$

Comme l'exposant $w/2 - \lfloor w/2 \rfloor$ vaut 0 ou 1/2, et comme les poids d'entrée et de sortie vérifient $w \leq m + Nk \leq N(m + k)$ et $d \leq m + Nn \leq N(m + n)$, ce facteur est minoré par la constante :

$$\frac{1}{(2m + k + n)^{1/2}}$$

ce qui permet d'écrire :

$$N^{\lfloor w/2 \rfloor} (w + d)^{\lceil w/2 \rceil} \leq (2m + k + n)^{1/2} N^{w/2} (w + d)^{w/2} = O(1) N^{w/2} (w + d)^{w/2}$$

□

Corollaire 4.4.2. *Soit $(n + m, k + m, \mathcal{E})$ un encodeur formel anti-récuratif.*

$$a_{in}^N(w, d) = O(1)^d \frac{N^{\frac{d}{2}} (w + d)^{\frac{d}{2}}}{d^d} \quad (\text{Borne } 2I)$$

Démonstration. Comme on l'a remarqué en début de section précédente, tous les résultats de la section précédente et de cette section s'appliquent également aux encodeurs formels inversés. Cela permet d'obtenir la borne (2I) directement à partir de la borne (1I) de la manière suivante. Soit $(k + m, n + m, \mathcal{E}')$ l'encodeur formel inversé réciproque de l'encodeur formel $(n + m, k + m, \mathcal{E})$. Notons a_{in}^N la répartition de poids liée à l'encodeur convolutif formel inversé présentée dans la définition 4.1.3. D'après le lemme 4.1.2, on a pour tout couple d'entiers (w, d) :

$$a_{in}^N(w, d) \leq a_{in}^N(d, w)$$

De plus, comme $(n + m, k + m, \mathcal{E})$ est anti-récuratif, alors d'après le lemme 3.3.1, $(k + m, n + m, \mathcal{E}')$ est récuratif. La borne (1I) s'applique donc à cet encodeur formel inversé, et l'on a :

$$a'_{in}^N(d, w) \leq a'_{in}^N(d, \leq w) = O(1)^d \frac{N^{\frac{d}{2}}(w + d)^{\frac{d}{2}}}{d^d}$$

ce qui prouve la borne (2I). $\square$

4.5 Preuve des théorèmes

On rappelle que l'on se propose d'établir des bornes asymptotiques en N sur la distance minimale d'un turbo-encodeur formel de taille $N = N_{in}$, basé sur un encodeur formel externe $(n_{ext}, k_{ext}, \mathcal{E}_{ext})$, un encodeur formel interne $(n_{in} + m_{in}, k_{in} + m_{in}, \mathcal{E}_{in})$ où m_{in} est implicitement connu, et un entrelaceur formel aléatoire de taille N , et que l'on désigne respectivement par d_{min} et d_{deg} les distances minimale et dégénérée de l'encodeur formel externe.

On consacre cette section à démontrer les théorèmes 3.4.1 et 3.4.2 en utilisant les différentes bornes établies aux sections précédentes. L'articulation de ces différentes bornes se fait conformément aux figures 4.1 et 4.2 qui se traduisent par le lemme suivant. Comme précisé auparavant, la borne D désigne tour à tour N^α où $\alpha < \frac{d_{deg}-2}{d_{deg}}$ lorsque l'on montre la borne polynomiale des théorèmes 3.4.1 et 3.4.2, et $\alpha \log N / (\log \log N)$ où $\alpha < d_{min} - 2$ lorsque l'on montre la borne sous-logarithmique du théorème 3.4.2.

Lemme 4.5.1. *Pour tout réel D et pour tout $x \in [0; n_{in}[$:*

$$p_N(\leq D) \leq \sum_{w=0}^{\lfloor D \rfloor + m_{in}} p_N(w, \leq D) + \sum_{w=\lfloor D \rfloor + m_{in} + 1}^{\lfloor xN \rfloor - 1} \sum_{d=0}^{\lfloor D \rfloor} p_N(w, d) + \sum_{w=\lfloor xN \rfloor}^{Nk_{in} + m_{in}} \sum_{d=0}^{\lfloor D \rfloor} p_N(w, d)$$

Démonstration. D'après la construction du turbo-encodeur formel, le poids w de la partie de mémoire et logique d'une erreur intermédiaire peut varier entre 0 et $Nk_{in} + m_{in}$. Le lemme s'ensuit immédiatement grâce à l'inégalité de Boole, en notant que les probabilités indiquées dans le côté droit de l'inégalité correspondent à des événements dont la réunion est qu'il existe une erreur d'entrée du turbo-encodeur formel dont l'image est une erreur nocive de poids inférieur ou égal à D . $\square$

La démonstration des théorèmes se fait donc en majorant séparément chacune des trois sommes par une expression négligeable devant 1 lorsque N tend vers l'infini. Il faut noter que, comme le montre la figure 4.1, le théorème 3.4.1 ne nécessite de majorer que la première somme, les deux autres étant nulles grâce au caractère systématique de l'encodeur convolutif formel interne. On adopte donc la notation suivante.

Définition 4.5.1. On appelle respectivement **première somme**, **deuxième somme** et **troisième somme** les expressions suivants :

$$\sum_{w=0}^{\lfloor D \rfloor + m_{in}} p_N(w, \leq D), \quad \sum_{w=\lfloor D \rfloor + m_{in} + 1}^{\lfloor xN \rfloor - 1} \sum_{d=0}^{\lfloor D \rfloor} p_N(w, d) \quad \text{et} \quad \sum_{w=\lfloor xN \rfloor}^{Nk_{in} + m_{in}} \sum_{d=0}^{\lfloor D \rfloor} p_N(w, d)$$

La liberté de choix laissée sur x permet de choisir cette valeur judicieusement afin de contrôler les deuxième et troisième sommes. On montrera en effet qu'il existe une valeur de x pour laquelle la deuxième somme devient négligeable devant 1 lorsque la taille N du turbo-encodeur formel tend vers l'infini, tandis que la troisième somme possède cette propriété quelque soit la valeur de x . On choisira ainsi x en fonction des contraintes imposées par la deuxième somme.

Avant de présenter la démonstration des théorèmes 3.4.1 et 3.4.2, intéressons-nous aux raisons qualitatives justifiant l'expression des bornes polynomiale et sous-logarithmique. Le détail des calculs est laissé aux lemmes 4.5.3 et 4.5.6 plus bas. Dans le cas où $d_{deg} > 2$, la première somme se comporte d'une manière similaire au cas classique présenté par Urbanke et Kahale. Elle est majorable, grâce aux bornes $1E$ et $1I$, par la somme sur w d'une expression dont le comportement est dicté par la forme suivante :

$$O(D)^{\frac{w}{2}} / N^{\frac{w}{2} \frac{d_{deg}-2}{d_{deg}}}$$

Cette forme est géométrique en w , et on constate alors que sa raison est strictement inférieure à 1 pour tout D se comportant comme N^α où $\alpha < \frac{d_{deg}-2}{d_{deg}}$.

Dans le cas où $d_{deg} = 2$, l'exposant de N dans la forme précédente est nul et ne permet pas de majorer la première somme. Néanmoins, la forme précédente ne tient compte, dans la borne $(1E)$, que de l'exposant en N qui est proportionnel à w . Une approche plus fine tenant compte de la forme exacte de la borne $(1E)$ donne dans ce cas la forme suivante :

$$O(D)^{\frac{w}{2}} / N^{\frac{d_{min}-2}{2}}$$

La somme de cette expression sur w est dominée par le terme où w atteint sa valeur maximale qui est de l'ordre de D . Ainsi, en choisissant une borne D telle que D^D est de l'ordre de N^α où $\alpha < d_{min} - 2$, on peut de nouveau majorer convenablement la première somme. La propriété vérifiée par D est importante pour la suite ; on l'exprime donc dans le lemme suivant.

Lemme 4.5.2. Notons $t = d_{min} - 2$. La solution de l'équation $x^x = N^t$ a pour équivalent $D = t \frac{\ln N}{\ln \ln N}$ lorsque N tend vers l'infini. De plus, $D^D \leq N^t$ à partir d'un certain rang.

Démonstration. Notons x_N la solution de l'équation $x^x = N^t$, et $u_N = x_N D^{-1}$.

Alors :

$$\begin{aligned}
t \ln N &= x_N \ln x_N = u_N D \ln(u_N D) \\
&= u_N t \frac{\ln N}{\ln \ln N} (\ln u_N + \ln t + \ln \ln N - \ln \ln \ln N) \\
&= u_N t \frac{\ln N}{\ln \ln N} (\ln u_N + \ln \ln N + o(\ln \ln N))
\end{aligned}$$

En simplifiant par $t \ln N$ on obtient :

$$\ln \ln N = u_N (\ln u_N + \ln \ln N + o(\ln \ln N)) \quad (4.5.1)$$

A partir d'un certain rang on a donc :

$$u_N \ln u_N \leq \ln \ln N$$

ce qui prouve que $u_N \leq \ln \ln N$. Donc $u_N = o(\ln N)$ et l'équation (4.5.1) devient :

$$\ln \ln N = u_N (\ln \ln N + o(\ln \ln N))$$

Cela prouve que $u_N \sim 1$ et donc $x_N \sim D$. D'autre part :

$$\begin{aligned}
D \ln D &= t \frac{\ln N}{\ln \ln N} (\ln t + \ln \ln N - \ln \ln \ln N) \\
&= t \ln N \frac{\ln t + \ln \ln N - \ln \ln \ln N}{\ln \ln N}
\end{aligned}$$

ce qui est inférieur à $t \ln N$ à partir d'un certain rang, impliquant alors que $D^D \leq N^t$. $\square$

On va à présent combiner les quatre bornes démontrées précédemment afin de prouver chacun des théorèmes 3.4.1 et 3.4.2. Comme on l'a remarqué au moment de dresser la table 4.1, ces théorèmes contiennent des hypothèses qui rendent valides les bornes utilisées pour leur preuve. Ces hypothèses sont même plus fortes que celles nécessaires à la validité de ces bornes. Ainsi, la distance minimale polynomiale dans les théorèmes 3.4.1 et 3.4.2 requiert l'hypothèse $d_{deg} > 2$, et la distance sous-logarithmique dans le théorème 3.4.2 requiert l'hypothèse $d_{min} > d_{deg}$. Ces hypothèses supplémentaires sont requises pour majorer la première et la deuxième sommes, comme cela apparaîtra dans la suite.

La démonstration des théorèmes 3.4.1 et 3.4.2 se fait sur six lemmes, qui majorent les différentes sommes du lemme 4.5.1 par des expressions négligeables devant 1 lorsque N tend vers l'infini. Les lemmes 4.5.3, 4.5.4 et 4.5.5 majorent respectivement les première, deuxième et troisième somme sous les hypothèses correspondant à la borne polynomiale du théorème 3.4.2, tandis que les lemmes 4.5.6, 4.5.7 et 4.5.8 majorent respectivement les première, deuxième et troisième somme sous les hypothèses correspondant à la borne sous-logarithmique du théorème 3.4.2. Cela démontre ainsi le théorème 3.4.2. En notant de plus que le lemme 4.5.3 s'applique également sous les hypothèses du théorème 3.4.1,

et que les deuxième et troisième sommes dans le cadre de ce théorème sont nulles de par le caractère systématique de l'encodeur convolutif formel interne, le théorème 3.4.1 s'en trouve également démontré.

Lemme 4.5.3. *On pose $D = N^\alpha$, où $\alpha < \frac{d_{deg}-2}{d_{deg}}$. Si $d_{deg} > 2$ et l'encodeur formel interne est récursif, la première somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. D'après le lemme 4.1.1, chaque terme de la première somme se majore de la manière suivante :

$$p_N(w, \leq \lfloor D \rfloor + m_{in}) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, \leq \lfloor D \rfloor + m_{in})}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \quad (4.5.2)$$

On applique les bornes 1E et 1I. En remarquant que $d_{min} \geq d_{deg}$, puis que $N_{ext} = O(N)$, la borne 1E donne :

$$\begin{aligned} a_{ext}^{N_{ext}}(w) &= O(1)^w \left(\frac{N_{ext}}{w} \right)^{\frac{w - d_{min}}{d_{deg}} + 1} \\ &= O(1)^w \left(\frac{N_{ext}}{w} \right)^{\frac{w}{d_{deg}}} \\ &= O(1)^w \left(\frac{N}{w} \right)^{\frac{w}{d_{deg}}} \end{aligned}$$

Et en remarquant que $w + \lfloor D \rfloor + m_{in} = O(D)$ car $w \leq \lfloor D \rfloor + m_{in}$, la borne 1I donne :

$$\begin{aligned} a_{in}^N(w, \leq \lfloor D \rfloor + m_{in}) &= O(1)^w \frac{N^{\frac{w}{2}} (w + \lfloor D \rfloor + m_{in})^{\frac{w}{2}}}{w^w} \\ &= O(1)^w \frac{N^{\frac{w}{2}} D^{\frac{w}{2}}}{w^w} \end{aligned}$$

Le dénominateur du membre de droite de l'inégalité (4.5.2) se minore quant à lui grâce à la propriété pour tout couple d'entiers $u \geq v$:

$$\binom{u}{v} \geq \left(\frac{u}{v} \right)^v$$

et grâce au fait que $N_{ext} n_{ext} \geq N$, ce qui donne :

$$(|G| - 1)^w \binom{N_{ext} n_{ext}}{w} \geq \left(\frac{N_{ext} n_{ext}}{w} \right)^w \geq \left(\frac{N}{w} \right)^w$$

L'inégalité (4.5.2) donne donc :

$$\begin{aligned}
p_N(w, \leq \lfloor D \rfloor + m_{in}) &= O(1)^w \left(\frac{N}{w} \right)^{\frac{w}{d_{deg}}} \frac{N^{\frac{w}{2}} D^{\frac{w}{2}}}{w^w} \left(\frac{w}{N} \right)^w \\
&= O(1)^w D^{\frac{w}{2}} N^{-\frac{w}{2} + \frac{w}{d_{deg}}} \\
&= \left(O(1) D^{\frac{1}{2}} N^{-\frac{d_{deg}-2}{2d_{deg}}} \right)^w \\
&= \left(O(1) N^{\frac{1}{2} \left(\alpha - \frac{d_{deg}-2}{d_{deg}} \right)} \right)^w
\end{aligned}$$

où l'on a remplacé D par sa valeur N^α dans la dernière ligne.

De plus, la borne 1E affirme que lorsque $w < d_{min}$, $a_{ext}^{N_{ext}}(w) = 0$, ce qui donne $p_N(w, \leq \lfloor D \rfloor + m_{in}) = 0$. Ainsi la première somme s'écrit :

$$\sum_{w=d_{min}}^{\lfloor D \rfloor + m_{in}} p_N(w, \leq D) \leq \sum_{w=d_{min}}^{\infty} \left(O(1) N^{\frac{1}{2} \left(\alpha - \frac{d_{deg}-2}{d_{deg}} \right)} \right)^w$$

où $O(1)$ est majorable par une constante $a > 0$ indépendante de w et N . Comme $\alpha < (d_{deg} - 2)/d_{deg}$, N vérifie à partir d'un certain rang :

$$aN^{\frac{1}{2} \left(\alpha - \frac{d_{deg}-2}{d_{deg}} \right)} < 1$$

ce qui implique :

$$\sum_{w=d_{min}}^{\lfloor D \rfloor + m_{in}} p_N(w, \leq D) \leq O(1) \left(aN^{\frac{1}{2} \left(\alpha - \frac{d_{deg}-2}{d_{deg}} \right)} \right)^{d_{min}}$$

Comme $N^{\frac{1}{2} \left(\alpha - \frac{d_{deg}-2}{d_{deg}} \right)}$ tend vers 0 lorsque N tend vers l'infini, la deuxième somme est donc bien négligeable devant 1 lorsque N tend vers l'infini. $\square$

Lemme 4.5.4. *On pose $D = N^\alpha$, où $\alpha < \frac{d_{deg}-2}{d_{deg}}$. Si $d_{deg} > 2$ et l'encodeur formel interne est anti-récuratif, il existe $x > 0$ tel que la deuxième somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. D'après le lemme 4.1.1, chaque terme de la deuxième somme se majore de la manière suivante :

$$p_N(w, d) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \quad (4.5.3)$$

On applique les bornes 1E et 2I. En remarquant que $d_{min} \geq d_{deg}$, puis que $N_{ext} = O(N)$, la borne 1E donne comme précédemment :

$$a_{ext}^{N_{ext}}(w) = O(1)^w \left(\frac{N}{w} \right)^{\frac{w}{d_{deg}}}$$

Et en remarquant que $w \geq d$ et donc $w + d = O(1)w$, la borne $2I$ donne :

$$\begin{aligned} a_{in}^N(w, d) &= O(1)^d \frac{N^{\frac{d}{2}} (w + d)^{\frac{d}{2}}}{d^d} \\ &= O(1)^w \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \end{aligned}$$

Le dénominateur du membre de droite de l'inégalité (4.5.3) est minoré, comme précédemment, par :

$$\left(\frac{N}{w}\right)^w$$

L'inégalité (4.5.3) donne donc :

$$\begin{aligned} p_N(w, d) &= O(1)^w \left(\frac{N}{w}\right)^{\frac{w}{d_{deg}}} \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \left(\frac{w}{N}\right)^w \\ &= \left(O(1) \left(\frac{w}{N}\right)^{1 - \frac{1}{d_{deg}}}\right)^w \left(\frac{Nw}{d^2}\right)^{\frac{d}{2}} \end{aligned}$$

Majorons $O(1)$ par une constante $a > 0$. Ecrivons le logarithme de la relation obtenue :

$$\ln p_N(w, d) \leq f(w, d)$$

où

$$f(w, d) = w \left(\ln a + \left(1 - \frac{1}{d_{deg}}\right) (\ln w - \ln N) \right) + \frac{d}{2} (\ln N + \ln w - 2 \ln d)$$

On étudie f sur le domaine $\llbracket [D]; [xN] - 1 \rrbracket \times \llbracket 0; [D] \rrbracket$, qui est plus large que le domaine des couples (w, d) intervenant dans la deuxième somme, et ce afin de simplifier les calculs. On se propose de démontrer qu'il existe $x > 0$ tel que le maximum de f est atteint en $([D], [D])$. Dérivons donc d'abord f par rapport à w :

$$\begin{aligned} \frac{\partial}{\partial w} f(w, d) &= \ln a + \left(1 - \frac{1}{d_{deg}}\right) (\ln w - \ln N) + 1 - \frac{1}{d_{deg}} + \frac{d}{2w} \\ &= \ln a + \left(1 - \frac{1}{d_{deg}}\right) (\ln w - \ln N + 1) + \frac{d}{2w} \end{aligned}$$

Or, $w \geq d$ donc $d/2w \leq 1/2$, et d'autre part $w \leq xN$ donc $\ln w - \ln N \leq \ln x$. Ainsi :

$$\frac{\partial}{\partial w} f(w, d) \leq \ln a + \left(1 - \frac{1}{d_{deg}}\right) (\ln x + 1) + \frac{1}{2}$$

Lorsque x tend vers 0, le membre de droite de l'inégalité précédente est équivalent à $(1 - 1/d_{deg}) \ln x$, qui tend vers $-\infty$ car, comme $d_{deg} > 1$, $1 - 1/d_{deg} > 0$. Il existe donc $x > 0$ tel que ce membre est strictement négatif, si bien que $\frac{\partial f}{\partial w} < 0$ dans le domaine considéré.

Supposons donc choisi un tel réel x . f atteint alors son maximum sur le segment $\{\lfloor D \rfloor\} \times \llbracket 0; \lfloor D \rfloor \rrbracket$ puisqu'il faut que w soit minimal. Dérivons à présent f par rapport à d sur ce segment, en remplaçant w par sa valeur $\lfloor D \rfloor = \lfloor N^\alpha \rfloor$:

$$\begin{aligned} \frac{\partial}{\partial d} f(w, d) &= \frac{1}{2} (\ln N + \ln w - 2 \ln d) - 1 \\ &= \frac{1}{2} (\ln N + \ln(\lfloor N^\alpha \rfloor) - 2 \ln d - 2) \end{aligned}$$

Cette expression est une fonction décroissante de d et atteint son minimum en $d = \lfloor N^\alpha \rfloor$. On a donc :

$$\begin{aligned} \frac{\partial}{\partial d} f(w, d) &\geq \frac{1}{2} (\ln N - \ln(\lfloor N^\alpha \rfloor) - 2) \\ &= \frac{1}{2} (\ln N - \alpha \ln N + o(1) - 2) \\ &\underset{N \rightarrow \infty}{\sim} \frac{1 - \alpha}{2} \ln N \end{aligned}$$

la deuxième ligne étant obtenue en remarquant que $\lfloor N^\alpha \rfloor \sim N^\alpha$. Comme $\alpha < 1$, on en déduit que lorsque N est suffisamment grand, $\frac{\partial f}{\partial d} < 0$ sur le segment $\{\lfloor D \rfloor\} \times \llbracket 0; \lfloor D \rfloor \rrbracket$, et f atteint son minimum en $w = d = \lfloor N^\alpha \rfloor$.

Ainsi, dans la deuxième somme, dont le nombre de termes est majorable simplement par $O(N^2)$, le logarithme de chaque terme est majoré par :

$$\begin{aligned} f(\lfloor N^\alpha \rfloor, \lfloor N^\alpha \rfloor) &= \lfloor N^\alpha \rfloor \left(\ln a + \left(1 - \frac{1}{d_{deg}} \right) (\ln \lfloor N^\alpha \rfloor - \ln N) \right) \\ &\quad + \frac{\lfloor N^\alpha \rfloor}{2} (\ln N + \ln \lfloor N^\alpha \rfloor - 2 \ln \lfloor N^\alpha \rfloor) \\ &= \lfloor N^\alpha \rfloor \left(\ln a + \left(\frac{1}{2} - \frac{1}{d_{deg}} \right) (\ln \lfloor N^\alpha \rfloor - \ln N) \right) \end{aligned}$$

Le logarithme de la deuxième somme est donc majoré par :

$$\begin{aligned} \ln(O(N^2)) + f(\lfloor N^\alpha \rfloor, \lfloor N^\alpha \rfloor) &= 2 \ln N + \lfloor N^\alpha \rfloor \left(\ln a + \left(\frac{1}{2} - \frac{1}{d_{deg}} \right) (\ln \lfloor N^\alpha \rfloor - \ln N) \right) \\ &\underset{N \rightarrow \infty}{\sim} N^\alpha \left(\frac{1}{2} - \frac{1}{d_{deg}} \right) (\alpha - 1) \ln N \end{aligned}$$

Comme $d_{deg} > 2$ et $\alpha < 1$, ce majorant tend vers $-\infty$ lorsque N tend vers l'infini, donc la deuxième somme tend vers 0 lorsque N tend vers l'infini. $\square$

Lemme 4.5.5. *On pose $D = N^\alpha$, où $\alpha < \frac{d_{deg}-2}{d_{deg}}$. Si $|G| \geq 4$, $d_{deg} \geq 2$ et l'encodeur formel interne est anti-récursif, alors pour tout $x > 0$, la troisième somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. D'après le lemme 4.1.1, chaque terme de la troisième somme se majore de la manière suivante :

$$p_N(w, d) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \quad (4.5.4)$$

On applique les bornes (2E) et (2I). La borne (2E) affirme qu'il existe une constante $c \in]0; 1[$ telle que :

$$\frac{a_{ext}^{N_{ext}}(w)}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \leq c^w$$

En notant que $w + d = O(N)$, la borne (2I) donne simplement :

$$a_{in}^N(w, d) = O(1)^d N^d$$

L'inégalité (4.5.4) implique alors :

$$p_N(w, d) = c^w O(1)^d N^d$$

Majorons $O(1)$ par une constante $a > 0$ et appliquons le logarithme :

$$\begin{aligned} \ln p_N(w, d) &\leq w \ln c + d(\ln a + \ln N) \\ &\leq \lfloor xN \rfloor \ln c + N^\alpha (\ln a + \ln N) \end{aligned}$$

où la dernière ligne s'obtient en notant que $\ln c < 0$, $w \geq \lfloor xN \rfloor$ et $d \leq N^\alpha$. Ceci est vrai pour tout couple (w, d) correspondant à un terme de la troisième somme. Comme le nombre de ces termes est majorable simplement par $O(N^2)$, on peut majorer le logarithme de la troisième somme par l'expression :

$$\ln(O(N^2)) + \lfloor xN \rfloor \ln c + N^\alpha (\ln a + \ln N)$$

Comme $\alpha < 1$, cette expression est équivalente à $xN \ln c$, et tend vers $-\infty$ lorsque N tend vers l'infini. Cela montre que la troisième somme est négligeable devant 1 lorsque N tend vers l'infini. $\square$

Lemme 4.5.6. *On pose $D = \alpha \log N / (\log \log N)$ où $\alpha < d_{min} - 2$. Si $d_{min} > d_{deg} = 2$ et l'encodeur formel interne est récursif, la première somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. Afin de simplifier la présentation, adoptons la notation :

$$\text{llog} N = \frac{\ln N}{\ln \ln N}$$

si bien que $D = \alpha \text{llog} N$. D'après le lemme 4.1.1, chaque terme de la première somme se majore de la manière suivante :

$$p_N(w, \leq \lfloor D \rfloor + m_{in}) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, \leq \lfloor D \rfloor + m_{in})}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \quad (4.5.5)$$

On applique les bornes $1E$ et $1I$. En remarquant que $d_{deg} = 2$ et $N_{ext} = O(N)$, la borne $1E$ donne :

$$\begin{aligned} a_{ext}^{N_{ext}}(w) &= O(1)^w \left(\frac{N_{ext}}{w} \right)^{\frac{w-d_{min}}{2}+1} \\ &= O(1)^w \left(\frac{N}{w} \right)^{\frac{w-d_{min}}{2}+1} \\ &= O(1)^w N^{\frac{w-d_{min}}{2}+1} \end{aligned}$$

Et en remarquant que $w + \lfloor D \rfloor + m_{in} = O(D)$ car $w \leq \lfloor D \rfloor + m_{in}$, la borne $1I$ donne :

$$\begin{aligned} a_{in}^N(w, \leq \lfloor D \rfloor + m_{in}) &= O(1)^w \frac{N^{\frac{w}{2}} (w + \lfloor D \rfloor + m_{in})^{\frac{w}{2}}}{w^w} \\ &= O(1)^w \frac{N^{\frac{w}{2}} D^{\frac{w}{2}}}{w^w} \end{aligned}$$

Le dénominateur du membre de droite de l'inégalité (4.5.5) se minore quant à lui, comme précédemment, par :

$$\left(\frac{N}{w} \right)^w$$

L'inégalité (4.5.5) donne donc :

$$\begin{aligned} p_N(w, \leq \lfloor D \rfloor + m_{in}) &= O(1)^w N^{\frac{w-d_{min}}{2}+1} \frac{N^{\frac{w}{2}} D^{\frac{w}{2}}}{w^w} \left(\frac{w}{N} \right)^w \\ &= O(1)^w D^{\frac{w}{2}} N^{-\frac{d_{min}}{2}+1} \end{aligned}$$

Il existe donc une constante $a > 1$ telle que :

$$p_N(w, \leq \lfloor D \rfloor + m_{in}) \leq (aD)^{\frac{w}{2}} N^{-\frac{d_{min}}{2}+1}$$

Le membre de droite de ce majorant est une fonction croissante de w . Son maximum est atteint en $w = \lfloor D \rfloor + m_{in}$ et vérifie :

$$\begin{aligned} (aD)^{\frac{\lfloor D \rfloor + m_{in}}{2}} N^{-\frac{d_{min}}{2}+1} &\leq a^{\frac{D+m_{in}}{2}} D^{\frac{D+m_{in}}{2}} N^{-\frac{d_{min}}{2}+1} \\ &= a^{\frac{m_{in}}{2}} a^{\frac{D}{2}} D^{\frac{m_{in}}{2}} D^{\frac{D}{2}} N^{-\frac{d_{min}}{2}+1} \end{aligned}$$

Remplaçons D par sa valeur $\text{allog} N$. Majorons dans l'ordre les facteurs de

l'expression ci-dessus. D'abord, $a^{\frac{m_{in}}{2}} = N^{o(1)}$. Ensuite :

$$\begin{aligned}
a^{\frac{\alpha \text{llog} N}{2}} &= \exp\left(\frac{\alpha \text{llog} N}{2} \ln a\right) \\
&= \exp\left(\frac{\alpha \ln N}{2 \ln \ln N} \ln a\right) \\
&= \exp\left(\ln N \frac{\alpha \ln a}{2 \ln \ln N}\right) \\
&= N^{\frac{\alpha \ln a}{2 \ln \ln N}} \\
&= N^{o(1)}
\end{aligned}$$

Par ailleurs :

$$(\alpha \text{llog} N)^{\frac{m_{in}}{2}} \leq (\alpha \log N)^{\frac{m_{in}}{2}} = N^{o(1)}$$

Finalement, d'après le lemme 4.5.2 :

$$(\alpha \text{llog} N)^{\frac{\alpha \text{llog} N}{2}} \leq N^{\frac{\alpha}{2}}$$

Par conséquent :

$$\begin{aligned}
p_N(w, \leq \lfloor D \rfloor + m_{in}) &\leq N^{o(1)} N^{\frac{\alpha}{2}} N^{-\frac{d_{min}}{2} + 1} \\
&= N^{\frac{\alpha - (d_{min} - 2)}{2} + o(1)}
\end{aligned}$$

La première somme, composée de $\lfloor \alpha \text{llog} N \rfloor + 1 = N^{o(1)}$ termes tous majorés par l'expression précédente, se majore donc également par :

$$N^{\frac{\alpha - (d_{min} - 2)}{2} + o(1)}$$

Comme l'exposant asymptotique en N est strictement négatif, la première somme est négligeable devant 1 lorsque N tend vers l'infini. $\square$

Lemme 4.5.7. *On pose $D = \alpha \ln N / (\ln \ln N)$ où $\alpha < d_{min} - 2$. Si $d_{min} > d_{deg} = 2$ et l'encodeur formel interne est anti-récursif, il existe $x > 0$ tel que la deuxième somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. On adopte toujours la notation :

$$\text{llog} N = \frac{\ln N}{\ln \ln N}$$

si bien que $D = \alpha \text{llog} N$. D'après le lemme 4.1.1, chaque terme de la deuxième somme se majore de la manière suivante :

$$p_N(w, d) \leq \frac{a_{ext}^{N_{ext}}(w) a_{in}^N(w, d)}{(|G| - 1)^w \binom{N_{ext} n_{ext}}{w}} \quad (4.5.6)$$

On applique les bornes (1E) et (2I). En remarquant que $d_{deg} = 2$ et $N_{ext} = O(N)$, la borne 1E donne :

$$\begin{aligned} a_{ext}^{N_{ext}}(w) &= O(1)^w \left(\frac{N_{ext}}{w} \right)^{\frac{w-d_{min}}{2}+1} \\ &= O(1)^w \left(\frac{N}{w} \right)^{\frac{w-d_{min}}{2}+1} \end{aligned}$$

Et en remarquant que $w \geq d$ et donc $w + d = O(1)w$, la borne 2I donne :

$$\begin{aligned} a_{in}^N(w, d) &= O(1)^d \frac{N^{\frac{d}{2}} (w + d)^{\frac{d}{2}}}{d^d} \\ &= O(1)^w \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \end{aligned}$$

Le dénominateur du membre de droite de l'inégalité (4.5.6) est minoré, comme précédemment, par :

$$\left(\frac{N}{w} \right)^w$$

L'inégalité (4.5.6) donne donc :

$$\begin{aligned} p_N(w, d) &= O(1)^w \left(\frac{N}{w} \right)^{\frac{w-d_{min}}{2}+1} \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \left(\frac{w}{N} \right)^w \\ &= O(1)^w \left(\frac{N}{w} \right)^{-\frac{d_{min}}{2}+1} \left(\frac{w}{N} \right)^{\frac{w}{2}} \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \\ &= O(1)^w N^{-\frac{d_{min}}{2}+1} \left(\frac{w}{N} \right)^{\frac{w}{2}} \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \\ &= N^{-\frac{d_{min}}{2}+1} \left(O(1) \left(\frac{w}{N} \right)^{\frac{1}{2}} \right)^w \frac{N^{\frac{d}{2}} w^{\frac{d}{2}}}{d^d} \end{aligned}$$

Majorons $O(1)$ par une constante $a > 0$. On peut alors écrire :

$$p_N(w, d) \leq N^{-\frac{d_{min}}{2}+1} \exp(f(w, d)) \quad (4.5.7)$$

où

$$f(w, d) = w(\ln a + \frac{1}{2}(\ln w - \ln N)) + \frac{d}{2}(\ln N + \ln w - 2 \ln d)$$

On étudie f sur le domaine $[[\lfloor D \rfloor; \lfloor xN \rfloor - 1]] \times [[0; \lfloor D \rfloor]]$, qui est plus large que le domaine des couples (w, d) intervenant dans la deuxième somme, et ce afin de simplifier les calculs. On se propose de démontrer qu'il existe $x > 0$ tel que pour tout couple (w, d) de ce domaine :

$$f(w, d) \leq f(\lfloor D \rfloor, \lfloor D \rfloor) - (w - \lfloor D \rfloor) \quad (4.5.8)$$

Dérivons d'abord f par rapport à w :

$$\frac{\partial}{\partial w} f(w, d) = \ln a + \frac{1}{2}(\ln w - \ln N) + \frac{1}{2} + \frac{d}{2w}$$

Or, $w \geq d$ donc $d/2w \leq 1/2$, et d'autre part $w \leq xN$ donc $\ln w - \ln N \leq \ln x$. Ainsi :

$$\frac{\partial}{\partial w} f(w, d) \leq \ln a + \frac{1}{2} \ln x + 1$$

Ce majorant tend vers $-\infty$ lorsque x tend vers 0. Il existe donc $x > 0$ tel que ce majorant soit inférieur à -1 . On suppose choisi un tel réel x . Le fait que $\frac{\partial f}{\partial w} \leq -1$ entraîne alors que pour tout (w, d) du domaine d'étude, l'accroissement de f entre les points $(\lfloor D \rfloor, d)$ et (w, d) est inférieur à -1 , ou encore :

$$f(w, d) \leq f(\lfloor D \rfloor, d) - (w - \lfloor D \rfloor)$$

Il ne reste qu'à montrer que, pour N suffisamment grand, le maximum de f sur le segment $\{\lfloor D \rfloor\} \times \llbracket 0; \lfloor D \rfloor \rrbracket$ est atteint en $(\lfloor D \rfloor, \lfloor D \rfloor)$. Pour cela, on dérive f par rapport à d sur ce segment :

$$\frac{\partial}{\partial d} f(\lfloor D \rfloor, d) = \frac{1}{2}(\ln N + \ln \lfloor D \rfloor - 2 \ln d) - 1$$

Cette dérivée est une fonction décroissante de d . En remarquant que $D = o(N)$, cela implique :

$$\begin{aligned} \frac{\partial}{\partial d} f(\lfloor D \rfloor, d) &\geq \frac{\partial}{\partial d} f(\lfloor D \rfloor, \lfloor D \rfloor) \\ &= \frac{1}{2}(\ln N - \ln \lfloor D \rfloor) - 1 \\ &\underset{N \rightarrow \infty}{\sim} \frac{1}{2} \ln N \end{aligned}$$

Ainsi, pour N suffisamment grand, la dérivée partielle de f par rapport à d est positive, ce qui montre la propriété (4.5.8).

De plus :

$$\begin{aligned} f(\lfloor D \rfloor, \lfloor D \rfloor) &= \lfloor D \rfloor (\ln a + \frac{1}{2}(\ln \lfloor D \rfloor - \ln N)) + \frac{\lfloor D \rfloor}{2}(\ln N + \ln \lfloor D \rfloor - 2 \ln \lfloor D \rfloor) \\ &= \lfloor D \rfloor (\ln a + \frac{1}{2}(\ln \lfloor D \rfloor - \ln N)) + \lfloor D \rfloor \frac{1}{2}(\ln N - \ln \lfloor D \rfloor) \\ &= \lfloor D \rfloor \ln a \end{aligned}$$

La propriété (4.5.8) se réécrit donc :

$$f(w, d) \leq \lfloor D \rfloor \ln a - (w - \lfloor D \rfloor)$$

En outre, comme

$$\exp(\lfloor D \rfloor \ln a) = \exp\left(\ln N \frac{\alpha \ln a}{\ln \ln N}\right) = N^{\frac{\alpha \ln a}{\ln \ln N}} = N^{o(1)}$$

la relation (4.5.7) donne :

$$p_N(w, d) \leq N^{-\frac{d_{min}}{2}+1} N^{o(1)} \exp(-(w - \lfloor D \rfloor))$$

Par conséquent, pour tout $d \in \llbracket 0; \lfloor D \rfloor \rrbracket$:

$$\begin{aligned} \sum_{w=\lfloor D \rfloor + m_{in} + 1}^{\lfloor xN \rfloor - 1} p_N(w, d) &\leq N^{-\frac{d_{min}}{2}+1+o(1)} \sum_{w=\lfloor D \rfloor + m_{in} + 1}^{\lfloor xN \rfloor - 1} \exp(-(w - \lfloor D \rfloor)) \\ &= O(1) N^{-\frac{d_{min}}{2}+1+o(1)} \end{aligned}$$

La somme de cette relation sur les $\lfloor D \rfloor + 1 = N^{o(1)}$ valeurs de d permet également de majorer la troisième somme par l'expression :

$$O(1) N^{-\frac{d_{min}}{2}+1+o(1)}$$

qui, en notant que $d_{min} > 2$, est négligeable devant 1 lorsque N tend vers l'infini. $\square$

Lemme 4.5.8. *On pose $D = \alpha \ln N / (\ln \ln N)$ où $\alpha < d_{min} - 2$. Si $|G| \geq 4$, $d_{deg} = 2$ et l'encodeur formel interne est anti-récuratif, alors pour tout $x > 0$, la troisième somme est négligeable devant 1 lorsque N tend vers l'infini.*

Démonstration. Le même raisonnement que celui du lemme 4.5.5 s'applique, en remplaçant simplement N^α par $\alpha \ln N / (\ln \ln N)$. $\square$

Bibliographie

- [1] I. Andriyanova, D. Maurice, and J.P. Tillich. Quantum ldpc codes obtained by non-binary constructions. In *Information Theory Proceedings (ISIT), 2012 IEEE International Symposium on*, pages 343–347. IEEE, 2012.
- [2] E. Arikan. Channel polarization : A method for constructing capacity-achieving codes for symmetric binary-input memoryless channels. *CoRR*, abs/0807.3917, 2008.
- [3] S. Benedetto, D. Divsalar, G. Montorsi, and F. Pollara. Serial concatenation of interleaved codes : Performance analysis, design, and iterative decoding. *Information Theory, IEEE Transactions on*, 44(3) :909–926, 1998.
- [4] S. Benedetto and G. Montorsi. Performance evaluation of turbo-codes. *Electronics letters*, 31(3) :163–165, 1995.
- [5] S. Benedetto and G. Montorsi. Design of parallel concatenated convolutional codes. *Communications, IEEE Transactions on*, 44(5) :591–600, 1996.
- [6] S. Benedetto and G. Montorsi. Iterative decoding of serially concatenated convolutional codes. *Electronics letters*, 32(13) :1186–1188, 1996.
- [7] S. Benedetto and G. Montorsi. Serial concatenation of block and convolutional codes. *Electronics Letters*, 32(10) :887–888, 1996.
- [8] S. Benedetto and G. Montorsi. Unveiling turbo codes : Some results on parallel concatenated coding schemes. *Information Theory, IEEE Transactions on*, 42(2) :409–428, 1996.
- [9] C.H. Bennett, G. Brassard, et al. Quantum cryptography : Public key distribution and coin tossing. In *Proceedings of IEEE International Conference on Computers, Systems and Signal Processing*, volume 175. Bangalore, India, 1984.
- [10] C. Berrou, A. Glavieux, and P. Thitimajshima. Near Shannon limit error-correcting coding and decoding : Turbo-codes. In *Communications, 1993. ICC 93. Geneva. Technical Program, Conference Record, IEEE International Conference on*, volume 2, pages 1064–1070. IEEE, 1993.

- [11] E. Boutillon, C. Douillard, and G. Montorsi. Iterative decoding of concatenated convolutional codes : Implementation issues. *Proceedings of the IEEE*, 95(6) :1201–1227, 2007.
- [12] M. Breiling. A logarithmic upper bound on the minimum distance of turbo codes. *Information Theory, IEEE Transactions on*, 50(8) :1692–1710, 2004.
- [13] K.L. Brown, W.J. Munro, and V.M. Kendon. Using quantum computers for quantum simulation. *Entropy*, 12(11) :2268–2307, 2010.
- [14] T.M. Cover and J.A. Thomas. *Elements of information theory*. Wiley-interscience, 2006.
- [15] I. Devetak. The private classical capacity and quantum capacity of a quantum channel. *Information Theory, IEEE Transactions on*, 51(1) :44–55, 2005.
- [16] E. Farhi, J. Goldstone, S. Gutmann, and M. Sipser. Quantum computation by adiabatic evolution. *arXiv preprint quant-ph/0001106*, 2000.
- [17] R.P. Feynman. Simulating physics with computers. *International journal of theoretical physics*, 21(6) :467–488, 1982.
- [18] R. Gallager. Low-density parity-check codes. *Information Theory, IRE Transactions on*, 8(1) :21–28, 1962.
- [19] D. Gottesman. Stabilizer codes and quantum error correction. *thèse, Arxiv preprint quant-ph/9705052*, 1997.
- [20] M. Grassl and T. Beth. Quantum bch codes. *arXiv preprint quant-ph/9910060*, 1999.
- [21] J. Hagenauer, E. Offer, and L. Papke. Iterative decoding of binary block and convolutional codes. *Information Theory, IEEE Transactions on*, 42(2) :429–445, 1996.
- [22] N. Kahale and R. Urbanke. On the minimum distance of parallel and serially concatenated codes. In *Information Theory, 1998. Proceedings. 1998 IEEE International Symposium on*, page 31. IEEE, 1998. Disponible à l’adresse lthcwww.epfl.ch/~ruediger/papers/weight.ps.
- [23] K. Kasai, M. Hagiwara, H. Imai, and K. Sakaniwa. Quantum error correction beyond the bounded distance decoding limit. *Information Theory, IEEE Transactions on*, 58(2) :1223–1230, 2012.
- [24] I.H. Kim. *Quantum codes on Hurwitz surfaces*. PhD thesis, Massachusetts Institute of Technology, 2007.
- [25] J.H. Kim and J. Pearl. A computational model for causal and diagnostic reasoning in inference systems. In *Proceedings of the 8th International Joint Conference on Artificial Intelligence*, pages 190–193. Citeseer, 1983.

- [26] F.R. Kschischang, B.J. Frey, and H.A. Loeliger. Factor graphs and the sum-product algorithm. *Information Theory, IEEE Transactions on*, 47(2) :498–519, 2001.
- [27] S. Le Goff, A. Glavieux, and C. Berrou. Turbo-codes and high spectral efficiency modulation. In *Communications, 1994. ICC'94, SUPER-COMM/ICC'94, Conference Record, 'Serving Humanity Through Communications. 'IEEE International Conference on*, pages 645–649. IEEE, 1994.
- [28] D. Lidar and K. Birgitta Whaley. Decoherence-free subspaces and subsystems. *Irreversible Quantum Dynamics*, pages 83–120, 2003.
- [29] S. Lloyd. Capacity of the noisy quantum channel. *Physical Review A*, 55(3) :1613, 1997.
- [30] S. Lloyd et al. Universal quantum simulators. *SCIENCE-NEW YORK THEN WASHINGTON-*, pages 1073–1077, 1996.
- [31] D.J.C. MacKay. Good error-correcting codes based on very sparse matrices. *Information Theory, IEEE Transactions on*, 45(2) :399–431, 1999.
- [32] D.J.C. MacKay and R.M. Neal. Near Shannon limit performance of low density parity check codes. *Electronics letters*, 32(18) :1645, 1996.
- [33] R. McEliece. *The theory of information and coding*. Cambridge University Press, 2002.
- [34] M.A. Nielsen and I.L. Chuang. *Quantum computation and quantum information*. Cambridge university press, 2010.
- [35] J. Pearl. *Reverend Bayes on inference engines : a distributed hierarchical approach*. Cognitive Systems Laboratory, School of Engineering and Applied Science, University of California, Los Angeles, 1982.
- [36] J. Pearl. Fusion, propagation, and structuring in belief networks. *Artificial intelligence*, 29(3) :241–288, 1986.
- [37] J. Pearl. *Probabilistic reasoning in intelligent systems : networks of plausible inference*. Morgan Kaufmann, 1988.
- [38] E. Pelchat and D. Poulin. Degenerate viterbi decoding. *arXiv preprint arXiv :1204.2439*, 2012.
- [39] L.C. Perez, J. Seghers, and D.J. Costello Jr. A distance spectrum interpretation of turbo codes. *Information Theory, IEEE Transactions on*, 42(6) :1698–1709, 1996.
- [40] D. Poulin and Y. Chung. On the iterative decoding of sparse quantum codes. *Quantum Information & Computation*, 8(10) :987–1000, 2008.

- [41] D. Poulin, J.P. Tillich, and H. Ollivier. Quantum serial turbo codes. *Information Theory, IEEE Transactions on*, 55(6) :2776–2798, 2009.
- [42] J. Preskill. Lecture notes for Physics 229 : Chapter 7 : Quantum Error Correction. *California Institute of Technology*, 1999.
- [43] J.M. Renes, F. Dupuis, and R. Renner. Efficient quantum polar coding. *arXiv preprint arXiv :1109.3195*, 2011.
- [44] P. Robertson. Illuminating the structure of code and decoder of parallel concatenated recursive systematic (turbo) codes. In *Global Telecommunications Conference, 1994. GLOBECOM'94. Communications : The Global Bridge., IEEE*, volume 3, pages 1298–1303. IEEE, 1994.
- [45] P.K. Sarvepalli, A. Klappenecker, and M. Rotteler. Asymmetric quantum ldpc codes. In *Information Theory, 2008. ISIT 2008. IEEE International Symposium on*, pages 305–309. IEEE, 2008.
- [46] P.W. Shor. Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer. *SIAM Journal on Computing*, 26(5) :1484–1509, 1997.
- [47] P.W. Shor. The quantum channel capacity and coherent information. In *lecture notes, MSRI Workshop on Quantum Computation*, 2002.
- [48] P.W. Shor and J.A. Smolin. Quantum error-correcting codes need not completely reveal the error syndrome. *arXiv preprint quant-ph/9604006*, 1996.
- [49] J.P. Tillich and G. Zémor. Quantum ldpc codes with positive rate and minimum distance proportional to $n^{1/2}$. In *Information Theory, 2009. ISIT 2009. IEEE International Symposium on*, pages 799–803. IEEE, 2009.
- [50] M. Wilde and S. Guha. Polar codes for classical-quantum channels. 2011.
- [51] M.M. Wilde. From classical to quantum shannon theory. *arXiv preprint arXiv :1106.1445*, 2011.
- [52] M.M. Wilde and M.H. Hsieh. Entanglement boosts quantum turbo codes. In *Information Theory Proceedings (ISIT), 2011 IEEE International Symposium on*, pages 445–449. IEEE, 2011.
- [53] G. Zémor. On cayley graphs, surface codes, and the limits of homological coding for quantum error correction. *Coding and Cryptology*, pages 259–273, 2009.