



Étude de la distance minimale des codes stabilisateurs locaux

Nicolas Saussay

► To cite this version:

Nicolas Saussay. Étude de la distance minimale des codes stabilisateurs locaux. Théorie de l'information et codage [math.IT]. Université de Limoges, 2025. Français. NNT : 2025LIMO0052 . tel-05422156

HAL Id: tel-05422156

<https://theses.hal.science/tel-05422156v1>

Submitted on 17 Dec 2025

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.



HAL Authorization

Université de Limoges

ED 653 : SCIENCES ET INGENIERIE (SI)

Faculté des Sciences et Techniques – Institut de Recherche XLIM

Thèse pour obtenir le grade de
Docteur de l'Université de Limoges
Mathématiques

Présentée et soutenue par

Nicolas SAUSSAY

Le 09 décembre 2025

ÉTUDE DE LA DISTANCE MINIMALE DES CODES STABILISATEURS LOCAUX

Thèse dirigée par François ARNAULT et Philippe GABORIT

JURY :

Président du jury

Gilles ZÉMOR, Professeur à l'Université de Bordeaux – IMB (UMR 5251) – Bordeaux

Rapporteurs

Xavier CARUSO, Directeur de recherche CNRS – IMB (UMR 5251) – Bordeaux

Alain COUVREUR, Directeur de recherche INRIA (GRACE) – LIX (CNRS UMR 7161) – Palaiseau

Examineurs

François ARNAULT, Maître de conférences à l'Université de Limoges – XLIM (UMR 7552) – Limoges

Philippe GABORIT, Professeur à l'Université de Limoges – XLIM (UMR 7552) – Limoges

Olivier RUATTA, Professeur à l'Université de Limoges – XLIM (UMR 7552) – Limoges



Dédicace à mes parents
et à celui que j'aspire à devenir.

Do not stop when you are tired.

Stop when you are done.

David Goggins

Remerciements

Tout d'abord, je tiens à remercier mes parents. Vous avez toujours été là, pour me soutenir. TOUJOURS ! Je ne serais jamais arrivé là sans votre aide et vos conseils. Cœur sur vous, la team de l'ombre !

Merci à François Arnault, avec qui j'ai eu de nombreuses discussions autour de mes preuves. François, parler de maths avec toi a toujours été un plaisir, surtout quand tu me présentais ta vision topologique des codes surfaces. Deux mots : simple et élégante.

Merci à Xavier Caruso, qui, par ses remarques constructives, m'a permis d'affiner certains résultats et d'améliorer ma compréhension globale des outils mathématiques sur lesquels mes recherches se sont appuyées.

Merci à Philippe Gaborit, qui m'a montré comment gagner en efficacité dans mon travail. À tes côtés, j'ai compris qu'être un bon chercheur ne se résume pas à rédiger des preuves : il faut savoir où chercher, quand abandonner ou accélérer un projet et il faut également savoir présenter ses travaux pour qu'ils soient acceptés. Les compétences que j'ai acquises pendant ces trois années me serviront tout au long de ma vie. Merci.

Merci à Wouter Rozendaal, avec qui j'ai eu énormément de plaisir à travailler et à échanger. Même si nos approches étaient différentes : toi avec ton intuition géométrique et moi avec mon formalisme algébrique, on a réussi à avancer et à produire des résultats que je trouve MA-GNI-FI-QUE. La preuve de la borne de Bravyi-Terhal améliorée : une pépite.

Enfin, merci à Gilles Zémor, que j'ai d'abord découvert à travers ses articles, avant d'avoir la chance de travailler directement avec lui. Gilles, c'est en te lisant que j'ai compris que, même dans ce domaine, je pouvais retrouver les mathématiques théoriques que j'aimais tant. À chacun de nos échanges, j'ai appris de ta culture et de ton intuition. Ce fut un vrai plaisir.

Droits d'auteurs

Cette création est mise à disposition selon le Contrat :

« Attribution-Pas d'Utilisation Commerciale-Pas de modification 3.0 France »

disponible en ligne : <http://creativecommons.org/licenses/by-nc-nd/3.0/fr/>



- COPYRIGHT Version française Cette création est mise à disposition selon les termes de la Licence Creative Commons **Attribution - Pas d'Utilisation Commerciale - Pas de Modification 4.0 International**.

Disponible en ligne: <https://creativecommons.org/licenses/by-nc-nd/4.0/>



En cas de doute, se référer à : <https://creativecommons.org/licenses/by-nc-nd/4.0/>

Table des Matières

1	Introduction générale	9
2	Guide des notations et conventions	12
2.1	Notations générales	13
2.2	Conventions d'écriture	14
3	Préliminaires	16
3.1	Rappels sur les réseaux	17
3.1.1	Sous-groupes discrets de $\mathbb{R}^D$	17
3.1.2	Les réseaux et leurs propriétés	18
3.2	Rappels sur les codes quantiques	21
3.2.1	Fondamentaux sur les qubits	21
3.2.2	Le groupe de Pauli sur n-qubits	22
3.2.3	Codes quantiques	24
4	Classification des GB(2,2) et comparaison avec les codes surfaces	28
4.1	Introduction	30
4.1.1	Motivation	30
4.1.2	Anciens résultats	31
4.1.3	État de l'art et contributions	31
4.1.4	Vue d'ensemble générale	34
4.1.5	Organisation du chapitre	35
4.2	Notions clés	36
4.2.1	Codes Bicycle généralisés (GB)	36
4.2.2	Code CSS associé au graphe de Cayley d'un groupe abélien	38
4.3	Borne inférieure sur la distance minimale des GB(2,2)	42
4.3.1	Propriétés des codes GB(2,2)	42
4.3.2	Interprétation des codes $GB(1+X^a, 1+X^b, n)$ à l'aide des graphes de Cayley	43

4.3.3	Borne inférieure sur la distance minimale des codes GB(2,2)	44
4.4	Preuve du théorème 4.3.2	45
4.4.1	Plan de la preuve :	45
4.4.2	Réinterprétation de la distance minimale	46
4.4.3	Des marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ vers les marches sur $\mathbb{Z}^2$	48
4.4.4	Des marches sur $\mathbb{Z}^2$ vers les marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$	52
4.4.5	Fin de la preuve du théorème 4.3.2	53
4.5	Application du théorème 4.3.2	56
4.5.1	Etude de $GB(1 + X, 1 + X^n, n^2)$	57
4.5.2	Etude de $GB(1 + X, 1 + X^{2r-1}, 2r^2)$	59
4.5.3	Etude de $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$	60
4.6	Relations d'équivalence sur les codes quantiques	63
4.6.1	Relation d'équivalence générale	63
4.6.2	Relation d'équivalence CGP	64
4.7	Comparaison des codes GB aux codes surface 2D de poids quatre	66
4.7.1	Théorème de Whitney	67
4.7.2	Vue d'ensemble de la démonstration de la non-équivalence	68
4.7.3	Comparaison avec le code surface 2D optimal pour la distance paire	69
4.7.4	Comparaison avec les codes de Kitaev standards	84
4.7.5	Comparaison avec les codes surface 2D optimaux pour la distance impaire	86
4.8	Classification des codes GB(2,2)	94
4.8.1	Méthode :	94
4.8.2	Classification des codes GB(2,2) extrémaux	95
4.8.3	Classification de tous les codes GB(2,2) (Longueur ≤ 200)	101
4.8.4	Résumé des meilleurs codes surface 2D à stabilisateurs de poids quatre	101
4.9	Conclusion du chapitre	105
5	Borne de Bravyi-Terhal généralisée	108
5.1	Introduction	110
5.1.1	Motivations	110
5.1.2	Résultats connus : borne de Bravyi-Terhal	111
5.1.3	Contributions	112
5.1.4	Vue d'ensemble sur le théorème 5.4.1	113
5.2	Organisation du chapitre	114

5.3	Résultats préliminaires	114
5.3.1	Mesure de Haar	115
5.3.2	Constantes de Hermite et de Rankin	115
5.4	Borne de Bravyi-Terhal généralisée	117
5.4.1	Borne de Bravyi-Terhal originale	117
5.4.2	Borne de Bravyi-Terhal généralisée	118
5.5	Esquisse de la preuve du théorème	119
5.5.1	Cas des réseaux de faible covolume : $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$. . .	119
5.5.2	Cas général : $\ell^{1/D} \geq 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$	120
5.6	Preuve du théorème 5.4.1	122
5.6.1	Construction d'une base du réseau adaptée au découpage . . .	122
5.6.2	Découpage du domaine fondamental	125
5.6.3	Application du lemme de nettoyage	129
5.6.4	Majoration du nombre de qubits par tranche	131
5.6.5	Preuve complète du théorème 5.4.1 :	136
5.7	Étude de la distance minimale des codes 2BGA abéliens	139
5.7.1	Présentation des codes 2BGA abéliens	140
5.7.2	Étude des codes 2BGA abéliens non triviaux	141
5.7.3	Démonstration du théorème 5.7.1	143
5.8	Conclusion du chapitre :	153
6	Conclusion générale	155
7	Bibliographie	157
	Références	158
	Liste des travaux	161

1

Introduction générale

Introduction générale

Bien qu'aujourd'hui de nombreux calculs soient irréalisables dans des délais raisonnables, la perspective d'un ordinateur quantique suffisamment puissant pourrait changer la donne en permettant de résoudre efficacement ces problèmes [1]. Cependant, la réalisation de calculs quantiques fiables et à grande échelle demeure un défi de taille. Le principal obstacle réside dans la fragilité inhérente de l'information quantique, qui se dégrade rapidement sous l'effet des perturbations environnementales. Sans des mécanismes de correction d'erreur quantique efficaces – incluant des codes quantiques robustes, des circuits de mesure de syndrome pour détecter les erreurs et des algorithmes de décodage avancés – les promesses de l'informatique quantique ne sauraient être tenues.

Les codes stabilisateurs quantiques LDPC (Low-Density Parity-Check) sont considérés comme essentiels pour le développement de l'informatique quantique à grande échelle, en raison de leur capacité à corriger efficacement les erreurs.

Nos travaux se concentrent sur la sous-classe des codes stabilisateurs locaux. Ces

codes, dont les qubits sont représentés sur des points de $\mathbb{R}^D$, limitent la propagation des erreurs en garantissant que chaque stabilisateur n'opère que sur un petit groupe de qubits voisins. Notre objectif principal est d'étudier comment l'agencement spatial des qubits dans l'espace influence la distance minimale de ces codes.

Structure du document

Ce document est organisé en quatre chapitres répartis en deux parties : la première présente les concepts clés et les outils, tandis que la seconde expose les contributions originales de cette thèse.

Fondations théoriques

Les deux premiers chapitres posent les bases de l'étude :

- Le Chapitre 2 présente les notations et conventions utilisées tout au long de la thèse.
- Le Chapitre 3 introduit les notions théoriques relatives aux réseaux et aux codes quantiques, qui ont été utilisées pour établir nos résultats.

Contributions et résultats

Les deux derniers chapitres sont dédiés à la présentation des travaux de recherche.

- Le Chapitre 4 est consacré à l'étude des codes Bicycle généralisés (GB) de poids quatre. Nous y détaillons la construction de plusieurs familles de codes inédites, dont les performances égalent celles des meilleurs codes surface 2D de poids quatre. Ce chapitre propose également deux classifications, chacune portant sur des codes de poids quatre et de longueur inférieure ou égale à 200 : l'une pour les meilleurs codes GB et l'autre pour les meilleurs codes surface 2D.
- Enfin, le Chapitre 5 est consacré à nos travaux, menés en collaboration avec Wouter Rozendaal et Gilles Zémor, sur les codes stabilisateurs locaux. Cette étude se concentre sur les codes stabilisateurs dont les qubits sont représentés sur les points à coordonnées entières de l'espace euclidien. Nous introduisons

une borne théorique sur la distance minimale de ces codes, qui généralise la célèbre borne de Bravyi-Terhal. L'application de ce nouveau résultat aux codes 2BGA abéliens nous a permis d'établir une borne sur leur distance minimale et d'affiner une borne existante pour la sous-famille des codes GB.

2

Guide des notations et conventions

Sommaire

2.1	Notations générales	13
2.2	Conventions d'écriture	14

Ce guide présente les notations et les conventions d'écriture utilisées pour assurer la cohérence de ce mémoire.

2.1 Notations générales

Ensembles fondamentaux

- $\mathbb{Z}$: Ensemble des entiers relatifs.
- $\mathbb{R}$: Ensemble des nombres réels.
- $\mathbb{C}$: Ensemble des nombres complexes.
- $\mathbb{R}^D$: Ensemble des vecteurs à $D \geq 1$ coordonnées réelles.
- $\mathbb{Z}/n\mathbb{Z}$: Groupe cyclique à $n \geq 1$ éléments, représentant le quotient de $\mathbb{Z}$ par l'idéal $n\mathbb{Z}$.

Algèbre linéaire et espaces vectoriels

- $[\mathbf{A} \mid \mathbf{B}]$: Représente la matrice obtenue en mettant les matrices $\mathbf{A}$ et $\mathbf{B}$ côte à côte.
- $\mathbf{M}_{n,m}(\mathbb{K})$: Désigne l'ensemble des matrices rectangulaires à n lignes et m colonnes, dont les coefficients appartiennent à l'ensemble $\mathbb{K}$.
- $GL_n(\mathbb{K})$: Ensemble des matrices de taille $n \times n$ à coefficients dans $\mathbb{K}$ qui sont inversibles.
- $\ker(\mathbf{A})$: Noyau à droite de la matrice $\mathbf{A}$, défini comme l'ensemble des vecteurs $\mathbf{x}$ tels que $\mathbf{A}\mathbf{x} = \mathbf{0}$.
- $\text{Vect}_{\mathbb{L}} \mathbf{A}$ ou $rs(\mathbf{A})$: Espace vectoriel engendré par les lignes de la matrice $\mathbf{A}$.
- $\langle \cdot, \cdot \rangle$: le produit scalaire euclidien usuel de $\mathbb{R}^D$.
- $\mathbb{A}^\perp$: Pour $\mathbb{A} \subset \mathbb{R}^D$, $\mathbb{A}^\perp$ désigne l'ensemble des éléments orthogonaux à $\mathbb{A}$ pour le produit scalaire euclidien usuel.
- $\|\cdot\|$: Représente la norme euclidienne sur $\mathbb{R}^D$.
- $\text{Im}\psi$: Image du morphisme $\psi : \mathbb{A} \rightarrow \mathbb{B}$, définie comme l'ensemble $\{\psi(x) \mid x \in \mathbb{A}\}$.

Théorie des codes et algèbre de groupe

- $\mathbb{F}_2[\mathbb{G}]$: Algèbre binaire du groupe $\mathbb{G}$, représentant l'ensemble des sommes formelles d'éléments de $\mathbb{G}$ à coefficients dans le corps $\mathbb{F}_2$.

- $\text{Circ}(A(X), n)$: Matrice circulante de taille $n \times n$ engendrée par le polynôme $A(X) \in \mathbb{F}_2[X]$. Ses lignes sont données par les décalages cycliques successifs des coefficients de $A(X)$.
- $\text{GB}(A(X), B(X), n)$: Code Bicycle Généralisé de longueur $2n$, engendré par les matrices $\text{Circ}(A(X), n)$ et $\text{Circ}(B(X), n)$.

Opérateurs quantiques et groupes de Pauli

- $\mathbb{P}_n$: Groupe de Pauli sur n qubits.
- X : Matrice de Pauli de type "bit-flip": $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$.
- Z : Matrice de Pauli de type "phase-flip": $\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$.
- Y : Matrice de Pauli de type "bit-phase-flip": $iXZ = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}$.
- S : Matrice de phase : $\begin{pmatrix} 1 & 0 \\ 0 & i \end{pmatrix}$.
- H : Matrice de Hadamard : $\frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$.

Théorie des graphes

- $(\mathbb{G}, g_1, g_2)$: Graphe de Cayley sur le groupe $\mathbb{G}$ engendré par les éléments g_1, g_2 .
- $C_0 \rightarrow C_1 \rightarrow \dots \rightarrow C_n$: Désigne une marche dans un graphe allant du sommet C_0 au sommet C_n .

2.2 Conventions d'écriture

Tout au long du mémoire, nous adopterons les conventions suivantes pour la notation des éléments, des ensembles et des codes :

Typographie des éléments mathématiques

- Les **ensembles** seront généralement notés avec des lettres capitales stylisées en double barre (ex : $\mathbb{R}, \mathbb{S}, \mathbb{Z}$).
- Les **entiers** seront représentés par des lettres minuscules (ex : n, ℓ).
- Les **intervalles d'entiers** seront représentés entre $\llbracket \cdot, \cdot \rrbracket$ (ex : $\llbracket 1, 4 \rrbracket = \{1, 2, 3, 4\}$).

- Les **polynômes** seront désignés par des lettres majuscules (ex : $A(X), B(X)$).
- Les **vecteurs** seront notés en gras avec des lettres minuscules (ex : $\mathbf{u}, \mathbf{v}$). Par convention, tous les vecteurs seront représentés en colonne.
- Les **matrices** seront notées en gras avec des lettres majuscules (ex : $\mathbf{A}, \mathbf{B}$).
- Les **éléments de quotients** de $\mathbb{Z}$ (comme $\mathbb{Z}/n\mathbb{Z}$) seront marqués d'une barre au-dessus (ex : $\bar{0}, \bar{1}$).
- Les **éléments de quotients** de $\mathbb{R}^D$ (comme $\mathbb{R}^D/\mathbb{L}$) seront notés par $\mathbf{x} \bmod \mathbb{L}$ ou $[\mathbf{x}]$.

Conventions spécifiques aux codes

- Les paramètres d'un **code quantique** seront encadrés par des doubles crochets, par exemple $[[N, k, d]]$.
- Les **codes quantiques** seront toujours notés avec la lettre Q en majuscule calligraphiée (ex: $\mathcal{Q}, \mathcal{Q}'$).
- Les **opérateurs** agissant sur les qubits seront notés en lettres majuscules calligraphiées (ex: $\mathcal{O}, \mathcal{S}$).

3

Préliminaires

Sommaire

3.1	Rappels sur les réseaux	17
3.1.1	Sous-groupes discrets de $\mathbb{R}^D$	17
3.1.2	Les réseaux et leurs propriétés	18
3.2	Rappels sur les codes quantiques	21
3.2.1	Fondamentaux sur les qubits	21
3.2.2	Le groupe de Pauli sur n-qubits	22
3.2.3	Codes quantiques	24

Ce chapitre a pour objectif de revoir les concepts mathématiques fondamentaux nécessaires à la compréhension des deux prochains chapitres. Nous commencerons par faire un rappel sur les réseaux, avant d'aborder les notions fondamentales des codes quantiques.

Tout au long de ce document, nous supposerons connues les définitions et propriétés de base des structures algébriques (groupes, anneaux, corps, modules sur un anneau, espaces vectoriels et quotients de groupes), ainsi que celles relatives aux produits scalaires, aux normes et à la compacité dans un espace euclidien. Les notions fondamentales des codes correcteurs classiques seront également considérées comme connues.

3.1 Rappels sur les réseaux

3.1.1 Sous-groupes discrets de $\mathbb{R}^D$

Définition 3.1.1 (Groupe discret). *Un sous-groupe additif $\mathbb{G} \subseteq \mathbb{R}^D$ est discret si tout compact de $\mathbb{R}^D$ ne contient qu'un nombre fini de points de $\mathbb{G}$.*

La proposition suivante caractérise les sous-groupes discrets : un sous-groupe $\mathbb{G}$ de $\mathbb{R}^D$ est discret si et seulement si l'origine est un point isolé de $\mathbb{G}$.

Proposition 3.1.1. *Un sous-groupe $\mathbb{G}$ de $\mathbb{R}^D$ est discret si et seulement s'il existe $\alpha > 0$ tel que la boule ouverte $B(0, \alpha)$ (centrée en 0 et de rayon α , pour la norme euclidienne $\|\cdot\|$) ne contienne aucun point de $\mathbb{G}$ autre que l'origine :*

$$B(0, \alpha) \cap \mathbb{G} = \{0\}$$

Exemple 3.1.1. $\mathbb{Z}^D$ et tous ses sous-groupes sont discrets.

Lorsque $\mathbb{G}$ est un sous-groupe discret de $\mathbb{R}^D$, le groupe quotient $\mathbb{R}^D/\mathbb{G}$ est naturellement muni d'une topologie métrisable, induite par celle de $\mathbb{R}^D$.

Proposition 3.1.2 (Distance dans $\mathbb{R}^D/\mathbb{G}$). *Soit $\mathbb{G} \subset \mathbb{R}^D$ un sous-groupe discret et $\|\cdot\|$ la norme euclidienne. Notons dist l'application définie par :*

$$\begin{aligned} \text{dist} : \mathbb{R}^D/\mathbb{G} \times \mathbb{R}^D/\mathbb{G} &\longrightarrow \mathbb{R}^+ \\ ([\mathbf{a}], [\mathbf{b}]) &\mapsto \inf_{\mathbf{g} \in \mathbb{G}} \|\mathbf{a} - (\mathbf{b} + \mathbf{g})\| \end{aligned} \quad (3.1)$$

où $[\mathbf{a}]$ et $[\mathbf{b}]$ désignent les classes modulo $\mathbb{G}$ des éléments $\mathbf{a}$ et $\mathbf{b}$.

L'espace $\mathbb{R}^D/\mathbb{G}$ muni de dist est un espace métrique. De plus, pour n'importe quelle paire d'éléments $(\mathbf{a}, \mathbf{b}) \in \mathbb{R}^D \times \mathbb{R}^D$, la distance entre $[\mathbf{a}]$ et $[\mathbf{b}]$ est atteinte :

$$\exists \mathbf{g}_0 \in \mathbb{G}, \text{dist}([\mathbf{a}], [\mathbf{b}]) = \|\mathbf{a} - (\mathbf{b} + \mathbf{g}_0)\|$$

3.1.2 Les réseaux et leurs propriétés

Les réseaux constituent un type important de groupes discrets. Il s'agit de $\mathbb{Z}$ -modules, libres engendrés par des vecteurs linéairement indépendants de $\mathbb{R}^D$.

Définition et structure des réseaux

Définition 3.1.2 (Réseaux). *Un réseau $\mathbb{L}$ est un sous-groupe discret de $\mathbb{R}^D$ qui est de la forme : $\mathbb{L} = \bigoplus_{i=1}^k \mathbb{Z}\mathbf{u}_i$ où k est un entier compris entre 1 et D et où $\mathbf{u}_1, \dots, \mathbf{u}_k \in \mathbb{R}^D$ sont des vecteurs linéairement indépendants sur $\mathbb{R}$. La famille $\{\mathbf{u}_1, \dots, \mathbf{u}_k\}$ est appelée une base de $\mathbb{L}$, et l'entier k son rang. Si $k = D$, le réseau est dit de rang plein. Dans le cas contraire, on le qualifie de réseau relatif.*

Remarque 3.1.1. *Tout réseau $\mathbb{L} \subset \mathbb{R}^D$ est un $\mathbb{Z}$ -module libre de rang $1 \leq k \leq D$. Lorsque son rang est D , on peut écrire $\mathbb{L} = \mathbf{M} \cdot \mathbb{Z}^D$ où $\mathbf{M} \in GL_D(\mathbb{R})$.*

La double nature algébrique et géométrique des réseaux permet de les étudier sous différents angles. Parmi leurs propriétés cruciales, celles issues de la géométrie sont souvent définies à l'aide des domaines fondamentaux des réseaux.

Invariants géométriques et algébriques des réseaux

Définition 3.1.3 (Domaine fondamental). *Soit $\mathbb{L} = \bigoplus_{i=1}^k \mathbb{Z}\mathbf{u}_i \subset \mathbb{R}^D$ un réseau engendré par k vecteurs linéairement indépendants $\mathbf{u}_1, \dots, \mathbf{u}_k$. Le domaine fondamental de $\mathbb{L}$*

associé à cette base, noté $\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_k)$, est l'ensemble défini par :

$$\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_k) = \left\{ \sum_{i=1}^k x_i \mathbf{u}_i \mid \forall i \in \llbracket 1, k \rrbracket, 0 \leq x_i < 1 \right\} \quad (3.2)$$

Pour un réseau de rang plein, le domaine fondamental permet de définir un pavage régulier de $\mathbb{R}^D$, où chaque tuile est obtenue en translatant le domaine fondamental par un vecteur du réseau :

$$\mathbb{R}^D = \coprod_{\mathbf{g} \in \mathbb{L}} (\mathbf{g} + \mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D))$$

La Figure 3.1 illustre cette décomposition.

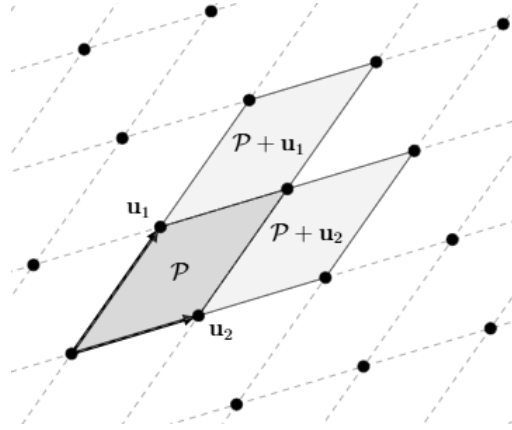


Figure 3.1: Domaine fondamental $\mathcal{P}$ d'un réseau 2D engendré par deux vecteurs $\{\mathbf{u}_1, \mathbf{u}_2\}$

Chaque réseau possède un invariant géométrique appelé covolume. En pratique, ce dernier s'avère particulièrement utile pour distinguer les réseaux entre eux.

Définition 3.1.4 (Covolume du réseau). Soit $\mathbb{L} = \bigoplus_{i=1}^k \mathbb{Z}\mathbf{u}_i$ un réseau engendré par k vecteurs linéairement indépendants de $\mathbb{R}^D$. On note $\langle \cdot, \cdot \rangle$ le produit scalaire usuel de $\mathbb{R}^D$.

Le déterminant de Gram $\Delta(\mathbf{u}_1, \dots, \mathbf{u}_k)$ est défini comme le déterminant de la matrice de Gram $(\langle \mathbf{u}_i, \mathbf{u}_j \rangle)_{1 \leq i, j \leq k}$ des produits scalaires des vecteurs $\mathbf{u}_1, \dots, \mathbf{u}_k$. Ce déterminant est strictement positif et ne dépend pas du choix de la base de $\mathbb{L}$.

Le covolume de $\mathbb{L}$, noté $\text{vol}(\mathbb{L})$, est alors défini comme la racine carrée du déterminant de Gram :

$$\text{vol}(\mathbb{L}) = \sqrt{\Delta(\mathbf{u}_1, \dots, \mathbf{u}_k)} > 0$$

Exemple 3.1.2. Le covolume de $\mathbb{Z}^D$ est 1 : $\text{vol}(\mathbb{Z}^D) = 1$.

Le covolume est un invariant géométrique que possèdent tous les réseaux, quel que soit leur rang. Pour les réseaux de rang plein, il existe également un invariant algébrique, appelé déterminant du réseau, qui coïncide avec le covolume.

Proposition 3.1.3 (Déterminant d'un réseau). *Soit $\mathbb{L} \subset \mathbb{R}^D$ un réseau de rang plein, muni d'une base $\mathbf{u}_1, \dots, \mathbf{u}_D$. Le déterminant de $\mathbb{L}$, noté $\det(\mathbb{L})$, est donné par la valeur absolue du déterminant de la matrice dont les colonnes sont formées par les vecteurs $\mathbf{u}_1, \dots, \mathbf{u}_D$:*

$$\det(\mathbb{L}) = |\det(\mathbf{u}_1, \dots, \mathbf{u}_D)|$$

Le déterminant de $\mathbb{L}$ ne dépend que du réseau lui-même et non du choix de sa base. De plus, il coïncide avec le covolume de $\mathbb{L}$ ainsi qu'avec le volume D -dimensionnel du domaine fondamental $\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D)$:

$$\det(\mathbb{L}) = \text{Leb}(\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D)) = \text{vol}(\mathbb{L})$$

Le covolume permet de différencier les réseaux : des covolumes distincts impliquent des réseaux différents. Cependant, dans le cas où leurs covolumes sont égaux, on ne peut généralement pas distinguer les réseaux sans l'ajout de conditions supplémentaires. Le théorème de décomposition des modules sur un anneau principal permet de déterminer des cas où deux réseaux de même covolume sont identiques.

Relations entre réseaux : le théorème des facteurs invariants

Théorème 3.1.4 (Facteurs invariants [2]). *Soient $\mathbb{L} \subset \mathbb{R}^D$ un réseau de rang r et $\mathbb{L}' \subseteq \mathbb{L}$ un sous-réseau de rang r' . Alors il existe une base $\{\epsilon_1, \dots, \epsilon_r\}$ de $\mathbb{L}$ et des entiers non nuls $q_1, \dots, q_{r'} \in \mathbb{Z} \setminus \{0\}$ tels que $\{q_1\epsilon_1, \dots, q_{r'}\epsilon_{r'}\}$ forme une base de $\mathbb{L}'$. Le groupe quotient $\mathbb{L}/\mathbb{L}'$ est donc isomorphe à $\mathbb{Z}^{r-r'} \times \prod_{i=1}^{r'} \mathbb{Z}/q_i\mathbb{Z}$.*

De plus, si $\mathbb{L}$ et $\mathbb{L}'$ ont le même rang (c'est-à-dire $r = r'$), alors :

$$|\mathbb{L}/\mathbb{L}'| = \prod_{i=1}^{r'} |q_i| = \text{vol}(\mathbb{L}')/\text{vol}(\mathbb{L})$$

En particulier, $\mathbb{L} = \mathbb{L}'$ si et seulement s'ils ont le même covolume.

À partir du théorème précédent, nous pouvons exprimer différemment le covolume des réseaux de rang plein inclus dans $\mathbb{Z}^D$.

Corollaire 3.1.1. *Le covolume d'un réseau de rang plein $\mathbb{L}$ qui est inclus dans $\mathbb{Z}^D$, est égal à la fois à l'ordre du groupe quotient $\mathbb{Z}^D/\mathbb{L}$ et au nombre de points à coordonnées entières contenus dans $\mathcal{P}$, le domaine fondamental associé à une base de $\mathbb{L}$:*

$$\text{vol}(\mathbb{L}) = |\mathbb{Z}^D/\mathbb{L}| = |\mathbb{Z}^D \cap \mathcal{P}|$$

3.2 Rappels sur les codes quantiques

3.2.1 Fondamentaux sur les qubits

Dans le monde de l'informatique quantique, le qubit est une représentation mathématique permettant de décrire l'état d'une particule quantique. Contrairement au bit qui ne peut valoir que 0 ou 1, un qubit peut exister dans une superposition de ces deux états simultanément.

Définition 3.2.1 (Qubits). *Soit $\mathbb{H}$ un espace de Hilbert complexe de dimension 2, de base orthonormale $|0\rangle, |1\rangle$. Un qubit $|\psi\rangle$ est un vecteur unitaire de $\mathbb{H}$:*

$$|\psi\rangle = \alpha |0\rangle + \beta |1\rangle \quad \text{avec } |\alpha|^2 + |\beta|^2 = 1.$$

Dans tout le reste du mémoire, nous prendrons $\mathbb{H} = \mathbb{C}^2$, avec les états de base $|0\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ et $|1\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$.

Transformations unitaires

Étant donné que les systèmes quantiques évoluent, les particules qui les composent subissent également des transformations. Pour les systèmes isolés (sans interactions avec l'extérieur), ces transformations se modélisent mathématiquement par l'action de matrices unitaires sur des qubits.

Définition 3.2.2 (Transformation unitaire sur un qubit). *Une transformation unitaire (ou une erreur) sur un qubit est une matrice U de taille 2×2 vérifiant $UU^\dagger = I_2$.*

Il est important de noter que l'application d'une matrice unitaire transforme un qubit en un autre qubit, garantissant ainsi que l'état de la particule est correctement modélisé à tout moment.

Exemple 3.2.1. $\mathbf{I}_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$, $\mathbf{X} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$, $\mathbf{Z} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$, $\mathbf{Y} = i\mathbf{XZ} = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}$
vérifient la relation $\mathbf{I}_2^2 = \mathbf{X}^2 = \mathbf{Y}^2 = \mathbf{Z}^2 = \mathbf{I}_2$

Le problème de la correction d'erreurs

Un problème courant en informatique quantique est celui de la correction d'erreurs sur un qubit. Ce problème consiste à retrouver l'état original $|\psi\rangle$ en ne connaissant que l'état corrompu $\mathbf{U}|\psi\rangle$, où $\mathbf{U} \in \mathbf{M}_2(\mathbb{C})$.

Pour résoudre ce problème, il n'est pas nécessaire de traiter chaque transformation individuellement. Il suffit de se concentrer sur la correction d'un ensemble restreint d'erreurs fondamentales, car toutes les autres peuvent être exprimées à partir de celles-ci.

Proposition 3.2.1. *Toute matrice complexe $\mathbf{A} \in \mathbf{M}_2(\mathbb{C})$ s'écrit comme combinaison linéaire complexe de $\mathbf{I}_2, \mathbf{X}, \mathbf{Y}, \mathbf{Z}$:*

$$\exists \alpha, \beta, \gamma, \delta \in \mathbb{C}, \quad \mathbf{A} = \alpha \mathbf{I}_2 + \beta \mathbf{X} + \gamma \mathbf{Y} + \delta \mathbf{Z}$$

Les matrices $\mathbf{X}, \mathbf{Y}, \mathbf{Z}$ sont appelées erreurs de Pauli de type x, y, z .

3.2.2 Le groupe de Pauli sur n-qubits

Pour modéliser mathématiquement l'état de n particules quantiques d'un système et leurs évolutions, on utilisera les produits de Kronecker de vecteurs ou de matrices unitaires.

Préliminaires d'algèbre linéaire : le produit de Kronecker

Définition 3.2.3 (Produit de Kronecker). *Le produit de Kronecker de deux matrices $\mathbf{A} = (a_{ij})$ de taille $m \times n$ et $\mathbf{B} = (b_{kl})$ de taille $p \times q$, est la matrice notée $\mathbf{A} \otimes \mathbf{B}$, de taille*

$mp \times nq$, définie par :

$$\mathbf{A} \otimes \mathbf{B} = \begin{pmatrix} a_{11}\mathbf{B} & \cdots & a_{1n}\mathbf{B} \\ \vdots & \ddots & \vdots \\ a_{m1}\mathbf{B} & \cdots & a_{mn}\mathbf{B} \end{pmatrix}$$

Proposition 3.2.2. *Le produit de Kronecker de matrices est bilinéaire, associatif et est compatible avec les opérations usuelles sur les matrices :*

- **Produit** : $(\mathbf{A} \otimes \mathbf{B})(\mathbf{C} \otimes \mathbf{D}) = (\mathbf{AC}) \otimes (\mathbf{BD})$
- **Inversion** : $(\mathbf{A} \otimes \mathbf{B})^{-1} = \mathbf{A}^{-1} \otimes \mathbf{B}^{-1}$
- **Transposition** : $(\mathbf{A} \otimes \mathbf{B})^{\top} = \mathbf{A}^{\top} \otimes \mathbf{B}^{\top}$

Systèmes à n particules

Pour décrire l'état de n particules quantiques au sein d'un système, nous utilisons les n -qubits. Il s'agit de vecteurs unitaires appartenant au $\mathbb{C}$ -espace vectoriel $(\mathbb{C}^2)^{\otimes n}$, engendré par l'ensemble des produits de Kronecker de n vecteurs de $\mathbb{C}^2$:

$$(\mathbb{C}^2)^{\otimes n} = \text{Vect}_{\mathbb{C}} \left(\{ \mathbf{v}_1 \otimes \cdots \otimes \mathbf{v}_n \mid \forall i \in \llbracket 1, n \rrbracket, \mathbf{v}_i \in \mathbb{C}^2 \} \right)$$

La base canonique de cet espace est constituée des vecteurs $|t_1\rangle \otimes \cdots \otimes |t_n\rangle$ où chaque $t_i \in \{0, 1\}$. En général, on préfère la notation simplifiée $|t_1 \dots t_n\rangle$ à $|t_1\rangle \otimes \cdots \otimes |t_n\rangle$.

Définition 3.2.4 (n -qubits). *Un n -qubit est un vecteur unitaire de l'espace $(\mathbb{C}^2)^{\otimes n}$. Une erreur sur un n -qubit est une transformation unitaire $\mathbf{U} \in \mathbb{U}_{2^n}(\mathbb{C})$.*

Exemple 3.2.2. $|0\rangle \otimes |1\rangle = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix}$ est un 2-qubit et $\mathbf{X} \otimes \mathbf{X} = \begin{pmatrix} 0 & \mathbf{X} \\ \mathbf{X} & 0 \end{pmatrix}$ est une erreur sur un 2-qubit.

Tout comme les transformations unitaires sur les qubits, celles agissant sur les n -qubits peuvent se décomposer en une combinaison linéaire d'opérateurs fondamentaux. En informatique quantique, la collection de ces opérateurs forme un groupe essentiel appelé groupe de Pauli.

Le groupe de Pauli sur n -qubits

Définition 3.2.5 (Groupe de Pauli). *Le groupe de Pauli sur n -qubits, noté $\mathbb{P}_n$, est l'ensemble de tous les opérateurs $\mathcal{O}$ de la forme $i^m \mathbf{O}_1 \otimes \cdots \otimes \mathbf{O}_n$ où $m \in \{0, 1, 2, 3\}$,*

représente un facteur de phase et où chaque $O_j \in \{I_2, X, Y, Z\}$ est l'une des quatre erreurs de base sur un qubit.

$$\mathbb{P}_n = \{\mathcal{O} = i^m O_1 \otimes \cdots \otimes O_n \mid m \in \llbracket 0, 3 \rrbracket, O_j \in \{I_2, X, Y, Z\}\}$$

Le poids d'un élément $\mathcal{O} = i^m O_1 \otimes \cdots \otimes O_n$ du groupe de Pauli, noté $\text{wt}(\mathcal{O})$, est le nombre d'opérateurs O_j qui sont différents de la matrice identité I_2 :

$$\text{wt}(\mathcal{O}) = |\{j \in \llbracket 1, n \rrbracket \mid O_j \neq I_2\}|$$

Représentation binaire des opérateurs

Tout élément $\mathcal{O}$ du groupe de Pauli $\mathbb{P}_n$ s'écrit de manière unique sous la forme :

$$\mathcal{O} = i^m X^u Z^v \quad \text{où } m \in \llbracket 0, 3 \rrbracket, u, v \in \{0, 1\}^n$$

et où les notations X^u et Z^v désignent :

$$X^u = X^{u_1} \otimes \cdots \otimes X^{u_n} \quad \text{et} \quad Z^v = Z^{v_1} \otimes \cdots \otimes Z^{v_n}$$

On représente alors $\mathcal{O}$ par le couple de vecteurs (u, v) .

Dans cette représentation, la multiplication de deux éléments du groupe de Pauli se traduit simplement par l'addition dans $\mathbb{F}_2$ des vecteurs correspondants.

3.2.3 Codes quantiques

Un code correcteur quantique $\mathcal{Q}$ est un sous-espace vectoriel de $(\mathbb{C}^2)^{\otimes n}$. À l'instar de leurs homologues classiques, les codes correcteurs quantiques ont été conçus pour protéger l'information quantique des erreurs auxquelles elle est constamment sujette.

Parmi les diverses approches de protection de l'information quantique, l'utilisation de codes correcteurs, dits stabilisateurs, occupe une place prépondérante. Ils ont été spécifiquement développés pour permettre la réalisation de calculs tolérants aux fautes. La particularité de ces codes réside dans le fait que les n -qubits encodés sont des points fixes pour un ensemble d'opérateurs formant un sous-groupe du groupe de Pauli. Cette propriété fondamentale simplifie considérablement la détection et la correction des erreurs, facilitant ainsi la protection de l'information encodée.

Codes stabilisateurs

Définition 3.2.6. *Un code stabilisateur $\mathcal{Q}$ de paramètres $[[n, k, d]]$ est un sous-espace vectoriel de $(\mathbb{C}^2)^{\otimes n}$, de dimension complexe 2^k , dont les éléments sont des points fixes pour un ensemble d'opérateurs $\mathbb{S}$ formant un sous-groupe commutatif du groupe de Pauli $\mathbb{P}_n$ et ne contenant pas l'opérateur $-\mathbf{I}_2^{\otimes n}$*

$$\mathcal{Q} = \{|\psi\rangle \in (\mathbb{C}^2)^{\otimes n} \mid \mathcal{S}|\psi\rangle = |\psi\rangle, \forall \mathcal{S} \in \mathbb{S}\}$$

Sa distance minimale d est définie comme le poids minimum d'un opérateur logique non trivial, c'est à dire le poids d'un opérateur de Pauli qui commute avec les éléments de $\mathbb{S}$, sans appartenir à $\mathbb{S}$:

$$d = \min\{\text{wt}(\mathcal{O}) \mid \mathcal{O} \in \mathbb{P}_n \setminus \mathbb{S} \text{ vérifiant } \mathcal{O}\mathcal{S} = \mathcal{S}\mathcal{O}, \forall \mathcal{S} \in \mathbb{S}\}$$

Ici n est appelée la longueur du code, k sa dimension et d sa distance minimale.

À l'image des codes correcteurs classiques, les codes stabilisateurs permettent systématiquement de corriger les erreurs survenant sur les n -qubits, à condition que leur poids soit relativement petit par rapport à la distance minimale.

Proposition 3.2.3. *Soit $\mathcal{Q}$ un code stabilisateur de paramètres $[[n, k, d]]$ associé à un groupe $\mathbb{S}$. Si un état encodé $|\psi\rangle$ est corrompu par un opérateur de Pauli $\mathcal{O}$ de poids $t < d/2$, alors on peut retrouver l'état original $|\psi\rangle$ à partir de la seule donnée de $\mathcal{O}|\psi\rangle$.*

Plus précisément, il est possible d'identifier un opérateur $\mathcal{O}' \in \mathbb{P}_n$ tel que le produit $\mathcal{O}\mathcal{O}'$ appartienne au groupe $\mathbb{S}$. Ainsi, en appliquant l'opérateur $\mathcal{O}'$ à l'état corrompu $\mathcal{O}|\psi\rangle$, on retrouve $|\psi\rangle$.

En tirant parti de la représentation binaire des opérateurs de Pauli, les codes stabilisateurs peuvent être décrits de façon plus compacte grâce à l'introduction de matrices génératrices.

Proposition 3.2.4. *Soit $\mathcal{Q}$ un code stabilisateur de paramètres $[[n, k, d]]$, associé à un groupe stabilisateur $\mathbb{S} = \langle \mathbf{G}_1, \dots, \mathbf{G}_{n-k} \rangle$, engendré par $n - k$ opérateurs de Pauli linéairement indépendants et commutant deux à deux.*

La matrice génératrice $\mathbf{H} = [\mathbf{A}_X \mid \mathbf{A}_Z]$ du code stabilisateur est la matrice dont les lignes sont données par les représentations binaires $(\mathbf{u}, \mathbf{v})$ des opérateurs $\mathbf{G}_j = i^m \mathbf{X}^{\mathbf{u}} \mathbf{Z}^{\mathbf{v}}$ qui

engendrent $\mathbb{S}$. La condition de commutativité des générateurs de $\mathbb{S}$ se traduit alors par la relation matricielle :

$$\mathbf{A}_X \mathbf{A}_Z^\top + \mathbf{A}_Z \mathbf{A}_X^\top = \mathbf{0} \pmod{2}$$

Les paramètres $[[n, k, d]]$ sont alors reliés à cette matrice $\mathbf{H}$ de la façon suivante :

- $\mathbf{H}$ est une matrice ayant $2n$ colonnes.
- La dimension du code est $k = n - \text{rang } \mathbf{H}$.
- La distance minimale d est définie comme le plus petit poids de Hamming symplectique d'un vecteur $\begin{pmatrix} u \\ v \end{pmatrix}$ qui appartient au noyau de $\mathbf{H}$ mais qui n'est pas dans l'espace vectoriel engendré par les lignes de $\mathbf{H}$:

$$d = \min\{\text{wt}_{\text{symp}}(\mathbf{u}, \mathbf{v}) \mid (\mathbf{u}, \mathbf{v}) \in \ker \mathbf{H} \setminus \text{Vect}_{\mathbf{L}} \mathbf{H}\}$$

où $\text{wt}_{\text{symp}}(\mathbf{u}, \mathbf{v}) = |\{i \in \{1, \dots, n\} \mid (u_i, v_i) \neq (0, 0)\}|$ et où $\text{Vect}_{\mathbf{L}} \mathbf{H}$ désigne l'espace vectoriel engendré par les lignes de $\mathbf{H}$.

Codes CSS

Les codes CSS, nommés d'après Calderbank, Shor et Steane [3], [4], forment une sous-classe importante des codes stabilisateurs. Leur particularité est que chaque stabilisateur est soit un produit d'opérateurs de type bit-flip ($\mathbf{X}$) uniquement, soit un produit d'opérateurs de type phase-flip ($\mathbf{Z}$) uniquement. Un code CSS se caractérise alors par une matrice génératrice de la forme :

$$\mathbf{H} = \begin{pmatrix} \mathbf{H}_X & 0 \\ 0 & \mathbf{H}_Z \end{pmatrix}$$

où $\mathbf{H}_X$ et $\mathbf{H}_Z$ sont deux matrices binaires dont les lignes sont orthogonales, c'est-à-dire $\mathbf{H}_X \cdot \mathbf{H}_Z^\top = \mathbf{0}$. Les paramètres du code sont alors définis comme suit :

- La longueur du code n (nombre de qubits physiques) est le nombre de colonnes de $\mathbf{H}_X$ (ou $\mathbf{H}_Z$).
- La dimension du code k (nombre de qubits logiques) est $k = n - \text{rang } \mathbf{H}_X - \text{rang } \mathbf{H}_Z$.
- La distance minimale d est $\min(d_X, d_Z)$ où d_X (resp. d_Z) est le plus petit poids de Hamming d'un mot de code de $\ker \mathbf{H}_X$ (resp. $\ker \mathbf{H}_Z$) n'appartenant pas à l'espace

engendré par les lignes de $\mathbf{H}_Z$ (resp $\mathbf{H}_X$).

Dans le cadre de cette thèse, notre attention s'est portée sur l'étude des codes stabilisateurs et, en particulier sur des codes CSS dont les sous-matrices génératrices $\mathbf{H}_X$ et $\mathbf{H}_Z$ ont une structure spécifique : $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$ où $\mathbf{A}, \mathbf{B}$ sont des matrices binaires qui commutent.

Le chapitre suivant sera consacré à la classification de codes de cette forme, où les matrices $\mathbf{A}$ et $\mathbf{B}$ sont des matrices circulantes, avec chacune deux éléments non-nuls par lignes.

4

Classification des GB(2,2) et comparaison avec les codes surfaces

Sommaire

4.1	Introduction	30
4.1.1	Motivation	30
4.1.2	Anciens résultats	31
4.1.3	État de l'art et contributions	31
4.1.4	Vue d'ensemble générale	34
4.1.5	Organisation du chapitre	35
4.2	Notions clés	36
4.2.1	Codes Bicycle généralisés (GB)	36
4.2.2	Code CSS associé au graphe de Cayley d'un groupe abélien	38
4.3	Borne inférieure sur la distance minimale des GB(2,2)	42
4.3.1	Propriétés des codes GB(2,2)	42
4.3.2	Interprétation des codes $GB(1 + X^a, 1 + X^b, n)$ à l'aide des graphes de Cayley	43
4.3.3	Borne inférieure sur la distance minimale des codes GB(2,2)	44
4.4	Preuve du théorème 4.3.2	45

4.4.1	Plan de la preuve :	45
4.4.2	Réinterprétation de la distance minimale	46
4.4.3	Des marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ vers les marches sur $\mathbb{Z}^2$. . .	48
4.4.4	Des marches sur $\mathbb{Z}^2$ vers les marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$. . .	52
4.4.5	Fin de la preuve du théorème 4.3.2	53
4.5	Application du théorème 4.3.2	56
4.5.1	Etude de $GB(1 + X, 1 + X^n, n^2)$	57
4.5.2	Etude de $GB(1 + X, 1 + X^{2r-1}, 2r^2)$	59
4.5.3	Etude de $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$	60
4.6	Relations d'équivalence sur les codes quantiques	63
4.6.1	Relation d'équivalence générale	63
4.6.2	Relation d'équivalence CGP	64
4.7	Comparaison des codes GB aux codes surface 2D de poids quatre .	66
4.7.1	Théorème de Whitney	67
4.7.2	Vue d'ensemble de la démonstration de la non-équivalence .	68
4.7.3	Comparaison avec le code surface 2D optimal pour la distance paire	69
4.7.4	Comparaison avec les codes de Kitaev standards	84
4.7.5	Comparaison avec les codes surface 2D optimaux pour la distance impaire	86
4.8	Classification des codes GB(2,2)	94
4.8.1	Méthode :	94
4.8.2	Classification des codes GB(2,2) extrémaux	95
4.8.3	Classification de tous les codes GB(2,2) (Longueur ≤ 200) . .	101
4.8.4	Résumé des meilleurs codes surface 2D à stabilisateurs de poids quatre	101
4.9	Conclusion du chapitre	105

4.1 Introduction

4.1.1 Motivation

L'informatique quantique promet de transformer en profondeur nos capacités de calcul actuelles, en rendant possible, en des temps raisonnables, la résolution de certains problèmes réputés intraitables par les ordinateurs classiques [1]. Cependant, cette nouvelle approche repose sur des principes physiques qui la rendent particulièrement vulnérable.

L'information quantique est, par nature, extrêmement fragile : le bruit environnemental et le phénomène de décohérence (entraînant la perte des propriétés quantiques des particules) exposent les ordinateurs quantiques à un risque permanent d'erreurs durant les opérations. À ces erreurs s'ajoutent les imprécisions liées aux mesures, pouvant mener à des résultats incorrects. Pour que les résultats de longues séquences de calculs quantiques soient fiables, il est donc indispensable de mettre en place des mécanismes de correction d'erreurs.

La correction d'erreurs quantique repose sur l'utilisation de codes quantiques, de circuits de mesure de syndrome pour détecter les erreurs et d'algorithmes de décodage chargés de les corriger dès qu'elles surviennent. Puisque des erreurs peuvent également se produire lors de la mesure ou lors du décodage, le processus de correction doit être à la fois rapide et précis. Cela permet de répéter régulièrement les cycles de correction et ainsi limiter la propagation des erreurs.

Pour limiter la propagation des erreurs et obtenir des résultats fiables grâce aux mécanismes de correction d'erreurs, il faut réduire le taux d'erreur physique — c'est-à-dire la probabilité qu'il y ait une erreur sur un qubit physique — en dessous d'un seuil critique [5].

4.1.2 Anciens résultats

Les **codes de Kitaev**, introduits en 2003, constituent les premiers codes **LDPC quantiques** (qLDPC) capables de protéger deux qubits logiques, tout en possédant une distance minimale croissant proportionnellement avec la racine carrée de la longueur du code [6]. Ce sont des codes stabilisateurs locaux : chaque générateur n'agit que sur un nombre fini de qubits et chaque qubit n'est affecté que par un nombre fini de générateurs.

Depuis leur introduction, de nombreuses généralisations ont été proposées dans le but d'améliorer leurs paramètres. Parmi les plus notables, on trouve :

- les **codes de Kitaev optimisés**, construits via des quotients du plan euclidien par des réseaux [7], offrant une meilleure distance minimale ;
- les **codes surface**, présentant de meilleurs taux.

Toutefois, aucun code issu de ces généralisations n'a de distance minimale dont l'ordre de croissance excède la racine carrée de la longueur du code.

Jusqu'à récemment, tous les codes qLDPC connus étaient contraints par cette limitation. Néanmoins, d'importants progrès ont été réalisés ces dernières années, avec la démonstration théorique de l'existence de codes qLDPC asymptotiquement bons [8], [9], dont la dimension et la distance minimale croissent linéairement avec la longueur du code.

Bien que ces constructions offrent un avantage théorique considérable en matière de correction d'erreurs, leur mise en œuvre pratique reste complexe, contrairement aux codes surface, qui bénéficient à la fois de seuils d'erreur élevés et d'algorithmes de décodage efficaces.

4.1.3 État de l'art et contributions

Les codes **Bicycle généralisés** (GB), notés $GB(a, b)$, sont construits à partir de paires de matrices circulantes ayant respectivement a et b éléments non nuls par ligne. Ils se distinguent comme l'une des alternatives les plus prometteuses aux codes surface pour une implémentation pratique. Une avancée récente de l'équipe de recherche d'IBM [10] a mis en évidence leur potentiel en introduisant une famille de codes GB(3,3) atteignant des seuils d'erreur proches de ceux des codes surface, tout en offrant de meilleurs taux.

Parmi ces codes, la classe des codes GB(2,2) qui généralisent les codes de Kitaev et leurs variantes optimisées [11], présente un intérêt particulier. Ces codes sont particulièrement bien adaptés pour l'implémentation physique grâce à leurs stabilisateurs

locaux de poids quatre. De plus, on peut les représenter géométriquement en plaçant leurs qubits sur les sommets d'un réseau bidimensionnel.

En se basant sur l'intuition de Haah, Hastings et O'Donnell [12], selon laquelle une modification de la géométrie du code torique pourrait permettre d'obtenir des codes avec des distances minimales supérieures à celles des codes de Kitaev, nous proposons une nouvelle approche pour construire des codes GB(2,2) performants.

Notre méthode repose sur une approche combinant l'algèbre et la théorie des graphes. Nous interprétons les codes GB(2,2) comme des codes CSS, dont les sous-matrices génératrices sont des matrices d'incidence de graphes de Cayley à deux générateurs définis sur des groupes abéliens. Spécifiquement, les codes GB(2,2) correspondent à des graphes de Cayley non orientés de la forme $(\mathbb{Z}/n\mathbb{Z}, \bar{a}, \bar{b})$.

Dans cette perspective, nous établissons une borne inférieure sur la distance minimale des codes étudiés en nous appuyant sur la plus petite norme L^1 des vecteurs non nuls de certains sous-réseaux de $\mathbb{Z}^2$ associés.

Cette borne s'est avérée indispensable, permettant la construction explicite de trois familles infinies de codes surface 2D, à stabilisateurs de poids quatre, optimaux :

- GB($1 + X, 1 + X^n, n^2$) de paramètres $[[2n^2, 2, n]]$
- GB($1 + X, 1 + X^{2r-1}, 2r^2$) de paramètres $[[4r^2, 2, 2r]]$
- GB($1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2$) de paramètres $[(2t + 1)^2 + 1, 2, 2t + 1]$

Kovalev et Pryadko [11] avaient identifié une famille de codes GB(2,2) de paramètres $[(2t + 1)^2 + 1, 2, 2t + 1]$, qui égalait les performances des meilleurs codes surface 2D à stabilisateurs de poids quatre pour les distances impaires. Cependant, la littérature était restée largement dépourvue de codes GB(2,2) dont la distance minimale croissait linéairement avec la racine carrée de la longueur du code, en particulier pour les distances paires. Notre construction de la famille $[[4r^2, 2, 2r]]$ a résolu ce problème en démontrant que les codes GB(2,2) pouvaient également égaler les performances des codes surface 2D à stabilisateurs de poids quatre optimaux pour les distances paires. Ce résultat majeur a remis en question une ancienne conjecture affirmant qu'il était impossible d'atteindre de telles performances avec des codes GB(2,2). Auparavant, les meilleurs paramètres connus pour des codes GB(2,2) à distance paire étaient de $[[4r^2 + 4, 2, 2r]]$, uniquement réalisables pour des valeurs paires de r .

En plus d'avoir de bons paramètres, les codes quantiques doivent également posséder de bonnes propriétés de décodage pour être utiles en pratique. En effet, deux

codes partageant les mêmes paramètres $[[n, k, d]]$ peuvent offrir des performances de décodage très différentes selon leur structure. À l’instar des problèmes équivalents que l’on regroupe en théorie de la complexité, les codes quantiques présentant des propriétés similaires en termes de correction d’erreurs sont regroupés dans des classes d’équivalence. Pour catégoriser et étudier efficacement ces codes, la recherche de notions d’équivalence adaptées revêt alors une importance capitale.

La relation d’équivalence usuelle sur les codes correcteurs stabilisateurs, selon laquelle deux codes $\mathcal{Q}$ et $\mathcal{Q}'$ sont équivalents s’il existe un opérateur unitaire U , s’écrivant comme un produit de matrices de permutation et de transformations locales de Clifford, tel que $\mathcal{Q}' = U\mathcal{Q}$, n’est pas adaptée à notre étude. Cette équivalence ne préserve ni la structure CSS, ni celle du graphe sous-jacent des codes CSS dérivés de graphes de Cayley que nous examinons. Or, préserver ces structures s’avère crucial pour les codes GB(2,2) car certains algorithmes de décodage établis (comme l’algorithme de renormalisation [13]) s’appuient sur les propriétés de ces dernières.

Afin d’évaluer l’originalité de nos codes, selon une relation qui préserve ces propriétés structurelles, nous introduisons une nouvelle notion d’équivalence spécifiquement adaptée.

Dans ce cadre, nous démontrons que les deux premières familles de codes que nous construisons ne sont équivalentes, ni aux codes de Kitaev standards, ni aux codes surface 2D optimaux à distance paire avec stabilisateurs de poids quatre (voir Propositions 4.7.4 et 4.7.7). Des simulations préliminaires indiquent même que les polynômes générateurs de poids de nos codes sont différents de ceux des codes de surface optimaux à distance paire et des codes de Kitaev. Cela suggère, de manière intuitive, que nos codes ne leur sont pas équivalents, même au sens usuel.

En revanche, notre troisième famille, $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$, constitue une nouvelle réalisation du code surface 2D de poids quatre, à distance impaire, optimal.

Ce travail s’inscrit dans la continuité de l’étude amorcée par Pryadko et Wang [14]. Dans cette étude, ils ont analysé les performances des codes GB en fonction du poids de leurs stabilisateurs, notamment ceux de la forme $GB(1 + X, 1 + X^a, n)$, de longueur $2n$ où n est un nombre premier choisi pour que 2 soit une racine primitive modulo n . Nous avons approfondi cette direction en menant une étude complète de tous les codes GB(2,2) de longueur inférieure ou égale à 200.

Enfin, nous avons classifié dans la table 4.1 les meilleurs codes GB(2,2) non équivalents, de longueur inférieure ou égale à 200. En outre, une table comparative des meilleurs

codes surface 2D de poids quatre connus est également incluse (table 4.2). L'ensemble constitue une base de données précieuse pour de futures applications pratiques, en fournissant des exemples de codes hautement structurés présentant de très bonnes propriétés de décodage.

4.1.4 Vue d'ensemble générale

Cette section détaille notre approche pour développer de nouveaux codes GB(2,2) optimaux en termes de distance minimale et pour justifier leur originalité. Nous détaillerons également les plans de la démonstration que nous avons suivis pour construire ces codes et pour prouver leur distinction par rapport aux codes surface 2D de poids quatre notables.

Construction

Notre méthode de construction des codes GB(2,2) performants exploite à la fois les propriétés algébriques et géométriques, induites par les structures de leurs graphes de Cayley sous-jacents. Cette approche se décline en quatre étapes clés :

- **Réinterprétation en termes de graphes de Cayley** : Nous interprétons les codes GB(2,2) comme des codes CSS dérivés de graphes de Cayley non orientés, à deux générateurs, de la forme : $(\mathbb{Z}/n\mathbb{Z}, \bar{a}, \bar{b})$. La distance minimale du code correspond alors à la longueur d'un certain type de cycle du graphe.
- **Association à un réseau** : Nous associons chacun de ces codes à un réseau 2D $\mathbb{L}$ et définissons une application qui envoie les cycles du graphe sur des éléments de ce réseau.
- **Minoration de la distance minimale** : Nous prouvons que la distance minimale du code est minorée par la plus petite norme L^1 d'un vecteur non nul du réseau associé.
- **Conception des codes** : En concevant des réseaux $\mathbb{L}$ dont les vecteurs non nuls possèdent une grande norme L^1 , nous construisons des codes GB(2,2) dotés de grandes distances minimales.

Comparaison

Afin de confirmer l'originalité de nos codes, nous avons établi un cadre de comparaison rigoureux. Plus précisément, nous avons ramené la question de leur équivalence avec les codes de surface 2D optimaux de distance paire, ou avec les codes de Kitaev, à celle de l'existence d'un isomorphisme de groupes. Notre démonstration a suivi le plan suivant :

- **Réinterprétation en termes de graphes de Cayley** : Nous interprétons les codes de surface 2D optimaux et les codes de Kitaev comme des codes CSS dérivés de graphes de Cayley non orientés de groupes abéliens.
- **Définition d'une relation d'équivalence** : Nous introduisons une relation d'équivalence spécifique, qui préserve à la fois la structure CSS et celle du graphe sous-jacent.
- **Etablissement d'un isomorphisme de graphes** : Nous démontrons que l'équivalence de nos codes GB optimaux avec les codes de surface optimaux à distance paire, ou avec les codes de Kitaev, impliquerait nécessairement qu'il existe un isomorphisme entre les graphes sous-jacents.
- **Implication sur les groupes abéliens** : Nous prouvons ensuite que cet isomorphisme de graphes induit nécessairement un isomorphisme entre les groupes abéliens correspondants.
- **Preuve de la distinction** : Enfin, nous montrons que les groupes abéliens sous-jacents à nos nouveaux codes ne sont pas isomorphes à ceux des codes surface 2D en question. Cela permet d'établir clairement la non-équivalence des codes, démontrant ainsi l'originalité de nos constructions.

4.1.5 Organisation du chapitre

Nous débutons par la Section 4.2 pour aborder les concepts clés des codes GB et des codes CSS dérivés des graphes de Cayley. Les Sections 4.3 et 4.4 présenteront ensuite la borne inférieure sur la distance minimale des codes GB(2,2). Puis, la Section 4.5 décrira explicitement les constructions des familles de codes GB(2,2) optimaux, dont les performances sont comparables à celles des meilleurs codes surfaces 2D à stabilisateurs de poids quatre.

La Section 4.6 introduit les notions sur les relations d'équivalence entre codes quantiques et établit le cadre méthodologique utilisé pour comparer nos codes aux codes surface de poids quatre optimaux. Ensuite, dans la Section 4.7, nous effectuons les comparaisons entre nos codes GB et ces codes surface.

Enfin, nous concluons en présentant, dans la Section 4.8, les résultats de notre classification exhaustive des codes GB(2,2), ainsi qu'une table récapitulant tous les meilleurs codes surface 2D de poids quatre connus.

4.2 Notions clés

4.2.1 Codes Bicycle généralisés (GB)

Matrices circulantes

Soit un entier $n \geq 1$ et un polynôme $A(X) = \sum_{i=0}^{n-1} a_i X^i \in \mathbb{F}_2[X]_{\leq n-1}$. La matrice circulante $Circ(A(X), n)$ est définie par :

$$Circ(A(X), n) = \begin{pmatrix} a_0 & a_{n-1} & \dots & a_1 \\ a_1 & a_0 & \ddots & \vdots \\ \vdots & \ddots & \ddots & a_{n-1} \\ a_{n-1} & \dots & a_1 & a_0 \end{pmatrix}.$$

L'ensemble des matrices circulantes $n \times n$ à coefficients dans $\mathbb{F}_2$ forme une algèbre, qui est isomorphe à celle des polynômes dans $\mathbb{F}_2[X]/(X^n - 1)$. Cet isomorphisme découle du fait que toute matrice circulante peut s'écrire comme un polynôme en la matrice $Circ(X, n)$.

Une première conséquence est que le produit de deux matrices circulantes est commutatif. Une autre est que le produit matriciel $Circ(A(X), n) \mathbf{b}$ de $Circ(A(X), n)$ avec un

vecteur $\mathbf{b} = \begin{pmatrix} b_0 \\ \vdots \\ b_{n-1} \end{pmatrix} \in \mathbb{F}_2^n$ peut être directement interprété comme une multiplication

de deux polynômes. Plus précisément, cette opération revient à calculer le produit des polynômes associés dans l'algèbre $\mathbb{F}_2[X]/(X^n - 1)$:

$$\left(\sum_{i=0}^{n-1} a_i X^i \right) \cdot \left(\sum_{i=0}^{n-1} b_i X^i \right) \mod X^n - 1$$

Codes Bicycle généralisés (GB)

Introduits par Pryadko et Kovalev en 2013 [11], les codes Bicycle généralisés (GB) étendent le concept des codes Bicycle de MacKay et al. [15]. Tandis que les codes Bicycle traditionnels sont des codes CSS dont les sous-matrices génératrices sont construites à partir d'une unique matrice circulante ($\mathbf{H}_X = \mathbf{H}_Z = [\mathbf{A} \mid \mathbf{A}^\top]$), les sous-matrices génératrices des GB utilisent deux matrices circulantes.

Un code GB est un code CSS associé à un entier $n \geq 2$ et à deux polynômes binaires $A(X), B(X) \in \mathbb{F}_2[X]_{\leq n-1}$. Ses sous-matrices génératrices sont données par $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$, où $\mathbf{A} = \text{Circ}(A(X), n)$ et $\mathbf{B} = \text{Circ}(B(X), n)$.

Ce code, noté $GB(A(X), B(X), n)$, est de longueur $2n$. Sa dimension est donnée par $2 \deg(\text{pgcd}(A(X), B(X), X^n - 1))$ où pgcd désigne le plus grand diviseur commun. Enfin, étant donné que $\mathbf{H}_Z$ peut être obtenue à partir de $\mathbf{H}_X$ en permutant des lignes et des colonnes, la distance minimale du code satisfait $d = d_X = d_Z$ [16].

L'isomorphisme entre l'algèbre des matrices circulantes et l'algèbre $\mathbb{F}_2[X]/(X^n - 1)$ permet d'interpréter la notion de distance minimale des codes GB à l'aide de polynômes.

Réinterprétation de la distance minimale des codes GB

Considérons $GB(A(X), B(X), n)$, le code GB associé à $A(X), B(X) \in \mathbb{F}_2[X]_{\leq n-1}$, dont les sous-matrices génératrices sont définies par : $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$, où $\mathbf{A} = \text{Circ}(A(X), n)$ et $\mathbf{B} = \text{Circ}(B(X), n)$.

La distance minimale de ce code, et plus généralement les propriétés de ses sous-matrices génératrices, peuvent être définies à l'aide de polynômes.

Soit un vecteur colonne $\mathbf{c} = \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} \in \mathbb{F}_2^{2n}$, avec $U(X)$ et $V(X) \in \mathbb{F}_2[X]_{\leq n-1}$ les polynômes associés respectivement à $\mathbf{u}$ et $\mathbf{v}$. Nous avons alors les équivalences suivantes :

- $\mathbf{c}$ appartient au noyau de $\mathbf{H}_X$ si et seulement si ses polynômes associés vérifient l'égalité :

$$A(X)U(X) + B(X)V(X) \equiv 0 \pmod{X^n - 1}$$

- c appartient au sous-espace vectoriel engendré par les lignes de H_Z si et seulement s'il existe un polynôme $H(X) \in \mathbb{F}_2[X]$ tel que :

$$\begin{cases} U(X) \equiv B(X)H(X) \mod X^n - 1 \\ V(X) \equiv A(X)H(X) \mod X^n - 1 \end{cases}$$

Dans leur étude de 2008 sur les codes homologiques, Bombin et Delgado [7] ont introduit une famille de codes surface optimaux qui surpassent significativement les codes toriques d'origine. Bien que leurs constructions diffèrent, les codes surface de Bombin et Delgado et les codes GB(2,2) peuvent être décrits à travers un point de vue unifié : celui des codes CSS dérivés de graphes de Cayley de groupes abéliens.

Dans la section suivante, nous commencerons par présenter cette notion, puis nous expliquerons comment les codes optimaux et les codes GB(2,2) sont construits avec cette approche.

4.2.2 Code CSS associé au graphe de Cayley d'un groupe abélien

Graphes de Cayley à deux générateurs d'un groupe abélien

Définition 4.2.1. Soit $\mathbb{G} \neq \{0\}$ un groupe abélien non-trivial et soient g_1, g_2 deux éléments non nuls de $\mathbb{G}$. Le graphe de Cayley associé à $\mathbb{G}, g_1$ et g_2 est le graphe, noté $(\mathbb{G}, g_1, g_2)$, dont les sommets sont les éléments de $\mathbb{G}$ et où chaque sommet $g \in \mathbb{G}$ est relié exactement à quatre voisins : $g + g_1, g - g_1, g + g_2$ et $g - g_2$. Ce nombre de quatre connexions est conservé même si elles mènent vers le même sommet.

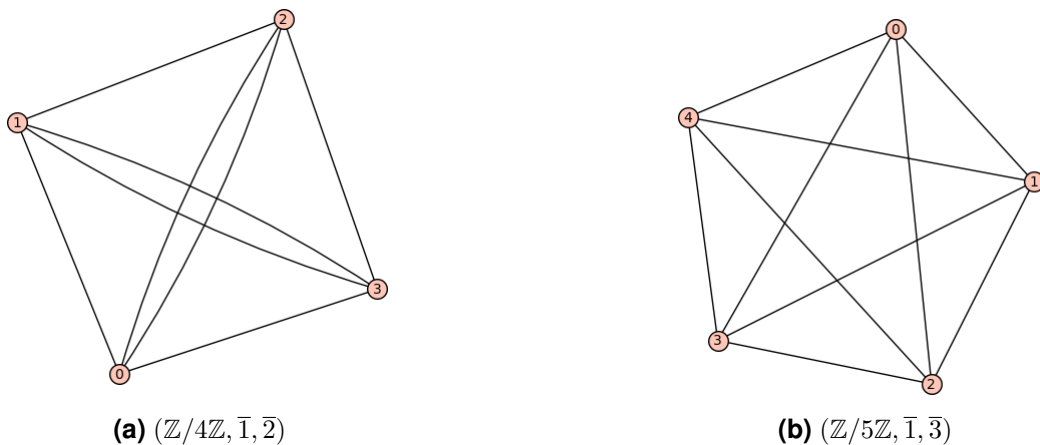


Figure 4.1: Exemples de graphes de Cayley à deux générateurs sur des groupes cycliques

Matrices d'incidence associées à $(\mathbb{G}, g_1, g_2)$

Considérons $(\mathbb{G}, g_1, g_2)$ le graphe de Cayley associé à un groupe abélien non trivial $\mathbb{G} \neq \{0\}$ et à g_1, g_2 deux éléments non nuls de $\mathbb{G}$.

De façon analogue aux constructions de codes de Kitaev, nous définissons sur ce graphe les notions de cocycles et de faces.

Définition 4.2.2 (Cocycles et faces de $(\mathbb{G}, g_1, g_2)$). *Pour tout sommet $g \in \mathbb{G}$, le cocycle et la face associés à g sont définis de la façon suivante :*

- Son **cocycle** est le multi-ensemble constitué des arêtes du graphe $(\mathbb{G}, g_1, g_2)$ reliant g à ses quatre voisins :

$$\{\{g, g + g_1\}, \{g, g - g_1\}, \{g, g + g_2\}, \{g, g - g_2\}\}$$

- Sa **face** $\mathcal{F}_g$ est le multi-ensemble formé par les arêtes du 4-cycle de $(\mathbb{G}, g_1, g_2)$ reliant successivement les sommets $g, g + g_1, g + g_1 + g_2$ et $g + g_2$:

$$\{\{g, g + g_1\}, \{g + g_1, g + g_1 + g_2\}, \{g + g_1 + g_2, g + g_2\}, \{g + g_2, g\}\}$$

Remarque 4.2.1. Dans les définitions des cocycles et des faces, on considère les arêtes comme distinctes même si elles aboutissent au même sommet. Ainsi, chaque multi-ensemble contient toujours quatre éléments.

On fixe un ordre (c'est-à-dire une numérotation) sur les sommets, les arêtes et les faces du graphe. À l'instar des codes de Kitaev [6], les matrices génératrices du code CSS associé à $(\mathbb{G}, g_1, g_2)$ sont construites à partir d'une paire de matrices d'incidence sur ce graphe.

Définition 4.2.3 (Matrices d'incidence associées à $(\mathbb{G}, g_1, g_2)$). *On définit sur $\mathbb{F}_2$ les matrices d'incidence suivantes :*

- $\mathbf{H}_X(\mathbb{G}, g_1, g_2)$: la matrice d'incidence sommet-arête du graphe $(\mathbb{G}, g_1, g_2)$.
- $\mathbf{H}_Z(\mathbb{G}, g_1, g_2)$: la matrice d'incidence face-arête du graphe $(\mathbb{G}, g_1, g_2)$.

dont les lignes sont respectivement indexées par les sommets g et les faces $\mathcal{F}_g$ du graphe $(\mathbb{G}, g_1, g_2)$, tandis que les colonnes sont indexées par les arêtes du graphe.

Construction de $CSS(\mathbb{G}, g_1, g_2)$

Une propriété essentielle du graphe $(\mathbb{G}, g_1, g_2)$ est que chaque face $\mathcal{F}$ forme un cycle, c'est-à-dire un multi-ensemble composé d'arêtes, tel que pour chaque sommet du graphe, il y a un nombre pair d'arêtes appartenant à $\mathcal{F}$ qui sont incidentes à ce sommet. Cette caractéristique conduit directement à la relation suivante entre les deux matrices d'incidence :

$$\mathbf{H}_X(\mathbb{G}, g_1, g_2) \cdot \mathbf{H}_Z(\mathbb{G}, g_1, g_2)^\top = \mathbf{0}$$

Définition 4.2.4. *Le code $CSS(\mathbb{G}, g_1, g_2)$ associé au graphe $(\mathbb{G}, g_1, g_2)$ se définit alors comme le code CSS dont les sous-matrices génératrices sont $\mathbf{H}_X(\mathbb{G}, g_1, g_2)$ et $\mathbf{H}_Z(\mathbb{G}, g_1, g_2)$.*

Remarque 4.2.2. *En numérotant différemment les sommets, les arêtes et les faces du graphe $(\mathbb{G}, g_1, g_2)$, nous aurions trouvé une paire de matrices d'incidence $(\mathbf{H}'_X, \mathbf{H}'_Z)$ différente. Toutefois, les matrices d'incidence sommet-arête $\mathbf{H}_X, \mathbf{H}'_X$ et face-arête $(\mathbf{H}_Z, \mathbf{H}'_Z)$ vérifieraient les relations suivantes :*

$$\mathbf{H}'_X = \mathbf{P}\mathbf{H}_X\mathbf{Q} \quad \text{et} \quad \mathbf{H}'_Z = \mathbf{R}\mathbf{H}_Z\mathbf{Q}.$$

où $\mathbf{P}, \mathbf{Q}, \mathbf{R}$ sont des matrices de permutation.

En conséquence, en dépit du changement d'ordre, les codes $CSS(\mathbf{H}_X, \mathbf{H}_Z)$ et $CSS(\mathbf{H}'_X, \mathbf{H}'_Z)$ ont les mêmes paramètres.

Interprétation des codes surface optimaux

Via la construction de codes homologiques issue de la topologie, Bombin et Delgado ont développé des codes surface de poids quatre optimaux [7]. Ces codes se caractérisent par des stabilisateurs locaux de poids quatre et des qubits positionnés sur les points à coordonnées entières du domaine fondamental d'un sous-réseau de $\mathbb{Z}^2$. Bien que les qubits du support d'un générateur du code stabilisateur ne soient pas nécessairement physiquement proches dans le plan, ils le deviennent lorsque l'on identifie les côtés opposés du domaine fondamental.

En utilisant le formalisme des codes CSS dérivés de graphes de Cayley, les codes surface 2D optimaux, de poids quatre, se décrivent de la manière suivante :

- $[(2t+1)^2+1, 2, 2t+1]$: Le code surface 2D optimal pour la distance $2t+1$, où $t \geq 1$, peut être exprimé comme :

$$CSS \left(\mathbb{Z}^2/\mathbb{L}_t, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}_t, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}_t \right) \text{ où } \mathbb{L}_t = \mathbb{Z} \begin{pmatrix} t+1 \\ t \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} -t \\ t+1 \end{pmatrix}$$

- $[4r^2, 2, 2r]$: Le code surface 2D optimal pour la distance $2r$, où $r \geq 1$, peut être exprimé comme :

$$CSS \left(\mathbb{Z}^2/\mathbb{L}_r, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}_r, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}_r \right) \text{ où } \mathbb{L}_r = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ -r \end{pmatrix}$$

Remarque 4.2.3. Avec ce même formalisme, les codes de Kitaev de paramètres $[[2n^2, 2, 2n]]$ correspondent aux CSS suivants :

$$CSS \left(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L} \right) \text{ où } \mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 0 \\ n \end{pmatrix}$$

Maintenant que les notions fondamentales requises pour ce chapitre ont été introduites, nous pouvons aborder l'étude des codes GB(2,2).

Un code GB(2,2) est un code de la forme $GB(A(X), B(X), n)$ où $A(X)$ et $B(X)$ sont deux polynômes à coefficients dans $\mathbb{F}_2$, ayant chacun exactement deux coefficients non nuls.

Notre objectif principal est d'établir une minoration de la distance minimale de ces codes et de classer ceux qui ont les meilleurs paramètres. Pour ce faire, ce chapitre est organisé en trois parties :

- Dans la première partie, nous nous concentrons sur l'étude de la distance minimale et nous établissons une borne inférieure pour une famille spécifique de codes GB(2,2).
- La deuxième partie exploite cette borne pour construire explicitement des familles de codes GB(2,2) dotés d'une grande distance minimale.
- Enfin, dans la dernière partie, nous établissons la classification de ces codes.

4.3 Borne inférieure sur la distance minimale des GB(2,2)

Pour l'étude de la distance minimale, il suffit de se concentrer sur certaines familles de codes GB(2,2). La proposition suivante précise les familles sur lesquelles nous pouvons nous focaliser.

4.3.1 Propriétés des codes GB(2,2)

Proposition 4.3.1. Soient $n > 0$ un entier et $A(X), B(X) \in \mathbb{F}_2[X]_{\leq n-1}$ deux polynômes.

- Si k est un entier premier avec n , alors les codes $GB(A(X), B(X), n)$ et $GB(A(X^k) \bmod X^n - 1, B(X^k) \bmod X^n - 1, n)$ ont les mêmes paramètres [16]. En particulier, pour des entiers $r, s \geq 1$, si r est premier avec n , alors

$$d(GB(1 + X^r, 1 + X^s, n)) = d(GB(1 + X, 1 + X^\alpha, n))$$

où $\alpha = st \bmod n$ avec $t = (r \bmod n)^{-1}$.

- Pour tous entiers naturels i, j , les codes $GB(A(X), B(X), n)$ et $GB(A(X)X^i \bmod X^n - 1, B(X)X^j \bmod X^n - 1, n)$ ont les mêmes paramètres.

Preuve : Pour le premier point, se référer à [16].

Pour le second point, considérons $R_i(X)$ et $S_j(X)$ les restes de $A(X)X^i$ et $B(X)X^j$ modulo $X^n - 1$. On pose : $\mathbf{A} = \text{Circ}(A(X), n)$, $\mathbf{B} = \text{Circ}(B(X), n)$, $\mathbf{C} = \text{Circ}(X, n)$ et on considère l'application h définie par :

$$h : \mathbb{F}_2^{2n} \rightarrow \mathbb{F}_2^{2n} \\ \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} \mapsto \begin{pmatrix} \mathbf{C}^i \mathbf{u} \\ \mathbf{C}^j \mathbf{v} \end{pmatrix}$$

Cette application induit des isométries (pour le poids de Hamming) entre :

- les noyaux à droite de $[\mathbf{A} \mid \mathbf{B}]$ et de $[\mathbf{C}^i \mathbf{A} \mid \mathbf{C}^j \mathbf{B}]$
- les espaces vectoriels engendrés par les lignes de $[\mathbf{B}^\top \mid \mathbf{A}^\top]$ et de $[(\mathbf{C}^j \mathbf{B})^\top \mid (\mathbf{C}^i \mathbf{A})^\top]$.

Puisque $\text{Circ}(R_i(X), n) = \mathbf{C}^i \mathbf{A}$ et $\text{Circ}(S_j(X), n) = \mathbf{C}^j \mathbf{B}$, les deux codes $GB(A(X), B(X), n)$ et $GB(R_i(X), S_j(X), n)$ ont les mêmes paramètres. $\square$

Une conséquence directe de cette proposition est que, pour l'étude de la distance minimale des codes GB(2,2), nous pouvons nous limiter aux codes de la forme $GB(1 + X^a, 1 + X^b, n)$ où a et b sont des entiers compris entre 1 et $\frac{n}{2}$.

Dans la section suivante, nous expliquerons comment chacun de ces codes peut être associé à un sous-réseau de $\mathbb{Z}^2$ et être exprimé comme un code CSS dérivé d'un graphe de Cayley.

4.3.2 Interprétation des codes $GB(1 + X^a, 1 + X^b, n)$ à l'aide des graphes de Cayley

Graphe associé à $GB(1 + X^a, 1 + X^b, n)$

Les matrices génératrices du code $GB(1 + X^a, 1 + X^b, n)$ peuvent être directement interprétées dans le cadre de la théorie des graphes, à l'aide des écarts a et b entre les positions des coefficients non nuls des polynômes associés.

Les sous-matrices génératrices du code, $\mathbf{H}_X$ et $\mathbf{H}_Z$, qui sont formées à partir des matrices circulantes $\mathbf{A} = \text{Circ}(1 + X^a, n)$ et $\mathbf{B} = \text{Circ}(1 + X^b, n)$, se réinterprètent alors de la façon suivante :

- $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$: matrice d'incidence sommet-arête du graphe de Cayley $(\mathbb{Z}/n\mathbb{Z}, \bar{a}, \bar{b})$.
- $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$: matrice d'incidence face-arête de ce même graphe.

Les i -ièmes lignes des matrices $\mathbf{H}_X$ et $\mathbf{H}_Z$ sont indexées respectivement par l'élément $\bar{i} \in \mathbb{Z}/n\mathbb{Z}$ et par la face $\mathcal{F}_i$ qui est définie par :

$$\mathcal{F}_i = \left\{ \{\bar{i}, \bar{i} + a\}, \{\bar{i} + a, \bar{i} + a + b\}, \{\bar{i} + a + b, \bar{i} + b\}, \{\bar{i} + b, \bar{i}\} \right\}$$

Pour $0 \leq k < n$, les k -ième et $(n + k)$ -ième colonnes de ces matrices sont indexées respectivement par les arêtes h_k reliant les sommets $\bar{k}$ et $\bar{k} + a$, et h_{n+k} reliant $\bar{k}$ et $\bar{k} + b$. Ainsi, la matrice $\mathbf{H}_X$ (resp. $\mathbf{H}_Z$) contient un 1 en position (i, j) si et seulement si l'arête h_j est incidente au sommet $\bar{i}$ (resp. à la face $\mathcal{F}_i$).

Par conséquent, le code $GB(1 + X^a, 1 + X^b, n)$ est le code CSS dérivé du graphe de Cayley $(\mathbb{Z}/n\mathbb{Z}, \bar{a}, \bar{b})$.

Réseau associé à $GB(1 + X^a, 1 + X^b, n)$

Le réseau 2D associé au code $GB(1 + X^a, 1 + X^b, n)$ se construit également à l'aide des écarts a et b entre les positions des coefficients non nuls des polynômes associés.

Définition 4.3.1. Soient a, b et n trois entiers tels que $1 \leq a, b \leq n - 1$.

Le réseau associé au code $GB(1 + X^a, 1 + X^b, n)$ est donné par le noyau de l'application $\mathbb{Z}$ -linéaire suivante :

$$\begin{aligned} \Phi : \mathbb{Z}^2 &\rightarrow \mathbb{Z}/n\mathbb{Z} \\ \begin{pmatrix} u \\ v \end{pmatrix} &\mapsto au + bv \pmod{n} \end{aligned}$$

Dans le cas où $a = 1$ et $b = \alpha$, le réseau associé est alors $\mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} \alpha \\ -1 \end{pmatrix}$.

Grâce à ces réseaux, nous pouvons établir une borne inférieure pour la distance minimale des codes GB associés.

4.3.3 Borne inférieure sur la distance minimale des codes GB(2,2)

Sous certaines conditions, la distance minimale des codes GB de la forme $GB(1 + X^a, 1 + X^b, n)$, excède la plus petite longueur d'un vecteur non nul du réseau associé.

Le théorème et le corollaire suivants précisent quelles sont les conditions requises.

Théorème 4.3.2. Soient $n \geq 2$ et $1 \leq \alpha \leq n - 1$ deux entiers. La distance minimale d_{min} du code Bicycle généralisé $GB(1 + X, 1 + X^\alpha, n)$ satisfait la borne :

$$d_{min} \geq \lambda(\mathbb{L}) = \min\{ \|\mathbf{u}\|_1 \mid \mathbf{u} \in \mathbb{L} \setminus \{0_{\mathbb{Z}^2}\} \}$$

où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} \alpha \\ -1 \end{pmatrix}$ est le sous-réseau de $\mathbb{Z}^2$ associé et $\|\cdot\|_1$ la norme L^1 .

Corollaire 4.3.1. Soient n, u, v trois entiers tels que $1 \leq u, v \leq n - 1$. Si u est premier avec n , alors la distance minimale du code $GB(1 + X^u, 1 + X^v, n)$ vérifie :

$$d(GB(1 + X^u, 1 + X^v, n)) \geq \min\{ \|\mathbf{u}\|_1 \mid \mathbf{u} \in \mathbb{L} \setminus \{0_{\mathbb{Z}^2}\} \}$$

où $\alpha = (v \pmod{n}) \cdot (u \pmod{n})^{-1}$ et $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} \alpha \\ -1 \end{pmatrix}$ est le réseau mentionné dans la définition 4.3.1, associé au code $GB(1 + X, 1 + X^\alpha, n)$.

4.4 Preuve du théorème 4.3.2

Pour les valeurs $n \in \llbracket 2, 5 \rrbracket$, on peut vérifier à la main que le théorème est vrai. Pour les autres valeurs, nous allons présenter une démonstration qui prouve que le théorème est toujours vrai.

Désormais, tout au long de cette section, nous adopterons les notations suivantes :

- Soit $n \geq 6$ un entier et $1 \leq \alpha \leq n - 1$.
- Posons $A(X) = 1 + X$ et notons $\mathbf{A} = \text{Circ}(A(X), n)$ la matrice circulante associée.
- Posons $B(X) = 1 + X^\alpha$ et notons $\mathbf{B} = \text{Circ}(B(X), n)$ la matrice circulante associée.
- Notons $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} \alpha \\ -1 \end{pmatrix}$ le réseau associé au code $GB(A(X), B(X), n)$.

Comme il a été mentionné dans la Proposition 4.3.1, la distance minimale du code $GB(A(X), B(X), n)$ reste inchangée si l'on réordonne les colonnes des matrices $\mathbf{A}$ ou $\mathbf{B}$. Cela nous permet de simplifier la preuve et de supposer que $1 \leq \alpha \leq \frac{n}{2}$.

Pour les cas particuliers où $\alpha = 1$ ou $\alpha = \frac{n}{2}$, le Théorème 4.3.2 est vérifié puisque la distance minimale du code $GB(1 + X, 1 + X^\alpha, n)$ est égale à $2 = \lambda(\mathbb{L})$. Nous considérons à présent le cas général où $1 < \alpha < \frac{n}{2}$.

4.4.1 Plan de la preuve :

Pour établir la minoration sur la distance minimale du code $GB(A(X), B(X), n)$ dans les cas non triviaux, nous allons utiliser une approche basée sur la théorie des graphes. Celle-ci se décompose en trois étapes :

- **Réinterprétation de la distance minimale** : D'abord, nous réinterpréterons les vecteurs de $\mathbb{F}_2^n$, comme des sous-ensembles d'arêtes du graphe. Puis, nous introduirons la notion de cycle *simple* et nous montrerons que la distance minimale du code correspond à la longueur d'un cycle *simple* ne pouvant être exprimé comme une somme de faces.
- **Utilisation du réseau associé** : Ensuite, en reproduisant les déplacements du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ dans $\mathbb{Z}^2$, nous associerons à chaque cycle simple du graphe un élément de $\mathbb{L}$ dont la norme L^1 est inférieure ou égale à la longueur du cycle.

- **Minoration de la longueur des cycles** : Enfin, nous démontrerons que pour les cycles simples qui ne sont pas des sommes de faces, l'élément de $\mathbb{L}$ associé est non nul. Cela prouvera alors que la distance minimale du code $GB(A(X), B(X), n)$ est supérieure à la plus petite norme L^1 d'un vecteur non nul de $\mathbb{L}$, le réseau associé à notre code GB.

Commençons par réinterpréter la distance minimale de $GB(A(X), B(X), n)$.

4.4.2 Réinterprétation de la distance minimale

En premier lieu, voyons comment les sous-ensembles d'arêtes de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ s'identifient à des vecteurs de $\mathbb{F}_2^n$.

Réinterprétation de l'ensemble des arêtes

On observe tout d'abord que, puisque $n \geq 6$ et $1 < \alpha < \frac{n}{2}$, les arêtes du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ sont deux à deux distinctes. En effet, comme chaque arête est soit de la forme $h_k = \{\bar{k}, \bar{k} + \bar{1}\}$, soit $h_{n+p} = \{\bar{p}, \bar{p} + \bar{\alpha}\}$, avec $k, p \in \llbracket 0, n-1 \rrbracket$, deux arêtes ne peuvent relier les mêmes sommets que si leurs indices sont identiques.

Ceci provient du fait que si deux arêtes d'indices distincts coïncidaient, cela impliquerait que l'une des trois égalités suivantes soit vérifiée dans $\mathbb{Z}/n\mathbb{Z}$:

$$\bar{2} = \bar{0}, \quad \bar{\alpha} - \bar{1} = \bar{0}, \quad \text{ou} \quad \bar{\alpha} + \bar{1} = \bar{0}.$$

Cependant, ce n'est pas possible car $n \geq 6$ et $0 < \alpha - 1 < \alpha + 1 < \frac{n}{2} + 1 < n$.

Ainsi, on peut identifier tout sous-ensemble d'arêtes du graphe à un vecteur de $\mathbb{F}_2^{2n}$ au moyen de la bijection $\mathbf{v} \in \mathbb{F}_2^{2n} \mapsto e_{\mathbf{v}} = \{h_j \mid 0 \leq j < 2n, v_j = 1\}$. Il est important de noter que $e_{\mathbf{v}}$ contient $\text{wt}(\mathbf{v})$ éléments, si $\text{wt}(\mathbf{v})$ désigne le poids de Hamming de $\mathbf{v}$ et que ses éléments sont 2 à 2 distincts.

Nous appellerons $\mathbf{v}$ le **vecteur caractéristique** de $e_{\mathbf{v}}$. Dans la suite, nous utiliserons le terme de **somme** pour désigner la différence symétrique de deux ensembles $e_{\mathbf{v}}$ et $e_{\mathbf{w}}$, puisque cette opération correspond à l'ensemble $e_{\mathbf{v}+\mathbf{w}}$ associé à la somme des vecteurs caractéristiques $\mathbf{v}$ et $\mathbf{w}$.

Grâce à cette identification, les lignes et les noyaux des matrices génératrices peuvent être réinterprétés, en utilisant les définitions de cocycles et de faces introduites dans la Définition 4.2.2.

Réinterprétation de $\text{Vect}_{\mathbf{L}}(\mathbf{H}_X)$ et de $\ker \mathbf{H}_X$:

Pour tout entier $p \in \llbracket 0, n-1 \rrbracket$, la $p^{\text{ème}}$ ligne de $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ s'identifie via la bijection au cocycle associé à $\bar{p}$, qui est l'ensemble des arêtes reliant le sommet $\bar{p}$ à ses quatre voisins :

$$\{ \{\bar{p}, \bar{p} + \bar{1}\}, \quad \{\bar{p}, \bar{p} - \bar{1}\}, \quad \{\bar{p}, \bar{p} + \bar{\alpha}\}, \quad \{\bar{p}, \bar{p} - \bar{\alpha}\} \}$$

Les éléments du noyau de $\mathbf{H}_X$ étant orthogonaux aux lignes de $\mathbf{H}_X$, sont caractérisés par la propriété suivante : $\mathbf{v} \in \ker \mathbf{H}_X$ si et seulement si l'ensemble $e_{\mathbf{v}}$ vérifie que $\forall x \in \mathbb{Z}/n\mathbb{Z}$, il y a un nombre pair d'arêtes de $e_{\mathbf{v}}$ incidentes en x .

Cette propriété fait de $e_{\mathbf{v}}$ un cycle du graphe. Cependant, contrairement aux cycles généraux du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, qui sont des multi-ensembles pouvant contenir plusieurs occurrences d'une même arête, $e_{\mathbf{v}}$ se distingue par le fait qu'il contienne au plus une fois chaque arête. Par exemple, le cycle correspondant à deux parcours successifs de $0 \rightarrow 1 \rightarrow \alpha + 1 \rightarrow \alpha \rightarrow 0$, n'est l'associé d'aucun élément du noyau.

Dans la suite, nous appellerons cycles **particuliers** les cycles qui, comme $e_{\mathbf{v}}$, ne font intervenir qu'au plus une fois chaque arête de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$.

Réinterprétation de $\text{Vect}_{\mathbf{L}}(\mathbf{H}_Z)$ et de $\ker \mathbf{H}_Z$:

Pour tout entier $p \in \llbracket 0, n-1 \rrbracket$, la $p^{\text{ème}}$ ligne de $\mathbf{H}_Z = [\mathbf{B}^{\top} \mid \mathbf{A}^{\top}]$ s'identifie via la bijection à la face associée à $\bar{p}$, qui est l'ensemble des arêtes du 4-cycle reliant successivement les sommets $\bar{p}, \bar{p} + \bar{1}, \bar{p} + \bar{\alpha} + \bar{1}$ et $\bar{p} + \bar{\alpha}$:

$$\{ \{\bar{p}, \bar{p} + \bar{1}\}, \quad \{\bar{p} + \bar{1}, \bar{p} + \bar{1} + \bar{\alpha}\}, \quad \{\bar{p} + \bar{1} + \bar{\alpha}, \bar{p} + \bar{\alpha}\}, \quad \{\bar{p} + \bar{\alpha}, \bar{p}\} \}$$

Réinterprétation de la distance minimale de $GB(1 + X, 1 + X^{\alpha}, n)$

La distance minimale de $GB(1 + X, 1 + X^{\alpha}, n)$ est définie comme le plus petit poids de Hamming d'un vecteur appartenant au noyau de $\mathbf{H}_X$ mais n'appartenant pas à $\text{Vect}_{\mathbf{L}}(\mathbf{H}_Z)$, l'espace vectoriel engendré par les lignes de $\mathbf{H}_Z$.

$$d = \min\{\text{wt}(\mathbf{u}) \mid \mathbf{u} \in \ker \mathbf{H}_X \setminus \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z)\}$$

se réinterprète alors comme la longueur du plus court cycle particulier du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, qui ne peut pas être exprimé comme une somme de faces.

Remarque 4.4.1. *Un plus court tel cycle est nécessairement non vide et **simple**, c'est-à-dire particulier et ne contenant aucun sous-cycle propre non vide. En particulier, un tel cycle est nécessairement **connexe**.*

Nous allons maintenant utiliser cette transcription de la distance minimale, en termes de cycles, pour établir une borne inférieure. Pour y parvenir, nous allons définir des marches sur $\mathbb{Z}^2$ qui reproduisent le comportement des marches que l'on peut effectuer dans le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$.

4.4.3 Des marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ vers les marches sur $\mathbb{Z}^2$

Marche et orientation des arêtes

Une marche $\mathcal{C}$ de longueur $r \geq 1$ sur un graphe est une suite d'arêtes $e_0, \dots, e_{r-1}$ reliant des sommets adjacents. Quand $e_i = \{C_i, C_{i+1}\}$ pour $0 \leq i \leq r-1$, on peut attribuer une orientation à $\mathcal{C}$ en la marquant de la façon suivante :

$$\mathcal{C} : C_0 \rightarrow \dots \rightarrow C_r$$

Dans les marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, on peut rencontrer exactement quatre types d'arêtes, selon l'orientation qu'on leur aura donnée :

$$\begin{aligned} \bar{k} \xrightarrow{+\bar{1}} \bar{k} + \bar{1} \quad \text{ou} \quad \bar{k} + \bar{1} \xrightarrow{-\bar{1}} \bar{k} \\ \bar{k} \xrightarrow{+\bar{\alpha}} \bar{k} + \bar{\alpha}, \quad \text{ou} \quad \bar{k} + \bar{\alpha} \xrightarrow{-\bar{\alpha}} \bar{k} \end{aligned}$$

Afin de reproduire les marches du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ dans $\mathbb{Z}^2$, nous associerons un déplacement plan à chaque type d'arête.

Construction de marches associées sur $\mathbb{Z}^2$

A toute marche $\mathcal{C} : C_0 \rightarrow \dots \rightarrow C_r$ sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, nous pouvons associer une marche correspondante sur $\mathbb{Z}^2$, notée $\gamma_{\mathcal{C}}$. Cette marche $\gamma_{\mathcal{C}}$ reproduit fidèlement le comportement de $\mathcal{C}$ en attribuant un type de déplacement unique à chacune des quatre orientations possibles d'arêtes. La table suivante décrit, de façon informelle, l'association que nous allons effectuer :

Déplacement dans le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$	Déplacement équivalent dans $\mathbb{Z}^2$
Avancer de $\bar{1}$: $\bar{k} \xrightarrow{+\bar{1}} \bar{k} + \bar{1}$	Faire un pas vers la droite : $+\begin{pmatrix} 1 \\ 0 \end{pmatrix}$
Reculer de $\bar{1}$: $\bar{k} \xrightarrow{-\bar{1}} \bar{k} - \bar{1}$	Faire un pas vers la gauche : $-\begin{pmatrix} 1 \\ 0 \end{pmatrix}$
Avancer de $\bar{\alpha}$: $\bar{k} \xrightarrow{+\bar{\alpha}} \bar{k} + \bar{\alpha}$	Faire un pas vers le haut : $+\begin{pmatrix} 0 \\ 1 \end{pmatrix}$
Reculer de $\bar{\alpha}$: $\bar{k} \xrightarrow{-\bar{\alpha}} \bar{k} - \bar{\alpha}$	Faire un pas vers le bas : $-\begin{pmatrix} 0 \\ 1 \end{pmatrix}$

Pour visualiser concrètement cette correspondance, voici un exemple (Figure 4.2). À gauche se trouve le chemin $\bar{0} \rightarrow \bar{3} \rightarrow \bar{4} \rightarrow \bar{5} \rightarrow \bar{6} \rightarrow \bar{1}$ du graphe $(\mathbb{Z}/8\mathbb{Z}, \bar{1}, \bar{3})$. Les arêtes de type $\bar{k} \rightarrow \bar{k} + \bar{\alpha}$, où $\alpha = 3$, sont représentées en bleu et les arêtes $\bar{k} \rightarrow \bar{k} + \bar{1}$ sont en rouge avec des pointillés. À droite se trouve la marche associée dans $\mathbb{Z}^2$.

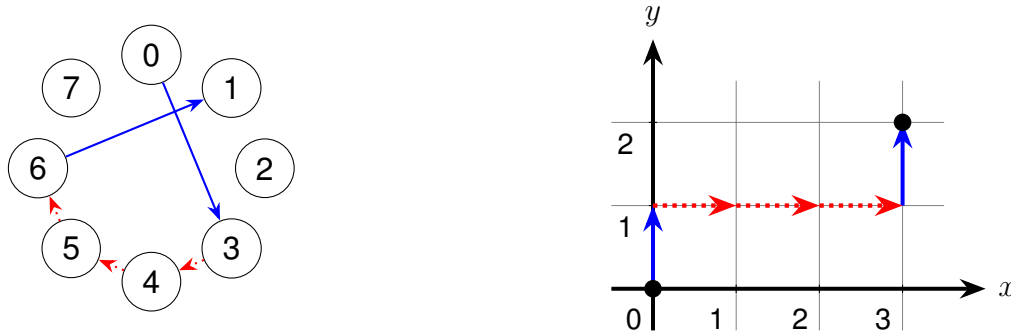


Figure 4.2: Marche sur $\mathbb{Z}^2$ associée à une marche sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$

Formellement, la marche $\gamma_{\mathcal{C}}$ se définit comme une suite d'arêtes $P_0 \rightarrow P_1 \rightarrow \dots \rightarrow P_r$ dans $\mathbb{Z}^2$, dont les arêtes successives sont construites en fonction du type d'arêtes rencontrées dans la marche originale $\mathcal{C} : C_0 \rightarrow C_1 \rightarrow \dots \rightarrow C_r$.

Le point de départ de $\gamma_{\mathcal{C}}$ est $P_0 = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$. Ensuite, pour chaque arête $C_i \rightarrow C_{i+1}$ de la marche $\mathcal{C}$, le point d'arrivée P_{i+1} de l'arête associée dans $\gamma_{\mathcal{C}}$ est obtenu en ajoutant un vecteur $\mathbf{v}$ au point précédent P_i . Ce vecteur $\mathbf{v}$, qui dépend de l'orientation de l'arête traversée dans $\mathcal{C}$, est choisi selon les règles suivantes :

- Si l'arête de $\mathcal{C}$ est $C_i \rightarrow C_{i+1} = C_i + \bar{1}$, on se déplace dans $\mathbb{Z}^2$ d'un pas vers la droite le long de l'axe des abscisses : $\mathbf{v} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$.

- Si l'arête de $\mathcal{C}$ est $C_i \rightarrow C_{i+1} = C_i - \bar{1}$, on se déplace dans $\mathbb{Z}^2$ d'un pas vers la gauche le long de l'axe des abscisses : $\mathbf{v} = -\begin{pmatrix} 1 \\ 0 \end{pmatrix}$.
- Si l'arête de $\mathcal{C}$ est $C_i \rightarrow C_{i+1} = C_i + \bar{\alpha}$, on se déplace dans $\mathbb{Z}^2$ d'un pas vers le haut le long de l'axe des ordonnées : $\mathbf{v} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$.
- Si l'arête de $\mathcal{C}$ est $C_i \rightarrow C_{i+1} = C_i - \bar{\alpha}$, on se déplace dans $\mathbb{Z}^2$ d'un pas vers le bas le long de l'axe des ordonnées : $\mathbf{v} = -\begin{pmatrix} 0 \\ 1 \end{pmatrix}$.

Lorsque la marche initiale $\mathcal{C}$ forme un cycle connexe dans le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, le point d'arrivée de la marche associée dans $\mathbb{Z}^2$ présente des caractéristiques intéressantes, que nous détaillerons dans la section suivante.

Point d'arrivée des marches sur $\mathbb{Z}^2$ associées aux cycles de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$

Considérons une marche $\mathcal{C}$ sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$. Les coordonnées du point d'arrivée de sa marche associée sur $\mathbb{Z}^2$ s'expriment en fonction du nombre d'arêtes de type $\{\pm\bar{1}, \pm\bar{\alpha}\}$ rencontrées durant le parcours de $\mathcal{C}$.

Plus précisément :

- L'abscisse du point d'arrivée est égale à la différence entre le nombre d'arêtes de type $+\bar{1}$ rencontrées et le nombre d'arêtes de type $-\bar{1}$ rencontrées.
- Son ordonnée est égale à la différence entre le nombre d'arêtes de type $+\bar{\alpha}$ rencontrées et le nombre d'arêtes de type $-\bar{\alpha}$ rencontrées.

Proposition 4.4.1. *Soit $\mathcal{C} : C_0 \rightarrow \dots \rightarrow C_r$ une marche sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ de longueur $r \geq 1$. Les coordonnées de $\mathbf{P}_r$, le point d'arrivée de la marche associée sur $\mathbb{Z}^2$, sont données par :*

$$\mathbf{P}_r = \begin{pmatrix} n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C}) \\ n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C}) \end{pmatrix}$$

où pour $\epsilon \in \{\pm\bar{1}, \pm\bar{\alpha}\}$, la notation $n_\epsilon(\mathcal{C})$ représente le nombre de fois où une arête est parcourue dans la direction correspondant à un pas de ϵ dans la marche $\mathcal{C}$:

$$n_\epsilon(\mathcal{C}) = |\{ i \in \llbracket 0, r-1 \rrbracket \mid C_{i+1} - C_i = \epsilon \}|$$

La norme L^1 de P_r est inférieure ou égale à la longueur de $\mathcal{C}$. De plus, si $\mathcal{C}$ est un cycle connexe, alors P_r appartient au réseau $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} -\alpha \\ 1 \end{pmatrix}$.

Preuve : Par définition, les coordonnées du point d'arrivée P_r de la marche $\gamma_{\mathcal{C}}$ sont données par $\begin{pmatrix} n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C}) \\ n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C}) \end{pmatrix}$. De plus, sa norme L^1 est toujours inférieure à la longueur de $\mathcal{C}$ qui est donnée par $n_{\bar{1}}(\mathcal{C}) + n_{-\bar{1}}(\mathcal{C}) + n_{\bar{\alpha}}(\mathcal{C}) + n_{-\bar{\alpha}}(\mathcal{C})$.

Considérons la somme $S = \sum_{i=0}^{r-1} C_{i+1} - C_i$. Si $\mathcal{C} : C_0 \rightarrow \dots \rightarrow C_r = C_0$ est un cycle connexe de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, alors S vaut $\bar{0}$. Or, S vérifie aussi l'égalité suivante :

$$S = (n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C})) \cdot \bar{1} + (n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C})) \cdot \bar{\alpha}$$

Donc si $\mathcal{C}$ est un cycle connexe, alors $P_r = \begin{pmatrix} n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C}) \\ n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C}) \end{pmatrix}$ appartient à $\mathbb{L}$, qui est constitué des vecteurs $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{Z}^2$ satisfaisant $x + \alpha y \equiv 0[n]$. □

Simplification de la preuve du théorème 4.3.2:

Comme établi dans la section 4.4.2, la distance minimale du code $GB(1 + X, 1 + X^\alpha, n)$ est égale la longueur du plus court cycle particulier du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ qui ne peut être exprimé comme une somme de faces. Ce cycle étant nécessairement simple et connexe (voir Remarque 4.4.1), la Proposition 4.4.1 fournit un minorant de sa longueur. Elle stipule en effet que P_r , le point d'arrivée de la marche associée à ce cycle dans $\mathbb{Z}^2$, appartient au réseau $\mathbb{L}$ associé au code $GB(1 + X, 1 + X^\alpha, n)$ et que sa norme L^1 est inférieure ou égale à la longueur du cycle.

Cette observation simplifie la preuve du Théorème 4.3.2. Pour démontrer que la distance minimale du code $GB(1 + X, 1 + X^\alpha, n)$ est minorée par la plus petite norme L^1 d'un vecteur non nul de $\mathbb{L}$, il suffit alors de prouver que P_r est non nul : $P_r \neq \begin{pmatrix} 0 \\ 0 \end{pmatrix}$.

Pour ce faire, nous adopterons le point de vue inverse de celui adopté précédemment sur les marches.

Alors que dans la section précédente, nous avons construit une marche sur $\mathbb{Z}^2$ qui imitait le comportement d'une marche du graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, dans la section suivante, nous montrerons comment construire des marches sur le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ à partir

d'une marche sur $\mathbb{Z}^2$. Cette correspondance inverse nous permettra de démontrer que, pour tout cycle simple non vide $\mathcal{C}$ qui n'est pas une somme de faces, le point d'arrivée P_r de la marche $\gamma_{\mathcal{C}}$ dans $\mathbb{Z}^2$ a nécessairement au moins une coordonnée non nulle.

4.4.4 Des marches sur $\mathbb{Z}^2$ vers les marches sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$

Considérons le graphe de Cayley infini de $\mathbb{Z}^2$, engendré par les éléments $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$ et $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$.

A partir de n'importe quelle marche $\gamma : \gamma_0 \rightarrow \dots \rightarrow \gamma_r$ sur ce graphe, nous pouvons créer une famille de marches $\{(C_\gamma)_{t_0} \mid t_0 \in \mathbb{Z}/n\mathbb{Z}\}$ sur le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, dont chaque membre reproduit le comportement de la marche de départ γ .

La marche $(C_\gamma)_{t_0} : t_0 \rightarrow \dots \rightarrow t_r$ commence à $t_0 \in \mathbb{Z}/n\mathbb{Z}$. Chaque arête $t_i \rightarrow t_{i+1}$ est déterminée par l'orientation qu'avait l'arête correspondante $\gamma_i \rightarrow \gamma_{i+1}$ de la marche initiale γ .

Précisément, pour chaque pas $\gamma_i \rightarrow \gamma_{i+1}$ de γ , le sommet t_{i+1} est obtenu en ajoutant un élément (modulo n) à t_i . L'élément qu'on ajoute est choisi selon les critères suivants :

- Si l'arête initiale dans $\mathbb{Z}^2$ était $\gamma_i \rightarrow \gamma_{i+1}$ avec $\gamma_{i+1} = \gamma_i + \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, alors on ajoute $\bar{1}$ à $t_i : t_{i+1} = t_i + \bar{1}$.
- Si l'arête initiale dans $\mathbb{Z}^2$ était $\gamma_i \rightarrow \gamma_{i+1}$ avec $\gamma_{i+1} = \gamma_i - \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, alors on soustrait $\bar{1}$ à $t_i : t_{i+1} = t_i - \bar{1}$.
- Si l'arête initiale dans $\mathbb{Z}^2$ était $\gamma_i \rightarrow \gamma_{i+1}$ avec $\gamma_{i+1} = \gamma_i + \begin{pmatrix} 0 \\ 1 \end{pmatrix}$, alors on ajoute $\bar{\alpha}$ à $t_i : t_{i+1} = t_i + \bar{\alpha}$.
- Si l'arête initiale dans $\mathbb{Z}^2$ était $\gamma_i \rightarrow \gamma_{i+1}$ avec $\gamma_{i+1} = \gamma_i - \begin{pmatrix} 0 \\ 1 \end{pmatrix}$, alors on soustrait $\bar{\alpha}$ à $t_i : t_{i+1} = t_i - \bar{\alpha}$.

Avec ce formalisme, toute marche $\mathcal{C}$ sur le graphe $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ s'exprime sous la forme $(\gamma_{\mathcal{C}})_{C_0}$, où C_0 est le point de départ de $\mathcal{C}$, et où $\gamma_{\mathcal{C}}$ est la marche sur $\mathbb{Z}^2$ associée à $\mathcal{C}$, introduite dans la section 4.4.3.

Dans la prochaine section, nous expliquerons comment cette écriture permet de conclure la preuve du théorème 4.3.2.

4.4.5 Fin de la preuve du théorème 4.3.2

Considérons un cycle simple non vide de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, que l'on note $\mathcal{C}$.

Notre objectif est de prouver que si $\mathcal{C}$ n'est pas somme de faces, alors P_r , le point d'arrivée de sa marche associée $\gamma_{\mathcal{C}}$ sur $\mathbb{Z}^2$, est non nul. Pour ce faire, nous allons utiliser un raisonnement par contraposition : nous montrerons que si P_r est nul alors $\mathcal{C}$ peut s'exprimer comme une somme de faces.

Notre argument s'appuiera sur les propriétés de la marche associée $\gamma_{\mathcal{C}}$.

Par construction, puisque $\mathcal{C}$ est un cycle simple non vide, $\gamma_{\mathcal{C}}$ est une marche de longueur $r \geq 2$ qui ne repasse ni par une arête, ni par un sommet, à l'exception peut-être des sommets aux extrémités. Ainsi, $\gamma_{\mathcal{C}}$ forme un cycle (simple) si et seulement si P_r , son point d'arrivée, est égal à son point de départ, qui est l'origine $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$.

Dans le lemme suivant, nous allons montrer que tout cycle sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, qui est la marche associée d'un cycle simple non vide de $\mathbb{Z}^2$ passant par l'origine, est une somme de faces.

Puisque, la condition $P_r = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ implique que $\gamma_{\mathcal{C}}$ est un cycle simple non vide de $\mathbb{Z}^2$ passant par l'origine, la démonstration de ce lemme prouvera que si $P_r = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ alors $\mathcal{C}$ est une somme de faces. Cela prouvera alors la contraposée du résultat annoncé.

Lemme 4.4.2. *Si $\gamma : \gamma_0 \rightarrow \dots \rightarrow \gamma_r$ est un cycle simple non vide sur $\mathbb{Z}^2$ commençant à l'origine, alors pour tout $t \in \mathbb{Z}/n\mathbb{Z}$, la marche $(C_{\gamma})_t$ associée sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ est une somme de faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$.*

Preuve : Par abus de langage, on confondra γ avec le contour qu'il définit dans le plan.

On appelle carrés élémentaires : les carrés de côté 1 de $\mathbb{R}^2$, dont les sommets sont des points à coordonnées entières.

Puisque γ est un contour de Jordan, $\mathbb{R}^2 \setminus \gamma$ est la réunion de deux ouverts connexes non vides, l'un étant borné et l'autre non. On notera $Int(\gamma)$ celui qui est borné et on l'appellera intérieur de γ .

Comme $Int(\gamma) \cup \gamma$ s'écrit comme une réunion finie et non vide de carrés élémentaires, la démonstration du lemme se fera par récurrence sur le nombre de carrés élémentaires distincts qui sont entourés par γ , c'est-à-dire inclus dans $Int(\gamma) \cup \gamma$.

Si γ est un cycle simple de $\mathbb{Z}^2$ passant par l'origine et entourant exactement un carré élémentaire de $\mathbb{Z}^2$, alors γ est égal au contour de ce carré. Donc, $\forall t \in \mathbb{Z}/n\mathbb{Z}$, $(C_\gamma)_t$ est l'une des faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ commençant en t :

- (DL) : $t \rightarrow t + \bar{1} \rightarrow t + \bar{1} + \bar{\alpha} \rightarrow t + \bar{\alpha} \rightarrow t$
- (DR) : $t \rightarrow t + \bar{\alpha} \rightarrow t - \bar{1} + \bar{\alpha} \rightarrow t - \bar{1} \rightarrow t$
- (UL) : $t \rightarrow t - \bar{\alpha} \rightarrow t + \bar{1} - \bar{\alpha} \rightarrow t + \bar{1} \rightarrow t$
- (UR) : $t \rightarrow t - \bar{1} \rightarrow t - \bar{1} - \bar{\alpha} \rightarrow t - \bar{\alpha} \rightarrow t$

Maintenant, supposons qu'il existe un entier $q \geq 1$, tel que pour n'importe quel cycle simple γ de $\mathbb{Z}^2$, démarrant à l'origine et entourant exactement $1 \leq k \leq q$ carrés élémentaires distincts, toutes les marches $(C_\gamma)_t$ associées sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ soient des sommes de faces de ce graphe.

Considérons Γ un cycle simple sur $\mathbb{Z}^2$, commençant à l'origine et entourant exactement $q + 1$ carrés élémentaires distincts. Nous allons montrer que pour n'importe quel $t \in \mathbb{Z}/n\mathbb{Z}$, Γ_t la marche associée sur $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ est une somme de faces.

Notons $Int(\Gamma)$ l'unique composante connexe bornée de $\mathbb{R}^2 \setminus \Gamma$. Comme Γ entoure $q + 1 \geq 2$ carrés élémentaires distincts, il y a au moins deux carrés élémentaires distincts contenus dans $Int(\Gamma) \cup \Gamma$ qui partagent des arêtes avec Γ .

On se retrouve donc dans l'un des deux scénarios exclusifs suivants :

- Soit tous les carrés élémentaires de $Int(\Gamma) \cup \Gamma$ qui partagent des arêtes avec Γ sont incidents à l'origine (c'est-à-dire que l'origine est l'un de leurs sommets). Dans ce cas, nous fixons l'un de ces carrés que nous appelons S . (voir Figure 4.3)
- Soit parmi les carrés élémentaires de $Int(\Gamma) \cup \Gamma$ qui partagent des arêtes avec Γ , il y en a au moins un qui n'est pas incident à l'origine. Si on est dans ce cas, nous fixons ce carré, que nous désignerons par S . (voir Figure 4.3)

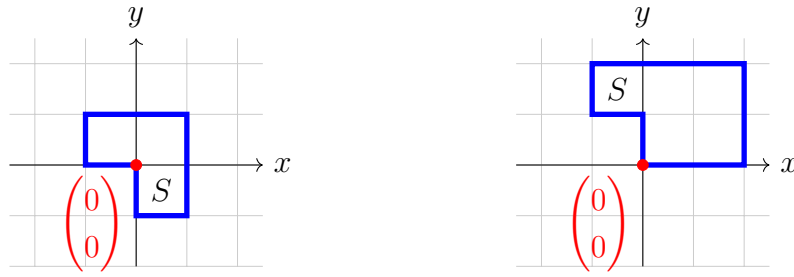


Figure 4.3: Deux configurations possibles dans l'intérieur de Γ . A gauche, le cas où tous les carrés élémentaires partageant des arêtes avec Γ sont incidents à l'origine. A droite, le cas où au moins un de ces carrés n'est pas incident à l'origine

Notons γ le cycle simple de $\mathbb{Z}^2$ qui entoure les mêmes carrés de $\mathbb{Z}^2$ que Γ , à l'exception de S . Puisque γ passe par l'origine et entoure exactement $q \geq 1$ carrés élémentaires distincts, alors par récurrence, nous avons que γ_t est une somme de faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$.

Notons $\begin{pmatrix} u \\ v \end{pmatrix}$ les coordonnées du coin inférieur gauche du carré S et posons $P = \overline{u + \alpha v}$.

Nous obtenons alors l'égalité : $\Gamma_t = \gamma_t + S(t + P)$, où $S(t + P)$ est l'une des faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ ayant une arête incidente au point $t + P$:

- (DL): $t + P \rightarrow t + P + \bar{1} \rightarrow t + P + \bar{1} + \bar{\alpha} \rightarrow t + P + \bar{\alpha} \rightarrow t + P$
- (DR): $t + P \rightarrow t + P + \bar{\alpha} \rightarrow t + P + \bar{\alpha} - \bar{1} \rightarrow t + P - \bar{1} \rightarrow t + P$
- (UL): $t + P \rightarrow t + P - \bar{\alpha} \rightarrow t + P - \bar{\alpha} + \bar{1} \rightarrow t + P + \bar{1} \rightarrow t + P$
- (UR): $t + P \rightarrow t + P - \bar{1} \rightarrow t + P - \bar{1} - \bar{\alpha} \rightarrow t + P - \bar{\alpha} \rightarrow t + P$

Ainsi, Γ_t est une somme de faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$. □

Nous pouvons maintenant conclure la preuve du Théorème 4.3.2.

Corollaire 4.4.1. *Soit $\mathcal{C} : C_0 \rightarrow \dots \rightarrow C_{r-1} \rightarrow C_0$ un cycle simple non vide de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$. Si $\mathcal{C}$ n'est pas une somme de faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$, alors au moins l'une des conditions suivantes doit être vraie : $n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C}) \neq 0$ ou $n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C}) \neq 0$.*

Preuve : Nous allons prouver ce corollaire par contraposée.

Considérons $\gamma_{\mathcal{C}}$ la marche associée à $\mathcal{C}$ sur $\mathbb{Z}^2$. Puisque $\mathcal{C}$ est un cycle simple non vide de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$ alors par construction $\gamma_{\mathcal{C}}$ est un marche non vide sur $\mathbb{Z}^2$ commençant à l'origine, ne repassant par aucun sommet ni aucune arête.

D'après la proposition 4.4.1 , les coordonnées du point d'arrivée de ce chemin sont :

$$\begin{pmatrix} n_{\bar{1}}(\mathcal{C}) - n_{-\bar{1}}(\mathcal{C}) \\ n_{\bar{\alpha}}(\mathcal{C}) - n_{-\bar{\alpha}}(\mathcal{C}) \end{pmatrix}$$

Si $n_{\bar{1}}(\mathcal{C}) = n_{-\bar{1}}(\mathcal{C})$ et $n_{\bar{\alpha}}(\mathcal{C}) = n_{-\bar{\alpha}}(\mathcal{C})$ alors $\gamma_{\mathcal{C}}$ est un cycle simple non vide commençant et finissant à l'origine. Donc , d'après le Lemme 4.4.2, $\mathcal{C} = (\gamma_{\mathcal{C}})_{\mathcal{C}_0}$ est une somme de faces de $(\mathbb{Z}/n\mathbb{Z}, \bar{1}, \bar{\alpha})$. $\square$

4.5 Application du théorème 4.3.2

Dans le domaine des codes quantiques, un objectif majeur est de concevoir des codes ayant une distance minimale la plus élevée possible, car celle-ci est directement liée aux capacités de correction d'erreurs du code. Pour les codes GB(2,2), la croissance de leur distance minimale est limitée. Elle ne peut pas croître plus vite que la racine carrée de la longueur du code.

Historiquement, les codes surface toriques étaient considérés comme les meilleurs codes de surface 2D avec des générateurs de poids quatre (voir Section 4.2.2), mais nos travaux marquent une avancée significative en la matière.

En nous appuyant sur le théorème 4.3.2, nous construisons des codes GB performants, ayant une distance au moins aussi grande que celle des codes de Kitaev. Certains de ces codes présentent même des paramètres aussi performants que les meilleurs codes surfaces de poids quatre précédemment cités.

Kovalev et Pryadko [11] avaient déjà démontré comment construire des codes surface optimaux à distance impaire en utilisant le formalisme GB (notamment avec $GB(1 + X^{2t^2+1}, X + X^{2t^2}, t^2 + (t+1)^2)$ de paramètres $[(2t+1)^2 + 1, 2, 2t+1]$). Avec notre construction, nous avons reproduit cet exploit en introduisant une nouvelle famille de codes GB qui a les mêmes paramètres.

Jusqu'à présent, la communauté scientifique considérait qu'il était impossible de construire des codes GB aussi performants que les codes surface 2D optimaux de distance minimale paire, soit avec les paramètres $[[4r^2, 2, 2r]]$ [16]. En effet, les constructions GB(2,2) les plus prometteuses n'atteignaient au mieux que des paramètres $[[4r^2 + 4, 2, 2r]]$ et uniquement pour des valeurs de r paires.

Notre travail a permis de surmonter cette limitation. Nous avons construit une nouvelle famille de codes GB(2,2) qui réalise pour la première fois les paramètres optimaux $[[4r^2, 2, 2r]]$ via le formalisme des codes GB.

La proposition qui suit présente les familles de codes GB(2,2) performants, obtenus à l'aide du théorème 4.3.2 :

Proposition 4.5.1. *Soient $n \geq 2$, $r \geq 1$ et $t \geq 1$ trois entiers.*

- $GB(1 + X, 1 + X^n, n^2)$ a pour paramètres $[[2n^2, 2, n]]$
- $GB(1 + X, 1 + X^{2r-1}, 2r^2)$ a pour paramètres $[[4r^2, 2, 2r]]$
- $GB(1 + X, 1 + X^{2t+1}, t^2 + (t+1)^2)$ a pour paramètres $[(2t+1)^2 + 1, 2, 2t+1]$

Nous avons décomposé la preuve de cette proposition en plusieurs parties, chacune étant dédiée à un type de construction spécifique. Comme le calcul des longueurs et de la dimension de ces codes est immédiat, nous nous concentrerons uniquement sur le calcul de leur distance minimale.

4.5.1 Etude de $GB(1 + X, 1 + X^n, n^2)$

Lemme 4.5.2. *Pour $n \geq 2$, la distance minimale du code $GB(1 + X, 1 + X^n, n^2)$ est égale à n .*

Preuve : Pour $n = 2$, nous avons prouvé l'égalité à l'aide de l'ordinateur. Pour les autres valeurs ($n \geq 3$), nous avons établi l'égalité en montrant que la distance minimale est à la fois majorée et minorée par n .

Majoration de la distance minimale

La distance minimale du code $GB(1 + X, 1 + X^n, n^2)$ correspond au plus petit poids de Hamming d'un couple de polynômes $(U(X), V(X)) \in (\mathbb{F}_2[X]_{\leq n^2-1})^2$ satisfaisant les propriétés suivantes :

- $(1 + X)U(X) + (1 + X^n)V(X) \equiv 0 \pmod{X^{n^2} - 1}$
- Il n'existe pas de polynôme $H(X) \in \mathbb{F}_2[X]$ tel que :

$$\begin{cases} U(X) \equiv (1 + X^n)H(X) & \pmod{X^{n^2} - 1} \\ V(X) \equiv (1 + X)H(X) & \pmod{X^{n^2} - 1} \end{cases}$$

Dans $\mathbb{F}_2[X]$, nous avons l'égalité : $(1+X^n) \sum_{k=0}^{n-1} X^{nk} = 1+X^{n^2}$. Cela implique alors que :

$$(1+X) \times 0 + (1+X^n) \sum_{k=0}^{n-1} X^{nk} \equiv 0 \pmod{X^{n^2} - 1}$$

Maintenant, supposons par l'absurde qu'il existe un polynôme $H(X) \in \mathbb{F}_2[X]$ tel que :

$$(1+X^n)H(X) \equiv 0 \pmod{X^{n^2} - 1} \quad \text{et} \quad (1+X)H(X) \equiv \sum_{k=0}^{n-1} X^{nk} \pmod{X^{n^2} - 1}$$

Nous avons donc $(\sum_{q=0}^{n-1} X^q)(1+X)H(X) \equiv (1+X^n)H(X) \equiv 0 \pmod{X^{n^2} - 1}$. Cependant, comme $(1+X)H(X) \equiv \sum_{k=0}^{n-1} X^{nk} \pmod{X^{n^2} - 1}$, il s'en suit que :

$$\left(\sum_{k=0}^{n-1} X^{nk}\right) \cdot \left(\sum_{q=0}^{n-1} X^q\right) \equiv 0 \pmod{X^{n^2} - 1}$$

Or, puisque le polynôme $(\sum_{k=0}^{n-1} X^{nk}) \cdot (\sum_{q=0}^{n-1} X^q)$ n'est pas nul et a un degré égal à $n^2 - 1$, il ne peut pas être divisible par $X^{n^2} - 1$.

En conséquence, la distance minimale de $GB(1+X, 1+X^n, n^2)$ est majorée par le poids de Hamming du couple de polynômes $(0, \sum_{k=0}^{n-1} X^{nk})$, qui vaut n .

Minoration de la distance minimale

Selon le Théorème 4.3.2, pour $n \geq 3$, la distance minimale du code $GB(1+X, 1+X^n, n^2)$ est minorée par la plus petite norme L^1 d'un vecteur non nul du réseau $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n^2 \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} n \\ -1 \end{pmatrix}$.

Soit $\mathbf{T}$ un élément non nul de $\mathbb{L}$, défini par : $\mathbf{T} = x \begin{pmatrix} n^2 \\ 0 \end{pmatrix} + y \begin{pmatrix} n \\ -1 \end{pmatrix}$, avec (x, y) un couple d'entiers non simultanément nuls. Sa norme L^1 vaut $\|\mathbf{T}\|_1 = |n(xn + y)| + |y|$.

Si $xn + y = 0$ alors : x est un entier non nul car $(x, y) \neq (0, 0)$. Donc $\|\mathbf{T}\|_1 = n|x| \geq n$.

Si $xn + y \neq 0$ alors : $|xn + y|$ est un entier non nul. La norme L^1 de $\mathbf{T}$ satisfait alors :

$$\|\mathbf{T}\|_1 = |n(xn + y)| + |y| \geq n|xn + y| \geq n$$

En conclusion, $d(GB(1+X, 1+X^n, n^2)) \geq n$. □

4.5.2 Etude de $GB(1 + X, 1 + X^{2r-1}, 2r^2)$

Lemme 4.5.3. *Pour $r \geq 1$, la distance minimale du code $GB(1+X, 1+X^{2r-1}, 2r^2)$ vaut $2r$.*

Preuve : Pour $r = 1$, l'égalité est trivialement vraie. Pour les entiers $r \geq 2$, nous prouvons à nouveau l'égalité en utilisant la méthode de la double inégalité.

Majoration de la distance minimale

Soit $\mathbf{c} = \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix}$ un vecteur binaire de poids de Hamming $2r$ composé de deux sous-vecteurs : $\mathbf{u} \in \mathbb{F}_2^{2r^2}$ qui contient des '1' sur ses $2r - 1$ premières coordonnées et des '0' ailleurs et $\mathbf{v} \in \mathbb{F}_2^{2r^2}$ qui ne contient qu'un seul '1' sur sa première coordonnée, le reste étant des '0'.

Le vecteur $\mathbf{c}$ appartient à l'ensemble $\ker \mathbf{H}_X \setminus \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z)$ où les matrices $\mathbf{H}_X$ et $\mathbf{H}_Z$ sont les suivantes :

- $\mathbf{H}_X = [\text{Circ}(1 + X, 2r^2) \mid \text{Circ}(1 + X^{2r-1}, 2r^2)]$
- $\mathbf{H}_Z = [\text{Circ}(1 + X^{2r-1}, 2r^2)^{\top} \mid \text{Circ}(1 + X, 2r^2)^{\top}]$.

En conséquence, $d(GB(1 + X, 1 + X^{2r-1}, 2r^2)) \leq 2r$.

Minoration de la distance minimale :

Selon le Théorème 4.3.2, pour $r \geq 2$, la distance minimale de $GB(1 + X, 1 + X^{2r-1}, 2r^2)$ est minorée par la plus petite norme L^1 d'un vecteur non nul de $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r^2 \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 2r - 1 \\ -1 \end{pmatrix}$.

Soit $\mathbf{T}$ un élément non nul de $\mathbb{L}$, défini par $\mathbf{T} = x \begin{pmatrix} 2r^2 \\ 0 \end{pmatrix} + y \begin{pmatrix} 2r - 1 \\ -1 \end{pmatrix}$ avec x, y des entiers non simultanément nuls.

Si $y + rx = 0$ alors : $x \neq 0$ car $(x, y) \neq (0, 0)$. Ainsi, la norme L^1 de $\mathbf{T}$ satisfait :

$$\|\mathbf{T}\|_1 = \left\| \begin{pmatrix} rx \\ rx \end{pmatrix} \right\|_1 = 2r|x| \geq 2r$$

Si $y + rx \neq 0$ alors : $|y + rx| \geq 1$ puisque $r, x, y \in \mathbb{Z}$. Nous obtenons donc :

$$\begin{aligned} \|\mathbf{T}\|_1 &= |2r^2x + (2r - 1)y| + |y| = |2r(rx + y) - y| + |y| \\ &\geq 2r|rx + y| - |y| + |y| \\ &\geq 2r \end{aligned}$$

En conclusion, $d(GB(1 + X, 1 + X^{2r-1}, 2r^2) \geq 2r$. □

4.5.3 Etude de $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$

Lemme 4.5.4. *Pour $t \geq 1$, la distance minimale $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ vaut $2t + 1$.*

Preuve : Pour $t = 1$, on prouve l'égalité en utilisant l'ordinateur. Pour les autres valeurs de t , on démontre l'égalité en montrant que la distance minimale est à la fois majorée et minorée par $2t + 1$.

Majoration de la distance minimale :

La distance minimale de $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ correspond au plus petit poids de Hamming d'une paire de polynômes $(U(X), V(X)) \in (\mathbb{F}_2[X]_{\leq t^2 + (t+1)^2 - 1})^2$ satisfaisant les conditions :

- $(1 + X)U(X) + (1 + X^{2t+1})V(X) \equiv 0 \pmod{X^{t^2 + (t+1)^2} - 1}$
- Il n'existe pas de polynôme $H(X) \in \mathbb{F}_2[X]$ tel que :

$$\begin{cases} U(X) \equiv (1 + X^{2t+1})H(X) & \pmod{X^{t^2 + (t+1)^2} - 1} \\ V(X) \equiv (1 + X)H(X) & \pmod{X^{t^2 + (t+1)^2} - 1} \end{cases}$$

Dans $\mathbb{F}_2[X]/(X^{t^2 + (t+1)^2} - 1)$, les relations suivantes sont vérifiées :

$$\begin{aligned} (1 + X^{2t+1})(1 + \sum_{k=0}^{t-1} X^{2t^2 - k(2t+1)}) &\equiv X^{2t+1} + X^{2t^2 - (t-1)(2t+1)} && \pmod{X^{t^2 + (t+1)^2} - 1} \\ &\equiv X^{2t+1} + X^{t+1} && \pmod{X^{t^2 + (t+1)^2} - 1} \\ &\equiv X^{t+1}(X^t + 1) && \pmod{X^{t^2 + (t+1)^2} - 1} \\ &\equiv X^{t+1}(\sum_{q=0}^{t-1} X^q)(1 + X) && \pmod{X^{t^2 + (t+1)^2} - 1} \end{aligned}$$

En conséquence, en posant $U(X) = X^{t+1}(\sum_{q=0}^{t-1} X^q)$ et $V(X) = (1 + \sum_{k=0}^{t-1} X^{2t-k(2t+1)})$, nous obtenons la congruence suivante modulo $X^{t^2+(t+1)^2} - 1$:

$$(1 + X)U(X) + (1 + X^{2t+1})V(X) \equiv 0 \pmod{X^{t^2+(t+1)^2} - 1}$$

Les degrés de $U(X)$ et $V(X)$ sont, tous deux, strictement inférieurs à $t^2 + (t + 1)^2$ et leurs poids de Hamming respectifs sont t et $t + 1$. Comme l'un de ces deux polynômes a forcément un poids de Hamming impair, on ne peut pas trouver de polynôme $H(X) \in \mathbb{F}_2[X]$ satisfaisant simultanément les deux congruences :

$$\begin{cases} U(X) \equiv (1 + X^{2t+1})H(X) & \pmod{X^{t^2+(t+1)^2} - 1} \\ V(X) \equiv (1 + X)H(X) & \pmod{X^{t^2+(t+1)^2} - 1} \end{cases}$$

En conséquence, la distance minimale du code $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ est majorée par $2t + 1$ qui est le poids de Hamming du couple $(U(X), V(X))$.

Minoration de la distance minimale :

D'après le Théorème 4.3.2, pour les valeurs de t qui sont supérieure ou égale à 2, la distance minimale du code $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ est minorée par la plus petite norme L^1 de vecteurs non nuls du réseau $\mathbb{Z} \begin{pmatrix} t^2 + (t + 1)^2 \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 2t + 1 \\ -1 \end{pmatrix}$.

Considérons $\mathbf{T}$ un élément non nul quelconque de ce réseau, défini par :

$$\mathbf{T} = x \begin{pmatrix} t^2 + (t + 1)^2 \\ 0 \end{pmatrix} + y \begin{pmatrix} 2t + 1 \\ -1 \end{pmatrix}$$

où x et y sont des entiers non-simultanément nuls. Nous allons minorer sa norme L^1 .

Commençons par traiter le cas où x est nul. Si $x = 0$ alors, y est un entier non nul. Donc $\|\mathbf{T}\|_1$ vérifie : $\|\mathbf{T}\|_1 \geq |y|(2t + 1) \geq 2t + 1$.

Maintenant supposons que $x \neq 0$. Nous allons distinguer deux cas : selon que $|y| \leq t|x|$ ou non.

Cas 1. $|y| \leq t|x|$:

$$\begin{aligned}
 \|\mathbf{T}\|_1 &= |x(t^2 + (t+1)^2) + y(2t+1)| + |y| \\
 &\geq |x|(t^2 + (t+1)^2) - |y|(2t+1) + |y| \\
 &= |x|(2t^2 + 2t + 1) - |y| \cdot 2t \\
 &= |x|(2t+1) + 2t(t|x| - |y|) \\
 &\geq |x|(2t+1) \\
 &\geq 2t+1
 \end{aligned}$$

Cas 2. $|y| > t|x|$: $\mathbf{T}$ vérifie les égalités suivantes :

$$\mathbf{T} = x \begin{pmatrix} t^2 + (t+1)^2 \\ 0 \end{pmatrix} + y \begin{pmatrix} 2t+1 \\ -1 \end{pmatrix} = \begin{pmatrix} (2t+1)((t+1)x + y) - tx \\ -y \end{pmatrix}$$

Si $|(t+1)x + y| = 0$, alors $\mathbf{T} = \begin{pmatrix} -tx \\ -(t+1)x \end{pmatrix}$ donc $\|\mathbf{T}\|_1 = (2t+1)|x| \geq 2t+1$.

Si $|(t+1)x + y| \neq 0$, alors la norme L^1 de $\mathbf{T}$ satisfait :

$$\begin{aligned}
 \|\mathbf{T}\|_1 &= |(2t+1)((t+1)x + y) - tx| + |y| \\
 &\geq |(2t+1)|((t+1)x + y) - t|x|| + |y| \\
 &\geq (2t+1)|((t+1)x + y) - t|x|| + |y| \\
 &> (2t+1)|((t+1)x + y)| \\
 &\geq 2t+1
 \end{aligned}$$

En conclusion, $d(GB(1+X, 1+X^{2t+1}, t^2 + (t+1)^2) \geq 2t+1$. □

Les familles de codes Bicycle généralisés que nous avons introduites dans la Proposition 4.5.1 atteignent des performances optimales, rivalisant avec celles des codes de Kitaev et des meilleurs codes surface 2D à stabilisateurs de poids quatre. Cependant, il n'y a aucune garantie que ces codes soient entièrement nouveaux. Il est tout à fait possible qu'ils soient simplement équivalents à des codes optimaux déjà connus.

Pour affirmer rigoureusement l'originalité de notre contribution, il est donc crucial de déterminer si ces codes représentent effectivement des structures inédites.

Pour y parvenir, nous allons examiner les différents types de relations d'équivalence qui permettent de classifier les codes quantiques.

4.6 Relations d'équivalence sur les codes quantiques

En théorie des codes quantiques, deux codes sont généralement considérés comme équivalents s'ils peuvent être transformés l'un en l'autre par des opérations qui préservent leurs capacités essentielles de correction d'erreurs. Cependant, comme nous le verrons, il existe plusieurs types de relations d'équivalence, chacune ayant des implications et des comportements distincts.

4.6.1 Relation d'équivalence générale

Définition

Deux codes quantiques, $\mathcal{Q}_1$ et $\mathcal{Q}_2$, sont dits équivalents s'il existe une opération unitaire U telle que $\mathcal{Q}_2 = U\mathcal{Q}_1$. Pour les codes stabilisateurs, on impose que U soit un produit de matrices de permutations et de transformations locales de Clifford, c'est-à-dire un produit tensoriel de portes de Hadamard $\frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$, Phase $\begin{pmatrix} 1 & 0 \\ 0 & i \end{pmatrix}$ et

$$\text{CNOT} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} \text{ [17].}$$

Limitations

Cette relation d'équivalence générale ne préserve pas toujours la structure CSS. Lorsqu'un code $\mathcal{Q}_1$ est stabilisé par un ensemble d'opérateurs $\mathbb{S}$, le code équivalent $\mathcal{Q}_2$ est stabilisé par sa classe conjugaison par U , soit $U\mathbb{S}U^\dagger$.

Sous cette action, certaines portes Clifford locales, comme la porte de Hadamard (H), transforment les opérateurs de Pauli X en Z ($HXH^\dagger = Z$) et inversement ($HZH^\dagger = X$). De même, la porte de Phase (S) transforme X en Y ($SXS^\dagger = Y = iXZ$). Ces transformations modifient la forme fondamentale des stabilisateurs, ce qui peut briser la structure CSS qui exige que les stabilisateurs soient composés exclusivement de produits d'opérateurs de type X ou de type Z .

Afin de comparer les codes GB optimaux (présentés dans la Section 4.5) aux meilleurs codes surface 2D à stabilisateurs de poids quatre, qui sont tous deux dérivés de graphes de Cayley de groupes abéliens, nous introduisons une relation d'équivalence, appelée *CGP* (pour *CSS-graph-preserving*), qui préserve à la fois la structure CSS et la structure de graphe sous-jacente.

4.6.2 Relation d'équivalence CGP

Définition :

Soit $\mathcal{Q}$ et $\mathcal{Q}'$ deux codes quantiques de longueur $n \geq 1$.

On dit que $\mathcal{Q}$ est CGP-équivalent à $\mathcal{Q}'$ et on note $\mathcal{Q} \sim \mathcal{Q}'$, si et seulement s'il existe une bijection ϕ sur l'ensemble d'indices des qubits $\{1, \dots, n\}$, telle que $\mathcal{Q}' = \mathbf{P}_\phi(\mathcal{Q})$, où $\mathbf{P}_\phi \in GL_{2^n}(\mathbb{C})$ est une matrice de permutation agissant sur les vecteurs de la base canonique de la façon suivante :

$$\forall t_1, \dots, t_n \in \{0, 1\}, \quad \mathbf{P}_\phi |t_1 \dots t_n\rangle = |t_{\phi(1)} \dots t_{\phi(n)}\rangle$$

Concrètement, cela signifie que les éléments de $\mathcal{Q}'$ sont obtenus en permutant les qubits des éléments de $\mathcal{Q}$:

$$\mathcal{Q}' = \left\{ \sum_{t \in \mathbb{F}_2^n} \alpha_t |t_{\phi(1)} \dots t_{\phi(n)}\rangle \mid \sum_{t \in \mathbb{F}_2^n} \alpha_t |t_1 \dots t_n\rangle \in \mathcal{Q} \right\}$$

Compatibilité avec la structure CSS :

Deux codes stabilisateurs CGP-équivalents $\mathcal{Q}$ et $\mathcal{Q}'$ (c'est-à-dire $\mathcal{Q}' = \mathbf{P}_\phi(\mathcal{Q})$), possèdent les mêmes paramètres et leurs groupes de stabilisateurs sont directement liés.

Plus précisément, un opérateur de Pauli $\mathbf{O}_1 \otimes \dots \otimes \mathbf{O}_n$ stabilise $\mathcal{Q}$ si et seulement si l'opérateur permuté $\mathbf{O}_{\phi(1)} \otimes \dots \otimes \mathbf{O}_{\phi(n)}$ stabilise $\mathcal{Q}'$. Cela implique que la structure CSS est préservée par cette relation d'équivalence : $\mathcal{Q}$ est un code CSS si et seulement si $\mathcal{Q}'$ l'est également.

Lien entre les matrices génératrices des codes CSS :

Les matrices génératrices de codes CSS CGP-équivalents sont elles aussi intrinsèquement liées.

Deux codes CSS $\mathcal{Q}$ et $\mathcal{Q}'$ sont CGP-équivalents (c'est-à-dire $\mathcal{Q}' = \mathbf{P}_\phi(\mathcal{Q})$) si et seulement si les sous-espaces vectoriels engendrés par les lignes de leurs sous-matrices génératrices respectives, $\mathbf{H} = \begin{pmatrix} \mathbf{H}_X & 0 \\ 0 & \mathbf{H}_Z \end{pmatrix}$ et $\mathbf{H}' = \begin{pmatrix} \mathbf{H}'_X & 0 \\ 0 & \mathbf{H}'_Z \end{pmatrix}$, satisfont les relations :

$$\begin{cases} \text{Vect}_{\mathbf{L}}(\mathbf{H}'_X) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_X \mathbf{Q}) \\ \text{Vect}_{\mathbf{L}}(\mathbf{H}'_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z \mathbf{Q}) \end{cases}$$

où $\mathbf{Q}$ est la matrice de permutation $n \times n$ associée à $\phi : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$.

Compatibilité avec la structure de graphes

Les codes GB(2,2), les codes de Kitaev et les codes de surface 2D à stabilisateurs de poids quatre sont tous dérivés de graphes de Cayley de groupes abéliens. Leurs sous-matrices génératrices $\mathbf{H}_X$ et $\mathbf{H}_Z$, correspondent aux matrices d'incidence sommet-arête et face-arête de leurs graphes de Cayley sous-jacents.

Pour deux codes de ce type, les conditions $\text{Vect}_{\mathbf{L}}(\mathbf{H}'_X) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_X \mathbf{Q})$ et $\text{Vect}_{\mathbf{L}}(\mathbf{H}'_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z \mathbf{Q})$, imposées par l'équivalence CGP, se traduisent en termes de graphes. Plus précisément, ces contraintes garantissent l'existence d'une bijection entre les arêtes des graphes sous-jacents, telle que les cycles simples d'un graphe sont directement transformés en cycles simples de l'autre. La proposition suivante formalise ce résultat.

Proposition 4.6.1. *Considérons deux codes CSS $\mathcal{Q}$ et $\mathcal{Q}'$, basés sur des graphes de Cayley ($\mathcal{G}$ et $\mathcal{G}'$), dont les sous-matrices génératrices satisfont les relations : $\text{Vect}_{\mathbf{L}}(\mathbf{H}'_X) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_X \mathbf{Q})$ et $\text{Vect}_{\mathbf{L}}(\mathbf{H}'_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z \mathbf{Q})$ où $\mathbf{Q}$ est une matrice de permutation.*

Alors, il existe une correspondance bijective entre les arêtes de $\mathcal{G}$ et $\mathcal{G}'$, induite par $\mathbf{Q}$, pour laquelle tout cycle simple de $\mathcal{G}$ correspond à un cycle simple de $\mathcal{G}'$.

Preuve : Soit N le nombre total d'arêtes dans $\mathcal{G}$ et $\mathcal{G}'$. On représente l'ensemble des parties de l'ensemble des arêtes de ces deux graphes par $\mathbb{F}_2^N$.

L'application $\mathbf{v} \in \mathbb{F}_2^N \rightarrow \mathbf{Q}^{-1}\mathbf{v} \in \mathbb{F}_2^N$ est une bijection de $\mathbb{F}_2^N$ sur lui-même. Cette bijection induit une correspondance entre les sous-ensembles d'arêtes de $\mathcal{G}$ et $\mathcal{G}'$, garantissant que chaque cycle simple de $\mathcal{G}$ est envoyé sur un cycle simple de $\mathcal{G}'$.

En effet, considérons un cycle simple $\mathcal{C}_v$ du graphe $\mathcal{G}$. Puisque, $\mathbf{H}_X$ est une matrice d'incidence sommet-arête du graphe $\mathcal{G}$, alors le cycle $\mathcal{C}_v$ est représenté par un vecteur $\mathbf{v} \in \ker(\mathbf{H}_X)$. En prenant son image par la transformation bijective, on trouve que $\mathbf{Q}^{-1}\mathbf{v} \in \ker(\mathbf{H}_X\mathbf{Q})$. Comme $\text{Vect}_{\mathbf{L}}(\mathbf{H}_X\mathbf{Q})^\perp = \text{Vect}_{\mathbf{L}}(\mathbf{H}'_X)^\perp = \ker(\mathbf{H}'_X)$, il s'en suit que $\mathbf{Q}^{-1}\mathbf{v} \in \ker(\mathbf{H}'_X)$. Cela signifie que $\mathcal{C}'_{\mathbf{Q}^{-1}\mathbf{v}}$, l'ensemble d'arête de $\mathcal{G}'$ correspondant, est un cycle du graphe $\mathcal{G}'$.

Maintenant, supposons par l'absurde que $\mathcal{C}'_{\mathbf{Q}^{-1}\mathbf{v}}$ ne soit pas un cycle simple du graphe $\mathcal{G}'$. Cela impliquerait qu'il existe un sous-cycle propre non vide $\emptyset \subsetneq \mathcal{C}'_{\mathbf{h}} \subsetneq \mathcal{C}'_{\mathbf{Q}^{-1}\mathbf{v}}$ de $\mathcal{C}'_{\mathbf{Q}^{-1}\mathbf{v}}$ dans $\mathcal{G}'$. Ce sous-cycle est représenté par un vecteur non nul $\mathbf{h} = \mathbf{Q}^{-1}\mathbf{w} \in \ker(\mathbf{H}'_X) = \ker(\mathbf{H}_X\mathbf{Q})$, dont le support est inclus strictement dans le support de $\mathbf{Q}^{-1}\mathbf{v}$.

Puisque $\mathbf{Q}$ est une matrice de permutation, alors $\mathbf{w} \in \ker(\mathbf{H}_X)$ et son support est un sous-ensemble strict et non vide du support de $\mathbf{v}$. En conséquence, le cycle $\mathcal{C}_{\mathbf{w}}$ de $\mathcal{G}$ associé à $\mathbf{w}$, est un sous-cycle propre non vide de $\mathcal{C}_v$, ce qui contredit la simplicité de $\mathcal{C}_v$. $\square$

Bien qu'elle soit plus restrictive que l'équivalence générale sur les codes quantiques, l'équivalence CGP permet de comparer les codes CSS dérivés de graphe de Cayley tout en préservant leur structure CSS et leur structure de graphes sous-jacente.

Dans la section suivante, nous verrons que pour les codes GB et les codes surface à stabilisateurs de poids quatre étudiés ici, la correspondance entre les cycles simples des sous-graphes sous-jacents induit un isomorphisme de graphes. Cela signifie que les relations d'incidence entre les arêtes sont intégralement préservées, offrant ainsi une forme encore plus forte d'équivalence structurelle.

4.7 Comparaison des codes GB aux codes surface 2D de poids quatre

Nos nouvelles familles de codes Bicycle généralisés, introduites dans la Section 4.5, ainsi que les codes de surface 2D décrits dans la Section 4.2.2, sont tous des codes CSS construits à partir de graphes de Cayley 3-connexes, c'est-à-dire des graphes qui demeurent connexes même après la suppression de deux sommets.

En exploitant cette structure, nous avons démontré que deux des familles de codes GB que nous avons construites sont différentes (non CGP-équivalentes) des codes

toriques de Kitaev et des codes surface 2D optimaux à distance paire. Cette preuve repose sur une idée clé en théorie des graphes : pour les graphes 3-connexes, tels que ceux qui sous-tendent nos codes, avoir une correspondance bijective entre les cycles simples des graphes induit un isomorphisme de graphes.

Plus précisément, notre raisonnement s'appuie sur le théorème de Whitney [18], qui énonce des conditions suffisantes pour garantir l'existence d'un isomorphisme de graphes entre deux graphes 3-connexes. Dans la section suivante, nous détaillerons la manière dont ce théorème sera utilisé pour prouver l'originalité de nos codes.

4.7.1 Théorème de Whitney

3-connexité

Un graphe est dit connexe s'il existe un chemin reliant n'importe quelle paire de sommets. Il est dit 3-connexe si même en retirant deux sommets distincts, il demeure connexe.

Définition 4.7.1 (3-connexité). *Un graphe $\mathcal{G}$ est dit 3-connexe s'il est connexe et que pour toute paire de sommets distincts (u, v) , le graphe obtenu en retirant ces deux sommets $\mathcal{G} - \{u, v\}$, reste connexe.*

Une condition suffisante pour qu'un graphe soit 3-connexe, est qu'il existe trois chemins indépendants (ne partageant aucun sommet, à l'exception de leurs extrémités) pour relier n'importe quelle paire de sommets distincts A et B .

Isomorphisme de graphes

L'isomorphisme de graphes, dont la définition est rappelée ci-dessous, est une propriété qui préserve les relations d'adjacence entre les sommets (et par extension la 3-connexité).

Définition 4.7.2 (Isomorphisme de graphes). *Deux graphes $\mathcal{G}$ et $\mathcal{H}$ sont dits isomorphes s'il existe une application bijective $\tau : \mathcal{G} \rightarrow \mathcal{H}$ entre les ensembles de sommets des graphes, qui est telle que : u et v sont voisins dans $\mathcal{G}$ si et seulement si $\tau(u)$ et $\tau(v)$ sont voisins dans $\mathcal{H}$.*

Un isomorphisme de graphes conserve les relations d'adjacence entre sommets et arêtes. Ainsi, si deux graphes sont isomorphes, leurs ensembles d'arêtes sont en bijection, et cette bijection induit à son tour une correspondance bijective entre leurs cycles simples. En revanche, la réciproque n'est pas vraie en général : deux graphes peuvent admettre une bijection entre leurs arêtes qui préserve les cycles simples, sans pour autant être isomorphes.

Théorème de Whitney

Le théorème de Whitney fournit une condition pour que cette réciproque soit vérifiée.

Théorème 4.7.1 (Théorème de Whitney ([19] et [20])). *Soient $\mathcal{G}_1$ et $\mathcal{G}_2$ deux graphes 3-connexes, ne contenant ni boucles, ni multi-arêtes. S'il existe une bijection entre leurs arêtes qui induit une correspondance entre leurs cycles simples, alors $\mathcal{G}_1$ et $\mathcal{G}_2$ sont isomorphes.*

Remarque 4.7.1. *Dans l'énoncé original du théorème 2 de [20], Whitney utilise le terme circuit. Cependant, d'après les définitions établies dans [19] et que Whitney a reprises dans [20], un circuit désigne un cycle connexe ne comportant ni sommet ni arête répétés. Cette définition coïncide donc avec la notion de cycle simple adoptée dans ce mémoire (voir remarque 4.4.1).*

4.7.2 Vue d'ensemble de la démonstration de la non-équivalence

Pour démontrer que les codes surface 2D mentionnés précédemment sont distincts (non CGP-équivalents) de nos codes GB, nous allons nous appuyer sur le fait que ces codes CSS sont dérivés de graphes de Cayley non orientés ($\mathcal{G}$ et $\mathcal{G}'$) de groupes abéliens ($\mathbb{G}$ et $\mathbb{G}'$).

Notre stratégie consiste à ramener la question de l'équivalence CGP des codes à celle de l'existence d'un isomorphisme de groupes, une question à laquelle nous pourrions répondre facilement. Plus précisément, pour prouver la non-équivalence entre les codes surface et nos codes, nous allons procéder comme suit :

1. Tout d'abord, nous établirons que les graphes sous-jacents $\mathcal{G}$ et $\mathcal{G}'$ des codes surfaces et de nos codes GB sont 3-connexes.

2. Ensuite, nous raisonnerons par l'absurde pour montrer que les codes ne sont pas CGP-équivalents. Si c'était le cas, il existerait une bijection entre les arêtes des graphes correspondants qui préserve les cycles simples. D'après le théorème de Whitney, cela impliquerait que les graphes sont isomorphes.
3. Enfin, nous démontrerons que cet isomorphisme de graphes induirait nécessairement un isomorphisme de groupes entre $\mathbb{G}$ et $\mathbb{G}'$. Cependant, nous montrerons que dans les cas étudiés, les groupes en question ne sont pas isomorphes. Cela prouvera ainsi que les codes sont distincts.

Dans la section suivante, nous démontrerons que le code $GB(1 + X, 1 + X^{2r-1}, 2r^2)$ n'est pas CGP-équivalent au code de surface 2D optimal de distance minimale $2r$, de paramètres $[[4r^2, 2, 2r]]$. Nous établirons ensuite que le code $GB(1 + X, 1 + X^n, n^2)$ n'est pas CGP-équivalent au code torique de Kitaev $[[2n^2, 2, n]]$. Enfin, nous prouverons que les codes $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ sont équivalents aux codes de surface à stabilisateurs de poids quatre optimaux, de paramètres $[(2t + 1)^2 + 1, 2, 2t + 1]$.

4.7.3 Comparaison avec le code surface 2D optimal pour la distance paire

Pour montrer que notre code $GB(1 + X, 1 + X^{2r-1}, 2r^2)$ n'est pas CGP-équivalent au code de surface 2D de distance minimale $2r$ optimal, nous allons d'abord établir que leurs graphes sous-jacents sont 3-connexes.

3-connexité des graphes sous-jacents

Lemme 4.7.2. *Pour $r \geq 3$, le graphe de Cayley $(\mathbb{Z}/2r^2\mathbb{Z}, \overline{1}, \overline{2r-1})$, sous-jacent au code $GB(1 + X, 1 + X^{2r-1}, 2r^2)$ est 3-connexe.*

Preuve : L'objectif est de prouver que même si on supprime deux sommets distincts, le graphe restant est toujours connexe.

Cas 1: On supprime $\bar{a}$ et $\overline{a+1}$

Le graphe original $(\mathbb{Z}/2r^2\mathbb{Z}, \overline{1}, \overline{2r-1})$ contient le cycle hamiltonien passant par tous les sommets $\bar{a} \rightarrow \overline{a+1} \rightarrow \dots \rightarrow \overline{a-1} \rightarrow \bar{a}$. Même après la suppression des sommets $\bar{a}$ et $\overline{a+1}$, les sommets restants demeurent connectés par le chemin $\overline{a+2} \rightarrow \overline{a+3} \rightarrow \dots \rightarrow \overline{a-1}$.

Cas 2 : On supprime $\bar{a}$ et $\bar{b}$ avec $0 \leq a < b \leq 2r^2 - 1$ et $\bar{a} \notin \{\bar{b} + 1, \bar{b} - 1\}$

Quand on retire les sommets $\bar{a}$ et $\bar{b}$ du graphe $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$ le cycle hamiltonien mentionné précédemment se scinde en deux sections distinctes :

- Une section s'étendant de $\overline{b+1}$ à $\overline{a-1}$

$$\overline{b+1} \rightarrow \overline{b+2} \rightarrow \dots \rightarrow \bar{0} \rightarrow \bar{1} \rightarrow \dots \rightarrow \overline{a-1}$$

- Une autre allant de $\overline{a+1}$ à $\overline{b-1}$

$$\overline{a+1} \rightarrow \overline{a+2} \rightarrow \dots \rightarrow \overline{b-1}$$

Si $\bar{a} = \overline{b+1 - (2r-1)} = \overline{b-2r+2}$, alors l'arête $\overline{a-2} \rightarrow \overline{a-2 + (2r-1)}$ assure la connexion entre les deux sections.

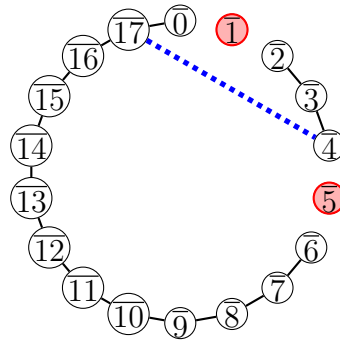


Figure 4.4: La connexité du graphe de Cayley $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$ où $r = 3$ est préservée après la suppression de $a = \bar{1}$ et $\overline{a+1 - (2r-1)} = \bar{5}$ grâce à l'arête reliant $\overline{a-2} = \bar{17}$ et $\overline{a-2 + (2r-1)} = \bar{4}$.

L'arête $\overline{a-2} \rightarrow \overline{a-2 + (2r-1)}$ est bien présente dans le graphe $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$. De plus, ses extrémités sont distinctes des sommets $\bar{a}$ et $\bar{b}$ qu'on a supprimés :

- $\overline{a-2} \neq \bar{a}$: ceci est vrai car $\bar{2} \neq \bar{0}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$ puisque $r \geq 3$.
- $\overline{a-2} \neq \bar{b}$: si c'était le cas, on aurait $\overline{b-2r} = \bar{b} \implies \overline{2r} = \bar{0}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$. Or, c'est impossible car $0 < 2r < 2r^2$.
- $\overline{a-2 + (2r-1)} \neq \bar{a}$: si c'était le cas, on aurait $\overline{2r-3} = \bar{0}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$. Or, c'est impossible car pour $r \geq 3$, nous avons que $0 < 2r-3 < 2r^2$.

- $\overline{a - 2 + (2r - 1)} \neq \bar{b}$: si c'était le cas, on aurait $\bar{b} - 1 = \bar{b}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$, ce qui entraînerait l'égalité $\bar{1} = \bar{0}$, qui est fausse.

Si $\bar{a} \neq \overline{b + 1 - (2r - 1)}$, alors l'arête $\overline{b + 1} \rightarrow \overline{b + 1 - (2r - 1)}$ assure la connexion entre les deux branches de $(\mathbb{Z}/2r^2\mathbb{Z} - \{\bar{a}, \bar{b}\}, \bar{1}, \overline{2r - 1})$ décrites plus haut.

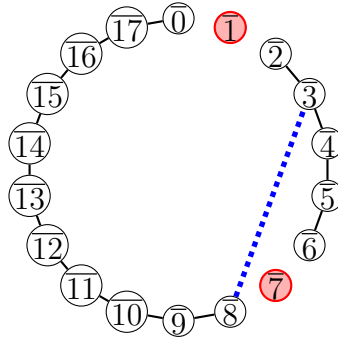


Figure 4.5: La connexité dans le graphe de Cayley $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r - 1})$ où $r = 3$ est préservée après la suppression des sommets $a = \bar{1}$ et $b = \bar{7} \neq \overline{a + 1 - (2r - 1)}$ grâce à l'arête entre les sommets $\overline{b + 1} = \bar{8}$ et $\overline{b + 1 - (2r - 1)} = \bar{3}$.

L'arête $\overline{b + 1} \rightarrow \overline{b + 1 - (2r - 1)}$ est présente dans le graphe original et ses deux extrémités sont distinctes des sommets $\bar{a}$ et $\bar{b}$:

- $\overline{b + 1} \neq \bar{a}$ par hypothèse.
- $\overline{b + 1} \neq \bar{b}$: si c'était le cas, cela impliquerait que $\bar{1} = \bar{0}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$, ce qui est impossible.
- $\overline{b + 1 - (2r - 1)} \neq \bar{a}$ par hypothèse.
- $\overline{b + 1 - (2r - 1)} \neq \bar{b}$: si c'était le cas, on aurait l'égalité $\bar{0} = \overline{2r - 2}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$. Cependant, elle est fausse car pour $r \geq 3$, on a $0 < 2r - 2 < 2r^2$.

En conséquence, le graphe $(\mathbb{Z}/2r^2\mathbb{Z} - \{\bar{a}, \bar{b}\}, \bar{1}, \overline{2r - 1})$ est connexe. Cela implique alors que le graphe $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r - 1})$ est 3-connexe. $\square$

Remarque 4.7.2. Plus généralement, en employant la même stratégie que dans cette preuve, il est possible de montrer que tout code GB(2,2) de la forme $GB(1+X, 1+X^\alpha, N)$, avec $2 \leq \alpha \leq N - 2$, possède un graphe sous-jacent $(\mathbb{Z}/N\mathbb{Z}, \bar{1}, \bar{\alpha})$ 3-connexe.

Nous allons maintenant démontrer que le graphe de Cayley sous-jacent au code surface optimal pour la distance paire (introduit en Section 4.2.2) est 3-connexe.

Lemme 4.7.3. *Soit $r \geq 3$ un entier. Posons $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix}$.*

Le graphe de Cayley $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ associé au code surface $2D$ à stabilisateurs de poids quatre optimal de distance $2r$ (introduit en Section 4.2.2) est 3-connexe.

Preuve : La démonstration se décompose en plusieurs étapes :

1. Tout d'abord, nous démontrerons que le graphe $\mathbb{Z}^2/\mathbb{L}$ demeure connexe lorsque l'on retire une paire quelconque de sommets non nuls distincts $P, Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$. Pour ce faire, notre stratégie consistera à construire trois chemins indépendants (c'est-à-dire des chemins qui ne partagent aucun sommet, hormis leurs extrémités) entre n'importe quel point non nul du groupe quotient $\mathbb{Z}^2/\mathbb{L}$ et $0_{\mathbb{Z}^2/\mathbb{L}}$. Cette construction garantit que toute paire de points non nuls distincts, différents de P et Q , pourra être reliée par un chemin n'empruntant ni P , ni Q .
2. Ensuite, nous établirons la connexité du graphe $\mathbb{Z}^2/\mathbb{L} - \{0_{\mathbb{Z}^2/\mathbb{L}}, Q\}$ pour tout point non nul $Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$. Pour y parvenir, nous présenterons un isomorphisme entre ce graphe et un graphe dont la connexité est déjà établie.

Étape 1 : Construction de chemins indépendants entre $R \neq 0_{\mathbb{Z}^2/\mathbb{L}}$ et $0_{\mathbb{Z}^2/\mathbb{L}}$

Les sommets du graphe sur lequel nous travaillons sont dans $\mathbb{Z}^2/\mathbb{L}$ où $\mathbb{L}$ est défini par $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix}$.

Dans le reste de la preuve, nous ferons référence à la classe d'équivalence $\begin{pmatrix} x \\ y \end{pmatrix} \bmod \mathbb{L}$ simplement en mentionnant l'un de ses représentants : $\begin{pmatrix} x \\ y \end{pmatrix}$.

Pour garantir que les chemins construits n'aient aucun sommet en commun (hormis leurs extrémités), il faut qu'ils passent par des représentants de classes d'équivalences distinctes.

Notons $\mathcal{P}$ le domaine fondamental de $\mathbb{L}$ associé à la base $\left\{ \begin{pmatrix} 2r \\ 0 \end{pmatrix}, \begin{pmatrix} r \\ r \end{pmatrix} \right\}$. Il est défini par :

$$\mathcal{P} = \left\{ \lambda \begin{pmatrix} 2r \\ 0 \end{pmatrix} + \mu \begin{pmatrix} r \\ r \end{pmatrix} \mid 0 \leq \lambda, \mu < 1 \right\}$$

Chaque élément de $\mathbb{Z}^2$ peut s'écrire de façon unique comme la somme d'un élément de $\mathbb{L}$ et d'un élément de $\mathbb{Z}^2 \cap \mathcal{P}$. Cet ensemble forme un ensemble complet de représentants pour les classes d'équivalences de $\mathbb{Z}^2 / \mathbb{L}$. Ainsi, construire des chemins indépendants revient à construire des chemins ne passant pas par les mêmes points de $\mathbb{Z}^2 \cap \mathcal{P}$.

Fixons $R \in \mathbb{Z}^2 \cap \mathcal{P}$ le représentant d'une classe d'équivalence non nulle. Comme les éléments de $\mathbb{Z}^2 \cap \mathcal{P}$ sont les vecteurs à coordonnées entières $\begin{pmatrix} x \\ y \end{pmatrix}$ satisfaisant les conditions $0 \leq y < r$ et $y \leq x < y + 2r$, on écrit $R = \begin{pmatrix} a \\ b \end{pmatrix}$ avec $0 \leq b < r$ et $b \leq a < b + 2r$.

Ici, il est important de noter que puisque $R \bmod \mathbb{L} \neq 0_{\mathbb{Z}^2 / \mathbb{L}}$ alors $R \neq \begin{pmatrix} 0 \\ 0 \end{pmatrix}$, donc en particulier $a > 0$ car $a \geq b \geq 0$.

Maintenant, nous allons détailler la construction de trois chemins indépendants reliant $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et R . Ces chemins prendront des formes différentes selon les valeurs spécifiques de a et b ; nous distinguerons donc différents cas. Afin d'en faciliter la compréhension, nous proposerons accompagnerons les descriptions précises des chemins d'illustrations permettant de visualiser leur construction.

Cas 1.1: $0 < a \leq r$ et $0 \leq b < r - 1$ (voir Figure 4.6)

- **Chemin 1 (chemin vert dans Figure 4.6):** Ce chemin part de l'origine et avance vers la droite le long de l'axe des abscisses jusqu'au sommet $\begin{pmatrix} a \\ 0 \end{pmatrix}$. Ensuite, il monte verticalement depuis $\begin{pmatrix} a \\ 0 \end{pmatrix}$ jusqu'à atteindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.
- **Chemin 2 (traits rouges pleins et pointillés sur la Figure 4.6) :** Ce chemin commence en $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} r \\ r \end{pmatrix}$ modulo $\mathbb{L}$. Il descend jusqu'à $\begin{pmatrix} r \\ b \end{pmatrix}$,

puis avance vers la gauche depuis ce point pour rejoindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.

- **Chemin 3 (pointillé bleu sur la Figure 4.6) :** Ce chemin démarre de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} 2r \\ 0 \end{pmatrix}$ modulo $\mathbb{L}$. Il monte ensuite jusqu'à $\begin{pmatrix} 2r \\ b+1 \end{pmatrix}$, puis avance vers la droite jusqu'à $\begin{pmatrix} a \\ b+1 \end{pmatrix}$ et enfin descend d'un cran pour atteindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.

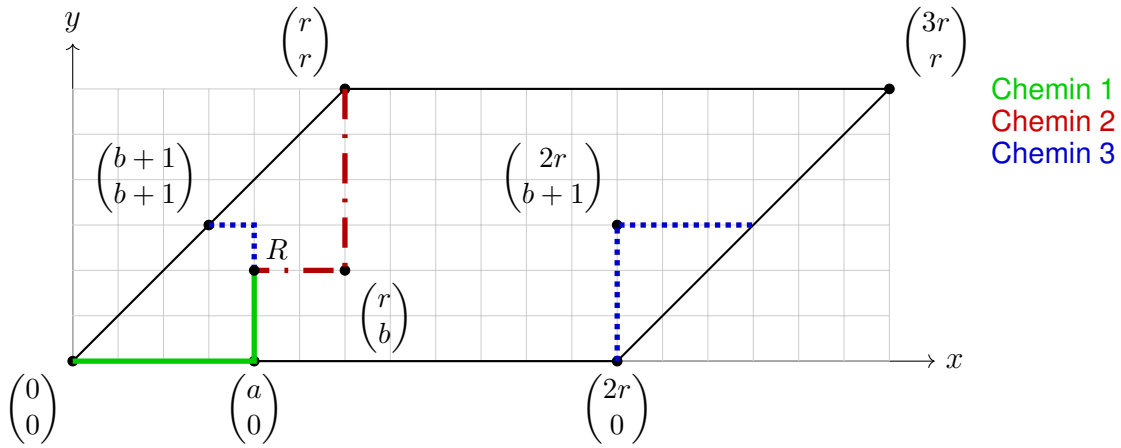


Figure 4.6: Trois chemins indépendants entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et $R = \begin{pmatrix} a \\ b \end{pmatrix}$ sur le tore $\mathbb{Z}^2/\mathbb{L}$. Ici $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix}$ avec $r = 6$, $a = 3$ et $b = 2$, respectant les conditions $0 < a \leq r$ et $0 \leq b < r - 1$.

Cas 1.2 : $0 < a \leq r$ et $b = r - 1$ (voir Figure 4.7)

- **Chemin 1 (trait vert sur la Figure 4.7) :** Ce chemin commence à l'origine $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, se déplace vers la droite le long de l'axe des abscisses jusqu'à $\begin{pmatrix} a \\ 0 \end{pmatrix}$, puis monte verticalement pour atteindre $R = \begin{pmatrix} a \\ r-1 \end{pmatrix}$.
- **Chemin 2 (traits rouges pleins et pointillés sur la Figure 4.7) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} r \\ r \end{pmatrix}$ modulo $\mathbb{L}$. Il descend ensuite de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ jusqu'à $\begin{pmatrix} r \\ r-1 \end{pmatrix}$, avant de se déplacer vers la gauche pour arriver à $R = \begin{pmatrix} a \\ r-1 \end{pmatrix}$.

- **Chemin 3 (pointillé bleu sur la Figure 4.7) :** Ce chemin commence à $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} 2r \\ 0 \end{pmatrix}$ modulo $\mathbb{L}$. Il monte ensuite jusqu'à $\begin{pmatrix} 2r \\ r-1 \end{pmatrix}$, puis avance vers la droite jusqu'à $\begin{pmatrix} a \\ r-1 \end{pmatrix}$.

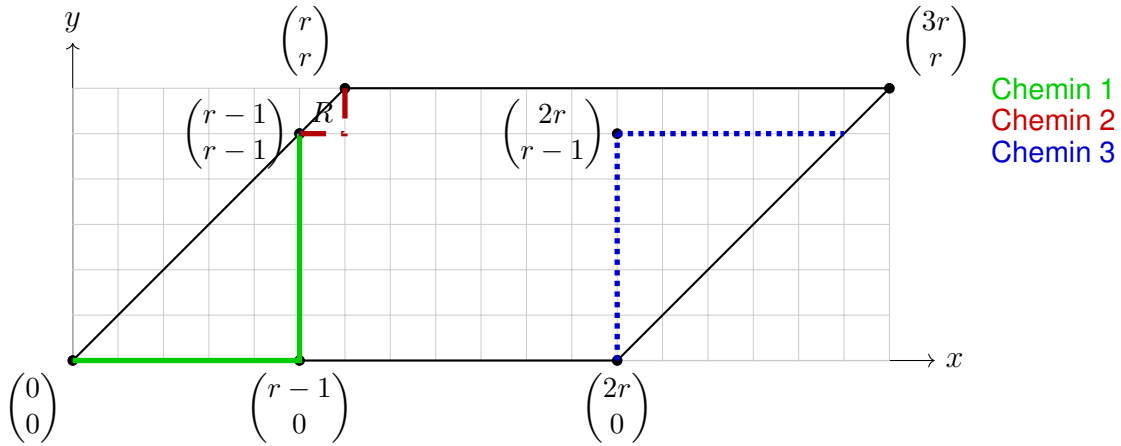


Figure 4.7: Trois chemins indépendants entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et $R = \begin{pmatrix} a \\ b \end{pmatrix}$ sur le tore $\mathbb{Z}^2/\mathbb{L}$. Ici $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix}$ avec $r = 6$ et $a = b = r - 1$, satisfaisant $r - 1 \leq a \leq r$ et $b = r - 1$.

Cas 2: $r < a < 2r$: (voir Figure 4.8)

- **Chemin 1 (trait vert sur la Figure 4.8) :** Ce chemin part de l'origine $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et avance vers la droite le long de l'axe des abscisses jusqu'à atteindre $\begin{pmatrix} a \\ 0 \end{pmatrix}$. De là, il monte verticalement pour atteindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.
- **Chemin 2 (traits rouges pleins et pointillés sur la Figure 4.8) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} r \\ r \end{pmatrix}$ modulo $\mathbb{L}$. Il descend d'abord jusqu'à $\begin{pmatrix} r \\ r-1 \end{pmatrix}$, puis se déplace vers la droite jusqu'à $\begin{pmatrix} a \\ r-1 \end{pmatrix}$. Enfin, il descend jusqu'à $R = \begin{pmatrix} a \\ b \end{pmatrix}$.
- **Chemin 3 (pointillé bleu sur la Figure 4.8) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est

équivalent à $\begin{pmatrix} 2r \\ 0 \end{pmatrix}$ modulo $\mathbb{L}$. Il monte ensuite verticalement jusqu'à $\begin{pmatrix} 2r \\ b \end{pmatrix}$. Puis, il se déplace vers la gauche jusqu'à atteindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.

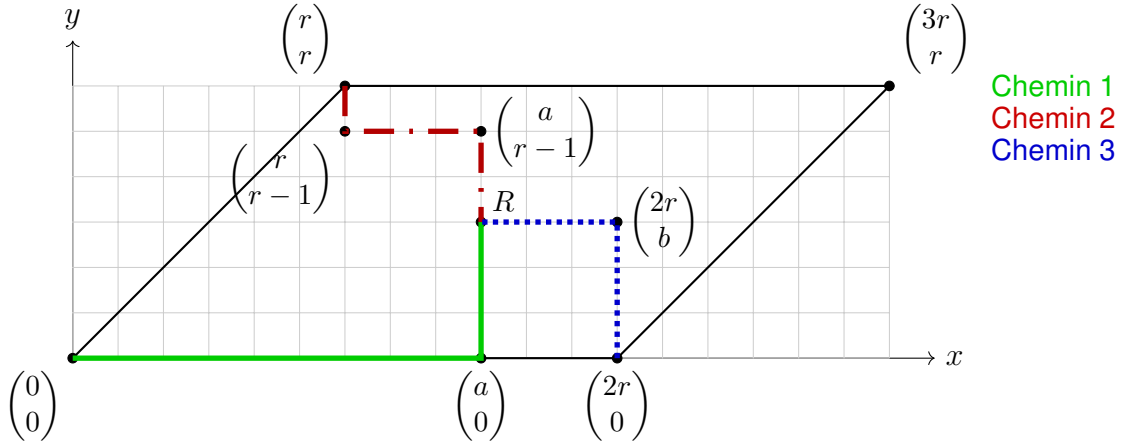


Figure 4.8: Trois chemins indépendants entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et $R = \begin{pmatrix} a \\ b \end{pmatrix}$ sur le tore $\mathbb{Z}^2/\mathbb{L}$. Ici $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix}$ avec $r = 6$, $a = 9$ et $b = 3$, satisfaisant les conditions $r < a < 2r$ et $0 \leq b < r$.

Cas 3 : $2r \leq a \leq 3r - 2$

Si a peut être égal à $2r$, il nous faut alors adapter la forme du premier et du troisième chemin. Par ailleurs, puisque $2r \leq a < b + 2r$ alors b doit être strictement positif.

Cas 3.1 : $2r \leq a \leq 3r - 2$ et $0 < b < r - 1$: (Figure 4.9)

- **Chemin 1 (chemin vert dans Figure 4.9):** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$. Ensuite, il se déplace vers la droite jusqu'à $\begin{pmatrix} 2r-1 \\ 0 \end{pmatrix}$. Puis, il monte ensuite verticalement jusqu'à $\begin{pmatrix} 2r-1 \\ b \end{pmatrix}$ avant de continuer vers la droite pour atteindre $R = \begin{pmatrix} a \\ b \end{pmatrix}$.
- **Chemin 2 (traits rouges pleins et pointillés dans la Figure 4.9) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} r \\ r \end{pmatrix}$ modulo $\mathbb{L}$. Il descend ensuite jusqu'à $\begin{pmatrix} r \\ r-1 \end{pmatrix}$, puis se déplace vers la droite jusqu'à $\begin{pmatrix} a \\ r-1 \end{pmatrix}$ et enfin descend verti-

calement jusqu'à $R = \begin{pmatrix} a \\ b \end{pmatrix}$.

- **Chemin 3 (pointillé bleu sur la Figure 4.9) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} 2r \\ 0 \end{pmatrix}$ modulo $\mathbb{L}$. Il monte ensuite verticalement jusqu'à $\begin{pmatrix} 2r \\ b-1 \end{pmatrix}$, puis se déplace vers la droite jusqu'à $\begin{pmatrix} a \\ b-1 \end{pmatrix}$ et puis effectue un pas vers le haut pour atteindre $\begin{pmatrix} a \\ b \end{pmatrix}$.

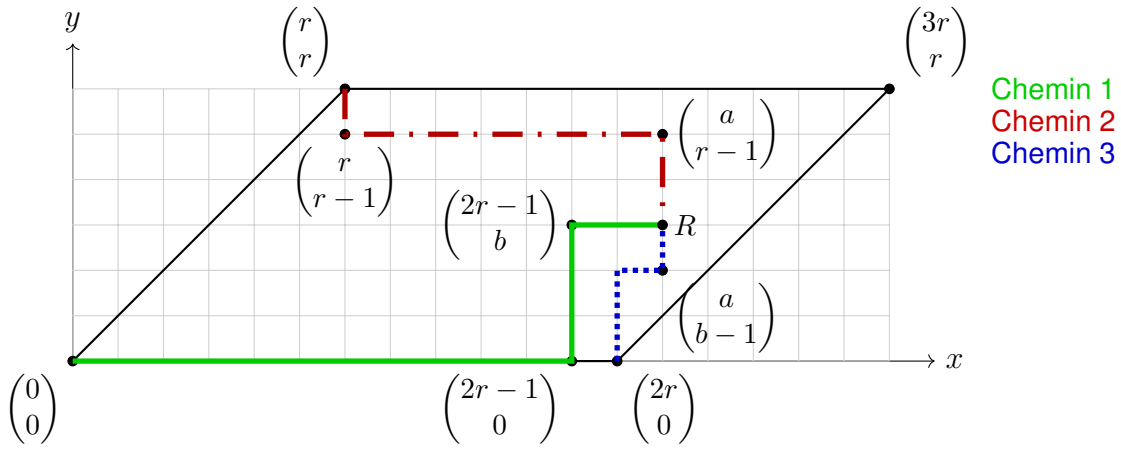


Figure 4.9: Trois chemins indépendants entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et $R = \begin{pmatrix} a \\ b \end{pmatrix}$ sur le tore $\mathbb{Z}^2/\mathbb{L}$. Ici $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix}$ avec $r = 6$, $a = 13$ et $b = 3$, satisfaisant les conditions $2r \leq a \leq 3r - 2$ et $0 < b < r - 1$.

Cas 3.2 : $2r \leq a \leq 3r - 2$ et $b = r - 1 \geq 2$

- **Chemin 1 (chemin vert dans la Figure 4.10) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et se déplace vers la droite jusqu'à $\begin{pmatrix} r-1 \\ 0 \end{pmatrix}$. Il monte ensuite verticalement jusqu'à $\begin{pmatrix} r-1 \\ r-1 \end{pmatrix}$, puis se déplace vers la gauche jusqu'à atteindre $R = \begin{pmatrix} a \\ r-1 \end{pmatrix}$.
- **Chemin 2 (traits rouges pleins et pointillés dans la Figure 4.10) :** Ce chemin

part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} r \\ r \end{pmatrix}$ modulo $\mathbb{L}$. Il descend verticalement d'abord jusqu'à $\begin{pmatrix} r \\ r-1 \end{pmatrix}$, puis continue vers la droite jusqu'à $R = \begin{pmatrix} a \\ r-1 \end{pmatrix}$.

- **Chemin 3 (pointillé bleu sur la Figure 4.10) :** Ce chemin part de $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$, qui est équivalent à $\begin{pmatrix} 2r \\ 0 \end{pmatrix}$ modulo $\mathbb{L}$. Il monte d'abord jusqu'à $\begin{pmatrix} 2r \\ r-2 \end{pmatrix}$, puis se déplace vers la droite jusqu'à $\begin{pmatrix} a \\ r-2 \end{pmatrix}$, avant de monter d'un cran pour atteindre $\begin{pmatrix} a \\ r-1 \end{pmatrix}$.

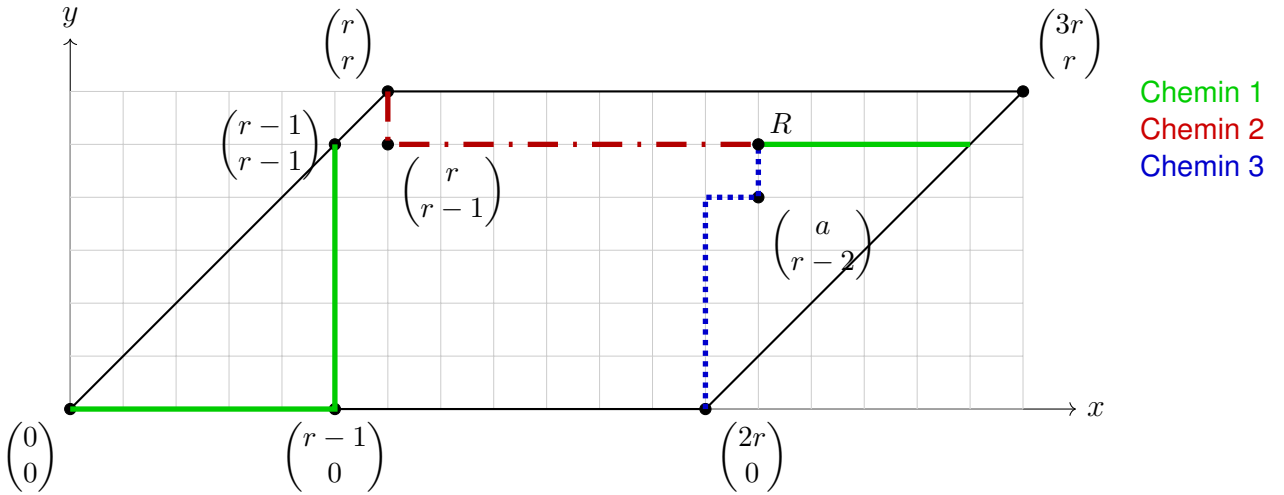


Figure 4.10: Trois chemins indépendants entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et $R = \begin{pmatrix} a \\ b \end{pmatrix}$ sur le tore $\mathbb{Z}^2/\mathbb{L}$. Ici $\mathbb{L} = \mathbb{Z} \begin{pmatrix} 2r \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix}$ avec $r = 6$, $a = 13$ et $b = 5$, satisfaisant les conditions $2r \leq a \leq 3r - 2$ et $b = r - 1$.

Dans tous les cas, en utilisant les points de $\mathbb{Z}^2 \cap \mathcal{P}$, nous avons construit entre $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ et R , trois chemins qui ne partagent aucun sommet, à l'exception de leurs extrémités.

Par conséquent, pour n'importe quel ensemble de quatre points non nuls distincts (P, Q, R, S) de $\mathbb{Z}^2/\mathbb{L}$, il existe toujours au moins un chemin entre R et S qui évite complètement P et Q . Cela signifie que le graphe $\mathbb{Z}^2/\mathbb{L} - \{P, Q\}$ est connexe.

Deuxième étape : Démontrer que $\mathbb{Z}^2/\mathbb{L} - \{0_{\mathbb{Z}^2/\mathbb{L}}, Q\}$ est connexe pour $Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$

Soit $Q \in \mathbb{Z}^2/\mathbb{L} \setminus \{0\}$. Puisque l'ordre du groupe $|\mathbb{Z}^2/\mathbb{L}|$ vaut $2r^2$ (avec $r \geq 3$), on peut toujours choisir un sommet $P \in \mathbb{Z}^2/\mathbb{L}$ tel que $P \neq 0_{\mathbb{Z}^2/\mathbb{L}}$ et $P \neq -Q$.

L'application $\psi : T \in \mathbb{Z}^2/\mathbb{L} \setminus \{0_{\mathbb{Z}^2/\mathbb{L}}, Q\} \mapsto T + P \in \mathbb{Z}^2/\mathbb{L} \setminus \{P, P + Q\}$ induit un isomorphisme de graphes.

Grâce à notre résultat précédent, nous savons que le graphe $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ privé de deux points non nuls distincts reste connexe. Comme $P \neq 0_{\mathbb{Z}^2/\mathbb{L}}$, $P + Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$ et $P + Q \neq P$, le graphe $\mathbb{Z}^2/\mathbb{L} - \{P, P + Q\}$ est connexe. Par conséquent, le graphe $\mathbb{Z}^2/\mathbb{L} - \{0_{\mathbb{Z}^2/\mathbb{L}}, Q\}$, qui lui est isomorphe, l'est également.

En résumé, nous avons montré que le graphe $\mathbb{Z}^2/\mathbb{L}$ reste connexe même après la suppression de n'importe quelle paire de sommets distincts. Cela couvre les deux scénarios possibles :

- la suppression de deux sommets P et Q qui sont distincts et non nuls.
- la suppression de $0_{\mathbb{Z}^2/\mathbb{L}}$ et d'un sommet non nul Q .

En conclusion, $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ est 3-connexe. □

Non-équivalence des codes :

Maintenant que nous avons montré que les graphes sous-jacents aux deux familles de codes sont 3-connexes, nous pouvons prouver que les codes ne sont pas CGP-équivalents en utilisant le théorème de Whitney.

Proposition 4.7.4. *Pour $r \geq 3$, le code surface 2D à stabilisateurs de poids quatre de distance minimale $2r$ optimal (dont les conditions aux bords sont déterminées par les vecteurs $\begin{pmatrix} r \\ r \end{pmatrix}, \begin{pmatrix} r \\ -r \end{pmatrix}$) [7], n'est pas CGP-équivalent au code GB($1 + X, 1 + X^{2r-1}, 2r^2$).*

Preuve : Raisonnons par l'absurde et supposons que ces deux codes soient CGP-équivalents. Ces deux codes sont des CSS construits à partir des graphes de Cayley suivants :

- $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$, où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ -r \end{pmatrix}$
- $(\mathbb{Z}/2r^2\mathbb{Z}, \overline{1}, \overline{2r-1})$.

Comme mentionné dans la section 4.6.2 et la proposition 4.6.1, leur CGP-équivalence impliquerait notamment l'existence d'une bijection entre leurs arêtes, préservant les cycles simples. Nous avons déjà démontré dans les Lemmes 4.7.2 et 4.7.3 que ces deux graphes sont 3-connexes. Puisqu'aucun de ces deux graphes ne contient de boucles ou de multi-arêtes, le théorème de Whitney implique ils sont isomorphes.

Nous allons maintenant expliquer comment l'existence de cet isomorphisme de graphe conduit à une contradiction, en suivant le plan ci-dessous :

- D'abord, nous prouverons que l'isomorphisme de graphes entre $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$ et $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ induit un isomorphisme de groupes entre $\mathbb{Z}/2r^2\mathbb{Z}$ et $\mathbb{Z}^2/\mathbb{L}$.
- Ensuite, nous démontrerons que ces deux groupes ne peuvent pas être isomorphes, étant donné que $\mathbb{Z}^2/\mathbb{L}$ ne contient pas d'éléments d'ordre $2r^2$.

Commençons par montrer que les groupes sous-jacents sont isomorphes.

Étape 1 : De l'isomorphisme de graphes à l'isomorphisme de groupes

Prenons $h : (\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1}) \rightarrow (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ un isomorphisme entre les deux graphes. Considérons ψ l'application définie, à partir de h , comme suit :

$$\begin{aligned} \psi : \mathbb{Z}/2r^2\mathbb{Z} &\rightarrow \mathbb{Z}^2/\mathbb{L} \\ \bar{T} &\mapsto h(\bar{T}) - h(\bar{0}) \end{aligned}$$

ψ est un isomorphisme de graphes :

Sa bijectivité est garantie par celle de h . De plus, ψ préserve les relations d'adjacence. En effet :

$\bar{x}$ et $\bar{y}$ sont voisins dans $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$

$$\iff h(\bar{x}) \text{ et } h(\bar{y}) \text{ sont voisins dans } (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$$

$$\iff h(\bar{x}) - h(\bar{y}) \in \left\{ \pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L} \right\}$$

$$\iff \psi(\bar{x}) - \psi(\bar{y}) \in \left\{ \pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L} \right\}$$

$$\iff \psi(\bar{x}) \text{ et } \psi(\bar{y}) \text{ sont voisins dans } (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$$

ψ est un morphisme de groupes :

Notons $\psi(\bar{1}) = U$.

On va démontrer que $\forall k \in \mathbb{Z}, \psi(\bar{k}) = kU$. On découpe la preuve en deux étapes. D'abord on va montrer par récurrence que cette égalité est vraie pour tous les entiers naturels. Ensuite, on va montrer que l'égalité s'étend aux entiers négatifs.

Commençons par montrer que $\forall k \in \mathbb{N}$, la proposition $P(k) : \psi(\bar{k}) = kU$ est vraie.

$P(0)$ et $P(1)$ sont vraies, puisque $\psi(\bar{0}) = 0_{\mathbb{Z}^2/\mathbb{L}}$ et $\psi(\bar{1}) = U$.

Supposons maintenant qu'il existe un entier $k \geq 1$ pour lequel la propriété $P(t)$ soit vérifiée pour chaque entier $0 \leq t \leq k$, c'est-à-dire $\psi(\bar{t}) = tU$. Nous voulons montrer que $\psi(\overline{k+1}) = (k+1)U$.

On sait que ψ est un isomorphisme de graphes entre $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$ et $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$. Comme les sommets $\bar{0}$ et $\bar{1}$ sont voisins dans $\mathbb{Z}/2r^2\mathbb{Z}$, il suit que $\psi(\bar{1})$ est l'un des quatre voisins de $\psi(\bar{0}) = 0_{\mathbb{Z}^2/\mathbb{L}}$ dans $\mathbb{Z}^2/\mathbb{L}$, ce qui se traduit par :

$$U = \psi(\bar{1}) \in \left\{ \pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L} \right\}$$

Maintenant posons $V \in \mathbb{Z}^2/\mathbb{L}$ tel que : $\left\{ \pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L} \right\} = \{\pm U, \pm V\}$

Pour $r \geq 3$, les éléments $0_{\mathbb{Z}^2/\mathbb{L}}, \pm U, \pm V$ sont deux à deux distincts. La raison en est que $\begin{pmatrix} \pm 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ \pm 1 \end{pmatrix}, \begin{pmatrix} \pm 1 \\ \pm 1 \end{pmatrix}, \begin{pmatrix} \pm 2 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ \pm 2 \end{pmatrix}$ n'appartiennent pas à $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ -r \end{pmatrix}$

Puisque $\psi : (\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1}) \rightarrow (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix})$ est un isomorphisme de graphes et que les sommets $\bar{k}$ et $\overline{k+1}$ sont voisins dans $(\mathbb{Z}/2r^2\mathbb{Z}, \bar{1}, \overline{2r-1})$, alors $\psi(\overline{k+1})$ est

un voisin de $\psi(\overline{k}) = kU$ dans $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix})$. Cela implique en particulier que $\psi(\overline{k+1}) \in \{kU \pm U, kU \pm V\}$.

Nous allons démontrer que la seule valeur possible pour $\psi(\overline{k+1})$ est $(k+1)U$. Ceci prouvera alors que la propriété $P(k+1)$ est vraie.

Cas 1. $\psi(\overline{k+1}) = (k-1)U$:

Par hypothèse de récurrence, $\psi(\overline{k+1}) = (k-1)U = \psi(\overline{k-1})$. Or comme ψ est une application bijective, cela impliquerait que $\overline{k+1} = \overline{k-1}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$ et en particulier que $\overline{2} = \overline{0}$ dans $\mathbb{Z}/2r^2\mathbb{Z}$. Or, c'est impossible car $0 < 2 < 2r^2$, pour $r \geq 3$.

Cas 2. $\psi(\overline{k+1}) = kU \pm V$:

$\psi(\overline{k+1})$ est l'un des sommets d'un des deux 4-cycles de $(\mathbb{Z}^2/\mathbb{L}, U, V)$ suivants :

- $kU \rightarrow kU + V \rightarrow (k-1)U + V \rightarrow (k-1)U \rightarrow kU$
- $kU \rightarrow kU - V \rightarrow (k-1)U - V \rightarrow (k-1)U \rightarrow kU$

Puisque ψ réalise une bijection entre les sommets $\mathbb{Z}/2r^2\mathbb{Z}$ et $\mathbb{Z}^2/\mathbb{L}$ et qu'elle vérifie $\psi(\overline{k-1}) = (k-1)U$ et $\psi(\overline{k}) = kU$, nous pouvons exprimer un des deux 4-cycles sous une forme standard :

$$\psi(\overline{k}) \rightarrow \psi(\overline{k+1}) \rightarrow \psi(\overline{T}) \rightarrow \psi(\overline{k-1}) \rightarrow \psi(\overline{k})$$

où $T \in \mathbb{Z}$ est un entier tel que :

- $\psi(\overline{T}) \in \{(k-1)U + V, (k-1)U - V\}$.
- $\overline{T}$ est un voisin de $\overline{k-1}$ dans graphe $(\mathbb{Z}/2r^2\mathbb{Z}, \overline{1}, \overline{2r-1})$:

$$\overline{T} \in \{\overline{k-2}, \overline{k}, \overline{k+2r-2}, \overline{k-2r}\}$$

- $\overline{T}$ est un voisin de $\overline{k+1}$ dans le graphe $(\mathbb{Z}/2r^2\mathbb{Z}, \overline{1}, \overline{2r-1})$:

$$\overline{T} \in \{\overline{k}, \overline{k+2}, \overline{k+2r}, \overline{k-2r+2}\}$$

Nous allons examiner chacune des valeurs possibles pour $\overline{T} : \{\overline{k-2}, \overline{k}, \overline{k+2r-2}, \overline{k-2r}\}$. Notre objectif est de montrer que chacune d'entre elles mène à une contradiction.

Nous ne pouvons pas avoir $\overline{T} = \overline{k}$. Si c'était le cas, on aurait $(k-1)U \pm V = \psi(\overline{T}) = \psi(\overline{k}) = kU$. Cela impliquerait que $U = \pm V$, mais nous avons déjà établi que $U \neq \pm V$.

Nous ne pouvons pas non plus avoir $\overline{T} = \overline{k-2}$. Si c'était le cas, en considérant que $\overline{T}$ doit appartenir à l'ensemble $\{\overline{k+2}, \overline{k-2r+2}, \overline{k+2r}\}$, l'un des scénarios suivants aurait lieu dans $\mathbb{Z}/2r^2\mathbb{Z}$:

- $\overline{k-2} = \overline{k+2} \implies \overline{4} = \overline{0}$
- $\overline{k-2} = \overline{k-2r+2} \implies \overline{2r-4} = \overline{0}$
- $\overline{k-2} = \overline{k+2r} \implies \overline{2r+2} = \overline{0}$

Cependant, pour $r \geq 3$, on a $0 < 4 < 2r^2$ et $0 < 2r-4 < 2r+2 < 2r^2$. Donc $\overline{T} \neq \overline{k-2}$.

Nous ne pouvons pas avoir $\overline{T} = \overline{k-2r}$ non plus. Si c'était le cas, en considérant que $\overline{T}$ doit appartenir à l'ensemble $\{\overline{k+2}, \overline{k-2r+2}, \overline{k+2r}\}$, l'un des scénarios suivants aurait lieu dans $\mathbb{Z}/2r^2\mathbb{Z}$:

- $\overline{k-2r} = \overline{k+2} \implies \overline{2r+2} = \overline{0}$
- $\overline{k-2r} = \overline{k-2r+2} \implies \overline{2} = \overline{0}$
- $\overline{k-2r} = \overline{k+2r} \implies \overline{4r} = \overline{0}$

Cependant pour $r \geq 3$, on a $0 < 2 < 2r+2 < 4r < 2r^2$. Donc $\overline{T} \neq \overline{k-2r}$.

Enfin, on n'a pas non plus $\overline{T} = \overline{k+2r-2}$. Si c'était le cas, en considérant que $\overline{T}$ doit appartenir à l'ensemble $\{\overline{k+2}, \overline{k-2r+2}, \overline{k+2r}\}$, l'un des scénarios suivants aurait lieu dans $\mathbb{Z}/2r^2\mathbb{Z}$:

- $\overline{k+2r-2} = \overline{k+2} \implies \overline{2r-4} = \overline{0}$
- $\overline{k+2r-2} = \overline{k-2r+2} \implies \overline{4r-4} = \overline{0}$
- $\overline{k+2r-2} = \overline{k+2r} \implies \overline{2} = \overline{0}$.

Cependant pour $r \geq 3$, on a $0 < 2 \leq 2r-4 < 4r-4 < 2r^2$, qui prouve que l'assertion ne peut être vraie. D'où $\overline{T} \neq \overline{k+2r-2}$.

En conséquence, on ne peut pas avoir $\psi(\overline{k+1}) = kU \pm V$ puisque tous les scénarios conduisent à une contradiction.

En résumé, $\psi(\overline{k+1})$ appartient à $\{(k+1)U, (k-1)U, kU+V, kU-V\}$, mais ne peut être égal ni à $(k-1)U$, ni à $kU+V$, ni à $kU-V$. Donc, nécessairement $\psi(\overline{k+1}) = (k+1)U$,

ce qui prouve que $P(k+1)$ est vraie. En conséquence, $\forall k \in \mathbb{N}$, $\psi(\bar{k}) = kU$.

Maintenant, prouvons que pour $q \in \mathbb{N}$, l'égalité $\psi(\overline{-q}) = -qU$ est vraie aussi.

Dans $\mathbb{Z}/2r^2\mathbb{Z}$, $\overline{-q} = \overline{-1} \cdot \bar{q} = \overline{(2r^2 - 1)} \cdot \bar{q} = \overline{(2r^2 - 1)q}$. Ainsi, nous avons l'égalité :

$$\psi(\overline{-q}) = \psi(\overline{(2r^2 - 1)q}) = (2r^2 - 1)qU$$

Puisque $U \in \{\pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}\}$, où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ -r \end{pmatrix}$, il suit que $2rU = 0_{\mathbb{Z}^2/\mathbb{L}}$. En conséquence, $\psi(\overline{-q}) = (2r^2 - 1)qU = -qU$.

Nous avons prouvé que $\forall k \in \mathbb{Z}$, $\psi(\bar{k}) = kU$, donc ψ est un morphisme de groupes entre $\mathbb{Z}/2r^2\mathbb{Z}$ et $\mathbb{Z}^2/\mathbb{L}$.

Étape 2 : Prouvons qu'il n'y a pas d'isomorphismes de groupes

Nous avons montré que pour $r \geq 3$, $\psi : \mathbb{Z}/2r^2\mathbb{Z} \rightarrow \mathbb{Z}^2/\mathbb{L}$ est un isomorphisme de groupes. Or, dans $\mathbb{Z}/2r^2\mathbb{Z}$, il y a des éléments d'ordre $2r^2$ ($\bar{1}$ par exemple), tandis que dans $\mathbb{Z}^2/\mathbb{L}$ où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} r \\ r \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} r \\ -r \end{pmatrix}$, l'ordre de chaque élément est majoré par $2r < 2r^2$. En effet :

$$\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{Z}^2 \implies 2r \begin{pmatrix} x \\ y \end{pmatrix} = x \begin{pmatrix} 2r \\ 0 \end{pmatrix} + y \begin{pmatrix} 0 \\ 2r \end{pmatrix} \in \mathbb{L}$$

En conclusion, même s'ils ont les mêmes paramètres $[[4r^2, 2, 2r]]$, le code surface 2D optimal de distance minimale $2r$ introduit dans la Section 4.2.2) et le code $GB(1+X, 1+X^{2r-1}, 2r^2)$ ne sont pas CGP-équivalents. $\square$

4.7.4 Comparaison avec les codes de Kitaev standards

Maintenant, démontrons que les codes toriques de Kitaev et les codes $GB(1+X, 1+X^n, n^2)$ ne sont pas CGP-équivalents.

La preuve de leur non-équivalence repose sur les mêmes principes que la démonstration précédente. Pour cette raison, nous nous limiterons à présenter les théorèmes clés et les grandes lignes de la preuve, sans en détailler chaque étape.

3-connexité des graphes sous-jacents

Lemme 4.7.5. *Pour tout entier $n \geq 3$, le graphe de Cayley $(\mathbb{Z}/n^2\mathbb{Z}, \bar{1}, \bar{n})$, qui est associé au code Bicycle généralisé $GB(1 + X, 1 + X^n, n^2)$ est 3-connexe.*

Preuve : D'après la remarque 4.7.2, il suffit de reprendre la même méthode que dans la démonstration du Lemme 4.7.2 pour démontrer que le graphe $(\mathbb{Z}/n^2\mathbb{Z}, \bar{1}, \bar{n})$ reste connexe même après la suppression de deux sommets distincts. Cela confirme alors la triple-connexité du graphe. $\square$

Lemme 4.7.6. *Soient $n \geq 3$ un entier et $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 0 \\ n \end{pmatrix}$ un sous-réseau de $\mathbb{Z}^2$.*

Alors le graphe de Cayley $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$, associé au code torique de Kitaev de paramètres $[[2n^2, 2, n]]$, est 3-connexe.

Preuve : Nous suivons la même stratégie de démonstration que celle du Lemme 4.7.3 :

1. D'abord, nous prouvons que $\mathbb{Z}^2/\mathbb{L}$ reste connexe même après la suppression de deux sommets non nuls distincts $P, Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$. Nous y parvenons en construisant trois chemins indépendants entre $0_{\mathbb{Z}^2/\mathbb{L}}$ et n'importe quel point non nul du quotient $\mathbb{Z}^2/\mathbb{L}$.
2. Puis, nous établissons la connexité du graphe $\mathbb{Z}^2/\mathbb{L} - \{0_{\mathbb{Z}^2/\mathbb{L}}, Q\}$ pour tout point non nul $Q \neq 0_{\mathbb{Z}^2/\mathbb{L}}$, en montrant qu'il est isomorphe à un graphe connexe déjà connu.

Ainsi, ceci montre que le graphe $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ est 3-connexe. $\square$

Proposition 4.7.7. *Pour $n \geq 4$, le code torique de Kitaev dont les conditions aux bords sont données par le réseau $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 0 \\ n \end{pmatrix}$ et le code Bicycle généralisé $GB(1 + X, 1 + X^n, n^2)$ ne sont pas CGP-équivalents.*

Preuve : Raisonnons par l'absurde et supposons que ces deux codes sont CGP-équivalents.

Ces codes sont des codes CSS construits à partir des graphes de Cayley suivants :

- $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$, où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 0 \\ n \end{pmatrix}$
- $(\mathbb{Z}/n^2\mathbb{Z}, \bar{1}, \bar{n})$

Comme aucun de ces graphes ne contient de boucles ou de multi-arêtes, le même raisonnement que celui de la démonstration de la Proposition 4.7.4, permet de montrer que si ces codes étaient équivalents, leurs graphes sous-jacents seraient isomorphes :

$$(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}) \cong (\mathbb{Z}/n^2\mathbb{Z}, \bar{1}, \bar{n})$$

Cet isomorphisme de graphes induirait à son tour un isomorphisme de groupes entre $\mathbb{Z}/n^2\mathbb{Z}$ et $\mathbb{Z}^2/\mathbb{L}$. Cependant, ces deux groupes ne sont pas isomorphes car $\mathbb{Z}^2/\mathbb{L}$ ne contient aucun élément d'ordre n^2 . En effet, tout élément de $\mathbb{Z}^2/\mathbb{L}$ a un ordre qui divise n , lequel est strictement inférieur à n^2 :

$$\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{Z}^2 \implies n \begin{pmatrix} x \\ y \end{pmatrix} = x \begin{pmatrix} n \\ 0 \end{pmatrix} + y \begin{pmatrix} 0 \\ n \end{pmatrix} \in \mathbb{L} = \mathbb{Z} \begin{pmatrix} n \\ 0 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} 0 \\ n \end{pmatrix}$$

En conclusion, le code torique de Kitaev de longueur $2n^2$ et le code Bicycle généralisé $GB(1+X, 1+X^n, n^2)$ ne sont pas CGP-équivalents. $\square$

Maintenant, comparons les codes surface 2D optimaux pour les distances impaires (présentés dans la Section 4.2.2) aux codes de la famille $GB(1+X, 1+X^{2t+1}, t^2+(t+1)^2)$.

4.7.5 Comparaison avec les codes surface 2D optimaux pour la distance impaire

Contrairement aux deux familles précédentes, les codes $GB(1+X, 1+X^{2t+1}, t^2+(t+1)^2)$ sont équivalents (au sens traditionnel des codes quantiques) aux codes surface 2D, décrits dans la Section 4.2.2, ayant les paramètres $[[(2t+1)^2 + 1, 2, 2t+1]]$.

La proposition suivante formalise cette équivalence :

Proposition 4.7.8. *Pour $t \geq 1$, le code surface 2D à distance $2t+1$ optimal, défini sur le tore $\mathbb{R}^2/\mathbb{L}$ avec $\mathbb{L} = \mathbb{Z} \begin{pmatrix} t \\ t+1 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} t+1 \\ -t \end{pmatrix}$ et le code Bicycle généralisé $GB(1+X, 1+X^{2t+1}, t^2+(t+1)^2)$ sont équivalents.*

Preuve : Le code Bicycle généralisé $GB(1+X, 1+X^{2t+1}, t^2+(t+1)^2)$ (décrit en Section 4.3.2) est construit à partir du graphe $(\mathbb{Z}/(t^2+(t+1)^2)\mathbb{Z}, \bar{1}, \overline{2t+1})$.

Le code de surface 2D optimal à distance $2t+1$ (présenté en Section 4.2.2) est construit à partir du graphe de Cayley $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$.

Notre objectif est de montrer que ces deux codes sont équivalents. Pour y parvenir, nous allons procéder en trois étapes :

- D'abord, nous construirons un isomorphisme entre les groupes sous-jacents, à savoir $\mathbb{Z}^2/\mathbb{L}$ et $\mathbb{Z}/(t^2+(t+1)^2)\mathbb{Z}$.
- Ensuite, nous démontrerons que cet isomorphisme de groupes induit un isomorphisme de graphes entre les deux graphes mentionnés ci-dessus.
- Enfin, grâce à cet isomorphisme, nous démontrerons que les sous-matrices génératrices des deux codes $(\mathbf{H}_X, \mathbf{H}_Z)$ et $(\mathbf{H}'_X, \mathbf{H}'_Z)$ satisfont les relations :

$$\begin{cases} \text{Vect}_{\mathbf{L}}(\mathbf{H}_X) = \text{Vect}_{\mathbf{L}}(\mathbf{H}'_X \mathbf{Q}) \\ \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}'_Z \mathbf{Q}) \end{cases}$$

où $\mathbf{Q}$ est une matrice de permutation. On en déduit alors que les codes sont CGP-équivalents, et donc équivalents au sens usuel pour les codes quantiques.

Étape 1 : Construction de l'isomorphisme de groupes

Soit η l'application définie par :

$$\begin{aligned} \eta : \quad \mathbb{Z}/(t^2+(t+1)^2)\mathbb{Z} &\rightarrow \mathbb{Z}^2/\mathbb{L} \\ \overline{T} := T \bmod (t^2+(t+1)^2) &\mapsto T \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} \end{aligned}$$

Commençons par montrer que l'application η est un isomorphisme de groupes qui envoie $\bar{1}$ sur $\begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}$ et $\overline{2t+1}$ sur $-\begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}$.

η est bien définie : Considérons deux entiers $k, q \in \mathbb{Z}$ tels que $\bar{k} = \bar{q}$ dans $\mathbb{Z}/(t^2+(t+1)^2)\mathbb{Z}$. Montrons que $\eta(\bar{k}) = \eta(\bar{q})$.

Par hypothèse, $\exists \lambda \in \mathbb{Z}$, $k - q = \lambda(t^2 + (t + 1)^2)$. Nous avons donc :

$$k \begin{pmatrix} 1 \\ 0 \end{pmatrix} - q \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \lambda \begin{pmatrix} t^2 + (t + 1)^2 \\ 0 \end{pmatrix} = \lambda \left((t + 1) \begin{pmatrix} t + 1 \\ -t \end{pmatrix} + t \begin{pmatrix} t \\ t + 1 \end{pmatrix} \right)$$

Donc $k \begin{pmatrix} 1 \\ 0 \end{pmatrix} - q \begin{pmatrix} 1 \\ 0 \end{pmatrix} \in \mathbb{L} = \mathbb{Z} \begin{pmatrix} t \\ t + 1 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} t + 1 \\ -t \end{pmatrix}$.

Ceci signifie que $\eta(\bar{k}) = k \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} = q \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} = \eta(\bar{q})$.

η est un morphisme de groupes : En effet, pour $k, q \in \mathbb{Z}$:

$$\begin{aligned} \eta(\bar{k} + \bar{q}) &= \eta(\overline{k + q}) = (k + q) \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} = k \begin{pmatrix} 1 \\ 0 \end{pmatrix} + q \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} \\ &= k \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} + q \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L} \\ &= \eta(\bar{k}) + \eta(\bar{q}) \end{aligned}$$

De plus, $\eta(\bar{1}) = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}$ et $\eta(\overline{2t + 1}) = - \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}$.

η est bijective : D'abord, notons que les ensembles $\mathbb{Z}^2/\mathbb{L}$ et $\mathbb{Z}/(t^2 + (t + 1)^2)\mathbb{Z}$ ont le même cardinal. En effet, puisque $\mathbb{L}$ est un sous-réseau de $\mathbb{Z}^2$, l'ordre du groupe $\mathbb{Z}^2/\mathbb{L}$ est donné par le déterminant de $\mathbb{L}$:

$$|\mathbb{Z}^2/\mathbb{L}| = \det(\mathbb{L}) = \det \begin{pmatrix} t + 1 & t \\ -t & t + 1 \end{pmatrix} = (t + 1)^2 + t^2 = |\mathbb{Z}/(t^2 + (t + 1)^2)\mathbb{Z}|$$

Ensuite, montrons que η est une application injective. Considérons un entier $k \in \mathbb{Z}$ tel que $\eta(\bar{k}) = 0_{\mathbb{Z}^2/\mathbb{L}}$. Montrons que $\bar{k} = \bar{0}$ dans $\mathbb{Z}/(t^2 + (t + 1)^2)\mathbb{Z}$.

Par définition de η , $k \begin{pmatrix} 1 \\ 0 \end{pmatrix} \in \mathbb{L} = \mathbb{Z} \begin{pmatrix} t \\ t + 1 \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} t + 1 \\ -t \end{pmatrix}$. Donc, il existe $\lambda, \mu \in \mathbb{Z}$ tels que $k \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \lambda \begin{pmatrix} t \\ t + 1 \end{pmatrix} + \mu \begin{pmatrix} t + 1 \\ -t \end{pmatrix}$. Ceci conduit aux égalités d'entiers suivantes :

$$k = \lambda t + \mu(t + 1) \quad \text{et} \quad \mu t = \lambda(t + 1)$$

Comme t et $t + 1$ sont premiers entre eux, on a que t divise λ et que $t + 1$ divise μ .

Donc, il existe des entiers p, q tels que $\lambda = pt$ et $\mu = q(t+1)$. En remplaçant par leur valeur dans les équations précédentes, on obtient que :

$$k = pt^2 + q(t+1)^2 \quad \text{et} \quad qt(t+1) = pt(t+1)$$

Or $t \geq 1$, donc en divisant par $t(t+1)$ dans la 2^{ème} égalité, on trouve que $p = q$. Cela implique que $k = p(t^2 + (t+1)^2)$, c'est-à-dire $\bar{k} = \bar{0}$ dans $\mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}$.

En conclusion, η étant un morphisme injectif entre deux groupes de même cardinal fini, c'est un isomorphisme de groupes.

Étape 2 : Preuve que η induit un isomorphisme de graphes :

L'application η définie par :

$$\begin{aligned} \eta : \mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z} &\rightarrow \mathbb{Z}^2/\mathbb{L} \\ \bar{T} &\mapsto T \begin{pmatrix} 1 \\ 0 \end{pmatrix} \pmod{\mathbb{L}} \end{aligned}$$

est un isomorphisme de groupes, envoyant $\bar{1}$ sur $\begin{pmatrix} 1 \\ 0 \end{pmatrix} \pmod{\mathbb{L}}$ et $\overline{2t+1}$ sur $-\begin{pmatrix} 0 \\ 1 \end{pmatrix} \pmod{\mathbb{L}}$.

De plus, η induit aussi un isomorphisme de graphes entre $(\mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}, \bar{1}, \overline{2t+1})$ et $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \pmod{\mathbb{L}}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \pmod{\mathbb{L}})$.

En effet, deux sommets $\bar{T}$ et $\bar{S}$ sont voisins dans $(\mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}, \bar{1}, \overline{2t+1})$ si et seulement si $\bar{T} - \bar{S} \in \{\pm\bar{1}, \pm\overline{2t+1}\}$. Puisque η est un isomorphisme, cette condition est équivalente à :

$$\eta(\bar{T} - \bar{S}) = \eta(\bar{T}) - \eta(\bar{S}) \in \left\{ \pm \begin{pmatrix} 1 \\ 0 \end{pmatrix} \pmod{\mathbb{L}}, \pm \begin{pmatrix} 0 \\ 1 \end{pmatrix} \pmod{\mathbb{L}} \right\}$$

prouvant ainsi que η préserve les relations d'adjacence.

Étape 3 : Preuve de l'équivalence des codes

Pour prouver l'équivalence des deux codes, au sens usuel, nous allons montrer qu'ils sont CGP-équivalents, ce qui est une condition suffisante.

Pour démontrer l'équivalence CGP de ces codes, il suffit de montrer que leurs sous-matrices génératrices respectives $(\mathbf{H}_X, \mathbf{H}_Z)$ et $(\mathbf{H}'_X, \mathbf{H}'_Z)$ sont reliées par les relations : $\text{Vect}_{\mathbf{L}}(\mathbf{H}_X) = \text{Vect}_{\mathbf{L}}(\mathbf{H}'_X \mathbf{Q})$ et $\text{Vect}_{\mathbf{L}}(\mathbf{H}_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}'_Z \mathbf{Q})$ où $\mathbf{Q}$ est une matrice de permutation.

Pour y arriver, nous procéderons de la manière suivante :

1. Nous commencerons par réinterpréter les sous-matrices génératrices comme des matrices d'incidence de graphes. Ainsi, $\mathbf{H}_X$ et $\mathbf{H}_Z$ seront respectivement les matrices sommet-arête et face-arête d'un graphe $\mathcal{G}$, tandis que $\mathbf{H}'_X$ et $\mathbf{H}'_Z$ joueront le même rôle pour un graphe $\mathcal{G}'$.
2. Nous renuméroterons ensuite les sommets et arêtes du graphe $\mathcal{G}'$ de sorte que les matrices d'incidence induites, $\mathbf{H}''_X$ et $\mathbf{H}''_Z$, vérifient $\mathbf{H}_X = \mathbf{H}''_X$ et $\text{Vect}_{\mathbf{L}}(\mathbf{H}_Z) = \text{Vect}_{\mathbf{L}}(\mathbf{H}''_Z)$.
3. Enfin, nous montrerons que les couples $(\mathbf{H}''_X, \mathbf{H}''_Z)$ et $(\mathbf{H}'_X, \mathbf{H}'_Z)$ satisfont les relations : $\mathbf{H}''_X = \mathbf{P} \mathbf{H}'_X \mathbf{Q}$ et $\mathbf{H}''_Z = \mathbf{R} \mathbf{H}'_Z \mathbf{Q}$, où $\mathbf{P}, \mathbf{Q}, \mathbf{R}$ sont des matrices de permutation. De ces relations, nous pourrions alors déduire l'équivalence annoncée.

Étape 3.1 : Réinterprétation des sous-matrices génératrices

Posons $\mathbf{A} = \text{Circ}(1 + X, t^2 + (t + 1)^2)$ et $\mathbf{B} = \text{Circ}(1 + X^{2t+1}, t^2 + (t + 1)^2)$.

D'après les travaux de la section 4.4.2, les sous-matrices génératrices $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$ du code $GB(1 + X, 1 + X^{2t+1}, t^2 + (t + 1)^2)$ correspondent respectivement aux matrices d'incidence sommet-arête et face-arête du graphe $\mathcal{G} = (\mathbb{Z}/(t^2 + (t + 1)^2)\mathbb{Z}, \overline{1}, \overline{2t+1})$, sur lequel les sommets et arêtes ont été numérotés comme suit.

Pour chaque $k \in \llbracket 0, t^2 + (t + 1)^2 - 1 \rrbracket$:

- Le k -ième sommet de $\mathcal{G}$ est $S_k := \overline{k}$.
- La k -ième arête de $\mathcal{G}$ est $h_k = \{\overline{k}, \overline{k+1}\}$.
- La $(n + k)$ -ième arête de $\mathcal{G}$ est $\{\overline{k}, \overline{k + (2t + 1)}\}$.

De même, les sous-matrices génératrices $\mathbf{H}'_X$ et $\mathbf{H}'_Z$ du code surface 2D optimal de paramètres $[[(2t+1)^2, 2, 2t+1]]$, introduit dans la section 4.2.2, peuvent être vues comme des matrices d'incidence sommet–arête et face–arête du graphe

$$\mathcal{G}' = (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$$

où $\mathbb{L} = \mathbb{Z} \begin{pmatrix} t+1 \\ t \end{pmatrix} \oplus \mathbb{Z} \begin{pmatrix} -t \\ t+1 \end{pmatrix}$. Toutefois, ici l'ordre choisi sur les sommets et les arêtes de $\mathcal{G}'$ n'est pas spécifié.

Étape 3.2 : Renumérotation des sommets, arêtes et faces de $\mathcal{G}'$

Comme nous l'avons vu précédemment, les graphes $\mathcal{G}$ et $\mathcal{G}'$ ont le même nombre de sommets et d'arêtes. Pour construire les matrices $\mathbf{H}''_X$ et $\mathbf{H}''_Z$, on renumérote les sommets et arêtes de $\mathcal{G}'$, au moyen de la bijection η .

Pour chaque $i \in \llbracket 0, t^2 + (t+1)^2 - 1 \rrbracket$, le i -ème sommet de $\mathcal{G}'$ sera $T_i = \eta(\bar{i})$, l'image par η du i -ième sommet de $\mathcal{G}$. De même, si la k -ième arête de $\mathcal{G}$ est $\{\bar{x}, \bar{y}\}$, alors la k -ième arête de $\mathcal{G}'$ sera $\{\eta(\bar{x}), \eta(\bar{y})\}$.

Notons les matrices d'incidence sommet–arête $\mathbf{H}''_X$ et face–arête $\mathbf{H}''_Z$ de $\mathcal{G}'$ correspondant à ce choix d'ordre.

Par construction, les matrices d'incidence sommet–arête des deux graphes sont identiques, c'est-à-dire $\mathbf{H}''_X = \mathbf{H}_X$. De plus, les lignes des matrices $\mathbf{H}_Z$ et $\mathbf{H}''_Z$ engendrent le même espace, c'est-à-dire $\text{Vect}_{\mathbb{L}}(\mathbf{H}''_Z) = \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z)$.

Comme ces espaces sont respectivement engendrés par les lignes des matrices $\mathbf{H}_Z$ et $\mathbf{H}''_Z$, il suffit de montrer que leurs ensembles de lignes coïncident.

Commençons par rappeler que, pour les ordres choisis sur les arêtes de $\mathcal{G}$ (respectivement $\mathcal{G}'$), les lignes de $\mathbf{H}_Z$ (respectivement $\mathbf{H}''_Z$) correspondent exactement aux vecteurs caractéristiques des faces du graphe $\mathcal{G}$ (respectivement $\mathcal{G}'$).

Maintenant, montrons que $\{\text{Lignes de } \mathbf{H}_Z\} \subseteq \{\text{Lignes de } \mathbf{H}''_Z\}$:

Soit $\mathbf{v} \in \mathbb{F}_2^{2(t^2+(t+1)^2)}$ une ligne de $\mathbf{H}_Z$. Puisque $\mathbf{H}_Z$ est la matrice d'incidence face-arête du graphe $(\mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}, \bar{1}, \overline{2t+1})$, il suit que $\mathbf{v}$ est le vecteur caractéristique d'une face de ce graphe. Une telle face $\mathcal{F}$ s'écrit alors :

$$\bar{T} \rightarrow \bar{T} + \bar{1} \rightarrow \bar{T} + \bar{1} + \overline{2t+1} \rightarrow \bar{T} + \overline{2t+1} \rightarrow \bar{T}$$

Pour l'ordre choisi sur les sommets et les arêtes de $\mathcal{G}'$, $\mathbf{v}$ correspond au vecteur caractéristique de $\eta(\mathcal{F})$, l'image de cette face par l'isomorphisme de graphes η :

$$\eta(\bar{T}) \rightarrow \eta(\bar{T} + \bar{1}) \rightarrow \eta(\bar{T} + \bar{1} + \overline{2t+1}) \rightarrow \eta(\bar{T} + \overline{2t+1}) \rightarrow \eta(\bar{T})$$

Puisque $\eta : q \bmod (t^2 + (t+1)^2) \mapsto \begin{pmatrix} q \\ 0 \end{pmatrix} \bmod \mathbb{L}$ est un morphisme de groupes envoyant $\overline{2t+1}$ sur $-\begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L}$, $\eta(\mathcal{F})$ se réécrit comme :

$$\begin{pmatrix} T \\ 0 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T+1 \\ 0 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T+1 \\ -1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T \\ -1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T \\ 0 \end{pmatrix} \bmod \mathbb{L}$$

Comme le graphe $\mathcal{G}' = (\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$ est non orienté, $\eta(\mathcal{F})$ correspond alors à la face :

$$\begin{pmatrix} T \\ -1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T+1 \\ -1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T+1 \\ 0 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T \\ 0 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} T \\ -1 \end{pmatrix} \bmod \mathbb{L}$$

Donc $\mathbf{v}$ est le vecteur caractéristique d'une face de $\mathcal{G}'$, ce qui implique que $\mathbf{v}$ est une ligne de $\mathbf{H}_Z''$.

Réciproquement, montrons que $\{\text{Lignes de } \mathbf{H}_Z''\} \subseteq \{\text{Lignes de } \mathbf{H}_Z\}$:

Soit $\mathbf{w}$ une ligne de $\mathbf{H}_Z''$. Par définition, $\mathbf{w} \in \mathbb{F}_2^{2(t^2+(t+1)^2)}$ est le vecteur caractéristique d'une face $\mathcal{F}_{\mathbb{L}}$ de $(\mathbb{Z}^2/\mathbb{L}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \bmod \mathbb{L}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \bmod \mathbb{L})$. Une telle face s'écrit comme :

$$\begin{pmatrix} x \\ y \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} x+1 \\ y \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} x+1 \\ y+1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} x \\ y+1 \end{pmatrix} \bmod \mathbb{L} \rightarrow \begin{pmatrix} x \\ y \end{pmatrix} \bmod \mathbb{L}$$

Puisque l'application $\eta : \mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z} \rightarrow \mathbb{Z}^2/\mathbb{L}$, $\bar{u} \mapsto \begin{pmatrix} u \\ 0 \end{pmatrix} \pmod{\mathbb{L}}$ est un isomor-

phisme, il existe un unique $\bar{k} \in \mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}$ tel que : $\eta(\bar{k}) = \begin{pmatrix} x \\ y \end{pmatrix} \pmod{\mathbb{L}}$.

De plus, comme $\eta(\overline{2t+1}) = -\begin{pmatrix} 0 \\ 1 \end{pmatrix} \pmod{\mathbb{L}}$, nous pouvons réécrire $\mathcal{F}_{\mathbb{L}}$ sous la forme :

$$\mathcal{F}_{\mathbb{L}} : \eta(\bar{k}) \rightarrow \eta(\overline{k+1}) \rightarrow \eta(\overline{k+1 - (2t+1)}) \rightarrow \eta(\overline{k - (2t+1)}) \rightarrow \eta(\bar{k})$$

Par construction de l'ordre choisi sur sommets et arêtes de $\mathcal{G}'$, $\mathbf{w}$ est également le vecteur caractéristique du sous-ensemble d'arêtes $\eta^{-1}(\mathcal{F}_{\mathbb{L}})$ de $\mathcal{G}$, défini par :

$$\bar{k} \rightarrow \overline{k+1} \rightarrow \overline{k+1 - (2t+1)} \rightarrow \overline{k - (2t+1)} \rightarrow \bar{k}$$

Puisqu'il s'agit d'une face du graphe $(\mathbb{Z}/(t^2 + (t+1)^2)\mathbb{Z}, \bar{1}, \overline{2t+1})$, alors $\mathbf{w}$ est une ligne de $\mathbf{H}_Z$.

En conséquence, les ensembles de lignes de $\mathbf{H}_Z$ et $\mathbf{H}_Z''$ coïncident. Ainsi, en particulier, $\text{Vect}_{\mathbb{L}}(\mathbf{H}_Z) = \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z'')$.

Étape 3.3 : Etablir les relations entre les couples $(\mathbf{H}_X'', \mathbf{H}_Z'')$ et $(\mathbf{H}_X', \mathbf{H}_Z')$:

Les couples de matrices d'incidence sommet-arête et face-arête : $(\mathbf{H}_X', \mathbf{H}_Z')$ et $(\mathbf{H}_X'', \mathbf{H}_Z'')$ sont obtenus en renumérotant différemment les sommets, arêtes et faces du même graphe $\mathcal{G}'$. D'après la remarque 4.2.2, ces matrices satisfont les relations :

$$\mathbf{H}_X'' = \mathbf{P}\mathbf{H}_X'\mathbf{Q} \quad \text{et} \quad \mathbf{H}_Z'' = \mathbf{R}\mathbf{H}_Z'\mathbf{Q}$$

où $\mathbf{P}, \mathbf{Q}, \mathbf{R}$ sont des matrices de permutations. Cela implique en particulier que :

$$\begin{aligned} \text{Vect}_{\mathbb{L}}(\mathbf{H}_X) &= \text{Vect}_{\mathbb{L}}(\mathbf{H}_X'') = \text{Vect}_{\mathbb{L}}(\mathbf{H}_X'\mathbf{Q}) \\ \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z) &= \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z'') = \text{Vect}_{\mathbb{L}}(\mathbf{H}_Z'\mathbf{Q}) \end{aligned}$$

En conclusion, pour tout $t \geq 1$, le code Bicycle généralisé $GB(1+X, 1+X^{2t+1}, t^2+(t+1)^2)$ est CGP-équivalent, et donc équivalent au sens usuel, au code surface 2D à distance $2t+1$ optimal, introduit dans la section 4.2.2. $\square$

4.8 Classification des codes GB(2,2)

Cette section présente notre classification des codes GB(2,2). Notre objectif principal est d'identifier systématiquement les meilleurs codes GB(2,2) (à CGP-équivalence près) pour toutes les longueurs inférieures à 200.

Cette classification est cruciale, car les capacités de correction et les performances de décodage de ces codes varient considérablement en fonction de leur classe d'équivalence. Comme nous l'avons mentionné, ces codes sont des codes CSS construits à partir de graphes de Cayley à deux générateurs de groupes cycliques, de la forme $(\mathbb{Z}/n\mathbb{Z}, a, b)$.

Nous nous sommes concentrés exclusivement sur la relation d'équivalence CGP, qui préserve les structures de codes CSS et celles de leurs graphes sous-jacents. Ceci est d'autant plus important que la performance de certains algorithmes de décodage, comme l'algorithme de renormalisation repose directement sur la structure des graphes sous-jacents [13].

D'après les travaux de [16], tous les codes GB peuvent être décomposés en une somme directe de codes GB élémentaires.

Proposition 4.8.1. *Le code $GB(A(X^\Delta), B(X^\Delta), n\Delta)$ est une somme directe de Δ copies de $GB(A(X), B(X), n)$. De plus, si les paramètres de $GB(A(X), B(X), n)$ sont $[[2n, k, d]]$, alors les paramètres du code original sont $[[2n\Delta, k\Delta, d]]$.*

En combinant ce résultat avec celui de la proposition 4.3.1, on peut montrer que tout code GB(2,2) se décompose en une somme directe de codes de la forme $GB(1 + X^a, 1 + X^b, n)$ où a, b et n sont trois entiers premiers entre eux dans leur ensemble.

Ce sont ces codes, que nous avons classifiés, car ils représentent les briques élémentaires dans la construction des codes GB(2,2) plus généraux.

4.8.1 Méthode :

Pour réaliser cette classification, nous avons procédé en plusieurs étapes :

1. **Identification des graphes non-isomorphes :** Nous avons d'abord recensé tous les graphes non-isomorphes de la forme $(\mathbb{Z}/n\mathbb{Z}, a, b)$ avec $\text{pgcd}(a, b, n) = 1$.

2. **Calcul des paramètres des codes** : Comme il peut exister une bijection préservant les cycles simples entre les ensembles d'arêtes de graphes non isomorphes, nous avons calculé les paramètres de tous les codes GB correspondants. Cette étape sert de filtre initial efficace, car deux codes CGP-équivalents ont nécessairement les mêmes paramètres.
3. **Vérification manuelle de la distinction des graphes** : Enfin, pour chaque ensemble de codes ayant les mêmes paramètres, nous avons effectué une vérification manuelle pour confirmer qu'aucune telle bijection n'existe. Ceci a été effectué à l'aide de deux méthodes :
 - **Vérification de la 3-connexité** : Nous avons vérifié si les graphes sous-jacents de tous les codes GB listés étaient 3-connexes. D'après le théorème de Whitney, pour les graphes 3-connexes, la non-isomorphie implique directement qu'il ne peut pas exister de bijection préservant les cycles simples entre les ensembles d'arêtes de ces graphes. Ce résultat est fondamental, car tous les codes de la forme $GB(1 + X^a, 1 + X^b, n)$ avec n compris entre 4 et 100, $1 \leq a < b \leq n - 2$ et $\text{pgcd}(a, b, n) = 1$, ont des graphes sous-jacents qui sont 3-connexes.
 - **Analyse d'équivalence par permutation** : Nous avons également vérifié, à l'aide de la méthode `.is_permutation_equivalent()` de Sage, qu'il était impossible de passer d'un espace vectoriel engendré par les lignes d'une matrice d'incidence sommet-arête d'un graphe à un autre par l'application d'une matrice de permutation. Cette méthode fournit un critère supplémentaire pour distinguer des codes non CGP-équivalents.

Les aspects calculatoires de cette classification ont été réalisés à l'aide des logiciels SageMath et Magma. Le premier a servi à déterminer les classes d'équivalence de graphes isomorphes et à évaluer leur 3-connexité, tandis que le deuxième a été utilisé pour calculer la distance minimale des codes.

4.8.2 Classification des codes GB(2,2) extrémaux

Nous avons dressé une table listant tous les codes GB(2,2) extrémaux, non CGP-équivalents, dont la longueur est inférieure à 200. Contrairement à la table de la Section 4.8.4, qui répertorie les meilleurs codes surface 2D à stabilisateurs de poids quatre connus, celle-ci se concentre exclusivement sur les codes GB(2,2). Nous l'avons structurée de la manière suivante : pour chaque longueur, nous indiquons les meilleurs

paramètres atteignables $[[2n, 2, d]]$ par des codes de la forme $GB(1 + X^a, 1 + X^b, n)$ avec a, b et n premiers dans leurs ensembles. Nous indiquons également le nombre total de codes non CGP-équivalents, les polynômes $A(X) = 1 + X^a$ et $B(X) = 1 + X^b$ des représentants, ainsi que leurs sources si ces codes ont déjà été documentés.

Table 4.1: Codes GB(2,2) extrémaux non CGP-équivalents $GB(A(X), B(X), n)$ de longueur inférieure à 200

n	Paramètres $[[2n, k, d]]$	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme $A(X)$	Polynôme $B(X)$
2	$[[4, 2, 2]]$	1	$1 + X$	$1 + X$ (Code surface optimal [7])
3	$[[6, 2, 2]]$	1	$1 + X$	$1 + X^2$
4	$[[8, 2, 2]]$	2	$1 + X$	$1 + X$ (code de Kitaev [6]), $1 + X^2$
5	$[[10, 2, 3]]$	1	$1 + X$	$1 + X^3$ (Code surface optimal [7])
6	$[[12, 2, 3]]$	1	$1 + X$	$1 + X^2$
7	$[[14, 2, 3]]$	1	$1 + X$	$1 + X^2$
8	$[[16, 2, 4]]$	1	$1 + X$	$1 + X^3$
9	$[[18, 2, 3]]$	2	$1 + X$	$1 + X^2, 1 + X^3$
10	$[[20, 2, 4]]$	2	$1 + X$	$1 + X^3, 1 + X^4$
11	$[[22, 2, 4]]$	1	$1 + X$	$1 + X^3$ [14]
12	$[[24, 2, 4]]$	3	$1 + X$	$1 + X^3, 1 + X^5$
			$1 + X^2$	$1 + X^3$
13	$[[26, 2, 5]]$	1	$1 + X$	$1 + X^5$ (Code surface optimal [7])
14	$[[28, 2, 5]]$	1	$1 + X$	$1 + X^4$
15	$[[30, 2, 5]]$	2	$1 + X$	$1 + X^4, 1 + X^6$
16	$[[32, 2, 5]]$	1	$1 + X$	$1 + X^6$
17	$[[34, 2, 5]]$	2	$1 + X$	$1 + X^4, 1 + X^5$
18	$[[36, 2, 6]]$	1	$1 + X$	$1 + X^5$
19	$[[38, 2, 5]]$	2	$1 + X$	$1 + X^4$ [14], $1 + X^7$

n	Paramètres [[$2n, k, d$]]	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme $A(X)$	Polynôme $B(X)$
20	[[40, 2, 5]]	3	$1 + X$	$1 + X^4, 1 + X^6, 1 + X^8$
21	[[42, 2, 6]]	2	$1 + X$	$1 + X^6, 1 + X^8$
22	[[44, 2, 6]]	2	$1 + X$	$1 + X^5, 1 + X^6$
23	[[46, 2, 6]]	1	$1 + X$	$1 + X^5$
24	[[48, 2, 6]]	5	$1 + X$ $1 + X^3$	$1 + X^5, 1 + X^7, 1 + X^9, 1 + X^{10}$ $1 + X^4$
25	[[50, 2, 7]]	1	$1 + X$	$1 + X^7$ (Code surface optimal [7])
26	[[52, 2, 7]]	1	$1 + X$	$1 + X^{10}$
27	[[54, 2, 7]]	1	$1 + X$	$1 + X^6$
28	[[56, 2, 7]]	2	$1 + X$	$1 + X^6, 1 + X^8$
29	[[58, 2, 7]]	2	$1 + X$	$1 + X^8$ [14], $1 + X^{12}$
30	[[60, 2, 7]]	1	$1 + X^2$	$1 + X^9$
31	[[62, 2, 7]]	2	$1 + X$	$1 + X^7, 1 + X^{12}$
32	[[64, 2, 8]]	1	$1 + X$	$1 + X^7$
33	[[66, 2, 7]]	3	$1 + X$	$1 + X^6, 1 + X^7, 1 + X^9$
34	[[68, 2, 8]]	1	$1 + X$	$1 + X^{13}$
35	[[70, 2, 7]]	4	$1 + X$	$1 + X^6, 1 + X^8, 1 + X^{10}, 1 + X^{15}$
36	[[72, 2, 8]]	3	$1 + X$	$1 + X^8, 1 + X^{10}, 1 + X^{15}$
37	[[74, 2, 8]]	1	$1 + X$	$1 + X^8$
38	[[76, 2, 8]]	2	$1 + X$	$1 + X^7, 1 + X^{16}$
39	[[78, 2, 8]]	2	$1 + X$	$1 + X^7, 1 + X^{15}$
40	[[80, 2, 8]]	5	$1 + X$ $1 + X^4$	$1 + X^7, 1 + X^9, 1 + X^{11}, 1 + X^{15}$ $1 + X^5$
41	[[82, 2, 9]]	1	$1 + X$	$1 + X^9$ (Code surface optimal [7])

n	Paramètres [[2n, k, d]]	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme A(X)	Polynôme B(X)
42	[[84, 2, 9]]	1	$1 + X$	$1 + X^{16}$
43	[[86, 2, 9]]	1	$1 + X$	$1 + X^{12}$
44	[[88, 2, 9]]	1	$1 + X$	$1 + X^8$
45	[[90, 2, 9]]	3	$1 + X$	$1 + X^8, 1 + X^{10}, 1 + X^{19}$
46	[[92, 2, 9]]	1	$1 + X$	$1 + X^{10}$
47	[[94, 2, 9]]	1	$1 + X$	$1 + X^{13}$
48	[[96, 2, 9]]	3	$1 + X$	$1 + X^{14}, 1 + X^{20}$
			$1 + X^2$	$1 + X^9$
49	[[98, 2, 9]]	1	$1 + X$	$1 + X^9$
50	[[100, 2, 10]]	2	$1 + X$	$1 + X^9, 1 + X^{19}$
51	[[102, 2, 9]]	4	$1 + X$	$1 + X^8, 1 + X^9, 1 + X^{11}, 1 + X^{15}$
52	[[104, 2, 9]]	3	$1 + X$	$1 + X^8, 1 + X^{20}, 1 + X^{22}$
53	[[106, 2, 9]]	3	$1 + X$	$1 + X^8, 1 + X^{12}, 1 + X^{23}$
54	[[108, 2, 10]]	1	$1 + X$	$1 + X^{15}$
55	[[110, 2, 10]]	3	$1 + X$	$1 + X^{10}, 1 + X^{12}, 1 + X^{21}$
56	[[112, 2, 10]]	1	$1 + X$	$1 + X^{10}$
57	[[114, 2, 10]]	1	$1 + X$	$1 + X^{24}$
58	[[116, 2, 10]]	4	$1 + X$	$1 + X^9, 1 + X^{16}, 1 + X^{17}, 1 + X^{22}$
59	[[118, 2, 10]]	1	$1 + X$	$1 + X^9$ [14]
60	[[120, 2, 10]]	6	$1 + X$	$1 + X^9, 1 + X^{11}, 1 + X^{13}, 1 + X^{25}$
			$1 + X^3$	$1 + X^8$
			$1 + X^5$	$1 + X^6$
61	[[122, 2, 11]]	1	$1 + X$	$1 + X^{11}$ (Code surface optimal [7])

n	Paramètres [[$2n, k, d$]]	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme $A(X)$	Polynôme $B(X)$
62	[[124, 2, 11]]	1	$1 + X$	$1 + X^{26}$
63	[[126, 2, 11]]	1	$1 + X$	$1 + X^{24}$
64	[[128, 2, 11]]	1	$1 + X$	$1 + X^{14}$
65	[[130, 2, 11]]	2	$1 + X$	$1 + X^{10}, 1 + X^{18}$
66	[[132, 2, 11]]	3	$1 + X$ $1 + X^2$	$1 + X^{10}, 1 + X^{12}$ $1 + X^9$
67	[[134, 2, 11]]	1	$1 + X$	$1 + X^{12}$ [14]
68	[[136, 2, 11]]	2	$1 + X$	$1 + X^{20}, 1 + X^{26}$
69	[[138, 2, 11]]	2	$1 + X$	$1 + X^{15}, 1 + X^{19}$
70	[[140, 2, 11]]	1	$1 + X^2$	$1 + X^{25}$
71	[[142, 2, 11]]	3	$1 + X$	$1 + X^{11}, 1 + X^{16}, 1 + X^{21}$
72	[[144, 2, 12]]	1	$1 + X$	$1 + X^{11}$
73	[[146, 2, 11]]	4	$1 + X$	$1 + X^{11}, 1 + X^{13}, 1 + X^{16}, 1 + X^{27}$
74	[[148, 2, 12]]	1	$1 + X$	$1 + X^{31}$
75	[[150, 2, 11]]	4	$1 + X$ $1 + X^5$	$1 + X^{10}, 1 + X^{53}, 1 + X^{33}$ $1 + X^{36}$
76	[[152, 2, 12]]	1	$1 + X$	$1 + X^{47}$
77	[[154, 2, 11]]	5	$1 + X$	$1 + X^{12}, 1 + X^{14}, 1 + X^{23},$ $1 + X^{43}, 1 + X^{56}$
78	[[156, 2, 12]]	4	$1 + X$ $1 + X^3$	$1 + X^{12}, 1 + X^{14}, 1 + X^{55}$ $1 + X^{62}$
79	[[158, 2, 12]]	1	$1 + X$	$1 + X^{12}$
80	[[160, 2, 12]]	3	$1 + X$	$1 + X^{45}, 1 + X^{58}$

n	Paramètres [[$2n, k, d$]]	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme A(X)	Polynôme B(X)
			$1 + X^4$	$1 + X^5$
81	[[162, 2, 12]]	2	$1 + X$	$1 + X^{24}, 1 + X^{31}$
82	[[164, 2, 12]]	2	$1 + X$	$1 + X^{36}, 1 + X^{67}$
83	[[166, 2, 12]]	2	$1 + X$	$1 + X^{11}$ [14], $1 + X^{23}$
84	[[168, 2, 12]]	8	$1 + X$ $1 + X^7$	$1 + X^{19}, 1 + X^{25}, 1 + X^{32}, 1 + X^{35}$ $1 + X^{69}, 1 + X^{71}, 1 + X^{73}$ $1 + X^{30}$
85	[[170, 2, 13]]	1	$1 + X$	$1 + X^{13}$ (Code surface optimal [7])
86	[[172, 2, 12]]	1	$1 + X$	$1 + X^{73}$
87	[[174, 2, 13]]	1	$1 + X$	$1 + X^{24}$
88	[[176, 2, 13]]	1	$1 + X$	$1 + X^{26}$
89	[[178, 2, 13]]	2	$1 + X$	$1 + X^{16}, 1 + X^{55}$
90	[[180, 2, 13]]	1	$1 + X$	$1 + X^{78}$
91	[[182, 2, 13]]	4	$1 + X$	$1 + X^{12}, 1 + X^{14}, 1 + X^{27}, 1 + X^{40}$
92	[[184, 2, 13]]	2	$1 + X$	$1 + X^{14}, 1 + X^{72}$
93	[[186, 2, 13]]	1	$1 + X$	$1 + X^{54}$
94	[[188, 2, 13]]	2	$1 + X$	$1 + X^{26}, 1 + X^{36}$
95	[[190, 2, 13]]	2	$1 + X$	$1 + X^{28}, 1 + X^{35}$
96	[[192, 2, 13]]	2	$1 + X$ $1 + X^4$	$1 + X^{42}$ $1 + X^{57}$
97	[[194, 2, 13]]	3	$1 + X$	$1 + X^{15}, 1 + X^{21}, 1 + X^{22}$
98	[[196, 2, 14]]	2	$1 + X$	$1 + X^{41}, 1 + X^{69}$
99	[[198, 2, 13]]	2	$1 + X$	$1 + X^{15}, 1 + X^{61}$
100	[[200, 2, 13]]	5	$1 + X$	$1 + X^{18}, 1 + X^{42}, 1 + X^{44}$

n	Paramètres [[2n, k, d]]	Nombre de codes GB(2,2) non CGP-Équivalents	Polynôme A(X)	Polynôme B(X)
			$1 + X^5$	$1 + X^{68}, 1 + X^{72}$

4.8.3 Classification de tous les codes GB(2,2) (Longueur ≤ 200)

Nous avons mis à disposition une classification complète de tous les codes GB(2,2) de la forme $GB(1 + X^a, 1 + X^b, n)$ avec $\gcd(a, b, n) = 1$. Cette classification, ainsi que le code source utilisé pour l'établir, sont accessibles sur notre dépôt GitHub [21] via ce lien : https://github.com/NicolasSaussay/weight-4_GB-Codes_Classification.

4.8.4 Résumé des meilleurs codes surface 2D à stabilisateurs de poids quatre

Cette section présente une table qui récapitule les paramètres des meilleurs codes surface de poids quatre connus, pour des longueurs inférieures à 200.

Contrairement à la table précédente, celle-ci rassemble les données des meilleurs codes surface 2D de poids quatre provenant de diverses sources, y compris des codes qui ne sont pas de type GB(2,2). Ces sources comprennent les codes toriques de Kitaev, les meilleurs codes surface 2D de poids quatre (présentés en Section 4.2.2), les codes issus du dépôt de Pryadko et Wang [14]. Sont également incluses nos deux familles infinies de codes GB (construites en Section 4.5), ainsi que l'ensemble plus large des codes GB extrémaux recensés dans la table 4.1, qui n'avaient pas été répertoriés jusqu'à présent.

Table 4.2: Récapitulatif des codes surface 2D de poids quatre optimaux de longueurs $2n$, pour n allant de 2 à 100

n	Paramètres optimaux [[2n, 2, d]]	Nombre de Codes Surfaces non CGP-Équivalents	Références des codes
2	[[4, 2, 2]]	1	Code surface optimal [7]
3	[[6, 2, 2]]	1	Table 4.1
4	[[8, 2, 2]]	2	Kitaev code Table 4.1

n	Paramètres optimaux [[2n, 2, d]]	Nombre de Codes Surfaces non CGP-Équivalents	Références des codes
5	[[10, 2, 3]]	1	Code surface optimal [7]
6	[[12, 2, 3]]	1	Table 4.1
7	[[14, 2, 3]]	1	Table 4.1
8	[[16, 2, 4]]	2	Code surface optimal [7] Notre code (voir Section 4.5.2)
9	[[18, 2, 3]]	3	Code torique de Kitaev Table 4.1
10	[[20, 2, 4]]	2	Table 4.1
11	[[22, 2, 4]]	1	Répertoire de Pryadko et Wang [14] (et Table 4.1)
12	[[24, 2, 4]]	3	Table 4.1
13	[[26, 2, 5]]	1	Code surface optimal [7]
14	[[28, 2, 5]]	1	Table 4.1
15	[[30, 2, 5]]	2	Table 4.1
16	[[32, 2, 5]]	1	Table 4.1
17	[[34, 2, 5]]	2	Table 4.1
18	[[36, 2, 6]]	2	Code surface optimal [7] Notre Code (voir Section 4.5.2)
19	[[38, 2, 5]]	2	Répertoire de Pryadko et Wang [14] Table 4.1
20	[[40, 2, 5]]	3	Table 4.1
21	[[42, 2, 6]]	2	Table 4.1
22	[[44, 2, 6]]	2	Table 4.1
23	[[46, 2, 6]]	1	Table 4.1
24	[[48, 2, 6]]	5	Table 4.1
25	[[50, 2, 7]]	1	Code surface optimal [7]
26	[[52, 2, 7]]	1	Table 4.1
27	[[54, 2, 7]]	1	Table 4.1
28	[[56, 2, 7]]	2	Table 4.1
29	[[58, 2, 7]]	2	[14], Table 4.1

n	Paramètres optimaux [[2n, 2, d]]	Nombre de Codes Surfaces non CGP-Équivalents	Références des codes
30	[[60, 2, 7]]	1	Table 4.1
31	[[62, 2, 7]]	2	Table 4.1
32	[[64, 2, 8]]	2	Code surface optimal [7] Notre Code (voir Section 4.5.2)
33	[[66, 2, 7]]	3	Table 4.1
34	[[68, 2, 8]]	1	Table 4.1
35	[[70, 2, 7]]	4	Table 4.1
36	[[72, 2, 8]]	3	Table 4.1
37	[[74, 2, 8]]	1	Répertoire de Pryadko et Wang [14] (et Table 4.1)
38	[[76, 2, 8]]	2	Table 4.1
39	[[78, 2, 8]]	2	Table 4.1
40	[[80, 2, 8]]	5	Table 4.1
41	[[82, 2, 9]]	1	Code surface optimal [7]
42	[[84, 2, 9]]	1	Table 4.1
43	[[86, 2, 9]]	1	Table 4.1
44	[[88, 2, 9]]	1	Table 4.1
45	[[90, 2, 9]]	3	Table 4.1
46	[[92, 2, 9]]	1	Table 4.1
47	[[94, 2, 9]]	1	Table 4.1
48	[[96, 2, 9]]	3	Table 4.1
49	[[98, 2, 9]]	1	Table 4.1
50	[[100, 2, 10]]	3	Table 4.1 Code surface optimal [7] Notre Code (voir Section 4.5.2)
51	[[102, 2, 9]]	4	Table 4.1
52	[[104, 2, 9]]	3	Table 4.1
53	[[106, 2, 9]]	3	Répertoire de Pryadko et Wang [14] Table 4.1
54	[[108, 2, 10]]	1	Table 4.1

n	Paramètres optimaux [[2n, 2, d]]	Nombre de Codes Surfaces non CGP-Équivalents	Références des codes
55	[[110, 2, 10]]	3	Table 4.1
56	[[112, 2, 10]]	1	Table 4.1
57	[[114, 2, 10]]	1	Table 4.1
58	[[116, 2, 10]]	4	Table 4.1
59	[[118, 2, 10]]	1	Répertoire de Pryadko et Wang [14]
60	[[120, 2, 10]]	6	Table 4.1
61	[[122, 2, 11]]	1	Code surface optimal [7]
62	[[124, 2, 11]]	1	Table 4.1
63	[[126, 2, 11]]	1	Table 4.1
64	[[128, 2, 11]]	1	Table 4.1
65	[[130, 2, 11]]	2	Table 4.1
66	[[132, 2, 11]]	3	Table 4.1
67	[[134, 2, 11]]	1	Répertoire de Pryadko et Wang [14] (et Table 4.1)
68	[[136, 2, 11]]	2	Table 4.1
69	[[138, 2, 11]]	2	Table 4.1
70	[[140, 2, 11]]	1	Table 4.1
71	[[142, 2, 11]]	3	Table 4.1
72	[[144, 2, 12]]	2	Code surface optimal [7] Notre Code (voir Section 4.5.2)
73	[[146, 2, 11]]	4	Table 4.1
74	[[148, 2, 12]]	1	Table 4.1
75	[[150, 2, 11]]	4	Table 4.1
76	[[152, 2, 12]]	1	Table 4.1
77	[[154, 2, 11]]	5	Table 4.1
78	[[156, 2, 12]]	4	Table 4.1
79	[[158, 2, 12]]	1	Table 4.1
80	[[160, 2, 12]]	2	Table 4.1
81	[[162, 2, 12]]	2	Table 4.1
82	[[164, 2, 12]]	2	Table 4.1

n	Paramètres optimaux [[2n, 2, d]]	Nombre de Codes Surfaces non CGP-Équivalents	Références des codes
83	[[166, 2, 12]]	2	Répertoire de Pryadko et Wang [14] Table 4.1
84	[[168, 2, 12]]	8	Table 4.1
85	[[170, 2, 13]]	1	Code surface optimal [7]
86	[[172, 2, 12]]	1	Table 4.1
87	[[174, 2, 13]]	1	Table 4.1
88	[[176, 2, 13]]	1	Table 4.1
89	[[178, 2, 13]]	2	Table 4.1
90	[[180, 2, 13]]	1	Table 4.1
91	[[182, 2, 13]]	4	Table 4.1
92	[[184, 2, 13]]	2	Table 4.1
93	[[186, 2, 13]]	1	Table 4.1
94	[[188, 2, 13]]	2	Table 4.1
95	[[190, 2, 13]]	2	Table 4.1
96	[[192, 2, 13]]	2	Table 4.1
97	[[194, 2, 13]]	3	Table 4.1
98	[[196, 2, 14]]	4	Code surface optimal [7] Table 4.1
99	[[198, 2, 13]]	2	Table 4.1
100	[[200, 2, 13]]	4	Table 4.1

4.9 Conclusion du chapitre

Dans ce travail, nous avons fait des progrès significatifs dans la compréhension et la conception de codes GB(2,2). Notre première contribution est d'avoir établi une borne inférieure sur leur distance minimale, directement liée à la plus petite norme L^1 des vecteurs non nuls de certains sous-réseaux de $\mathbb{Z}^2$. Cette borne constitue un outil théorique essentiel pour construire explicitement des codes GB(2,2) performants, une tâche qui aurait été très difficile, voire impossible, sans elle.

En utilisant cette borne, nous avons développé trois nouvelles familles infinies de codes GB(2,2) performants : $[[2n^2, 2, n]]$, $[[4r^2, 2, 2r]]$ et $[(2t+1)^2+1, 2, 2t+1]$. Il convient de souligner que la famille de paramètres $[[4r^2, 2, 2r]]$ comble un vide existant parmi les codes GB(2,2) optimaux, car ces codes n'étaient auparavant pas considérés comme réalisables.

En nous appuyant sur la relation d'équivalence spécifique aux codes CSS dérivés de graphes de Cayley (équivalence CGP), nous avons prouvé rigoureusement que deux de nos familles ($[[2n^2, 2, n]]$ et $[[4r^2, 2, 2r]]$) ne sont pas CGP-équivalentes aux codes surface 2D bien connus, tels que le code de Kitaev standard ou le code de surface 2D optimal pour la distance paire. La troisième famille, $[(2t+1)^2+1, 2, 2t+1]$, s'est avérée équivalente (au sens usuel des codes quantiques) au meilleur code surface 2D à distance impaire, offrant ainsi une nouvelle méthode de construction.

Enfin, nous avons proposé une classification exhaustive de tous les codes GB(2,2) optimaux de la forme $GB(1+X^a, 1+X^b, n)$ avec $\text{pgcd}(a, b, n) = 1$, de longueur inférieure à 200. Grâce à sa structure organisée, cette classification, disponible sur notre dépôt GitHub [21], constitue une ressource essentielle pour les chercheurs à la recherche de codes aux propriétés de décodage remarquables.

Par ailleurs, nous avons établi une table regroupant les meilleurs codes de surface 2D de poids quatre connus, qu'ils soient GB(2,2) ou non, pour des longueurs inférieures ou égales à 200. Cette table constitue un véritable annuaire des codes surface 2D performants à stabilisateurs de poids quatre, offrant une ressource précieuse pour les chercheurs s'intéressant aux codes stabilisateurs efficaces, de petites longueurs.

Ces contributions ouvrent plusieurs pistes prometteuses. Une extension naturelle serait de classer les codes GB, ou plus généralement les codes CSS, sous des relations d'équivalence plus larges, afin d'identifier les classes de codes CSS les plus adaptées pour l'implémentation pratique. Nous envisageons également d'adapter notre approche basée sur la théorie des graphes pour créer des codes GB performants avec des générateurs de poids plus élevés. Pour cela, nous élargirions les notions de graphes introduites aux hypergraphes construits à partir de graphes de Cayley ayant plus de générateurs.

Bien que nous ayons montré que les codes GB(2,2) et les codes de surface 2D pouvaient être représentés sur des réseaux bidimensionnels, l'analyse de ce chapitre s'est focalisée sur les qubits sur lesquels agissent les stabilisateurs, plutôt que sur leur disposition

géométrique. Dans le chapitre suivant, nous inverserons cette perspective pour nous concentrer sur l'agencement spatial des qubits. Nous examinerons en particulier les codes stabilisateurs locaux, qui généralisent les codes étudiés dans ce chapitre. Ce sont des codes dont les qubits sont représentés par des points de $\mathbb{R}^D$ de manière à ce que les stabilisateurs n'agissent que sur des qubits voisins. Notre objectif sera d'analyser comment cette contrainte géométrique influence la distance minimale des codes.

5

Borne de Bravyi-Terhal généralisée

Sommaire

5.1	Introduction	110
5.1.1	Motivations	110
5.1.2	Résultats connus : borne de Bravyi-Terhal	111
5.1.3	Contributions	112
5.1.4	Vue d'ensemble sur le théorème 5.4.1	113
5.2	Organisation du chapitre	114
5.3	Résultats préliminaires	114
5.3.1	Mesure de Haar	115
5.3.2	Constantes de Hermite et de Rankin	115
5.4	Borne de Bravyi-Terhal généralisée	117
5.4.1	Borne de Bravyi-Terhal originale	117
5.4.2	Borne de Bravyi-Terhal généralisée	118
5.5	Esquisse de la preuve du théorème	119
5.5.1	Cas des réseaux de faible covolume : $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$	119
5.5.2	Cas général : $\ell^{1/D} \geq 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$	120
5.6	Preuve du théorème 5.4.1	122
5.6.1	Construction d'une base du réseau adaptée au découpage	122

5.6.2	Découpage du domaine fondamental	125
5.6.3	Application du lemme de nettoyage	129
5.6.4	Majoration du nombre de qubits par tranche	131
5.6.5	Preuve complète du théorème 5.4.1 :	136
5.7	Étude de la distance minimale des codes 2BGA abéliens	139
5.7.1	Présentation des codes 2BGA abéliens	140
5.7.2	Étude des codes 2BGA abéliens non triviaux	141
5.7.3	Démonstration du théorème 5.7.1	143
5.8	Conclusion du chapitre :	153

5.1 Introduction

5.1.1 Motivations

Ces dernières années, les codes stabilisateurs quantiques LDPC (qLDPC) ont suscité un intérêt croissant pour leur capacité à corriger efficacement les erreurs quantiques.

Alors que la construction de bons codes LDPC classiques repose généralement sur des matrices aléatoires, celle de bons codes stabilisateurs est bien plus complexe, car ces derniers doivent satisfaire des contraintes structurelles qui rendent cette approche inapplicable.

Malgré tout, des avancées majeures ont récemment démontré l'existence de codes qLDPC asymptotiquement bons, dont la dimension et la distance minimale croissent linéairement avec la longueur du code [8], [9]. Bien que ces codes soient théoriquement prometteurs, en pratique leur efficacité ne se manifeste qu'à partir d'un grand nombre de qubits. Ainsi, pour les besoins pratiques, la priorité est de concevoir des codes performants de petite longueur adaptés aux capacités des ordinateurs quantiques actuels.

Pour répondre à cette demande, les recherches se sont tournées vers les codes 2-blocs. Ce sont des codes CSS définis par des sous-matrices génératrices de la forme $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$, où $\mathbf{A}$ et $\mathbf{B}$ sont des matrices qui commutent.

Parmi ces codes, on retrouve les codes Bicycle généralisés [11], qui ont été abordés dans le Chapitre 4, ainsi que leur généralisation, les codes 2BGA (Two-Block Group Algebra) [22], qui sont construits à partir d'algèbres de groupe de groupes abéliens. Ces derniers ont fait l'objet de nombreuses études ([11], [16], [23]) et présentent des capacités de correction d'erreurs très prometteuses. Dans les cas les plus favorables, ils peuvent atteindre un seuil d'erreur comparable à celui des codes surface tout en offrant une meilleure dimension pour un nombre inférieur de qubits [10]. Toutefois, le comportement de leur distance minimale est encore incompris à ce jour.

Malgré cette inconnue, les codes 2BGA abéliens demeurent très attrayants. À distance minimale égale, ils permettent de construire des codes de petite longueur avec un taux supérieur à celui des codes toriques de Kitaev. De plus, tout comme ces derniers, leurs qubits peuvent être représentés sur les sommets de réseaux euclidiens, de sorte que son groupe stabilisateur soit engendré par une famille de générateurs agissant uniquement sur des qubits voisins, ce qui est très avantageux pour une implémentation physique.

5.1.2 Résultats connus : borne de Bravyi-Terhal

En raison de leur intérêt pratique, une grande partie de la recherche s’est tournée vers les codes stabilisateurs dits locaux, dont le groupe stabilisateur est engendré par une famille de générateurs agissant chacun sur un petit nombre de qubits situés à proximité les uns des autres. Cependant, cette localité impose des contraintes sur leurs paramètres.

En 2009, Bravyi et Terhal ont établi une borne supérieure fondamentale sur la distance minimale de ces codes [24]. Leur travail a notamment montré que pour un code stabilisateur de longueur ℓ , avec des qubits disposés sur une partie de $\mathbb{Z}^D$, et dont le groupe stabilisateur est engendré par une famille de générateurs dont le support est inclus dans un cube de ρ^D sommets, la distance minimale vérifie l’inégalité : $d \leq \rho \ell^{\frac{D-1}{D}}$.

Cette borne a eu un impact majeur en définissant les limites théoriques des codes quantiques stabilisateurs locaux. Depuis, les recherches se sont poursuivies selon deux directions complémentaires.

D’une part, plusieurs travaux ont cherché à construire des codes atteignant cette borne. En dimension deux, le code de Kitaev en constitue un exemple emblématique. En dimensions supérieures ($D \geq 3$), des constructions récentes ont permis d’obtenir des codes locaux optimaux [25], [26], [27], [28], dont la distance minimale et le nombre de qubits logiques respectent $d = O(\ell^{(D-1)/D})$ et $k = O(\ell^{(D-2)/D})$, en accord avec la borne de Bravyi-Poulin-Terhal [29].

D’autre part, des travaux plus récents se sont attachés à étendre la portée de la borne de Bravyi-Terhal. Il a notamment été montré [30], [31] que, sous certaines conditions, la borne $d = O(\ell^{\frac{D-1}{D}})$ reste valable lorsque les qubits sont représentés sur les sommets de graphes inclus dans $\mathbb{R}^D$, et non plus uniquement sur les points d’un réseau euclidien. Dans la même lignée, Pryadko et Wang ont également démontré que la distance minimale des codes Bicycle généralisés est limitée à une croissance sous-linéaire en la

longueur du code. Plus précisément, sans préciser la constante, ils ont établi que pour des codes GB non triviaux de longueur ℓ , dont le groupe des stabilisateurs est engendré par des générateurs de poids w , la distance minimale est de l'ordre de $O(\ell^{\frac{D-1}{D}})$, avec $D \leq w - 1$ (ou $D \leq w - 2$ dans certains cas) [16].

5.1.3 Contributions

Enfin, ce chapitre est consacré à nos travaux, réalisés en collaboration avec Gilles Zémor et Wouter Rozendaal, sur les codes stabilisateurs locaux.

Ces travaux visent à étendre la portée de la borne de Bravyi-Terhal.

Nous présentons pour cela une généralisation qui s'applique aux codes stabilisateurs LDPC locaux, dont les qubits sont représentés sur un quotient de $\mathbb{Z}^D$ par un réseau de rang D arbitraire. Plus précisément, nous établissons une nouvelle borne sur la distance minimale d des codes stabilisateurs de longueur $n = m\ell$, dont les qubits sont indexés par les sommets du quotient $\mathbb{Z}^D/\Lambda$ (chaque sommet indexant m qubits) et dont le groupe stabilisateur est engendré par des générateurs dont le support est inclus dans une boule de rayon ρ du quotient (voir Théorème 5.4.1).

Alors que la borne de Bravyi-Terhal s'écrit $d \leq \rho \ell^{\frac{D-1}{D}}$, notre borne, plus générale, s'exprime sous la forme :

$$d < m \cdot \max(1, \sqrt{\gamma_{D,D-1}(\Lambda)}) \cdot (\sqrt{D} + 4\rho) \ell^{\frac{D-1}{D}}$$

où $\gamma_{D,D-1}(\Lambda)$ est une constante ne dépendant que du réseau Λ et de la dimension D . Cette constante peut être majorée par la constante de Hermite D -dimensionnelle $\gamma_D \geq 1$, qui ne dépend que de D . On obtient ainsi la borne simplifiée :

$$d < m\sqrt{\gamma_D} \cdot (\sqrt{D} + 4\rho) \ell^{\frac{D-1}{D}}$$

Tandis que la borne de Bravyi-Terhal se limite aux réseaux cubiques réguliers de la forme $(\mathbb{Z}/L\mathbb{Z})^D$, la nôtre s'applique à des quotients de $\mathbb{Z}^D$ bien plus généraux.

Cette généralisation n'entraîne qu'un coût négligeable : appliquée aux mêmes réseaux que ceux de Bravyi et Terhal, notre borne n'est moins précise que d'un facteur constant (inférieur ou égal à 6), et prend alors la forme $d_{\min} < m(4\rho + 2)\ell^{\frac{D-1}{D}}$ où $\ell = L^D$. Comme nous l'expliquerons plus tard dans le chapitre, cette perte de précision est inhérente à notre méthode de démonstration, qui doit rester valable pour tout choix de Λ .

En outre, en appliquant la borne du Théorème 5.4.1 aux codes 2BGA abéliens, nous établissons une borne sur leur distance minimale. Plus précisément, nous montrons que pour tout code de cette famille de longueur $n = 2\ell$, dont le groupe stabilisateur est engendré par des éléments de poids w , sa distance minimale d satisfait :

$$d < 2\sqrt{\gamma_D}(\sqrt{D} + 4)\ell^{\frac{D-1}{D}}$$

où $D = w - 2$ et γ_D est la constante d'Hermite D -dimensionnelle. Puisque $\gamma_D = O(D)$, cela signifie que la distance minimale de ces codes est majorée par de $O(w \ell^{(w-3)/(w-2)})$.

Cette borne sur la distance minimale s'applique également aux codes Bicycle généralisés. Elle est à la fois plus explicite et plus générale que celle donnée dans [16], qui était de $O(D \ell^{\frac{D-1}{D}})$ avec $D \leq w - 1$ où w est le poids des générateurs du groupe stabilisateur.

5.1.4 Vue d'ensemble sur le théorème 5.4.1

Notre stratégie de preuve s'inspire directement des travaux de Bravyi et Terhal, mais est adaptée pour les réseaux plus généraux. L'idée générale est de partitionner l'espace où sont situés les qubits en régions plus petites, puis de montrer qu'il existe un opérateur logique non trivial dont le support est entièrement contenu dans l'une de ces régions. De ce fait, on peut déduire que la distance minimale du code est bornée par le nombre de qubits dans cette région, qu'il ne reste plus qu'à majorer.

La preuve se divise en plusieurs étapes.

1. **Choix d'une base optimisée** : Nous commençons par choisir une base spécifique pour le réseau Λ , qui nous permet de diviser de manière optimale l'espace où sont situés les qubits, à savoir le domaine fondamental de Λ associé. Dans notre cas, il s'agit d'une base $(\mathbf{u}_1, \dots, \mathbf{u}_D)$ de Λ dont la norme $\|\mathbf{u}_D^*\|$ du dernier vecteur dans la base de Gram-Schmidt associée, $(\mathbf{u}_1^*, \dots, \mathbf{u}_D^*)$, soit maximisée.
2. **Découpage de l'espace en tranches** : Ensuite, nous découpons l'espace des qubits en tranches identiques, qui sont translatées les unes des autres. Elles sont conçues pour être suffisamment larges pour que les générateurs du groupe stabilisateur n'agissent que localement, sans jamais interagir avec des qubits situés dans des tranches non adjacentes.

3. **Existence d'un opérateur logique** : Ensuite, nous utilisons un lemme de nettoyage (introduit par Bravyi et Terhal) pour prouver qu'il existe forcément un opérateur logique non trivial dont le support est entièrement contenu dans l'une de ces tranches. Ceci confirme que la distance minimale du code ne peut être supérieure au nombre total de qubits contenus dans cette tranche, à savoir m fois le nombre de points à coordonnées entières contenus dans cette tranche.
4. **Majoration du nombre de qubits dans la tranche** :
Il est souvent complexe de compter précisément la quantité de points à coordonnées entières contenus dans une tranche, car les tranches et la grille du réseau $\mathbb{Z}^D$ ne sont pas nécessairement bien alignées. Nous utilisons donc une estimation en majorant ce nombre par le volume d'une région plus grande que la tranche. Cette approche est plus générale, mais elle nous impose une légère pénalité : notre borne est moins précise que celle de Bravyi-Terhal, d'un facteur pouvant aller jusqu'à $\sqrt{D}$. Ce facteur provient de l'approximation que nous avons dû faire pour estimer le nombre de qubits contenus dans une tranche du domaine fondamental d'un réseau général. Bravyi et Terhal ont pu éviter cette approximation, puisque leurs calculs étaient limités aux réseaux cubiques réguliers, pour lesquels le nombre de points entiers contenus dans une tranche correspond directement à son volume.

5.2 Organisation du chapitre

Ce chapitre est structuré en plusieurs sections. Nous commencerons par présenter les outils nécessaires à la démonstration du théorème 5.4.1 dans la section 5.3. Ensuite, les sections 5.4 à 5.6 seront consacrées à la présentation de notre borne sur la distance minimale. Pour finir, la section 5.7 détaillera comment une borne sur la distance minimale des codes 2BGA abéliens se déduit du théorème 5.4.1.

5.3 Résultats préliminaires

Outre les notions fondamentales sur les réseaux évoquées dans la section 3.1 du Chapitre 3, nous introduisons deux concepts plus avancés, essentiels pour la démonstration du théorème 5.4.1.

- **Mesure de Haar** : Nous utiliserons la mesure de Haar sur $\mathbb{R}^D/\Lambda$, induite par la mesure de Lebesgue. Cet outil servira à estimer et à majorer le nombre de qubits présents dans les tranches.
- **Constantes de Rankin** : Ces constantes joueront un rôle crucial. Elles nous guideront dans le choix d'une base particulière du Λ , grâce à laquelle nous pourrons effectuer le découpage de l'espace des qubits décrit dans la section précédente.

5.3.1 Mesure de Haar

Soit Λ un sous-réseau de $\mathbb{Z}^D$. La projection canonique $\pi : \mathbb{R}^D \rightarrow \mathbb{R}^D/\Lambda$ est une application continue pour la topologie induite par la distance dist , introduite dans la proposition 3.1.2 du Chapitre 3.

Tribu des boréliens sur $\mathbb{R}^D/\Lambda$

L'espace quotient $\mathbb{R}^D/\Lambda$ est équipé d'une tribu de boréliens, que nous noterons $\text{Bor}(\mathbb{R}^D/\Lambda)$. Cette tribu est composée de tous les sous-ensembles $\mathbb{A}$ de $\mathbb{R}^D/\Lambda$ dont l'image réciproque par π est un borélien de $\mathbb{R}^D$, soit $\pi^{-1}(\mathbb{A}) \in \text{Bor}(\mathbb{R}^D)$.

Puisque la projection $\pi : \mathbb{R}^D \rightarrow \mathbb{R}^D/\Lambda$ est surjective, la tribu des boréliens de $\mathbb{R}^D/\Lambda$ est également composée de toutes les images directes par π de boréliens de $\mathbb{R}^D$.

Définition de la mesure

Considérons $\mathcal{P}$ un domaine fondamental de Λ . La mesure de Haar $\mathcal{M}$ sur l'espace quotient $\mathbb{R}^D/\Lambda$ est définie comme la mesure image de la mesure de Lebesgue restreinte à $\mathcal{P}$, par l'application π :

$$\mathbb{A} \in \text{Bor}(\mathbb{R}^D/\Lambda), \mathcal{M}(\mathbb{A}) = \text{Leb}(\pi^{-1}(\mathbb{A}) \cap \mathcal{P})$$

5.3.2 Constantes de Hermite et de Rankin

Dans le chapitre 3, nous avons vu que les réseaux sont des groupes discrets. Chacun d'eux possède un vecteur non nul de longueur minimale à partir duquel nous définirons les constantes de Hermite et de Rankin.

Constante de Hermite

Historiquement, la constante de Hermite, introduite en 1850 par Charles Hermite [32], a été initialement utilisée dans l'étude des formes quadratiques, afin de déterminer la densité maximale d'hypersphères pouvant être empilées dans un volume donné. Dans ce chapitre, nous adoptons plutôt une approche basée sur la théorie des réseaux pour définir la constante de Hermite.

Définition 5.3.1 (Constante de Hermite). *Pour un réseau $\Lambda \subset \mathbb{R}^D$, on définit la quantité auxiliaire $\gamma_D(\Lambda) = \min \left\{ \left(\frac{\|x\|}{\text{vol}(\Lambda)^{1/D}} \right)^2 \mid x \in \Lambda \setminus \{0\} \right\}$.*

La constante de Hermite D -dimensionnelle, notée γ_D , se définit comme la valeur maximale de ces quantités prises sur tous les réseaux de rang plein:

$$\gamma_D = \max \left\{ \gamma_D(\Lambda) \mid \Lambda \subset \mathbb{R}^D \text{ réseau de rang plein} \right\}$$

Théorème de Minkowski

Calculer directement la longueur minimale d'un vecteur non nul d'un réseau est un problème complexe. Cependant, le théorème de Minkowski nous fournit un moyen de la majorer.

Théorème 5.3.1 (Minkowski [2]). *Soit Λ un réseau de $\mathbb{R}^D$ de rang plein et soit $\mathbb{A} \subseteq \mathbb{R}^D$ un ensemble Lebesgue-mesurable, convexe, symétrique tel que $\text{Leb}(\mathbb{A}) > 2^D \text{vol}(\Lambda)$. Alors $\mathbb{A}$ contient un élément non nul de Λ .*

Ce théorème garantit en particulier que toute boule de rayon suffisamment grand contient au moins un vecteur non nul du réseau. Cela permet de déduire une majoration de la constante de Hermite [33], à savoir :

$$0 < \gamma_D \leq 1 + \frac{D}{4} \tag{5.1}$$

Constante de Rankin

Les constantes de Rankin généralisent le concept introduit par Hermite. Alors que la constante de Hermite se définit à partir de la longueur du plus petit vecteur non nul d'un réseau Λ , la m -ième constante de Rankin s'appuie sur le covolume minimal d'un sous-réseau de Λ de rang m .

Définition 5.3.2 (Constante de Rankin). Soient $\Lambda \subset \mathbb{R}^D$ un réseau de rang plein et $1 \leq m < D$ un entier. La m -ième constante de Rankin se définit à l'aide de $\gamma_{D,m}(\Lambda)$, un invariant de Λ qui est défini par le rapport suivant :

$$\gamma_{D,m}(\Lambda) = \min \left\{ \left(\frac{\text{vol}(\{\mathbf{x}_1, \dots, \mathbf{x}_m\})}{\text{vol}(\Lambda)^{\frac{m}{D}}} \right)^2 \mid \mathbf{x}_1, \dots, \mathbf{x}_m \in \Lambda \text{ avec } \text{vol}(\{\mathbf{x}_1, \dots, \mathbf{x}_m\}) \neq 0 \right\}$$

où $\text{vol}(\{\mathbf{x}_1, \dots, \mathbf{x}_m\})$ correspond au covolume du sous-réseau de Λ engendré par la famille libre $\{\mathbf{x}_1, \dots, \mathbf{x}_m\}$ (pour la définition, se référer à 3.1.4).

La m -ième constante de Rankin $\gamma_{D,m}$ est la borne supérieure de ces quantités, prises sur tous les réseaux de rang D :

$$\gamma_{D,m} = \max \{ \gamma_{D,m}(\Lambda) \mid \Lambda \subset \mathbb{R}^D \text{ réseau de rang plein} \}$$

Les constantes de Hermite et Rankin sont reliées par les relations suivantes :

Proposition 5.3.2 (Relation avec la constante de Hermite [33]). Pour deux entiers m et D tels que $1 \leq m < D$ les constantes de Rankin et Hermite vérifient les relations :

- $\gamma_{D,m} = \gamma_{D,D-m}$
- $\gamma_{D,1}$ coïncide avec γ_D , la constante de Hermite D -dimensionnelle.

Cette dernière relation confirme que la constante de Rankin est une généralisation de la constante de Hermite.

5.4 Borne de Bravyi-Terhal généralisée

Avant de présenter notre résultat, à savoir la borne sur la distance minimale des codes stabilisateurs locaux, nous rappelons brièvement la borne de Bravyi-Terhal établie par nos prédécesseurs.

5.4.1 Borne de Bravyi-Terhal originale

En 2009, Bravyi et Terhal ont établi une borne supérieure sur la distance minimale des codes stabilisateurs géométriquement locaux [24]. Dans leur approche, les qubits sont placés sur une partie de $\mathbb{Z}^D$, de sorte que le groupe stabilisateur soit engendré par une famille de générateurs dont le support est inclus dans une boule de rayon ρ . Plus

précisément, ils ont prouvé que la distance minimale d de ces codes satisfait $d \leq \rho \ell^{\frac{D-1}{D}}$ où ℓ représente la longueur du code et ρ le paramètre de localité.

Nous généralisons ce résultat en étendant cette borne aux codes dont les qubits sont placés sur les points entiers du domaine fondamental d'un réseau $\Lambda \subset \mathbb{Z}^D$ de rang plein. Cette approche permet de couvrir une classe de codes beaucoup plus vaste, mais introduit un facteur correctif qui dépend de la dimension D et des propriétés du réseau.

5.4.2 Borne de Bravyi-Terhal généralisée

Théorème 5.4.1. *Soit $\mathcal{Q}$ un code stabilisateur de longueur $n = m\ell$ où m et ℓ sont deux entiers supérieur ou égaux à 1.*

Supposons qu'il existe un entier $D \geq 2$ et un sous-réseau de rang plein $\Lambda \subseteq \mathbb{Z}^D$, tels que :

- *Le covolume de Λ est égal à ℓ , c'est-à-dire $\text{vol}(\Lambda) = |\mathbb{Z}^D / \Lambda| = \ell$.*
- *Les qubits peuvent être représentés sur l'ensemble des sommets du réseau quotient $\mathbb{Z}^D / \Lambda$, de telle sorte que chaque sommet indexe m qubits et que le groupe stabilisateur soit engendré par une famille de générateurs dont le support est contenu dans une boule de rayon ρ de ce quotient (pour la distance induite par la norme euclidienne).*

Alors la distance minimale est majorée par :

$$d < m \cdot \max(1, \sqrt{\gamma_{D,D-1}(\Lambda)}) \cdot (\sqrt{D} + 4\rho) \ell^{\frac{D-1}{D}}$$

où $\gamma_{D,D-1}(\Lambda)$ désigne la constante auxiliaire introduite dans la définition 5.3.2 de la $(D-1)$ -ième constante de Rankin.

À partir de cette borne, il est possible d'obtenir, à l'aide de la constante de Hermite D -dimensionnelle, une majoration indépendante du choix de Λ .

Corollaire 5.4.1. *Sous les hypothèses du théorème 5.4.1, la distance minimale de $\mathcal{Q}$ vérifie :*

$$d_{\min} < m \sqrt{\gamma_D} (\sqrt{D} + 4\rho) \cdot \ell^{\frac{D-1}{D}}$$

où γ_D est la constante d'Hermite D -dimensionnelle.

Preuve : Pour établir le corollaire à partir du théorème, il suffit de montrer que la constante de Hermite D -dimensionnelle vérifie : $\max(1, \gamma_{D,D-1}(\Lambda)) \leq \gamma_D$.

D'après la définition 5.3.1, la constante de Hermite γ_D est définie comme le maximum des constantes $\gamma_D(\mathbb{L})$ prises sur tout les réseaux $\mathbb{L} \subset \mathbb{R}^D$ de rang plein. En particulier, nous avons que $\gamma_D \geq \gamma_D(\mathbb{Z}^D) = 1$.

De plus, d'après la proposition 5.3.2, la constante de Hermite γ_D coïncide avec le maximum, pris sur tous les réseaux $\mathbb{L}$ de dimension D , des constantes $\gamma_{D,D-1}(\mathbb{L})$ introduites dans la définition 5.3.2. On en déduit donc que $0 < \gamma_{D,D-1}(\Lambda) \leq \gamma_D$. $\square$

Remarque 5.4.1 (Abus de langage). *Pour simplifier les notations, nous identifierons désormais les qubits aux points du quotient $\mathbb{Z}^D / \Lambda$ qui les indexent. Ainsi, au lieu de dire que les qubits sont indexés par des points du quotient, nous dirons qu'ils sont situés sur ces points.*

5.5 Esquisse de la preuve du théorème

Pour établir la démonstration du théorème, nous distinguons deux cas : selon que $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$ ou non.

5.5.1 Cas des réseaux de faible covolume : $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$

Pour les réseaux Λ dont le covolume ℓ vérifie $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$, la preuve du Théorème 5.4.1 est directe.

D'après la borne Singleton, la distance minimale d'un code quantique de paramètres $[[n, k, d]]$ satisfait $n - k \geq 2(d - 1)$. Ainsi, en particulier, $2d \leq n - k + 2 \leq n + 2$.

Si $\ell^{1/D} < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$, alors nous avons la majoration suivante :

$$n = m\ell = m \cdot \ell^{\frac{1}{D}} \cdot \ell^{\frac{D-1}{D}} < 8m \cdot \rho\sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{\frac{D-1}{D}}$$

Comme $\ell, m \geq 1$ alors $m \cdot \ell^{\frac{D-1}{D}} \geq 1$, nous en déduisons que :

$$d \leq \frac{n+2}{2} < 4m \cdot \rho\sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{\frac{D-1}{D}} + m \cdot \ell^{\frac{D-1}{D}} = m \cdot (4\rho\sqrt{\gamma_{D,D-1}(\Lambda)} + 1) \cdot \ell^{\frac{D-1}{D}}$$

En considérant séparément les cas $\gamma_{D,D-1}(\Lambda) \leq 1$ et $\gamma_{D,D-1}(\Lambda) > 1$, et en utilisant le fait que $D \geq 2$, on déduit que :

$$d < m \cdot \max(1, \sqrt{\gamma_{D,D-1}(\Lambda)}) \cdot (1 + 4\rho) \ell^{\frac{D-1}{D}} < m \cdot \max(1, \sqrt{\gamma_{D,D-1}(\Lambda)}) \cdot (\sqrt{D} + 4\rho) \ell^{\frac{D-1}{D}}$$

5.5.2 Cas général : $\ell^{1/D} \geq 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$

Notre démonstration s'appuie sur des résultats d'algèbre que nous présenterons ultérieurement, ainsi que sur le lemme de nettoyage de Bravyi et Terhal [24].

Lemme de nettoyage

Avant d'énoncer le lemme de nettoyage, nous introduisons la notion d'opérateur logique.

Définition 5.5.1 (Opérateur logique). *Soit $\mathcal{Q}$ un code stabilisateur dont le groupe de stabilisateur $\mathbb{S}$ est un sous-groupe commutatif de $\mathbb{P}_n$ ne contenant pas $-\mathbf{I}^{\otimes n}$.*

On appelle centralisateur de $\mathbb{S}$ dans $\mathbb{P}_n$, noté $C_{\mathbb{P}_n}(\mathbb{S})$, le sous-groupe de $\mathbb{P}_n$ formé des opérateurs qui commutent avec tous les éléments de $\mathbb{S}$:

$$C_{\mathbb{P}_n}(\mathbb{S}) = \left\{ \mathcal{O} \in \mathbb{P}_n \mid \forall \mathcal{S} \in \mathbb{S}, \mathcal{O}\mathcal{S} = \mathcal{S}\mathcal{O} \right\}.$$

Comme $\mathbb{S}$ est abélien, on a $\mathbb{S} \subseteq C_{\mathbb{P}_n}(\mathbb{S})$. Les éléments de $C_{\mathbb{P}_n}(\mathbb{S})$ sont appelés opérateurs logiques, et ceux de $C_{\mathbb{P}_n}(\mathbb{S}) \setminus (\mathbb{S} \cup \{-\mathbf{I}^{\otimes n}, \pm i\mathbf{I}^{\otimes n}\})$ sont appelés opérateurs logiques non triviaux.

Lemme 5.5.1 (Lemme de nettoyage [24]). *Pour un code stabilisateur $\mathcal{Q}$ de groupe stabilisateur $\mathbb{S}$ et V un ensemble de qubits du code, il n'existe que deux possibilités :*

- *Soit il existe un opérateur logique non trivial, c'est-à-dire un élément de $C_{\mathbb{P}_n}(\mathbb{S}) \setminus \mathbb{S}$, dont le support est entièrement contenu dans V .*
- *Soit tout opérateur logique peut être "nettoyé" sur V , c'est-à-dire qu'il peut être transformé, en le multipliant par des générateurs du groupe stabilisateur, en un opérateur équivalent qui agit trivialement sur V .*

Plan de la preuve :

Notre stratégie se base sur une approche géométrique. Nous allons partitionner l'espace des qubits en petites régions contenant chacune moins de $m\sqrt{\gamma_{D,D-1}(\Lambda)}(\sqrt{D}+4\rho) \cdot \ell^{\frac{D-1}{D}}$ qubits. Puis, nous utiliserons le lemme de nettoyage pour prouver qu'il existe nécessairement un opérateur logique non trivial dont le support est entièrement contenu dans l'une de ces régions. La distance minimale du code sera alors limitée par le nombre de qubits contenus dans cette région.

Plus précisément, nous suivrons les étapes suivantes pour démontrer le théorème 5.4.1 dans le cas général :

- **1. Choix de la base de Λ :** Nous commençons par choisir une base de notre réseau Λ qui est optimisée pour le découpage. Il s'agit d'une base $\mathbf{u}_1, \dots, \mathbf{u}_D$ de Λ qui une fois orthogonalisée par l'algorithme de Gram-Schmidt, fournira un dernier vecteur $\mathbf{u}_D^*$ dont la longueur est maximale.
- **2. Partitionnement du domaine fondamental :** Nous découpons ensuite le domaine fondamental du réseau correspondant en un nombre pair de tranches similaires (identiques à translation près). Ce découpage garantit qu'aucun générateur du groupe stabilisateur respectant la condition de localité ne puisse agir simultanément sur des qubits de tranches non adjacentes.
- **3. Application du lemme de nettoyage :** Nous appliquons ce lemme à plusieurs reprises afin de garantir l'existence d'un opérateur logique non trivial dont le support est entièrement contenu dans l'une de ces tranches. La distance minimale sera alors majorée par le nombre de qubits dans cette tranche.
- **4. Majoration du nombre de qubits contenu dans une tranche :** Nous interprétons ensuite le nombre de qubits dans cette tranche comme le volume d'un ensemble, que nous majorons par $\sqrt{\gamma_{D,D-1}(\Lambda)}(\sqrt{D}+4\rho) \cdot \ell^{\frac{D-1}{D}}$.
- **5. Majoration de la distance minimale :** En combinant les résultats des étapes précédentes, nous obtenons la borne annoncée sur la distance minimale :

$$d < m\sqrt{\gamma_{D,D-1}(\Lambda)}(\sqrt{D}+4\rho) \cdot \ell^{\frac{D-1}{D}}$$

5.6 Preuve du théorème 5.4.1

Soient $\mathcal{Q}$ un code stabilisateur de longueur $n = m\ell$ et Λ un sous-réseau de $\mathbb{Z}^D$ de rang plein et de covolume $\text{vol}(\Lambda) = |\mathbb{Z}^D/\Lambda| = \ell$.

Dans toute la suite de cette section, nous supposons que les qubits de $\mathcal{Q}$ sont disposés sur les sommets de $\mathbb{Z}^D/\Lambda$, avec m qubits par sommet, et qu'il existe un réel $\rho > 0$ tel que le groupe stabilisateur soit engendré par une famille de générateurs dont le support est contenu dans une boule de $\mathbb{Z}^D/\Lambda$ de rayon ρ pour la distance induite par la norme euclidienne. Enfin, nous supposons que le covolume de Λ satisfait : $\ell^{1/D} \geq 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$.

5.6.1 Construction d'une base du réseau adaptée au découpage

La démonstration du Théorème 5.4.1 repose sur l'application répétée du lemme de nettoyage. Pour que ce lemme soit applicable, la première étape est de diviser l'espace des qubits en tranches, de sorte qu'aucun générateur du groupe stabilisateur ne puisse agir simultanément sur des qubits situés dans des tranches non adjacentes.

La clé pour réaliser ce découpage est de choisir une base particulière pour notre réseau Λ . L'objectif est de construire une base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ dont le dernier vecteur de Gram-Schmidt, $\mathbf{u}_D^*$, a une longueur maximale. Ce choix est crucial, car la longueur de $\mathbf{u}_D^*$ détermine directement l'épaisseur des tranches, qui doit être suffisamment grande pour assurer que la condition sur les générateurs du groupe stabilisateur soit satisfaite.

Même s'il n'est pas toujours possible de compléter une famille libre du réseau en une base, la base que nous recherchons peut être construite de cette manière. Précisément, la construction de la base se fera à partir d'une famille de $D - 1$ vecteurs $\mathbf{u}_1, \dots, \mathbf{u}_{D-1} \in \Lambda$ qui sont linéairement indépendants et qui engendrent un sous-réseau de covolume minimal parmi tous les sous-réseaux de Λ de rang $D - 1$.

Le lemme qui suit détaille la construction de cette base particulière.

Choix de la base

Lemme 5.6.1. *Il existe une famille de $D - 1$ vecteurs linéairement indépendants $\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\} \subset \Lambda$ qui peut être complétée en une base de Λ et qui satisfait la propriété suivante :*

$$0 < \text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\}) = \sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{\frac{D-1}{D}}$$

où :

- $\text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\})$ est le covolume du réseau engendré par $\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\}$.
- $\sqrt{\gamma_{D,D-1}(\Lambda)}$ est le minimum des rapports $\frac{\text{vol}(\mathbb{L})}{\ell^{(D-1)/D}}$ pris sur tous les sous-réseaux $\mathbb{L} \subset \Lambda$ de rang $D - 1$.

Preuve : Puisque Λ est inclus dans $\mathbb{Z}^D$ (par hypothèse), l'ensemble $\mathbb{V} = \{\text{vol}(\mathbb{L})^2 \mid \mathbb{L} \subset \Lambda \text{ sous-réseau de rang } D - 1\}$ est un sous-ensemble non vide d'entiers naturels non nuls. Il admet donc un plus petit élément : $\text{vol}(\mathbb{L}_0)^2$ obtenu pour $\mathbb{L}_0$ un sous-réseau de Λ de rang $D - 1$.

Par définition, nous avons donc $0 < \text{vol}(\mathbb{L}_0) = \sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{(D-1)/D}$. D'après le théorème des facteurs invariants présenté dans le Chapitre 3 (voir Théorème 3.1.4), il existe $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ une base de Λ et des entiers non nuls $a_1, \dots, a_{D-1} \in \mathbb{Z} \setminus \{0\}$ tels que $\{a_1 \mathbf{u}_1, \dots, a_{D-1} \mathbf{u}_{D-1}\}$ forme une base de $\mathbb{L}_0$.

Nous allons montrer que la famille $\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\}$ satisfait :

$$\text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\}) = \text{vol}(\mathbb{L}_0) = \sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{(D-1)/D}$$

Notons $\mathbb{T}$, le sous-réseau de Λ de rang $D - 1$ engendré par $\mathbf{u}_1, \dots, \mathbf{u}_{D-1}$. Comme $\mathbb{L}_0$ est un $\mathbb{Z}$ -module engendré par $a_1 \mathbf{u}_1, \dots, a_{D-1} \mathbf{u}_{D-1} \in \mathbb{T}$, nous avons la suite d'inclusions :

$$\mathbb{L}_0 \subseteq \mathbb{T} \subseteq \mathbb{Z}^D$$

En conséquence, leurs covolumes satisfont : $\text{vol}(\mathbb{L}_0) = [\mathbb{T} : \mathbb{L}_0] \cdot \text{vol}(\mathbb{T})$. Puisque $[\mathbb{T} : \mathbb{L}_0] = \prod_{i=1}^{D-1} |a_i| \geq 1$, alors $\text{vol}(\mathbb{L}_0) \geq \text{vol}(\mathbb{T})$. Or $\mathbb{T}$ étant un sous-réseau Λ de rang $D - 1$, son covolume vérifie également $\text{vol}(\mathbb{T}) \geq \text{vol}(\mathbb{L}_0)$ (par minimalité de $\text{vol}(\mathbb{L}_0)$). D'où $\text{vol}(\mathbb{L}_0) = \text{vol}(\mathbb{T})$. $\square$

Remarque 5.6.1. Outre l'égalité de leurs covolumes, les deux réseaux, $\mathbb{L}_0$ et $\mathbb{T}$ de la démonstration précédente, sont en réalité égaux. L'égalité provient du fait que nous avons déjà l'inclusion $\mathbb{L}_0 \subseteq \mathbb{T}$ et que les deux réseaux ont même rang et même covolume.

Nous allons maintenant examiner les propriétés de la base de Gram-Schmidt associée à la base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ de Λ .

Propriétés de la base de Gram-Schmidt associée :

Étant donné que $\Lambda \subseteq \mathbb{Z}^D$ est un réseau de rang plein, sa base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ est également une base de $\mathbb{R}^D$. Cela résulte du fait que le déterminant de la matrice formée par ces vecteurs est non nul :

$$|\det(\mathbf{u}_1, \dots, \mathbf{u}_D)| = \det(\Lambda) = |\mathbb{Z}^D / \Lambda| \geq 1$$

La base de Gram-Schmidt associée à $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ est une base orthogonale de $\mathbb{R}^D$ construite par récurrence :

- $\mathbf{u}_1^* = \mathbf{u}_1$
- $\forall i \in \llbracket 1, D-1 \rrbracket, \mathbf{u}_{i+1}^* := \mathbf{u}_{i+1} - \sum_{k=1}^i \frac{(\mathbf{u}_{i+1} | \mathbf{u}_k^*)}{\|\mathbf{u}_k^*\|^2} \mathbf{u}_k^*.$

Pour chaque entier $j \in \llbracket 1, D \rrbracket$, les j premiers vecteurs de la base de Gram-Schmidt ont deux propriétés importantes :

- L'espace vectoriel engendré par les j premiers vecteurs de la base originale, $\{\mathbf{u}_1, \dots, \mathbf{u}_j\}$, est le même que celui engendré par les j premiers vecteurs orthogonalisés, $\{\mathbf{u}_1^*, \dots, \mathbf{u}_j^*\}$:

$$\text{Vect}_{\mathbb{R}}(\mathbf{u}_1, \dots, \mathbf{u}_j) = \text{Vect}_{\mathbb{R}}(\mathbf{u}_1^*, \dots, \mathbf{u}_j^*)$$

- Le covolume du sous-réseau engendré par les j premiers vecteurs de la base originale est égal au produit des normes des j premiers vecteurs de la base de Gram-Schmidt :

$$\text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_j\}) = \prod_{k=1}^j \|\mathbf{u}_k^*\|$$

Cette dernière propriété nous permet d'établir une borne inférieure pour la norme de $\mathbf{u}_D^*$.

Lemme 5.6.2. $\|\mathbf{u}_D^*\| \geq 8\rho.$

Preuve : Par hypothèse sur ℓ , l'inégalité $\frac{\ell^{1/D}}{\sqrt{\gamma_{D,D-1}(\Lambda)}} \geq 8\rho$ est vraie.

Puisque le covolume ℓ de Λ vérifie $\ell = \prod_{i=1}^D \|\mathbf{u}_i^*\| = \text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\}) \|\mathbf{u}_D^*\|$ alors :

$$\begin{aligned}
 \|\mathbf{u}_D^*\| &= \frac{\ell}{\text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\})} \\
 &= \frac{\ell}{\sqrt{\gamma_{D,D-1}(\Lambda)} \cdot \ell^{(D-1)/D}} \quad (\text{par construction de } \mathbf{u}_1, \dots, \mathbf{u}_D) \\
 &= \frac{\ell^{1/D}}{\sqrt{\gamma_{D,D-1}(\Lambda)}} \\
 &\geq 8\rho
 \end{aligned}$$

□

Maintenant que nous avons démontré que la norme de $\mathbf{u}_D^*$ est supérieure à 8ρ (où ρ est le paramètre de localité) nous pouvons passer à l'étape suivante : le découpage du domaine fondamental.

5.6.2 Découpage du domaine fondamental

Pour appliquer le lemme de nettoyage, notre stratégie consiste à partitionner l'espace des qubits, $\mathbb{R}^D/\Lambda$, en plusieurs tranches. L'espace $\mathbb{R}^D$ étant pavé par des translatés du domaine fondamental $\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D)$ du réseau Λ , il nous suffit de découper ce domaine fondamental pour obtenir une partition de l'espace des qubits.

Définition des tranches

Nous allons décomposer $\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D)$ en μ tranches identiques, que nous noterons T_q . L'entier μ , que nous expliciterons dans la section suivante, est un nombre pair dont la valeur exacte dépend du réseau Λ et du paramètre de localité ρ .

Plus précisément, pour q allant de 0 à $\mu - 1$, la tranche T_q est définie par :

$$T_q = \left\{ \sum_{i=1}^{D-1} x_i \mathbf{u}_i + \frac{x_D}{\mu} \mathbf{u}_D \mid \forall i \in \llbracket 1, D-1 \rrbracket, 0 \leq x_i < 1 \text{ et } q \leq x_D < q+1 \right\}$$

L'espace $\mathbb{R}^D/\Lambda$ est alors partitionné en μ tranches disjointes $\pi(T_q)$, qui sont les images des tranches T_q par la projection canonique (voir Figure 5.1).

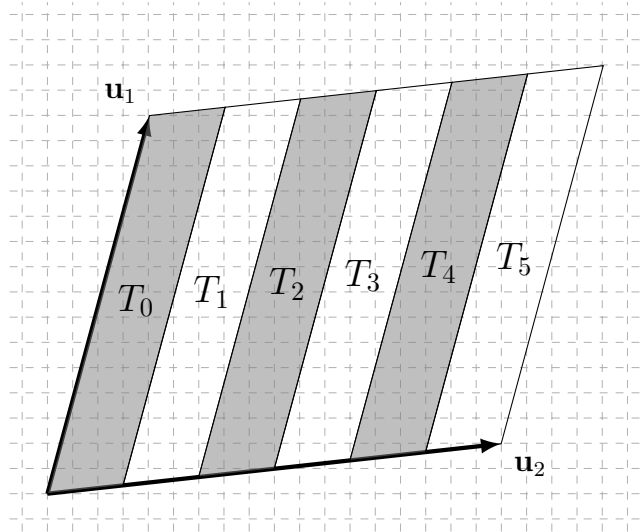


Figure 5.1: Découpage du domaine fondamental en 6 tranches.

Contraintes sur μ

La précision de notre borne dépend directement du nombre de tranches μ . Pour obtenir la meilleure borne possible sur la distance minimale, nous cherchons à minimiser le nombre de qubits par tranche, ce qui nécessite de maximiser μ , le nombre de tranches.

Cependant, le lemme de nettoyage que nous utilisons pour établir notre borne impose une contrainte de localité : les tranches doivent être suffisamment larges pour qu'aucun des générateurs du groupe stabilisateur, satisfaisant la condition de localité, ne puisse agir simultanément sur des qubits situés dans des tranches non adjacentes. Cette condition interdit de choisir des tranches trop fines, ce qui limite la valeur maximale que peut prendre μ .

La prochaine étape de la preuve consiste à choisir l'entier μ de manière à ce qu'il maximise le nombre de tranches tout en respectant cette contrainte de localité.

Construction de μ

Pour définir la valeur de μ , nous utilisons la borne que nous avons établie sur la norme de $\mathbf{u}_D^*$ dans le lemme 5.6.2, à savoir $\|\mathbf{u}_D^*\| \geq 8\rho$. Précisément, μ est choisi de la façon suivante :

$$\mu = \begin{cases} \left\lfloor \frac{\|\mathbf{u}_D^*\|}{2\rho} \right\rfloor & \text{si c'est un entier pair,} \\ \left\lfloor \frac{\|\mathbf{u}_D^*\|}{2\rho} \right\rfloor - 1 & \text{sinon.} \end{cases}$$

où pour $x \in \mathbb{R}$, la notation $\lfloor x \rfloor$ désigne la partie entière inférieure de x , c'est-à-dire l'unique entier $k \in \mathbb{Z}$ tel que $k \leq x < k + 1$.

Ce choix garantit deux choses :

1. μ est un entier pair supérieur ou égal à 4.
2. L'épaisseur (commune) des tranches, $\lambda = \frac{\|\mathbf{u}_D^*\|}{\mu}$, satisfait l'inégalité : $2\rho \leq \lambda < 4\rho$.

Le deuxième point découle des inégalités : $0 < \frac{\|\mathbf{u}_D^*\|}{2\rho} - 2 < \mu \leq \frac{\|\mathbf{u}_D^*\|}{2\rho}$, qui impliquent :

$$2\rho \leq \frac{\|\mathbf{u}_D^*\|}{\mu} = \lambda < \frac{\|\mathbf{u}_D^*\| \cdot 2\rho}{\|\mathbf{u}_D^*\| - 4\rho} \leq 4\rho$$

Localité des générateurs du groupe stabilisateur

La construction de nos tranches garantit que leur épaisseur est supérieure ou égale à 2ρ . Le lemme suivant prouve que cette épaisseur est suffisante pour s'assurer qu'aucun générateur du groupe stabilisateur, satisfaisant la condition de localité, n'agisse simultanément sur des qubits situés dans des tranches non adjacentes.

Lemme 5.6.3. *Si $\lambda \geq 2\rho$, alors aucun générateur du groupe stabilisateur, satisfaisant la condition de localité du théorème 5.4.1, n'agit simultanément sur des qubits de tranches non adjacentes.*

Preuve : Raisonnons par l'absurde et supposons qu'il existe $\mathcal{O}$ un générateur du groupe stabilisateur, respectant la condition de localité du théorème 5.4.1, qui agit simultanément sur des qubits situés dans deux tranches non adjacentes. Cela signifie qu'il existe deux indices distincts i et j dans $\llbracket 0, \mu - 1 \rrbracket$ tels que :

- $1 < j - i < \mu - 1$
- $\mathcal{O}$ agit simultanément sur des qubits des tranches $\pi(T_i)$ et $\pi(T_j)$ (qui ne sont ni confondues ni adjacentes), où ici π désigne la projection canonique $\mathbb{R}^D \rightarrow \mathbb{R}^D / \Lambda$.

Dans $\mathbb{R}^D / \Lambda$, la distance entre les ensembles $\pi(\mathbb{Z}^D \cap T_i)$ et $\pi(\mathbb{Z}^D \cap T_j)$ est donnée par :

$$\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) = \inf \{ \text{dist}([\mathbf{x}], [\mathbf{y}]) \mid ([\mathbf{x}], [\mathbf{y}]) \in \pi(\mathbb{Z}^D \cap T_i) \times \pi(\mathbb{Z}^D \cap T_j) \}$$

où $[\mathbf{x}]$ et $[\mathbf{y}]$ désignent les classes modulo Λ des vecteurs $\mathbf{x}$ et $\mathbf{y}$, et dist la distance induite par la norme euclidienne introduite dans la proposition 3.1.2 du Chapitre 3.

Pour aboutir à une contradiction, nous allons montrer que $\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j))$ est à la fois inférieure ou égale à 2ρ et strictement supérieure à 2ρ . Pour y arriver, nous utiliserons l'hypothèse sur $\mathcal{O}$, qui garantit que son support est inclus dans une boule de rayon ρ de $\mathbb{R}^D/\Lambda$, pour la distance dist .

$$\underline{\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) \leq 2\rho :}$$

Puisque $\mathcal{O}$ agit simultanément sur les deux tranches, il existe des points $\mathbf{q}_i \in \mathbb{Z}^D \cap T_i$ et $\mathbf{q}_j \in \mathbb{Z}^D \cap T_j$ tels que les qubits $\pi(\mathbf{q}_i) := [\mathbf{q}_i]$ et $\pi(\mathbf{q}_j) := [\mathbf{q}_j]$ appartiennent à son support. Comme par hypothèse, le support de $\mathcal{O}$ est inclus dans une boule de rayon ρ pour la distance dist , alors $\text{dist}([\mathbf{q}_i], [\mathbf{q}_j]) \leq 2\rho$. D'où $\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) \leq 2\rho$.

$$\underline{\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) > 2\rho :}$$

Puisque les ensembles $\pi(\mathbb{Z}^D \cap T_i)$ et $\pi(\mathbb{Z}^D \cap T_j)$ sont inclus dans $\pi(\mathbb{Z}^D) = \mathbb{Z}^D/\Lambda$, qui est fini, ils le sont aussi. Par conséquent, la distance entre ces deux ensembles est atteinte, c'est-à-dire qu'il existe un couple $(\mathbf{a}, \mathbf{b}) \in \mathbb{Z}^D \cap T_i \times \mathbb{Z}^D \cap T_j$ tel que :

$$\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) = \text{dist}([\mathbf{a}], [\mathbf{b}])$$

où $[\mathbf{a}]$ et $[\mathbf{b}]$ désignent les classes modulo Λ de $\mathbf{a}$ et $\mathbf{b}$.

D'après les propriétés de la distance dist , qui ont été rappelées dans la proposition 3.1.2 du Chapitre 3, il existe $\mathbf{g} \in \Lambda$ tels que $\text{dist}([\mathbf{a}], [\mathbf{b}]) = \|\mathbf{a} - (\mathbf{b} + \mathbf{g})\|$.

En utilisant le fait que $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ est une base de Λ et les définitions des tranches T_i et T_j de la section 5.6.2, nous obtenons les décompositions suivantes :

- $\mathbf{a} \in T_i$ s'écrit $\mathbf{a} = \sum_{k=1}^{D-1} a_k \mathbf{u}_k + \frac{a_D}{\mu} \mathbf{u}_D$ avec $a_1, \dots, a_{D-1} \in [0, 1[$ et $a_D \in [i, i+1[$
- $\mathbf{b} \in T_j$ s'écrit $\mathbf{b} = \sum_{k=1}^{D-1} b_k \mathbf{u}_k + \frac{b_D}{\mu} \mathbf{u}_D$ avec $b_1, \dots, b_{D-1} \in [0, 1[$ et $b_D \in [j, j+1[$
- $\mathbf{g} \in \Lambda$ s'écrit $\mathbf{g} = \sum_{k=1}^D g_k \mathbf{u}_k$ avec $g_k \in \mathbb{Z}$

Ce qui se réécrit de la façon suivante dans la base de Gram-Schmidt associée $\mathbf{u}_1^*, \dots, \mathbf{u}_D^*$:

- $\mathbf{a} = \sum_{k=1}^{D-1} \tilde{a}_k \mathbf{u}_k^* + \frac{a_D}{\mu} \mathbf{u}_D^*$ avec $\forall k \in \llbracket 1, D-1 \rrbracket, \tilde{a}_k \in \mathbb{R}$ et $a_D \in [i, i+1[$

- $\mathbf{b} = \sum_{k=1}^{D-1} \tilde{b}_k \mathbf{u}_k^* + \frac{b_D}{\mu} \mathbf{u}_D^*$ avec $\forall k \in \llbracket 1, D-1 \rrbracket$, $\tilde{b}_k \in \mathbb{R}$ et $b_D \in [j, j+1[$
- $\mathbf{g} = \sum_{k=1}^{D-1} \tilde{g}_k \mathbf{u}_k^* + g_D \mathbf{u}_D^*$ avec $\forall k \in \llbracket 1, D-1 \rrbracket$, $\tilde{g}_k \in \mathbb{R}$ et $g_D \in \mathbb{Z}$.

En appliquant le théorème de Pythagore, nous observons que :

$$\|\mathbf{a} - (\mathbf{b} + \mathbf{g})\|^2 \geq \left| \frac{a_D - (b_D + \mu g_D)}{\mu} \right|^2 \cdot \|\mathbf{u}_D^*\|^2 = |a_D - (b_D + \mu g_D)|^2 \cdot \lambda^2$$

Comme $\lambda \geq 2\rho$, pour montrer que $\text{dist}([\mathbf{a}], [\mathbf{b}]) = \|\mathbf{a} - (\mathbf{b} + \mathbf{g})\| > 2\rho$ il suffit de voir que l'inégalité $|a_D - (b_D + \mu g_D)| > 1$ est vraie pour tous $a_D \in [i, i+1[$, $b_D \in [j, j+1[$ et $g_D \in \mathbb{Z}$ (où i et j sont des entiers tels que $1 < j - i < \mu - 1$).

Pour $g_D = 0$, nous avons que $|a_D - (b_D + \mu g_D)| = |a_D - b_D| > j - (i+1) \geq 1$. Pour $g_D \neq 0$, on procède comme suit. D'abord, on remarque que :

$$|a_D - (b_D + \mu g_D)| = |\mu g_D - (a_D - b_D)| \geq |\mu |g_D| - |a_D - b_D||$$

Ensuite, puisque $\mu |g_D| \geq \mu$ et que $|b_D - a_D| < j+1 - i \leq \mu - 1$, nous obtenons que :

$$|\mu |g_D| - |b_D - a_D|| > \mu - (\mu - 1) \geq 1$$

En conclusion, $\text{dist}(\pi(\mathbb{Z}^D \cap T_i), \pi(\mathbb{Z}^D \cap T_j)) = \text{dist}([\mathbf{a}], [\mathbf{b}]) > 2\rho$. □

Ce lemme est crucial. Il nous permet d'appliquer le lemme de nettoyage, qui garantira l'existence d'un opérateur logique non trivial dont le support est confiné à une seule tranche. Nous en déduirons alors une majoration de la distance minimale du code.

5.6.3 Application du lemme de nettoyage

Maintenant que notre domaine fondamental est découpé en un nombre pair de tranches et que nous avons la garantie qu'aucun générateur du code stabilisateur, satisfaisant la condition de localité du théorème 5.4.1, n'agit simultanément sur des qubits situés dans des tranches non adjacentes, nous pouvons appliquer le lemme de nettoyage de Bravyi-Terhal.

Le lemme suivant montre qu'il existe un opérateur logique non trivial dont le support est entièrement contenu dans une tranche $\pi(T_j)$.

Lemme 5.6.4. *Il existe un opérateur logique non trivial dont le support est entièrement contenu dans une tranche $\pi(T_j)$. Par conséquent, la distance minimale du code $\mathcal{Q}$ est majorée par le nombre de qubits de cette tranche :*

$$d \leq m|\pi(\mathbb{Z}^D \cap T_j)|$$

où m est le nombre de qubits par point à coordonnées entières du domaine fondamental $\mathcal{P}(\mathbf{u}_1, \dots, \mathbf{u}_D)$ du réseau Λ .

Preuve : D'après le lemme 5.5.1 (lemme de nettoyage de Bravyi-Terhal), il n'existe que deux cas possibles.

Soit il existe un opérateur logique non trivial dont le support est entièrement inclus dans une tranche d'indice impair, $\pi(T_{2i+1})$. Dans ce cas, la distance minimale du code est immédiatement bornée par le nombre de qubits de cette tranche : $d \leq m|\pi(\mathbb{Z}^D \cap T_{2i+1})|$.

Soit tous les opérateurs logiques peuvent être "nettoyés" sur chaque tranche d'indice impair. Nous allons montrer que dans ce cas, il existe un opérateur logique non trivial dont le support est entièrement contenu dans une tranche d'indice pair.

Considérons un opérateur logique non trivial $\mathcal{O}$. Étant donné que le groupe stabilisateur est engendré par une famille de générateurs satisfaisant la condition de localité du théorème 5.4.1 et qu'aucun de ces générateurs n'agit non trivialement sur plusieurs tranches d'indice impair à la fois, nous pouvons appliquer le lemme de nettoyage de manière répétée. Cette application (qui consiste à multiplier $\mathcal{O}$ successivement par des générateurs du groupe stabilisateur de $\mathcal{Q}$ agissant non trivialement sur les tranches d'indice impair) permet d'obtenir un opérateur logique équivalent $\tilde{\mathcal{O}}$, qui agit trivialement sur toutes les tranches d'indice impair. Cet opérateur s'écrit

$$\tilde{\mathcal{O}} = \tilde{\mathcal{O}}_0 \times \tilde{\mathcal{O}}_2 \times \dots \times \tilde{\mathcal{O}}_{\mu-2}$$

où chaque opérateur $\tilde{\mathcal{O}}_{2i}$ a son support inclus dans la tranche d'indice pair $\pi(T_{2i})$.

Chacun des opérateurs $\tilde{\mathcal{O}}_{2i}$ commute avec tous stabilisateurs, et donc est opérateur logique.

Puisque le groupe stabilisateur $\mathbb{S}$ est engendré par une famille de générateurs satisfaisant la condition de localité du théorème 5.4.1, il suffit de vérifier que $\tilde{\mathcal{O}}_{2i}$ commute

avec ces générateurs pour confirmer qu'il s'agit bien d'un opérateur logique.

Considérons $\mathcal{S} \in \mathbb{S}$ un générateur de $\mathbb{S}$ satisfaisant la condition de localité.

- Si $\mathcal{S}$ n'agit pas sur $\pi(T_{2i})$ alors il commute avec $\tilde{\mathcal{O}}_{2i}$, puisque leurs supports sont disjoints.
- Si $\mathcal{S}$ agit sur $\pi(T_{2i})$, notre découpage garantit qu'il n'agit sur aucune autre tranche d'indice pair. Il commute donc avec tous les autres opérateurs $\tilde{\mathcal{O}}_{2k}$. Puisqu'il commute également avec $\mathcal{O}$ qui est un opérateur logique, il commute nécessairement avec $\tilde{\mathcal{O}}_{2i}$.

De plus, comme $\tilde{\mathcal{O}}$ n'appartient pas au groupe stabilisateur $\mathbb{S}$, au moins l'un des opérateurs $\tilde{\mathcal{O}}_{2i}$ n'appartient pas à ce groupe non plus. En conséquence, il existe un opérateur logique non trivial dont le support est entièrement inclus dans une tranche d'indice pair $\pi(T_{2i})$.

D'où $d \leq m|\pi(\mathbb{Z}^D \cap T_{2i})|$. □

5.6.4 Majoration du nombre de qubits par tranche

Pour achever la démonstration du Théorème 5.4.1, il ne nous reste qu'à majorer le nombre de qubits contenus dans une tranche $\pi(T_j)$. Puisque ce nombre est un multiple du nombre de points entiers qui sont contenus dans cette tranche, soit $|\pi(\mathbb{Z}^D \cap T_j)|$, il nous suffit de majorer ce dernier.

Le lemme suivant nous fournit cette estimation.

Lemme 5.6.5. *Pour $j \in \llbracket 0, \mu - 1 \rrbracket$, $|\pi(\mathbb{Z}^D \cap T_j)| \leq (\lambda + \sqrt{D}) \frac{\ell}{\|\mathbf{u}_D^*\|}$.*

Esquisse de la preuve :

La stratégie pour majorer le cardinal de $\pi(\mathbb{Z}^D \cap T_j)$ consiste à l'interpréter comme un volume, puis à utiliser les outils de la théorie de la mesure pour majorer ce volume. Voici les trois étapes principales de cette démonstration :

1. **Approximation par le volume** : Nous allons montrer que le cardinal de l'ensemble $|\pi(\mathbb{Z}^D \cap T_j)|$ est égal au volume de l'union des cubes unitaires centrés sur les points à coordonnées entières de tranche $\pi(T_j)$.
2. **Inclusion dans une région plus grande** : Nous allons trouver un ensemble, noté $\pi(\Gamma)$, qui est une version légèrement agrandie de la tranche $\pi(T_j)$ et qui contient cette union de cubes. Ainsi, le cardinal de $\pi(\mathbb{Z}^D \cap T_j)$ sera inférieur ou égal au volume de l'ensemble $\pi(\Gamma)$.
3. **Calcul du volume** : Nous allons calculer le volume de $\pi(\Gamma)$ et nous montrerons qu'il est égal à $\frac{\ell}{\|\mathbf{u}_D^*\|}(\lambda + \sqrt{D})$, ce qui nous donnera la majoration finale.

La figure 5.2 illustre comment la région de la tranche est étendue pour inclure tous les cubes. Sur celle-ci est représentée (à gauche) la réunion de tous les cubes unitaires (en gris) centrés en des points à coordonnées entières (en noirs) d'une tranche $\pi(T_j)$. Un zoom (à droite) met en évidence que cette union peut dépasser légèrement la tranche initiale. L'ensemble Γ correspondra à une extension de la tranche suffisamment large pour contenir tout ces cubes.

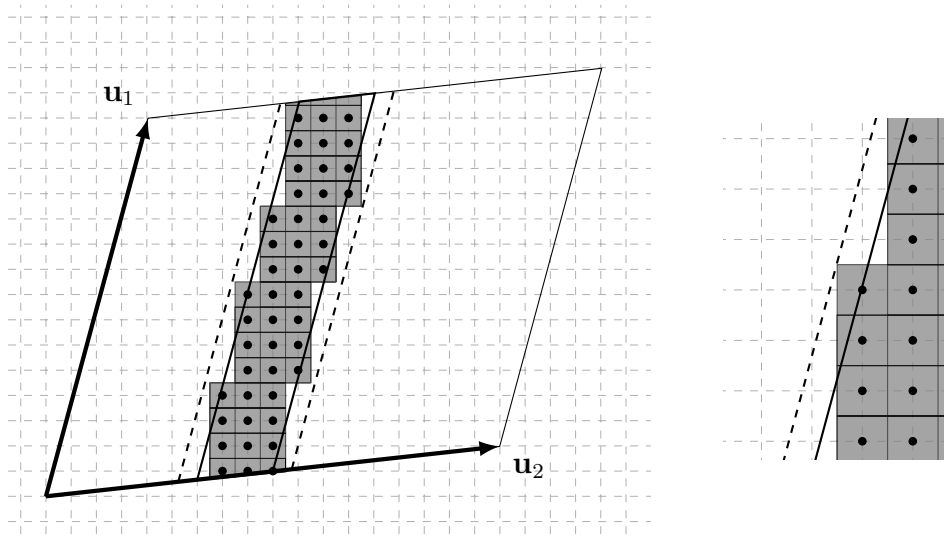


Figure 5.2: Représentation schématique des trois étapes de la preuve.

Démonstration :

Preuve : Afin d'alléger les notations, dans toute cette démonstration, nous dénoterons par $\mathcal{P}$ le domaine fondamental de Λ associé à la base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$.

Étape 1 : Réinterprétation de $|\pi(\mathbb{Z}^D \cap T_j)|$

La restriction à $\mathcal{P}$ de la projection canonique $\pi : \mathbb{R}^D \rightarrow \mathbb{R}^D/\Lambda$ est une application injective. En conséquence, pour $T_j \cap \mathbb{Z}^D$ qui est inclus dans $\mathcal{P}$ (par définition de T_j), nous avons que $|\pi(\mathbb{Z}^D \cap T_j)| = |\mathbb{Z}^D \cap T_j| < +\infty$.

Pour $\mathbf{x} \in \mathbb{Z}^D \cap T_j$, notons $C_{\mathbf{x}} = \{\mathbf{y} \in \mathbb{R}^D \mid \|\mathbf{y} - \mathbf{x}\|_{\infty} < \frac{1}{2}\}$ le cube unitaire centré en $\mathbf{x}$.

D'une part, les projections de ces cubes dans le quotient sont deux à deux disjointes. En effet, s'il existe deux éléments $\mathbf{x}, \mathbf{y} \in \mathbb{Z}^D \cap T_j$ tels que les ensembles $\pi(C_{\mathbf{x}})$ et $\pi(C_{\mathbf{y}})$ se rencontrent, alors il existe un couple $(\mathbf{t}_{\mathbf{x}}, \mathbf{t}_{\mathbf{y}}) \in C_{\mathbf{x}} \times C_{\mathbf{y}}$ tels que $\pi(\mathbf{t}_{\mathbf{x}}) = \pi(\mathbf{t}_{\mathbf{y}})$. Cela implique en particulier que :

$$\|\mathbf{x} - \mathbf{y} - (\mathbf{t}_{\mathbf{x}} - \mathbf{t}_{\mathbf{y}})\|_{\infty} \leq \|\mathbf{x} - \mathbf{t}_{\mathbf{x}}\|_{\infty} + \|\mathbf{t}_{\mathbf{y}} - \mathbf{y}\|_{\infty} < 1$$

Comme $\mathbf{x}, \mathbf{y}$ sont des entiers et que $\mathbf{t}_{\mathbf{x}} - \mathbf{t}_{\mathbf{y}} \in \Lambda \subset \mathbb{Z}^D$ alors $\mathbf{x} - \mathbf{y} - (\mathbf{t}_{\mathbf{x}} - \mathbf{t}_{\mathbf{y}}) = \mathbf{0}$. En conséquence, $\pi(\mathbf{x}) = \pi(\mathbf{y})$. Or, comme $\mathbf{x}$ et $\mathbf{y}$ appartiennent à $T_j \subset \mathcal{P}$, alors $\mathbf{x} = \mathbf{y}$ (par injectivité de $\pi|_{\mathcal{P}}$).

D'autre part, pour tout $\mathbf{x} \in \mathbb{Z}^D \cap T_j$, le volume de $\pi(C_{\mathbf{x}})$ vaut 1 (pour la mesure de Haar introduite dans la section 5.3.1). Cela repose sur le fait que le volume de $C_{\mathbf{x}}$ est 1 et que les ensembles $(C_{\mathbf{x}} + \mathbf{g})_{\mathbf{g} \in \Lambda}$ sont deux à deux disjointes. En effet, si $\exists \mathbf{t} \in (C_{\mathbf{x}} + \mathbf{g}) \cap (C_{\mathbf{x}} + \mathbf{h})$ où $\mathbf{g}$ et $\mathbf{h}$ sont des éléments de Λ , alors $\mathbf{t} - \mathbf{g} \in C_{\mathbf{x}}$ et $\mathbf{t} - \mathbf{h} \in C_{\mathbf{x}}$. Comme $\Lambda \subset \mathbb{Z}^D$, $\mathbf{g}$ et $\mathbf{h}$ sont des entiers. De plus, ils vérifient l'inégalité :

$$\|\mathbf{g} - \mathbf{h}\|_{\infty} = \|\mathbf{x} - (\mathbf{t} - \mathbf{g}) - (\mathbf{x} - (\mathbf{t} - \mathbf{h}))\|_{\infty} \leq \|\mathbf{x} - (\mathbf{t} - \mathbf{g})\|_{\infty} + \|\mathbf{x} - (\mathbf{t} - \mathbf{h})\|_{\infty} < 1$$

Par conséquent, $\mathbf{g}$ et $\mathbf{h}$ sont égaux.

En utilisant les propriétés de la mesure de Lebesgue (additivité et invariance par translation), celles du sous-groupe discret Λ (qui est dénombrable et symétrique) et

celles de son domaine fondamental ($\mathbb{R}^D = \coprod_{g \in \Lambda} \mathcal{P} + g$), nous obtenons alors :

$$\begin{aligned}
 \mathcal{M}(\pi(C_{\mathbf{x}})) &= \text{Leb}(\pi^{-1}(C_{\mathbf{x}}) \cap \mathcal{P}) = \text{Leb}\left(\bigcup_{g \in \Lambda} (C_{\mathbf{x}} + g) \cap \mathcal{P}\right) = \sum_{g \in \Lambda} \text{Leb}(C_{\mathbf{x}} + g \cap \mathcal{P}) \\
 &= \sum_{g \in \Lambda} \text{Leb}(C_{\mathbf{x}} \cap \mathcal{P} - g) \\
 &= \sum_{g \in \Lambda} \text{Leb}(C_{\mathbf{x}} \cap \mathcal{P} + g) \\
 &= \text{Leb}(C_{\mathbf{x}} \cap (\bigcup_{g \in \Lambda} \mathcal{P} + g)) \\
 &= \text{Leb}(C_{\mathbf{x}}) \\
 &= 1
 \end{aligned}$$

En conséquence, le cardinal de $\pi(\mathbb{Z}^D \cap T_j)$ se réinterprète comme :

$$|\pi(\mathbb{Z}^D \cap T_j)| = \sum_{\mathbf{x} \in \mathbb{Z}^D \cap T_j} 1 = \sum_{\mathbf{x} \in \mathbb{Z}^D \cap T_j} \mathcal{M}(\pi(C_{\mathbf{x}})) = \mathcal{M}(\pi(\bigcup_{\mathbf{x} \in \mathbb{Z}^D \cap T_j} C_{\mathbf{x}}))$$

Étape 2 : Construction de $\pi(\Gamma)$ contenant $\pi(\bigcup_{\mathbf{x} \in \mathbb{Z}^D \cap T_j} C_{\mathbf{x}})$

Notons $C = \bigcup_{\mathbf{x} \in \mathbb{Z}^D \cap T_j} C_{\mathbf{x}}$. D'après la Figure 5.2, $\pi(C)$, la projection de l'union de tous les cubes unitaires centrés en des points entiers de $\mathbb{Z}^D \cap T_j$, pourrait ne pas être entièrement contenue dans la tranche $\pi(T_j)$. Toutefois, les parties des cubes qui dépassent se trouvent à une distance d'au plus $\frac{\sqrt{D}}{2}$ d'un point entier de la tranche.

Afin de majorer le volume de $\pi(C)$, nous l'incluons dans une version légèrement agrandie de la tranche $\pi(T_j)$, contenant toutes ces parties débordantes.

Notons $\Gamma \subset \mathcal{P}$, l'ensemble défini par :

$$\Gamma = \left\{ \sum_{i=1}^D x_i \mathbf{u}_i \mid \forall i \in \llbracket 1, D-1 \rrbracket, 0 \leq x_i < 1 \text{ et } -\frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j}{\mu} \leq x_D < \frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j+1}{\mu} \right\}$$

Son image par π contient $\pi(C)$, c'est-à-dire $\pi(C) \subset \pi(\Gamma)$. En conséquence, son volume satisfait :

$$|\pi(\mathbb{Z}^D \cap T_j)| \leq \mathcal{M}(\pi(C)) \leq \mathcal{M}(\pi(\Gamma))$$

Étape 3 : Calcul de volume de $\pi(\Gamma)$

Comme Γ est une région du domaine fondamental $\mathcal{P}$, son volume est égal à celui de $\pi(\Gamma)$, sa projection dans le quotient. Autrement dit, $\mathcal{M}(\pi(\Gamma)) = \text{Leb}(\Gamma)$. Cette égalité se démontre en utilisant les mêmes arguments que ceux utilisés pour le calcul du volume de $\pi(C_x)$. En conséquence, le cardinal de $\pi(\mathbb{Z}^D \cap T_j)$ est majoré par le volume de Γ :

$$|\pi(\mathbb{Z}^D \cap T_j)| \leq \mathcal{M}(\pi(C)) \leq \mathcal{M}(\pi(\Gamma)) = \text{Leb}(\Gamma)$$

Pour calculer le volume de Γ , nous utilisons un changement de variables. Notons ϕ l'application linéaire définie par :

$$\begin{aligned} \phi : \mathbb{R}^D &\rightarrow \mathbb{R}^D \\ \begin{pmatrix} a_1 \\ \vdots \\ a_D \end{pmatrix} &\mapsto \sum_{i=1}^D a_i \mathbf{u}_i \end{aligned}$$

Puisque $\mathbf{u}_1, \dots, \mathbf{u}_D$ est une base du réseau Λ (et de $\mathbb{R}^D$ en tant qu'espace vectoriel), ϕ est un automorphisme. Son Jacobien est constant et est égal au covolume du réseau :

$$\forall \mathbf{y} \in \mathbb{R}^D, \quad |\det Jac_{\mathbf{y}} \phi| = |\det(\mathbf{u}_1, \dots, \mathbf{u}_D)| = \det(\Lambda) = \text{vol}(\Lambda) = \ell$$

De plus, l'ensemble Γ est donné par $\phi(\mathcal{U})$, l'image par ϕ de l'ensemble $\mathcal{U}$ défini par :

$$\mathcal{U} = [0, 1[^{D-1} \times \left[-\frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j}{\mu}, \frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j+1}{\mu} \right[$$

Par conséquent, en utilisant le théorème du changement de variables, nous obtenons donc :

$$\begin{aligned} \text{Leb}(\Gamma) &= \int_{\phi(\mathcal{U})} dy = \int_{\mathcal{U}} |\det Jac_{\mathbf{x}} \phi| d\mathbf{x} = \text{Leb}(\mathcal{U}) \ell \\ &= \left(\frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j+1}{\mu} - \left(-\frac{\sqrt{D}}{2\|\mathbf{u}_D^*\|} + \frac{j}{\mu} \right) \right) \ell \\ &= \left(\frac{\sqrt{D}}{\|\mathbf{u}_D^*\|} + \frac{1}{\mu} \right) \ell \end{aligned}$$

Puisque $\|\mathbf{u}_D^*\| = \lambda\mu$, alors $|\pi(\mathbb{Z}^D \cap T_j)| \leq \left(\frac{\sqrt{D}}{\|\mathbf{u}_D^*\|} + \frac{1}{\mu} \right) \ell = \left(\sqrt{D} + \lambda \right) \frac{\ell}{\|\mathbf{u}_D^*\|}$. □

5.6.5 Preuve complète du théorème 5.4.1 :

En combinant tous les éléments introduits dans les précédentes sections, nous pouvons fournir une preuve du théorème 5.4.1.

Preuve : Soient $m \geq 1$ un entier et $\mathcal{Q}$ un code stabilisateur de longueur $m\ell$ tel qu'il existe Λ un sous-réseau de $\mathbb{Z}^D$ de rang plein dont le covolume est égal à $\text{vol}(\Lambda) = |\mathbb{Z}^D/\Lambda| = \ell$.

Supposons que chaque sommet de $\mathbb{Z}^D/\Lambda$ indexe m qubits et qu'il existe $\rho > 0$ tel que le groupe stabilisateur soit engendré par une famille de générateurs dont le support est contenu dans une boule de $\mathbb{R}^D/\Lambda$ de rayon ρ pour la distance dist . Supposons de plus que $\ell^{\frac{1}{D}} \geq 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$ où $\gamma_{D,D-1}(\Lambda)$ est la constante auxiliaire introduite dans la définition 5.3.2 de la $D - 1$ -ième constante de Rankin.

Alors, avec l'approche de la section 5.6.1, on choisit une base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ du réseau Λ de telle sorte que le dernier vecteur de la base de Gram-Schmidt associée, $\mathbf{u}_D^*$, ait une norme suffisamment grande pour satisfaire l'inégalité : $\|\mathbf{u}_D^*\| = \frac{\ell^{1/D}}{\sqrt{\gamma_{D,D-1}(\Lambda)}} \geq 8\rho$.

Ensuite, en suivant la méthode de la section 5.6.2, on découpe $\mathbb{R}^D/\Lambda$ en μ tranches $(\pi(T_i))_{i \in \llbracket 0, \mu-1 \rrbracket}$, où μ est un entier pair supérieur ou égal à 4, construit avec la méthode de la section 5.6.2. Pour ce découpage, l'épaisseur de chaque tranche, $\lambda = \frac{\|\mathbf{u}_D^*\|}{\mu}$, satisfait la condition : $2\rho \leq \lambda < 4\rho$.

D'après le lemme 5.6.3, ce découpage garantit qu'aucun générateur du groupe stabilisateur, respectant la condition de localité, n'agit simultanément de manière non triviale sur des qubits situés dans des tranches non adjacentes. Le lemme 5.6.4 assure alors qu'il existe un opérateur logique non trivial dont le support est entièrement contenu dans une tranche $\pi(T_j)$, où $j \in \llbracket 0, \mu - 1 \rrbracket$. La distance minimale du code est donc majorée par le nombre de qubits dans cette tranche : $d \leq m|\pi(\mathbb{Z}^D \cap T_j)|$.

La dernière étape consiste à majorer le nombre de qubits par tranche en utilisant une approximation par le volume (lemme 5.6.5). Cela nous permet d'obtenir la borne finale pour la distance minimale d :

$$\begin{aligned}
 d &= m(\lambda + \sqrt{D}) \frac{\ell}{\|\mathbf{u}_D^*\|} \\
 &= m\sqrt{\gamma_{D,D-1}(\Lambda)}(\lambda + \sqrt{D})\ell^{\frac{D-1}{D}} \quad (\text{car } \|\mathbf{u}_D^*\| = \frac{\ell^{1/D}}{\sqrt{\gamma_{D,D-1}(\Lambda)}}) \\
 &< m\sqrt{\gamma_{D,D-1}(\Lambda)}(4\rho + \sqrt{D})\ell^{\frac{D-1}{D}} \quad (\text{Car } \lambda < 4\rho)
 \end{aligned}$$

□

Le théorème 5.4.1 s'applique à tout réseau Λ . Cependant, une connaissance plus fine de la structure de Λ permet d'obtenir une borne supérieure plus précise. En particulier, lorsque $\Lambda = \bigoplus_{i=1}^D \mathbb{Z}N\epsilon_i$ est le réseau cubique engendré par des multiples des vecteurs de la base canonique $(\epsilon_i)_{1 \leq i \leq D}$ de $\mathbb{Z}^D$, on retrouve — à un facteur constant près (en l'occurrence, 6) — la borne originale de Bravyi et Terhal.

Théorème 5.6.6. *Sous les hypothèses du théorème 5.4.1, si le réseau sur lequel est défini le code $\mathcal{Q}$ est $\Lambda = \bigoplus_{i=1}^D \mathbb{Z}N\epsilon_i$, alors la distance minimale du code satisfait l'inégalité :*

$$d_{\min} < m(4\rho + 2) \cdot N^{D-1} \leq 6m \cdot \max(1, \rho) \cdot N^{D-1}$$

Preuve : Commençons par remarquer que $\gamma_{D,D-1}(\Lambda) = 1$.

En effet, puisque $\Lambda = \bigoplus_{i=1}^D \mathbb{Z}N\epsilon_i$, toute base $\{\mathbf{u}_1, \dots, \mathbf{u}_D\}$ de Λ est constituée de vecteurs de la forme $\mathbf{u}_i = N\mathbf{v}_i$ où $\mathbf{v}_i$ est un vecteur à coordonnées entières. Ainsi, pour toute telle base on a :

$$\text{vol}(\{\mathbf{u}_1, \dots, \mathbf{u}_{D-1}\})^2 = N^{2(D-1)} \text{vol}(\{\mathbf{v}_1, \dots, \mathbf{v}_{D-1}\})^2$$

Comme les quantités $\text{vol}(\mathbf{u}_1, \dots, \mathbf{u}_{D-1})^2$ et $\text{vol}(\{\mathbf{v}_1, \dots, \mathbf{v}_{D-1}\})^2$ sont deux entiers non nuls, il en résulte que :

$$\frac{\text{vol}(\mathbf{u}_1, \dots, \mathbf{u}_{D-1})^2}{N^{2(D-1)}} \geq 1$$

Par conséquent, nous avons donc :

$$\gamma_{D,D-1}(\Lambda) = \min\left\{\frac{\text{vol}(\mathbb{L})^2}{N^{2(D-1)}} \mid \mathbb{L} \subset \Lambda \text{ réseau de rang } D-1\right\} \geq 1$$

Réciproquement, on peut exhiber un sous-réseau de Λ de covolume N^{D-1} . Pour $\mathbb{L} = \bigoplus_{i=1}^{D-1} \mathbb{Z}N\epsilon_i$, on obtient :

$$\sqrt{\gamma_{D,D-1}(\Lambda)} \leq \frac{\text{vol}(\mathbb{L})}{N^{D-1}} = \frac{N^{D-1}}{N^{D-1}} = 1$$

Ce qui permet donc de conclure de $\gamma_{D,D-1}(\Lambda) = 1$.

Prouvons à présent l'inégalité sur la distance minimale.

Comme dans la preuve du théorème, nous distinguons deux cas : selon que $\text{vol}(\Lambda)^{\frac{1}{D}} = N < 8\rho\sqrt{\gamma_{D,D-1}(\Lambda)}$ ou non.

si $N = \text{vol}(\Lambda)^{\frac{1}{D}} < 8\rho$ alors : en appliquant les résultats de la section 5.5.1, on obtient :

$$d_{\min} < m(4\rho\sqrt{\gamma_{D,D-1}(\Lambda)} + 1) \cdot N^{D-1} = m(4\rho + 1) \cdot N^{D-1}$$

Si $N = \text{vol}(\Lambda)^{\frac{1}{D}} \geq 8\rho$ alors : on choisit pour base de Λ les vecteurs $\mathbf{u}_i := N\epsilon_i$ pour tout $i \in \llbracket 1, D \rrbracket$. Le D -ième vecteur de la base de Gram-Schmidt associé, $\mathbf{u}_D^*$, est alors $\mathbf{u}_D^* = N\epsilon_D$ et sa norme vérifie $\|\mathbf{u}_D^*\| = N \geq 8\rho$.

En appliquant successivement procédures de découpage et de nettoyage des sections 5.6.2 et 5.6.3, nous prouvons l'existence d'un entier $q \in \mathbb{N}$ tel que $d_{\min} \leq m|\mathbb{Z}^D \cap T_q|$ où T_q est une tranche du domaine fondamental définie par :

$$T_q = \left\{ \sum_{i=1}^{D-1} x_i \mathbf{u}_i + \frac{x_D}{\mu} \mathbf{u}_D \mid \forall i \in \llbracket 1, D-1 \rrbracket, 0 \leq x_i < 1 \text{ et } q \leq x_D < q+1 \right\}$$

Dans notre cas, $T_q = [0, N[^{D-1} \times [\frac{Nq}{\mu}, \frac{N(q+1)}{\mu}[$ et μ est un entier choisi comme dans la section 5.6.2, qui vérifie $\mu \geq \lfloor \frac{N}{2\rho} \rfloor - 1$.

Ainsi la distance minimale satisfait :

$$\begin{aligned} d_{\min} &\leq m|\mathbb{Z}^D \cap T_q| \leq mN^{D-1} \left(\left\lfloor \frac{N(q+1)}{\mu} \right\rfloor - \left\lfloor \frac{Nq}{\mu} \right\rfloor + 1 \right) \\ &< mN^{D-1} \left(\frac{N(q+1)}{\mu} - \left(\frac{Nq}{\mu} - 1 \right) + 1 \right) \quad (\text{car } x - 1 < \lfloor x \rfloor \leq x) \\ &= mN^{D-1} \left(\frac{N}{\mu} + 2 \right) \end{aligned}$$

Or puisque $\mu \geq \lfloor \frac{N}{2\rho} \rfloor - 1 > \frac{N}{2\rho} - 2 = \frac{N-4\rho}{2\rho}$, nous obtenons que $\frac{N}{\mu} < 2\rho \frac{N}{N-4\rho} \leq 4\rho$. La dernière inégalité découlant de la suite d'équivalence :

$$\begin{aligned} 2\rho \cdot \frac{N}{N-4\rho} \leq 4\rho &\iff 2\rho \cdot N \leq 4\rho(N-4\rho) && (\text{car } N \geq 8\rho > 0) \\ &\iff 0 \leq 2\rho \cdot (N-8\rho) \end{aligned}$$

En conclusion, en combinant tous ces résultats nous obtenons alors :

$$d_{\min} < mN^{D-1} \cdot \left(\frac{N}{\mu} + 2\right) < mN^{D-1} \cdot (4\rho + 2)$$

□

Dans la section suivante, nous allons appliquer la borne trouvée dans le théorème 5.4.1 aux codes 2BGA abéliens. Pour ce faire, nous procéderons en trois étapes :

1. Nous décrivons la construction de ces codes.
2. Nous montrerons comment leurs qubits peuvent être représentés par paires sur les sommets d'un quotient $\mathbb{Z}^D/\Lambda$, de sorte que leurs groupes stabilisateurs soient engendrés par des générateurs dont le support est inclus dans une boule de rayon 1 de $\mathbb{R}^D/\Lambda$.
3. Enfin, nous appliquerons le Théorème 5.4.1 pour établir une borne sur leur distance minimale.

5.7 Étude de la distance minimale des codes 2BGA abéliens

Les *codes 2BGA abéliens* (Abelian Two-Block Group Algebra) sont des codes CSS construits à partir d'algèbres de groupe de groupes abéliens.

Inspirés des codes Bicycle déjà connus [11], les *codes 2BGA abéliens* ont été proposés par Pryadko et Lin [22], comme une généralisation des codes GB. Alors qu'un code GB se construit à partir d'une paire de matrices circulantes, un code 2BGA utilise une paire de matrices commutantes issues de l'algèbre d'un groupe abélien fini, lequel n'est pas nécessairement cyclique.

5.7.1 Présentation des codes 2BGA abéliens

Les codes 2BGA abéliens sont des codes CSS dont la construction repose sur l'algèbre $\mathbb{F}_2[\mathbb{G}]$ d'un groupe abélien. Dans tout le reste de ce chapitre, $\mathbb{G} = \{g_1, \dots, g_\ell\}$ sera un groupe abélien fini d'ordre ℓ , dont la loi sera notée multiplicativement.

Algèbre de groupe : définition et propriétés :

L'ensemble $\mathbb{F}_2[\mathbb{G}]$ est constitué des combinaisons linéaires formelles d'éléments de $\mathbb{G}$ à coefficients dans $\mathbb{F}_2$:

$$\mathbb{F}_2[\mathbb{G}] = \left\{ \sum_{g \in \mathbb{G}} a_g g \mid \forall g \in \mathbb{G}, a_g \in \mathbb{F}_2 \right\}$$

Puisque le groupe $\mathbb{G}$ est commutatif, l'ensemble $\mathbb{F}_2[\mathbb{G}]$ muni de l'addition terme à terme et du produit de convolution, hérite d'une structure d'algèbre commutative. Le produit de deux éléments $\mathbf{a} = (\sum_{g \in \mathbb{G}} a_g g)$ et $\mathbf{b} = (\sum_{g \in \mathbb{G}} b_g g)$ est donné par :

$$\mathbf{a} \cdot \mathbf{b} = \mathbf{b} \cdot \mathbf{a} = \sum_{g \in \mathbb{G}} \left(\sum_{\substack{u, v \in \mathbb{G} \\ uv = g}} a_u b_v \right) g$$

Remarque 5.7.1. Le produit d'un élément $\mathbf{a}$ de l'algèbre par une somme formelle contenant un unique élément $h \in \mathbb{G}$ (notée $1h$), s'obtient directement en permutant les coefficients de $\mathbf{a}$:

$$(1h) \cdot \mathbf{a} = \sum_{g \in \mathbb{G}} a_{h^{-1}g} g = \sum_{p \in \mathbb{G}} a_p hp$$

Matrices de permutation associées

A chaque élément, g , du groupe abélien $\mathbb{G} = \{g_1, \dots, g_\ell\}$, on associe une matrice de permutation $\mathbb{B}(g)$ de taille $\ell \times \ell$. Le coefficient en position (i, j) de cette matrice, $\mathbb{B}(g)_{i,j}$, est égal à 1 si et seulement si l'élément $g_i = gg_j$.

Exemple 5.7.1. La matrice associée à l'élément neutre e est l'identité, soit $\mathbb{B}(e) = \mathbf{I}_\ell$.

Comme $\mathbb{G}$ est commutatif, le produit de deux telles matrices vérifie :

$$\forall g, h \in \mathbb{G}, \mathbb{B}(gh) = \mathbb{B}(g)\mathbb{B}(h) = \mathbb{B}(h)\mathbb{B}(g) = \mathbb{B}(hg)$$

Par conséquent, pour tout $g \in \mathbb{G}$, la matrice $\mathbb{B}(g)$ est inversible et son inverse est donné par $\mathbb{B}(g)^{-1} = \mathbb{B}(g^{-1})$.

Par extension, tout élément de l'algèbre $\mathbb{F}_2[\mathbb{G}]$ peut être représenté par une matrice. Pour $\mathbf{a} = \sum_{g \in \mathbb{G}} a_g g$, la matrice correspondante est $\mathbf{A} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g)$. Ainsi, deux matrices $\mathbf{A}$ et $\mathbf{B}$, associées aux éléments $\mathbf{a}$ et $\mathbf{b}$ de l'algèbre $\mathbb{F}_2[\mathbb{G}]$, commutent entre elles : $\mathbf{AB} = \mathbf{BA}$.

Construction des codes 2BGA

Soient $\mathbb{G}$ un groupe abélien d'ordre ℓ et soient $\mathbf{a} = \sum_{g \in \mathbb{G}} a_g g$ et $\mathbf{b} = \sum_{g \in \mathbb{G}} b_g g$ deux éléments de l'algèbre $\mathbb{F}_2[\mathbb{G}]$. Un code 2BGA abélien est un code CSS dont les sous-matrices génératrices $\mathbf{H}_X$ et $\mathbf{H}_Z$ sont données par :

$$\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}] \quad \text{et} \quad \mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$$

où $\mathbf{A} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g)$ et $\mathbf{B} = \sum_{g \in \mathbb{G}} b_g \mathbb{B}(g)$ sont les matrices associées à $\mathbf{a}$ et $\mathbf{b}$.

Les paramètres de ces codes sont :

- **Longueur** : $n = 2\ell$
- **Dimension** : $k = 2\ell - 2\text{rang}\mathbf{H}_X$ car $\text{rang}\mathbf{H}_X = \text{rang}\mathbf{H}_Z$ [34].
- **Distance minimale** : $d = d_X = d_Z$ d'après [22].

Généralisation d'autres constructions

Les codes 2BGA généralisent plusieurs classes de codes bien connues :

- Lorsque $\mathbb{G}$ est cyclique, on retrouve les codes Bicycle généralisés [11].
- Si de plus $\mathbf{A} = \mathbf{B}$, on retrouve les codes Bicycle [15].
- Si le groupe $\mathbb{G}$ est un produit de deux groupes cycliques, $\mathbb{G} = \mathbb{Z}/r\mathbb{Z} \times \mathbb{Z}/s\mathbb{Z}$, la construction donne les codes Bicycle bivariés [10].

5.7.2 Étude des codes 2BGA abéliens non triviaux

L'objectif de cette section est de majorer la distance minimale des codes 2BGA abéliens non triviaux, qui sont ceux pour lesquels chacune des matrices $\mathbf{A}$ et $\mathbf{B}$ est non nulle. Pour y parvenir, nous allons appliquer le Théorème 5.4.1 à ces codes.

Majoration de la distance minimale des codes 2BGA non triviaux

Théorème 5.7.1. *Soit $\mathbb{G} = \{g_1, \dots, g_\ell\}$ un groupe abélien fini d'ordre ℓ . On considère deux éléments $\mathbf{a} = \sum_{g \in \mathbb{G}} a_g g$ et $\mathbf{b} = \sum_{g \in \mathbb{G}} b_g g$ de $\mathbb{F}_2[\mathbb{G}]$, de poids respectifs $r + 1$ et $s + 1$, avec $r, s \in \mathbb{N}$.*

Considérons un code 2BGA abélien non trivial de longueur $n = 2\ell$, dont les sous-matrices génératrices binaires sont données par $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$, où $\mathbf{A} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g)$ et $\mathbf{B} = \sum_{g \in \mathbb{G}} b_g \mathbb{B}(g)$ possèdent respectivement $r + 1$ et $s + 1$ éléments non nuls par ligne.

Alors, la distance minimale du code est bornée par :

$$d < 2\sqrt{\gamma_D}(\sqrt{D} + 4) \ell^{\frac{D-1}{D}}$$

où $D = r + s = w - 2$ et γ_D est la constante de Hermite D -dimensionnelle.

Esquisse de la preuve :

La preuve du théorème se déroule en plusieurs étapes :

- **Simplification de la preuve : Simplification de la preuve :** Nous commençons par montrer que tout code 2BGA abélien non trivial a la même distance minimale qu'un code 2BGA non trivial construit à partir des matrices $\tilde{\mathbf{A}} = \sum_{g \in \mathbb{G}} \tilde{a}_g \mathbb{B}(g)$ et $\tilde{\mathbf{B}} = \sum_{g \in \mathbb{G}} \tilde{b}_g \mathbb{B}(g)$ avec $\tilde{a}_e = \tilde{b}_e = 1$, où ici e désigne l'élément neutre du groupe. Ainsi, il suffit de prouver la borne uniquement pour ce type de code.
- **Représentation des qubits sur un quotient :** Ensuite, nous représentons les qubits de ce nouveau code sur les points d'un quotient de $\mathbb{Z}^D$ par un réseau Λ . Cette représentation est conçue de manière à placer deux qubits sur chaque sommet, de sorte que le groupe stabilisateur de ce code soit engendré par des générateurs dont le support est inclus dans une boule de rayon 1 de $\mathbb{R}^D / \Lambda$.
- **Application du théorème :** Après avoir vérifié ces conditions, nous appliquons directement le théorème 5.4.1. En fixant les paramètres à $m = 2$ (pour les paires de qubits par sommet) et $\rho = 1$ (pour le rayon des boules), nous obtenons la borne annoncée pour la distance minimale du code.

5.7.3 Démonstration du théorème 5.7.1

Soit $\mathbb{G} = \{g_1, \dots, g_\ell\}$ un groupe abélien fini d'ordre ℓ , d'élément neutre e .

Considérons $\mathcal{Q}$ un code 2BGA abélien non trivial de longueur $n = 2\ell$, dont les sous-matrices génératrices $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$ sont définies à l'aide des matrices $\mathbf{A} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g)$ et $\mathbf{B} = \sum_{g \in \mathbb{G}} b_g \mathbb{B}(g)$ ayant respectivement $r + 1$ et $s + 1$ éléments non nuls par ligne.

Commençons par montrer que ce code 2BGA a la même distance minimale qu'un code 2BGA non trivial engendré par des matrices $\tilde{\mathbf{A}} = \sum_{g \in \mathbb{G}} \tilde{a}_g \mathbb{B}(g)$ et $\tilde{\mathbf{B}} = \sum_{g \in \mathbb{G}} \tilde{b}_g \mathbb{B}(g)$ ayant chacune un coefficient 1 en position (e, e) .

Première simplification : construction du code simplifié

Puisque $\mathbb{B}(e)$ est la matrice identité, la présence d'un '1' en position (e, e) dans $\tilde{\mathbf{A}}$ et $\tilde{\mathbf{B}}$ signifie que leurs éléments dans l'algèbre de groupe correspondants, $\tilde{\mathbf{a}}$ et $\tilde{\mathbf{b}}$, ont un coefficient de '1' pour l'élément neutre, c'est-à-dire :

$$\tilde{\mathbf{a}} = \sum_{g \in \mathbb{G}} \tilde{a}_g g \quad \text{et} \quad \tilde{\mathbf{b}} = \sum_{g \in \mathbb{G}} \tilde{b}_g g \quad \text{avec} \quad \tilde{a}_e = \tilde{b}_e = 1$$

Le lemme suivant garantit l'existence d'un tel code.

Lemme 5.7.2. *Il existe un code 2BGA abélien non trivial sur $\mathbb{G}$ qui a la même distance minimale que $\mathcal{Q}$ et dont les sous-matrices génératrices $\tilde{\mathbf{H}}_X = [\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$ et $\tilde{\mathbf{H}}_Z = [\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$ sont définies à l'aide de $\tilde{\mathbf{A}}$ et $\tilde{\mathbf{B}}$ qui ont un '1' en position (e, e) .*

Preuve : La preuve se déroule en deux étapes : la construction des matrices $\tilde{\mathbf{A}}$ et $\tilde{\mathbf{B}}$ et la démonstration que les deux codes ont la même distance minimale.

Étape 1 : Construction de $\tilde{\mathbf{A}}$ et $\tilde{\mathbf{B}}$

Partons des éléments $\mathbf{a} = \sum_{g \in \mathbb{G}} a_g g$ et $\mathbf{b} = \sum_{g \in \mathbb{G}} b_g g$ de $\mathbb{F}_2[\mathbb{G}]$ qui définissent les matrices $\mathbf{A}$ et $\mathbf{B}$ du code $\mathcal{Q}$. Par hypothèse, $\mathbf{a}$ et $\mathbf{b}$ ont respectivement $r + 1$ et $s + 1$ coefficients non nuls. Nous choisissons donc deux indices g_a et g_b pour lesquels $a_{g_a} = b_{g_b} = 1$.

En multipliant $\mathbf{a}$ et $\mathbf{b}$ par des sommes formelles composées d'un seul terme, $1g_a^{-1}$ et $1g_b^{-1}$, nous définissons les nouveaux éléments $\tilde{\mathbf{a}} = (1g_a^{-1}) \cdot \mathbf{a}$ et $\tilde{\mathbf{b}} = (1g_b^{-1}) \cdot \mathbf{b}$. D'après

la remarque 5.7.1, leurs expressions sont données par :

$$\tilde{\mathbf{a}} = \sum_{g \in \mathbb{G}} a_g (g_a^{-1} g) \quad \text{et} \quad \tilde{\mathbf{b}} = \sum_{g \in \mathbb{G}} b_g (g_b^{-1} g)$$

Nous pouvons exprimer les matrices associées, $\tilde{\mathbf{A}}$ et $\tilde{\mathbf{B}}$, en utilisant les propriétés des matrices $\mathbb{B}(\cdot)$, rappelées en section 5.7.1 :

$$\tilde{\mathbf{A}} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g_a^{-1} g) = \mathbb{B}(g_a^{-1}) \left(\sum_{g \in \mathbb{G}} a_g \mathbb{B}(g) \right) = \mathbb{B}(g_a^{-1}) \mathbf{A}$$

De même, nous obtenons : $\tilde{\mathbf{B}} = \mathbb{B}(g_b^{-1}) \mathbf{B}$.

Étape 2 : Calcul de la distance minimale

D'après les propriétés des codes abéliens 2BGA (voir section 5.7.1), les distances minimales d et $\tilde{d}$ des codes 2BGA abéliens $\mathcal{Q}$ et $\tilde{\mathcal{Q}}$, associés respectivement aux couples $(\mathbf{A}, \mathbf{B})$ et $(\tilde{\mathbf{A}}, \tilde{\mathbf{B}})$, sont données par :

$$\begin{cases} d = \min\{\text{wt}(\mathbf{c}) \mid \mathbf{c} \in \ker[\mathbf{A} \mid \mathbf{B}] \setminus \text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]\} \\ \tilde{d} = \min\{\text{wt}(\mathbf{c}) \mid \mathbf{c} \in \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}] \setminus \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]\} \end{cases}$$

où $\text{Vect}_{\mathbf{L}} \mathbf{M}$ désigne l'espace vectoriel engendré par les lignes de la matrice $\mathbf{M}$.

Pour montrer que d et $\tilde{d}$ sont égaux, il nous suffit de montrer qu'il existe une isométrie sur $\mathbb{F}_2^{2\ell}$ qui induit des bijections entre $\ker[\mathbf{A} \mid \mathbf{B}]$ et $\ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$, et entre $\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]$ et $\text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$.

L'isométrie que nous construisons est l'application linéaire η définie par :

$$\eta : \mathbb{F}_2^{2\ell} \rightarrow \mathbb{F}_2^{2\ell} \\ \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} \mapsto \begin{pmatrix} \mathbb{B}(g_a) \mathbf{u} \\ \mathbb{B}(g_b) \mathbf{v} \end{pmatrix}$$

Puisque les matrices $\mathbb{B}(g_a)$ et $\mathbb{B}(g_b)$, sont des matrices de permutation alors η est un automorphisme qui préserve le poids de Hamming. Son application inverse η^{-1}

est donnée par :

$$\eta^{-1} : \mathbb{F}_2^{2\ell} \rightarrow \mathbb{F}_2^{2\ell}$$

$$\begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} \mapsto \begin{pmatrix} \mathbb{B}(g_a)^{-1} \mathbf{u} \\ \mathbb{B}(g_b)^{-1} \mathbf{v} \end{pmatrix}$$

Maintenant, nous allons démontrer que η induit les bijections souhaitées.

Montrons que $\eta(\ker[\mathbf{A} \mid \mathbf{B}]) = \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$:

On procède par double inclusion.

Commençons par montrer que $\eta(\ker[\mathbf{A} \mid \mathbf{B}]) \subseteq \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$.

Soit $\mathbf{c} = \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} \in \ker[\mathbf{A} \mid \mathbf{B}]$. Les matrices $\mathbb{B}(g_a)^{-1}$ et $\mathbb{B}(g_b)^{-1}$ commutent avec $\mathbf{A}$ et $\mathbf{B}$ (voir section 5.7.1), donc nous avons l'égalité :

$$\tilde{\mathbf{A}}\mathbb{B}(g_a)\mathbf{u} + \tilde{\mathbf{B}}\mathbb{B}(g_b)\mathbf{v} = \mathbb{B}(g_a)^{-1}\mathbf{A}\mathbb{B}(g_a)\mathbf{u} + \mathbb{B}(g_b)^{-1}\mathbf{B}\mathbb{B}(g_b)\mathbf{v} = \mathbf{A}\mathbf{u} + \mathbf{B}\mathbf{v} = \mathbf{0}$$

Donc $\eta(\mathbf{c}) = \begin{pmatrix} \mathbb{B}(g_a)\mathbf{u} \\ \mathbb{B}(g_b)\mathbf{v} \end{pmatrix} \in \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$, ce qui implique $\eta(\ker[\mathbf{A} \mid \mathbf{B}]) \subseteq \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$.

Réciproquement, si $\tilde{\mathbf{c}} = \begin{pmatrix} \tilde{\mathbf{u}} \\ \tilde{\mathbf{v}} \end{pmatrix} \in \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$ alors :

$$\mathbf{0} = \tilde{\mathbf{A}}\tilde{\mathbf{u}} + \tilde{\mathbf{B}}\tilde{\mathbf{v}} = \mathbb{B}(g_a)^{-1}\mathbf{A}\tilde{\mathbf{u}} + \mathbb{B}(g_b)^{-1}\mathbf{B}\tilde{\mathbf{v}} = \mathbf{A}\mathbb{B}(g_a)^{-1}\tilde{\mathbf{u}} + \mathbf{B}\mathbb{B}(g_b)^{-1}\tilde{\mathbf{v}}$$

Donc $\eta^{-1}(\tilde{\mathbf{c}}) = \begin{pmatrix} \mathbb{B}(g_a)^{-1}\tilde{\mathbf{u}} \\ \mathbb{B}(g_b)^{-1}\tilde{\mathbf{v}} \end{pmatrix} \in \ker[\mathbf{A} \mid \mathbf{B}]$.

D'où $\eta^{-1}(\ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]) \subseteq \ker[\mathbf{A} \mid \mathbf{B}]$, ce qui équivaut à $\ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}] \subseteq \eta(\ker[\mathbf{A} \mid \mathbf{B}])$.

En conclusion : $\eta(\ker[\mathbf{A} \mid \mathbf{B}]) = \ker[\tilde{\mathbf{A}} \mid \tilde{\mathbf{B}}]$.

Montrons que $\eta(\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]) = \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$

On va également procéder par double inclusion.

Commençons par montrer que $\eta(\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]) \subseteq \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$.

Soit $\mathbf{c} \in \text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]$. Nous pouvons écrire $\mathbf{c} = \sum_{i \in \mathcal{J}} \begin{pmatrix} \mathbf{B}\epsilon_i \\ \mathbf{A}\epsilon_i \end{pmatrix}$ où $\mathcal{J} \subseteq \llbracket 1, \ell \rrbracket$ et $(\epsilon_i)_{i \in \llbracket 1, \ell \rrbracket}$ est la base canonique de $\mathbb{F}_2^\ell$. En vertu des propriétés des matrices construites à partir d'éléments de $\mathbb{F}_2[\mathbb{G}]$ (voir section 5.7.1), nous avons la suite d'égalités suivante :

$$\begin{aligned} \eta(\mathbf{c}) &= \sum_{i \in \mathcal{J}} \begin{pmatrix} \mathbb{B}(g_a)\mathbf{B}\epsilon_i \\ \mathbb{B}(g_b)\mathbf{A}\epsilon_i \end{pmatrix} = \sum_{i \in \mathcal{J}} \begin{pmatrix} \mathbb{B}(g_b)^{-1}\mathbb{B}(g_b)\mathbb{B}(g_a)\mathbf{B}\epsilon_i \\ \mathbb{B}(g_a)^{-1}\mathbb{B}(g_a)\mathbb{B}(g_b)\mathbf{A}\epsilon_i \end{pmatrix} = \sum_{i \in \mathcal{J}} \begin{pmatrix} \mathbb{B}(g_b)^{-1}\mathbb{B}(g_b g_a)\mathbf{B}\epsilon_i \\ \mathbb{B}(g_a)^{-1}\mathbb{B}(g_a g_b)\mathbf{A}\epsilon_i \end{pmatrix} \\ &= \sum_{i \in \mathcal{J}} \begin{pmatrix} \mathbb{B}(g_b)^{-1}\mathbf{B}\mathbb{B}(g_b g_a)\epsilon_i \\ \mathbb{B}(g_a)^{-1}\mathbf{A}\mathbb{B}(g_a g_b)\epsilon_i \end{pmatrix} \\ &= \sum_{i \in \mathcal{J}} \begin{pmatrix} \tilde{\mathbf{B}}\mathbb{B}(g_b g_a)\epsilon_i \\ \tilde{\mathbf{A}}\mathbb{B}(g_a g_b)\epsilon_i \end{pmatrix} \end{aligned}$$

Notons σ la permutation sur $\llbracket 1, \ell \rrbracket$ induite par la matrice de permutation $\mathbb{B}(g_a g_b)$.

Comme $\mathbb{B}(g_b g_a) = \mathbb{B}(g_a g_b)$, il suit que $\eta(\mathbf{c}) = \sum_{i \in \mathcal{J}} \begin{pmatrix} \tilde{\mathbf{B}}\epsilon_{\sigma(i)} \\ \tilde{\mathbf{A}}\epsilon_{\sigma(i)} \end{pmatrix} \in \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$.

D'où $\eta(\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]) \subseteq \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$.

Réciproquement, en partant d'un vecteur $\tilde{\mathbf{c}} = \sum_{i \in \tilde{\mathcal{J}}} \begin{pmatrix} \tilde{\mathbf{B}}\epsilon_i \\ \tilde{\mathbf{A}}\epsilon_i \end{pmatrix}$ où $\tilde{\mathcal{J}} \subseteq \llbracket 1, \ell \rrbracket$, on trouve :

$$\eta^{-1}(\tilde{\mathbf{c}}) = \sum_{i \in \tilde{\mathcal{J}}} \begin{pmatrix} \mathbb{B}(g_a^{-1})\tilde{\mathbf{B}}\epsilon_i \\ \mathbb{B}(g_b^{-1})\tilde{\mathbf{A}}\epsilon_i \end{pmatrix} = \sum_{i \in \tilde{\mathcal{J}}} \begin{pmatrix} \mathbb{B}(g_a)^{-1}\mathbb{B}(g_b)^{-1}\mathbf{B}\epsilon_i \\ \mathbb{B}(g_b)^{-1}\mathbb{B}(g_a)^{-1}\mathbf{A}\epsilon_i \end{pmatrix} = \sum_{i \in \tilde{\mathcal{J}}} \begin{pmatrix} \mathbf{B}\mathbb{B}(g_a^{-1}g_b^{-1})\epsilon_i \\ \mathbf{A}\mathbb{B}(g_b^{-1}g_a^{-1})\epsilon_i \end{pmatrix}$$

Comme $\mathbb{B}(g_a^{-1}g_b^{-1}) = \mathbb{B}(g_b^{-1}g_a^{-1}) = \mathbb{B}(g_a g_b)^{-1}$, $\eta^{-1}(\tilde{\mathbf{c}})$ se réécrit de la façon suivante à l'aide de σ : $\eta^{-1}(\tilde{\mathbf{c}}) = \sum_{i \in \tilde{\mathcal{J}}} \begin{pmatrix} \mathbf{B}\epsilon_{\sigma^{-1}(i)} \\ \mathbf{A}\epsilon_{\sigma^{-1}(i)} \end{pmatrix} \in \text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]$.

Ainsi, $\eta^{-1}(\text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]) \subseteq \text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]$. D'où $\text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top] \subseteq \eta(\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top])$.

En conclusion : $\eta(\text{Vect}_{\mathbf{L}}[\mathbf{B}^\top \mid \mathbf{A}^\top]) = \text{Vect}_{\mathbf{L}}[\tilde{\mathbf{B}}^\top \mid \tilde{\mathbf{A}}^\top]$. □

Représentation des qubits :

Grâce au lemme 5.7.2, nous pouvons supposer sans perte de généralité que le code 2BGA abélien $\mathcal{Q}$ a déjà été simplifié. Concrètement, cela signifie que l'élément neutre e

est inclus dans les supports des éléments $\mathbf{a} = \sum_{g \in \mathbb{G}} a_g g$ et $\mathbf{b} = \sum_{g \in \mathbb{G}} b_g g$ qui permettent de définir les matrices $\mathbf{A}$ et $\mathbf{B}$. Les supports sont donc de la forme :

- $\text{Supp}(\mathbf{a}) = \{g \in \mathbb{G} \mid a_g \neq 0\} = \{e, g_{a_1}, \dots, g_{a_r}\}.$
- $\text{Supp}(\mathbf{b}) = \{g \in \mathbb{G} \mid b_g \neq 0\} = \{e, g_{b_1}, \dots, g_{b_s}\}.$

Pour appliquer le théorème 5.4.1, nous représentons les qubits du code $\mathcal{Q}$ sur les sommets d'un quotient de $\mathbb{Z}^D$ par le noyau d'une application $\mathbb{Z}$ -linéaire. L'application considérée se définit à l'aide des supports de $\mathbf{a}$ et $\mathbf{b}$ de la façon suivante :

$$\Psi : \mathbb{Z}^{r+s} \rightarrow \mathbb{G}, \quad \epsilon_i \mapsto \begin{cases} g_{a_i} & \text{si } 1 \leq i \leq r, \\ g_{b_{i-r}} & \text{si } r+1 \leq i \leq r+s \end{cases} \quad (5.2)$$

où $\{\epsilon_i\}_{1 \leq i \leq r+s}$ est la base canonique de $\mathbb{Z}^{r+s}$.

Pour prouver que le groupe stabilisateur est engendré par une famille de générateurs locaux, nous avons besoin que cette application soit surjective. Cependant, celle-ci ne l'est pas toujours. Pour que l'application soit surjective, il faudrait que les éléments du support de $\mathbf{a}$ et $\mathbf{b}$ engendrent le groupe $\mathbb{G}$.

Deuxième simplification : assurer la surjectivité de Ψ :

Pour résoudre ce problème de surjectivité, nous allons montrer que notre code $\mathcal{Q}$ a la même distance qu'un code 2BGA abélien, noté $\mathcal{Q}_{\mathbb{H}}$, qui satisfait les mêmes propriétés que $\mathcal{Q}$. Ce nouveau code est construit sur le sous-groupe, $\mathbb{H} \subseteq \mathbb{G}$, engendré par les supports de $\mathbf{a}$ et $\mathbf{b}$. En réalisant la construction avec ce sous-groupe plutôt qu'avec $\mathbb{G}$, l'application correspondante $\Psi_{\mathbb{H}} : \mathbb{Z}^{r+s} \rightarrow \mathbb{H}$ sera surjective.

Plus précisément, nous allons démontrer l'existence d'un code $\mathcal{Q}_{\mathbb{H}}$ vérifiant les propriétés suivantes :

- Les supports des éléments $\mathbf{a}(\mathbb{H})$ et $\mathbf{b}(\mathbb{H})$ de l'algèbre $\mathbb{F}_2[\mathbb{H}]$, qui permettent de construire $\mathcal{Q}_{\mathbb{H}}$, ont les mêmes tailles que les supports d'origine, à savoir $r+1$ et $s+1$. De plus, ils contiennent tous deux l'élément neutre e .
- L'application $\Psi_{\mathbb{H}} : \mathbb{Z}^{r+s} \rightarrow \mathbb{H}$ qui envoie les éléments de la base canonique sur les éléments non nuls (différents de e) des supports de $\mathbf{a}(\mathbb{H})$ et $\mathbf{b}(\mathbb{H})$, est surjective.

- La distance minimale de $\mathcal{Q}_{\mathbb{H}}$ est identique à celle de $\mathcal{Q}$.

Le lemme ci-dessous assure que nous pouvons réaliser cette simplification.

Lemme 5.7.3. *Pour prouver la borne du théorème 5.7.1, sur la distance minimale du code $\mathcal{Q}$, on peut supposer sans perte de généralité, que Ψ est surjective.*

Preuve : Supposons que Ψ ne soit pas surjective. Son image, que nous notons $\mathbb{H}$, est un sous-groupe propre de $\mathbb{G}$ composé de $|\mathbb{H}|$ éléments : $\mathbb{H} = \{h_1, \dots, h_{|\mathbb{H}|}\}$.

Construction des matrices $\mathbf{A}_{\mathbb{H}}$ et $\mathbf{B}_{\mathbb{H}}$: Le groupe $\mathbb{G}$ s'écrit comme la réunion disjointe des classes modulo $\mathbb{H}$. Il y a $[\mathbb{G} : \mathbb{H}] > 1$ classes et chacune a pour cardinal $|\mathbb{H}|$:

$$\mathbb{G} = \coprod_{i=1}^{[\mathbb{G}:\mathbb{H}]} t_i \mathbb{H} \quad \text{avec } t_1 = e$$

Pour simplifier la représentation, nous numérotions les éléments de $\mathbb{G}$ de manière spécifique. Nous ordonnons d'abord les classes consécutivement, puis à l'intérieur de chaque classe $t_j \mathbb{H}$, nous ordonnons les éléments en suivant la numérotation des éléments de $\mathbb{H}$, c'est-à-dire $t_j \mathbb{H} = \{t_j h_1, t_j h_2, \dots, t_j h_{|\mathbb{H}|}\}$.

En utilisant cette numérotation, les matrices $\mathbf{A}$ et $\mathbf{B}$ du code d'origine $\mathcal{Q}$ prennent une forme diagonale par blocs. Les blocs diagonaux de chaque matrice sont identiques. Nous noterons $\mathbf{A}_{\mathbb{H}}$, ceux de $\mathbf{A}$ et $\mathbf{B}_{\mathbb{H}}$, ceux de $\mathbf{B}$.

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_{\mathbb{H}} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{A}_{\mathbb{H}} & \ddots & \vdots \\ \vdots & \ddots & \ddots & \mathbf{0} \\ \mathbf{0} & \cdots & \mathbf{0} & \mathbf{A}_{\mathbb{H}} \end{pmatrix} \quad \text{et} \quad \mathbf{B} = \begin{pmatrix} \mathbf{B}_{\mathbb{H}} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{B}_{\mathbb{H}} & \ddots & \vdots \\ \vdots & \ddots & \ddots & \mathbf{0} \\ \mathbf{0} & \cdots & \mathbf{0} & \mathbf{B}_{\mathbb{H}} \end{pmatrix}$$

Les matrices $\mathbf{A}_{\mathbb{H}}$ et $\mathbf{B}_{\mathbb{H}}$ sont associées aux éléments de l'algèbre de groupe $\mathbf{a}_{\mathbb{H}} = \sum_{h \in \mathbb{H}} a_h h$ et $\mathbf{b}_{\mathbb{H}} = \sum_{h \in \mathbb{H}} b_h h$. Elles ont chacune un coefficient 1 en position (e, e) , puisque $a_e = b_e = 1$.

Surjectivité de $\Psi_{\mathbb{H}}$: Par définition, le sous-groupe $\mathbb{H}$ est engendré par les éléments non nuls des supports de $\mathbf{A}$ et $\mathbf{B}$: $\{g_{a_1}, \dots, g_{a_r}, g_{b_1}, \dots, g_{b_s}\}$.

Étant donné que les supports des matrices $\mathbf{A}_{\mathbb{H}}$ et $\mathbf{B}_{\mathbb{H}}$ sont les mêmes que ceux des matrices originales, cela garantit la surjectivité de l'application $\Psi_{\mathbb{H}}$ (associée aux matrices $\mathbf{A}_{\mathbb{H}}$ et $\mathbf{B}_{\mathbb{H}}$), définie par :

$$\Psi_{\mathbb{H}} : \mathbb{Z}^{r+s} \rightarrow \mathbb{H}, \epsilon_i \mapsto \begin{cases} g_{a_i} & \text{si } 1 \leq i \leq r, \\ g_{b_{i-r}} & \text{si } r+1 \leq i \leq r+s \end{cases}$$

Preuve que $d(\mathcal{Q}_{\mathbb{H}}) = d(\mathcal{Q})$: Les noyaux des matrices $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$ et $\mathbf{H}_X(\mathbb{H}) = [\mathbf{A}_{\mathbb{H}} \mid \mathbf{B}_{\mathbb{H}}]$ sont liés par la relation :

$$\begin{pmatrix} \mathbf{u}_1 \\ \vdots \\ \mathbf{u}_{[\mathbb{G}:\mathbb{H}]} \\ \mathbf{v}_1 \\ \vdots \\ \mathbf{v}_{[\mathbb{G}:\mathbb{H}]} \end{pmatrix} \in \ker \mathbf{H}_X \iff \forall i \in \llbracket 1, [\mathbb{G} : \mathbb{H}] \rrbracket, \begin{pmatrix} \mathbf{u}_i \\ \mathbf{v}_i \end{pmatrix} \in \ker \mathbf{H}_X(\mathbb{H})$$

De même, les espaces vectoriels engendrés par les lignes des matrices $\mathbf{H}_Z = [\mathbf{B}^{\top} \mid \mathbf{A}^{\top}]$ et $\mathbf{H}_Z(\mathbb{H}) = [\mathbf{B}_{\mathbb{H}}^{\top} \mid \mathbf{A}_{\mathbb{H}}^{\top}]$ sont liés par la relation :

$$\begin{pmatrix} \mathbf{u}_1 \\ \vdots \\ \mathbf{u}_{[\mathbb{G}:\mathbb{H}]} \\ \mathbf{v}_1 \\ \vdots \\ \mathbf{v}_{[\mathbb{G}:\mathbb{H}]} \end{pmatrix} \in \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z) \iff \forall i \in \llbracket 1, [\mathbb{G} : \mathbb{H}] \rrbracket, \begin{pmatrix} \mathbf{u}_i \\ \mathbf{v}_i \end{pmatrix} \in \text{Vect}_{\mathbf{L}} \mathbf{H}_Z(\mathbb{H})$$

En combinant ces deux faits, nous arrivons à la conclusion que les distances minimales d_X et $d_X(\mathbb{H})$ suivantes sont égales :

- $d_X = \min\{\text{wt}(\mathbf{c}) \mid \mathbf{c} \in \ker \mathbf{H}_X \setminus \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z)\}$
- $d_X(\mathbb{H}) = \min\{\text{wt}(\mathbf{c}_{\mathbb{H}}) \mid \mathbf{c}_{\mathbb{H}} \in \ker \mathbf{H}_X(\mathbb{H}) \setminus \text{Vect}_{\mathbf{L}}(\mathbf{H}_Z(\mathbb{H}))\}$

En conséquence, les distances minimales des codes 2BGA abéliens $\mathcal{Q}$ (associé à $\mathbf{A}, \mathbf{B}$) et $\mathcal{Q}_{\mathbb{H}}$ (associé $\mathbf{A}_{\mathbb{H}}$ et $\mathbf{B}_{\mathbb{H}}$), sont égales. $\square$

Le lemme 5.7.3 nous permet de simplifier la preuve du théorème 5.7.1. Nous pouvons désormais faire l'hypothèse que le groupe $\mathbb{G}$ est engendré par les éléments non nuls

des supports de $\mathbf{a}$ et $\mathbf{b}$, soit $\{g_{a_1}, \dots, g_{a_r}, g_{b_1}, \dots, g_{b_s}\}$.

Sous cette hypothèse, l'application Ψ (voir équation (5.2)) induit un isomorphisme $\bar{\Psi}^{-1} : \mathbb{G} \rightarrow \mathbb{Z}^{r+s} / \ker \Psi$. Cet isomorphisme nous sert à représenter chaque qubit du code $\mathcal{Q}$ (c'est-à-dire les colonnes des matrices $\mathbf{H}_X$ et $\mathbf{H}_Z$) sur un point du quotient $\mathbb{Z}^{r+s} / \ker \Psi$. Les supports des stabilisateurs associés aux lignes de ces matrices correspondent alors à des sous-ensembles de ce quotient.

Définition 5.7.1 (Représentation des qubits et stabilisateurs). *Pour $j \in \llbracket 1, \ell \rrbracket$, on représente le j -ième et le $(\ell + j)$ -ème qubit par le point $\bar{\Psi}^{-1}(g_j) \in \mathbb{Z}^{r+s} / \ker \Psi$.*

De même, pour chaque $i \in \llbracket 1, \ell \rrbracket$, les supports du i -ème stabilisateur en X (noté S_i^X) et du i -ème stabilisateur en Z (noté S_i^Z) sont représentés par les ensembles de points correspondants aux qubits de leur support :

$$\begin{aligned} \text{Supp } S_i^X &= \{\bar{\Psi}^{-1}(g_j) \mid j \in \llbracket 1, \ell \rrbracket \text{ tel que } (\mathbf{H}_X)_{i,j} = 1 \text{ ou } (\mathbf{H}_X)_{i,\ell+j} = 1\} \subseteq \mathbb{Z}^{r+s} / \ker \Psi \\ \text{Supp } S_i^Z &= \{\bar{\Psi}^{-1}(g_j) \mid j \in \llbracket 1, \ell \rrbracket \text{ tel que } (\mathbf{H}_Z)_{i,j} = 1 \text{ ou } (\mathbf{H}_Z)_{i,\ell+j} = 1\} \subseteq \mathbb{Z}^{r+s} / \ker \Psi \end{aligned}$$

Localité des stabilisateurs S_i^X et S_i^Z

Pour appliquer le Théorème 5.4.1, il faut que le code $\mathcal{Q}$ soit géométriquement local. Autrement dit, chaque stabilisateur en X et Z ne doit agir que sur un ensemble de qubits voisins dans l'espace quotient $\mathbb{Z}^{r+s} / \ker \Psi$.

Dans le lemme suivant, nous allons démontrer que le support de chaque stabilisateur en X et Z , défini plus haut, est contenu dans une boule de rayon 1 de cet espace.

Lemme 5.7.4. *Pour la représentation choisie, le code $\mathcal{Q}$ est géométriquement local. Plus précisément, le support de chaque stabilisateur S_i^X et S_i^Z est inclus dans une boule fermée de rayon 1, pour la distance induite par la norme euclidienne sur le quotient $\mathbb{Z}^{r+s} / \ker \Psi$.*

Preuve : Fixons $i \in \llbracket 1, \ell \rrbracket$. Notons $B_f(\bar{\Psi}^{-1}(g_i), 1)$ la boule fermée de centre $\bar{\Psi}^{-1}(g_i)$ et de rayon 1 de l'espace métrique $(\mathbb{R}^{r+s} / \ker \Psi, \text{dist})$:

$$B_f(\bar{\Psi}^{-1}(g_i), 1) = \{[\mathbf{x}] \in \mathbb{R}^{r+s} / \ker \Psi \mid \text{dist}(\bar{\Psi}^{-1}(g_i), [\mathbf{x}]) \leq 1\}$$

Localité de S_i^X : Le support du i -ème stabilisateur en X est déterminé par les positions des coefficients non nuls des i -èmes lignes des matrices A et B .

Dans la matrice A , ces positions correspondent à :

$$\{g_j \mid j \in \llbracket 1, \ell \rrbracket, (\mathbf{H}_X)_{i,j} = 1\} = \{g_i, g_{a_1}^{-1}g_i, \dots, g_{a_r}^{-1}g_i\}$$

Dans la matrice B , ces positions sont :

$$\{g_j \mid j \in \llbracket 1, \ell \rrbracket, (\mathbf{H}_X)_{i,\ell+j} = 1\} = \{g_i, g_{b_1}^{-1}g_i, \dots, g_{b_s}^{-1}g_i\}$$

Notons $(\epsilon_k)_{1 \leq k \leq r+s}$ la base canonique de $\mathbb{Z}^{r+s}$.

On sait que $\bar{\Psi}^{-1} : \mathbb{G} \rightarrow \mathbb{Z}^{r+s} / \ker \Psi$ est un isomorphisme de groupe satisfaisant $\bar{\Psi}^{-1}(g_{a_m}) = \epsilon_m$ et $\bar{\Psi}^{-1}(g_{b_n}) = \epsilon_{r+n}$ pour tout couple $(m, n) \in \llbracket 1, r \rrbracket \times \llbracket 1, s \rrbracket$. Ainsi, puisque le support de S_i^X s'écrit :

$$\text{Supp } S_i^X = \{\bar{\Psi}^{-1}(g_j) \mid (\mathbf{H}_X)_{i,j} = 1 \text{ ou } (\mathbf{H}_X)_{i,\ell+j} = 1\} \subset \mathbb{Z}^{r+s} / \ker \Psi$$

on en déduit que :

$$\text{Supp } S_i^X = \{\bar{\Psi}^{-1}(g_i), \bar{\Psi}^{-1}(g_i) - [\epsilon_1], \dots, \bar{\Psi}^{-1}(g_i) - [\epsilon_{r+s}]\} \subset B_f(\bar{\Psi}^{-1}(g_i), 1)$$

Localité de S_i^Z : Le support de S_i^Z , le i -ème stabilisateur en Z , dépend des positions des coefficients non nuls des i -èmes colonnes des matrices B et A .

Dans la matrice B , ces positions sont données par

$$\{g_j \mid j \in \llbracket 1, \ell \rrbracket, (\mathbf{H}_Z)_{i,j} = 1\} = \{g_i, g_{b_1}g_i, \dots, g_{b_s}g_i\}$$

Dans la matrice A , ces positions sont :

$$\{g_j \mid j \in \llbracket 1, \ell \rrbracket, (\mathbf{H}_Z)_{i,\ell+j} = 1\} = \{g_i, g_{a_1}g_i, \dots, g_{a_r}g_i\}$$

On sait que $\bar{\Psi}^{-1} : \mathbb{G} \rightarrow \mathbb{Z}^{r+s} / \ker \Psi$ est un isomorphisme de groupe satisfaisant $\bar{\Psi}^{-1}(g_{a_m}) = \epsilon_m$ et $\bar{\Psi}^{-1}(g_{b_n}) = \epsilon_{r+n}$ pour tout couple $(m, n) \in \llbracket 1, r \rrbracket \times \llbracket 1, s \rrbracket$. Ainsi, puisque le support de S_i^Z s'écrit

$$\text{Supp } S_i^Z = \{\bar{\Psi}^{-1}(g_j) \mid (\mathbf{H}_Z)_{i,j} = 1 \text{ ou } (\mathbf{H}_Z)_{i,\ell+j} = 1\} \subseteq \mathbb{Z}^{r+s} / \ker \Psi$$

on en déduit que :

$$\text{Supp } S_i^Z = \{\bar{\Psi}^{-1}(g_i), \bar{\Psi}^{-1}(g_i) + [\epsilon_1], \dots, \bar{\Psi}^{-1}(g_i) + [\epsilon_{r+s}]\} \subset B_f(\bar{\Psi}^{-1}(g_i), 1)$$

En conclusion, le code $\mathcal{Q}$ est géométriquement local et les supports des générateurs S_i^X et S_i^Z de son groupe stabilisateur sont contenus dans des boules fermées de rayon 1 de $\mathbb{Z}^{r+s} / \ker \Psi$. $\square$

Fin de la preuve du théorème :

Nous avons à présent tous les outils nécessaires pour appliquer le théorème 5.4.1. Nous allons maintenant montrer comment les combiner pour en déduire la majoration sur la distance minimale des codes 2BGA non triviaux.

Preuve du Théorème 5.7.1: Considérons un code 2BGA abélien non trivial de longueur $n = 2\ell$, dont les sous-matrices génératrices sont $\mathbf{H}_X = [\mathbf{A} \mid \mathbf{B}]$, $\mathbf{H}_Z = [\mathbf{B}^\top \mid \mathbf{A}^\top]$ où $\mathbf{A} = \sum_{g \in \mathbb{G}} a_g \mathbb{B}(g)$ et $\mathbf{B} = \sum_{g \in \mathbb{G}} b_g \mathbb{B}(g)$, ont respectivement $r + 1$ et $s + 1$ éléments non nuls par ligne.

En combinant les lemmes 5.7.2 et 5.7.3, la preuve se simplifie : ils garantissent qu'on peut, sans perte de généralités, supposer que l'élément neutre e est dans les supports de $\mathbf{A}$ et $\mathbf{B}$ et que le groupe $\mathbb{G}$ est engendré par leurs supports.

D'après les lemmes 5.7.3 et 5.7.4, il existe Λ un sous-réseau de $\mathbb{Z}^{r+s}$ (le noyau de l'application Ψ de l'équation 5.2) tel que :

- Chaque point du quotient $\mathbb{Z}^{r+s} / \Lambda$ indexe deux qubits.
- Le code $\mathcal{Q}$ est géométriquement local : son groupe stabilisateur est engendré par une famille de générateurs dont le support est inclus dans une boule de rayon 1 de ce quotient.

Ces conditions nous permettent d'appliquer le théorème 5.4.1 avec les paramètres $m = 2$ et $\rho = 1$. En conséquence, la distance minimale du code $\mathcal{Q}$ est majorée par :

$$d < 2\sqrt{\gamma_D}(\sqrt{D} + 4) \ell^{\frac{D-1}{D}}$$

où γ_D est la constante de Hermite D -dimensionnelle pour $D = r + s$. $\square$

Cette borne garantit que la distance minimale des codes 2BGA non triviaux croît de façon sous-linéaire avec leur longueur, selon la relation $d = O(\ell^{(D-1)/D})$ où D dépend du poids des générateurs locaux du groupe stabilisateur. Comme les codes GB sont une sous-famille des codes 2BGA, le théorème 5.7.1 s'applique également à eux.

Remarque 5.7.2. *Lorsqu'on applique cette borne aux codes GB de longueur 2ℓ et de poids w , nous obtenons que leur distance minimale est bornée par $O(\ell^{(D-1)/D})$ où $D = w - 2$. Cela constitue une amélioration par rapport à la borne de Pryadko et Wang [16], qui proposaient $d = O(\ell^{(D-1)/D})$ où D valait $w - 1$ en général et $D = w - 2$ uniquement quand ℓ était un nombre premier.*

5.8 Conclusion du chapitre :

Les paramètres des codes stabilisateurs locaux ont été largement étudiés, notamment par Bravyi et Terhal [24]. Ils ont démontré que lorsque les qubits de ces codes sont placés sur le tore discret $(\mathbb{Z}/L\mathbb{Z})^D$, de manière à ce que le groupe stabilisateur soit engendré par des générateurs dont le support est inclus dans un cube à ρ^D sommets, la distance minimale croît de façon sous-linéaire avec sa longueur ℓ : $d \leq \rho \ell^{(D-1)/D}$.

Notre travail généralise ce résultat. Le Théorème 5.4.1 établit une borne supérieure sur la distance minimale de tous les codes stabilisateurs dont les qubits peuvent être représentés sur un quotient de $\mathbb{Z}^D$ par un réseau Λ de rang plein. La distance minimale d d'un tel code de longueur ℓ est alors bornée : $d \leq \kappa_D \ell^{(D-1)/D}$, où $\kappa_D = O(D)$ est défini explicitement dans le théorème.

Notre borne est un peu moins précise que celle de Bravyi–Terhal lorsqu'on l'applique directement au même réseau, car sa formulation plus générale introduit un facteur constant supplémentaire. En contrepartie, elle est bien plus adaptable à une grande variété de réseaux.

En appliquant ce théorème à la classe des codes 2BGA abéliens de longueur ℓ , dont le code stabilisateur est engendré par des générateurs de poids w , nous obtenons une nouvelle borne sur leur distance minimale (Théorème 5.7.1) : $d < 2\sqrt{\gamma_D}(\sqrt{D} + 4) \ell^{\frac{D-1}{D}}$ où $D = w - 2$ et γ_D est la constante de Hermite D -dimensionnelle. Cette borne constitue une amélioration par rapport au résultat de Pryadko et Wang pour la sous-classe des codes GB de poids w [16], pour lesquels ils proposaient une borne de $d = O(\ell^{(D-1)/D})$ avec $D = w - 1$ (ou $w - 2$ si ℓ est un nombre premier).

Plusieurs pistes de recherche découlent de ce travail :

- **Affinement de la borne** : Il serait intéressant de déterminer des familles de codes ou des réseaux Λ spécifiques pour lesquels la borne pourrait être affinée, offrant ainsi une estimation plus précise de la distance minimale.
- **Implémentation pratique** : D'un point de vue pratique, il serait pertinent d'étudier l'implémentation de portes quantiques sur les codes 2BGA tels que nous les avons représentés (voir section 5.7.3), en particulier pour les dimensions $D = 2$ et $D = 3$. Cette analyse permettrait d'évaluer les performances de ces codes, et notamment leur tolérance aux fautes.



Conclusion générale

Dans cette thèse, nous avons exploré comment la représentation géométrique des qubits influence les performances des codes stabilisateurs locaux. Notre démarche a consisté à montrer comment l'espace sur lequel les qubits sont représentés et la façon dont ils sont positionnés peuvent être exploités pour construire des codes performants ou pour évaluer leurs capacités de correction d'erreurs.

Synthèse des contributions

Nous avons étudié les codes stabilisateurs locaux dont les qubits sont disposés sur des quotients de $\mathbb{Z}^D$ en adoptant deux approches distinctes.

Dans le Chapitre 4, notre démarche est allée du particulier au général. Nous avons mené une étude approfondie des codes GB(2,2), qui sont des codes locaux en deux dimensions. En utilisant le formalisme des graphes de Cayley, nous avons établi une borne inférieure sur leur distance minimale, qui nous a ensuite permis de construire trois nouvelles familles de codes GB(2,2) aux performances optimales.

En outre, en considérant une relation d'équivalence spécifique qui préserve les structures CSS, nous avons réalisé deux classifications, chacune portant sur des codes de poids quatre et de longueur inférieure ou égale à 200 : l'une pour les meilleurs codes

GB(2,2) et l'autre pour les meilleurs codes surface 2D connus.

À l'inverse, dans le chapitre 5, nous avons adopté une perspective plus générale, en allant du général au particulier. Nous avons élargi le champ d'application de la borne de Bravyi et Terhal à tous les codes stabilisateurs locaux dont les qubits sont placés sur les sommets de quotients de $\mathbb{Z}^D$ par des réseaux arbitraires. Cette approche nous a permis de montrer que la distance minimale de ces codes a une croissance sous-linéaire en leur longueur. De plus, en appliquant cette nouvelle borne à la sous-famille des codes GB, nous avons pu améliorer une borne existante sur leur distance minimale, démontrant ainsi l'efficacité de notre approche théorique sur des cas concrets.

Perspectives de recherche

Cette thèse ouvre de nombreuses pistes de recherche tant théoriques que pratiques pour la conception de futurs ordinateurs quantiques.

Vers une classification des codes quantiques

Un premier axe de recherche serait de s'intéresser à la classification des codes quantiques, un aspect encore peu exploré. Il serait pertinent d'identifier, pour chaque famille de codes, les relations d'équivalence qui préservent leurs structures sous-jacentes et leurs propriétés de correction d'erreurs. Une telle démarche permettrait aux chercheurs de déterminer si un nouveau code est véritablement inédit ou s'il s'agit d'une variation d'un code déjà connu. Cela permettrait de concentrer l'étude et l'optimisation des performances sur un ensemble plus restreint de codes.

De la théorie à la pratique : l'optimisation des représentations

Un second axe de recherche, davantage orienté vers l'implémentation, consiste à mieux définir ce qui fait un "bon" code stabilisateur local. En nous appuyant sur notre approche de représentation des qubits, il serait particulièrement pertinent de déterminer quelles configurations géométriques permettent d'implémenter les portes quantiques de manière optimale dans la pratique. La conception de codes stabilisateurs locaux facilement implémentables pourrait ainsi contribuer au développement de systèmes fiables de correction d'erreurs et à la démocratisation des ordinateurs quantiques.

7

Bibliographie

Sommaire

Références	158
Liste des travaux	161

Références

- [1] P. W. Shor, “Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer,” *SIAM Review*, vol. 41, no. 2, pp. 303–332, 1999. eprint: <https://doi.org/10.1137/S0036144598347011>. [Online]. Available: <https://doi.org/10.1137/S0036144598347011>.
- [2] G. Gras and M.-N. Gras, *Algèbre fondamentale - Arithmétique*. ELLIPSES, Jun. 2004, Niveau L3 et M1.
- [3] A. R. Calderbank and P. W. Shor, “Good quantum error-correcting codes exist,” *Physical Review A*, vol. 54, no. 2, pp. 1098–1105, Aug. 1996.
- [4] A. M. Steane, “Error correcting codes in quantum theory,” *Physical Review Letters*, vol. 77, pp. 793–797, 5 Jul. 1996.
- [5] D. Aharonov and M. Ben-Or, *Fault-tolerant quantum computation with constant error rate*, 1999. arXiv: [quant-ph/9906129](https://arxiv.org/abs/quant-ph/9906129) [quant-ph]. [Online]. Available: <https://arxiv.org/abs/quant-ph/9906129>.
- [6] A. Kitaev, “Fault-tolerant quantum computation by anyons,” *Annals of Physics*, vol. 303, no. 1, pp. 2–30, Jan. 2003. arXiv: [quant-ph/9707021](https://arxiv.org/abs/quant-ph/9707021).
- [7] H. Bombin and M. A. Martin-Delgado, “Homological error correction: Classical and quantum codes,” *Journal of Mathematical Physics*, vol. 48, no. 5, p. 052 105, May 2007. eprint: https://pubs.aip.org/aip/jmp/article-pdf/doi/10.1063/1.2731356/14848728/052105_1_online.pdf. [Online]. Available: <https://doi.org/10.1063/1.2731356>.
- [8] A. Leverrier and G. Zémor, “Quantum Tanner codes,” in *2022 IEEE 63rd Annual Symposium on Foundations of Computer Science (FOCS)*, IEEE, 2022, pp. 872–883.
- [9] P. Panteleev and G. Kalachev, “Asymptotically good quantum and locally testable classical LDPC codes,” in *Proceedings of the 54th Annual ACM SIGACT Symposium on Theory of Computing*, 2022, pp. 375–388.

- [10] S. Bravyi, A. W. Cross, J. M. Gambetta, D. Maslov, P. Rall, and T. J. Yoder, “High-threshold and low-overhead fault-tolerant quantum memory,” *Nature*, vol. 627, no. 8005, pp. 778–782, 2024.
- [11] A. A. Kovalev and L. P. Pryadko, “Quantum kronecker sum-product low-density parity-check codes with finite rate,” *Phys. Rev. A*, vol. 88, p. 012311, 1 Jul. 2013. [Online]. Available: <https://link.aps.org/doi/10.1103/PhysRevA.88.012311>.
- [12] M. B. Hastings, J. Haah, and R. O’Donnell, *Fiber bundle codes: breaking the $n^{1/2}\text{polylog}(n)$ barrier for Quantum LDPC codes* (STOC 2021). Virtual, Italy: Association for Computing Machinery, 2021, pp. 1276–1288. [Online]. Available: <https://doi.org/10.1145/3406325.3451005>.
- [13] W. Rozendaal and G. Zémor, *Analysis of the error-correcting radius of a renormalisation decoder for kitaev’s toric code*, 2023. arXiv: 2309.12165 [quant-ph]. [Online]. Available: <https://arxiv.org/abs/2309.12165>.
- [14] R. Wang and L. P. Pryadko, *Collection of codes constructed for “distance bounds for generalized bicycle codes”*, GitHub repository; updated on 2022-03-30, 2022. [Online]. Available: <https://github.com/QEC-pages/GB-codes>.
- [15] D. MacKay, G. Mitchison, and P. McFadden, “Sparse-graph codes for quantum error correction,” *IEEE Transactions on Information Theory*, vol. 50, no. 10, pp. 2315–2330, 2004.
- [16] R. Wang and L. P. Pryadko, “Distance bounds for Generalized Bicycle codes,” *Symmetry*, vol. 14, no. 7, p. 1348, 2022.
- [17] D. Gottesman, “Theory of fault-tolerant quantum computation,” *Physical Review A*, vol. 57, no. 1, pp. 127–137, Jan. 1998. [Online]. Available: <http://dx.doi.org/10.1103/PhysRevA.57.127>.
- [18] H. Whitney, “2-isomorphic graphs,” *American Journal of Mathematics*, vol. 55, no. 1, pp. 245–254, 1933. Accessed: Aug. 19, 2025. [Online]. Available: <http://www.jstor.org/stable/2371127>.
- [19] H. Whitney, “Non-separable and planar graphs,” *Transactions of the American Mathematical Society*, vol. 34, no. 2, pp. 339–362, 1932. [Online]. Available: <https://doi.org/10.1090/S0002-9947-1932-1501641-2>.
- [20] H. Whitney, “Congruent graphs and the connectivity of graphs,” *American Journal of Mathematics*, vol. 54, no. 1, pp. 150–168, 1932.
- [21] F. Arnault, P. Gaborit, and N. Saussay, *Classification of (2,2)-generalized bicycle (gb) codes*, GitHub repository; updated on 2025-07-13, 2025. [Online]. Available: https://github.com/NicolasSaussay/weight-4_GB-Codes_Classification.

- [22] H.-K. Lin and L. P. Pryadko, “Quantum Two-Block Group Algebra codes,” *Physical Review A*, vol. 109, no. 2, p. 022 407, 2024. arXiv: 2306.16400 [quant-ph].
- [23] P. Panteleev and G. Kalachev, “Degenerate quantum LDPC codes with good finite length performance,” *Quantum*, vol. 5, p. 585, Nov. 2021. [Online]. Available: <http://dx.doi.org/10.22331/q-2021-11-22-585>.
- [24] S. Bravyi and B. Terhal, “A no-go theorem for a two-dimensional self-correcting quantum memory based on stabilizer codes,” *New Journal of Physics*, vol. 11, no. 4, p. 043 029, Apr. 2009. [Online]. Available: <http://dx.doi.org/10.1088/1367-2630/11/4/043029>.
- [25] E. Portnoy, *Local quantum codes from subdivided manifolds*, 2023. arXiv: 2303.06755 [quant-ph]. [Online]. Available: <https://arxiv.org/abs/2303.06755>.
- [26] D. J. Williamson and N. Baspın, “Layer codes,” *Nature Communications*, vol. 15, no. 1, p. 9528, 2024.
- [27] T.-C. Lin, A. Wills, and M.-H. Hsieh, *Geometrically local quantum and classical codes from subdivision*, 2024. arXiv: 2309.16104 [quant-ph]. [Online]. Available: <https://arxiv.org/abs/2309.16104>.
- [28] X. Li, T.-C. Lin, and M.-H. Hsieh, *Transform arbitrary good quantum LDPC codes into good geometrically local codes in any dimension*, 2024. arXiv: 2408.01769 [quant-ph]. [Online]. Available: <https://arxiv.org/abs/2408.01769>.
- [29] S. Bravyi, D. Poulin, and B. Terhal, “Tradeoffs for reliable quantum information storage in 2D systems,” *Physical Review Letters*, vol. 104, no. 5, Feb. 2010. [Online]. Available: <http://dx.doi.org/10.1103/PhysRevLett.104.050503>.
- [30] N. Baspın and A. Krishna, “Connectivity constrains quantum codes,” *Quantum*, vol. 6, p. 711, May 2022.
- [31] N. Baspın, V. Guruswami, A. Krishna, and R. Li, “Improved rate-distance trade-offs for quantum codes with restricted connectivity,” *Quantum Science and Technology*, vol. 10, no. 1, p. 015 021, 2024.
- [32] C. Hermite, “Extraits de lettres de M. Ch. Hermite à M. Jacobi sur différents objets de la théorie des nombres,” fre, *Journal für die reine und angewandte Mathematik*, vol. 40, pp. 261–277, 1850.
- [33] J. Martinet, *Perfect lattices in Euclidean spaces* (Grundlehren der mathematischen Wissenschaften). Springer Berlin Heidelberg, 2002.
- [34] P. Panteleev and G. Kalachev, “On the minimum distance in one class of quantum LDPC codes,” *Intelligent systems. Theory and applications*, vol. 24, pp. 87–117, 4 2020.

Liste des Travaux

1. F. Arnault, P. Gaborit, W. Rozendaal, N. Saussay, and G. Zémor, "Upper Bounds on the Minimum Distance of Structured LDPC Codes," arXiv, 2025. [Online]. Available: <https://arxiv.org/abs/2501.19125>
2. F. Arnault, P. Gaborit, W. Rozendaal, N. Saussay, and G. Zémor, "A Variant of the Bravyi-Terhal Bound for Arbitrary Boundary Conditions," arXiv, 2025. [Online]. Available: <https://arxiv.org/abs/2502.04995>
3. F. Arnault, P. Gaborit, and N. Saussay, "(2,2)-GB Codes: Classification and Comparison with weight-4 Surface Codes", arXiv, 2025. [Online]. Available: <https://arxiv.org/abs/2507.21237>
4. F. Arnault, P. Gaborit, and N. Saussay, "Classification of (2,2)-Generalized Bicycle (GB) Codes", GitHub repository, 2025. [Online]. Available: https://github.com/NicolasSaussay/weight-4_GB-Codes_Classification (updated on 2025-07-13)

Étude de la distance minimale des codes stabilisateurs locaux

Résumé : Cette thèse étudie l'influence de l'agencement spatial des qubits sur les performances des codes stabilisateurs locaux. Nous montrons en particulier que leur disposition sur des sommets de $\mathbb{R}^D$ affecte directement la distance minimale des codes.

Nos contributions se répartissent en deux axes principaux. D'abord, nous avons exclusivement étudié les codes locaux en deux dimensions ($D = 2$). À partir d'une borne inférieure sur la distance minimale que nous avons établie, nous avons construit trois nouvelles familles de codes Bicycle généralisés (GB) de poids quatre, qui atteignent les meilleures performances possibles pour des codes surface 2D de même poids. Nous avons également établi deux classifications, chacune portant sur des codes de poids quatre et de longueur inférieure ou égale à 200 : l'une pour les meilleurs codes GB et l'autre pour les meilleurs codes surface 2D.

Ensuite, nous avons étendu notre étude aux codes locaux en dimension $D \geq 2$ quelconque. Nous avons démontré une borne théorique sur la distance minimale des codes stabilisateurs locaux définis sur des quotients de $\mathbb{Z}^D$ par des réseaux. Cette borne, qui élargit le champ d'application du célèbre théorème de Bravyi–Terhal, nous a permis d'établir que la distance minimale des codes 2BGA abéliens de poids w croît de manière sous-linéaire avec leur longueur, $d = O(\ell^{(D-1)/D})$ où $D = w - 2$. Dans le cas particulier des codes GB, qui constituent une sous-famille des codes 2BGA abéliens, notre résultat affine une borne déjà connue, soulignant ainsi la portée de notre approche théorique.

Mots clés : Codes LDPC quantiques (qLDPC), Codes stabilisateurs locaux, Codes Bicycle généralisés (GB), Codes 2BGA Abéliens, Borne sur la distance minimale, Classification de codes.

Study of the minimum distance of local stabilizer codes

Abstract: This thesis investigates how the spatial arrangement of qubits influences the performance of local stabilizer codes. In particular, we show that their placement on vertices of $\mathbb{R}^D$ directly affects the minimum distance of these codes.

Our contributions are organized along two main directions. First, we focused exclusively on local codes in two dimensions ($D = 2$).

Based on a lower bound on the minimum distance that we established, we constructed three new families of Generalized Bicycle (GB) codes of weight four, which attain the theoretically optimal performance for 2D surface codes of the same weight. We also established two classifications, each concerning codes of weight four and length less than or equal to 200: one for the best GB codes and the other for the best 2D surface codes.

Second, we extended our study to arbitrary dimensions $D \geq 2$. We derived a theoretical bound on the minimum distance of local stabilizer codes whose qubits are represented on quotients of $\mathbb{Z}^D$ by lattices. This bound generalizes and extends the well-known Bravyi–Terhal bound. As an application, we proved that abelian 2BGA codes of weight w have a minimum distance that grows sublinearly with their length, $d = O(\ell^{(D-1)/D})$ where $D = w - 2$. In the particular case of GB codes, which form a subclass of abelian 2BGA codes, our result sharpens a previously known bound, thereby highlighting the significance of our theoretical approach.

Keywords: Quantum LPDC codes, Local stabilizer Codes, Generalized Bicycle Codes (GB), 2BGA Abelian Codes, Upper Bound on the Minimum Distance, Quantum Codes Classification.