



Topological Quantum Error-Correcting Codes beyond dimension 2

Vivien Londe

► To cite this version:

Vivien Londe. Topological Quantum Error-Correcting Codes beyond dimension 2. Information Theory [math.IT]. Université de Bordeaux, 2019. English. NNT : 2019BORD0304 . tel-02429868v2

HAL Id: tel-02429868

<https://inria.hal.science/tel-02429868v2>

Submitted on 16 Nov 2021

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

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

THÈSE DE DOCTORAT DE L'UNIVERSITÉ DE BORDEAUX

Specialité

MATHEMATIQUES

École doctorale de Mathématiques et d'Informatique de Bordeaux

Codes correcteurs d'erreur quantique topologiques au
delà de la dimension 2

Auteur:

Vivien Londe

Sous la direction de:

Anthony Leverrier and Gilles Zémor

soutenue le 6 December 2019

Membres du Jury:

Anthony LEVERRIER	Inria de Paris	Directeur de Thèse
Gilles ZÉMOR	Université de Bordeaux	Directeur de Thèse
Benjamin AUDOUX	Aix-Marseille Université	Rapporteur
Dan BROWNE	University College London	Rapporteur
Valentin SAVIN	CEA-Leti	Membre du Jury
Elham KASHEFI	LIP6 Sorbonne Université	Membre du Jury
Christine BACHOC	Université de Bordeaux	Membre du Jury
Philippe GRANGIER	Institut d'Optique, Un. Paris Sud	Membre du Jury
Jean-Pierre TILLICH	Inria de Paris	Membre du Jury

Rapporteurs:

Benjamin AUDOUX	Aix-Marseille Université
Dan BROWNE	University College London

Financée par
AMX, Ecole Polytechnique

Équipe-Projet
SECRET
Inria de Paris,
2 rue Simone IFF,
75 012 Paris

Acknowledgment

Je tiens avant tout à remercier mes directeurs de thèse Anthony Leverrier et Gilles Zémor. En premier lieu, un immense merci à Anthony de m'avoir accompagné au quotidien et d'avoir répondu avec précision et pédagogie à mes nombreuses questions sur les sciences quantiques. Je le remercie d'avoir su me faire profiter de l'étendue de ses connaissances. Je lui suis particulièrement reconnaissant de m'avoir encouragé à me confronter à des problèmes ambitieux et de m'avoir si bien conseillé, soutenu, orienté et accompagné dans cette tâche exaltante. Merci pour les nombreux bons moments passés ensemble et les précieuses discussions scientifiques. Toute ma gratitude va à Gilles pour le temps passé avec lui à Bordeaux. Je le remercie de s'être autant impliqué dans les projets de recherche sur lesquels il m'a guidé. Ces journées de réflexion intense et collective sont un des souvenirs heureux de ces trois années.

I am also very grateful to my reviewers, Professor Benjamin Audoux and Professor Dan Browne, for their accurate and constructive reading of my manuscript. I am thankful to Benjamin for his precious advices during Tocq meetings as well. I also would like to thank Professors Valentin Savin, Elham Kashefi, Christine Bachoc, Philippe Grangier and Jean-Pierre Tillich for taking part in my thesis committee. Thank you to Jean-Pierre for both academic advices and good times during these three years at Inria as well.

Merci à Anne Canteaut et à tous les membres de l'équipe Secret pour les bonnes conditions de travail et la convivialité de l'environnement à l'Inria. Merci à Kaushik, André S., Rémi, Anirudh, Lucien, Christophe et à mon frère académique Antoine d'avoir fait du bureau quantique un endroit où j'avais envie de me rendre. Merci à Yann, Kevin, Thomas, Rodolfo, Andrea, Xavier, Valentin, Ferdinand, Matthieu, Daniel, Léo, Ivan, Simon, Christina, Maria, Gaëtan, Pascale, André C., Christelle et Nicolas pour les bons moments passés ensemble. Merci à Nikolas Breuckmann pour une collaboration agréable et productive.

Merci à mes parents pour leurs regards inconditionnellement bienveillants et encourageants. Merci à Mona pour son intérêt et les discussions sur l'histoire des sciences. Enfin et surtout merci à Marianne pour son écoute de tous les aspects de ma thèse, des plus techniques au plus psychologiques, pour son soutien quotidien et pour sa participation indirecte à cette aventure avec moi.

Contents

Introduction

Résumé

1	Quantum computing concepts	1
1.1	Quantum memory: qubits	1
1.2	Quantum processing: Unitary matrices	2
1.3	Quantum measurements: Hermitian operators	3
1.4	Pauli operators	3
1.5	Quantum Error Correcting Codes	4
1.6	Decoding and logical errors	5
1.7	Homological quantum codes	7
1.8	The orthogonal decomposition of $\mathbb{C}^{2^n}$ defined by a CSS code	8
1.9	$\mathbb{F}^n$ and $\mathbb{C}^{2^n}$ dualities	11
2	Golden codes	15
2.1	A variant of Guth and Lubotzky's construction based on a Regular Tessellation of Hyperbolic Space	16
2.2	Hyperbolic 4-space and its Regular Tessellation by Hypercubes	17
2.2.1	Hyperbolic space	18
2.2.2	Regular polytopes and tessellations	18
2.2.3	Isometry group of the tessellation	20
2.2.4	Coxeter group approach	22
2.3	Compact Manifolds equipped with a $\{4, 3, 3, 5\}$ Tessellation	23
2.3.1	Number fields and rings of integers	23
2.3.2	Arithmetic subgroups $\Gamma^{\{4,3,3,5\}}(I)$	24
2.3.3	Lower bound on the rate of the quantum codes	27
2.3.4	Lower bound on the minimum distances of the quantum codes	28
2.3.5	Estimates of the number of physical qubits	31
2.4	Local Decoders	32
2.4.1	Decoding Z -errors	33
2.4.2	Decoding X -errors	34
2.4.3	Proof of the Z -decoder Lemma	34
2.4.4	Proof of the X -decoder Lemma	38
2.5	Conclusion of the chapter	40
3	The decoding problem for Hyperbolic 4D codes	41
3.1	Arithmetic construction of $\{5,3,3,5\}$ quantum codes	41
3.1.1	Canonical Representation of Coxeter groups	41
3.1.2	Representation of the $\{5,3,3,5\}$ string Coxeter group	43
3.1.3	Abstract regular polytopes and string C-groups	45

3.1.4	Normal quotients of Coxeter groups are C-groups	46
3.1.5	The non-local subgroup condition	47
3.1.6	Reducing Γ modulo a prime ideal I of $\mathbb{Z}[\phi]$	48
3.1.7	Subgroups of $G_{\mathbb{F}_5}$ satisfying the non-local condition.	51
3.1.8	A note on normal subgroups and regular codes	52
3.2	Numerical search for normal subgroups of the Coxeter group	53
3.3	Decoding with Belief Propagation	53
3.3.1	Belief propagation on a tree	53
3.4	Numerical Results	56
4	Cube-quotient codes	61
4.1	The Long Exact Sequence	62
4.2	Number of logical qubits in a hemicube code	62
4.3	Cycle minimum distance for the hemicube code	63
4.4	Cocycle minimum distance for the hemicube code	65
4.5	Number of logical qubits in a cube-quotient code	66
4.6	Cohomology basis and short exact sequence in cohomology in a cube-quotient code	67
4.6.1	Adjunction and short exact sequence in homology in a cube-quotient code	68
4.7	Product cycles in a cube-quotient code	69
4.8	Homology basis in a cube-quotient code	71
4.9	Cocycle minimum distance in a cube-quotient code	71
4.10	Cycle minimum distance in a cube-quotient code	72
4.11	The non-local subgroup condition on cube-quotient codes	74
	Conclusion	77

Introduction

Error correction is the set of techniques used in order to store, process and transmit information reliably in a noisy context. The classical theory of error correction is based on encoding classical information redundantly. A major endeavor of the theory is to find optimal trade-offs between redundancy, which we try to minimize, and noise tolerance, which we try to maximize.

The quantum theory of error correction cannot directly imitate the redundant schemes of the classical theory because it has to cope with the no-cloning theorem: quantum information cannot be copied. Quantum error correction is nonetheless possible by spreading the information on more quantum memory elements than would be necessary. In quantum information theory, dilution of the information replaces redundancy since copying is forbidden by the laws of quantum mechanics.

Besides this conceptual difference, quantum error correction inherits a lot from its classical counterpart. In this PhD thesis we are concerned with a class of quantum error correcting codes whose classical counterpart was defined in 1961 by Gallager [Gal62]. At that time quantum information was not even a research domain yet. This class is the family of low density parity check (LDPC) codes. Informally a code is said to be LDPC if the constraints imposed to ensure redundancy in the classical setting or dilution in the quantum setting are local.

More precisely this PhD thesis focuses on a subset of the LDPC quantum error correcting codes: the homological quantum error correcting codes. These codes take their name from the mathematical field of homology, whose objects of study are sequences of linear maps such that the kernel of a map contains the image of its left neighbour. Originally introduced to study the topology of geometric shapes, homology theory now encompasses more algebraic branches as well, where the focus is more abstract and combinatorial. The same is true of homological codes: they were introduced in 1997 by Kitaev [Kit03] with a quantum code that has the shape of a torus. They now form a vast family of quantum LDPC codes, some more inspired from geometry than others. Homological quantum codes were designed from spherical, Euclidean and hyperbolic geometries, from 2-dimensional, 3-dimensional and 4-dimensional objects, from objects with increasing and unbounded dimension and from hypergraph or homological products.

After we introduce some general quantum information concepts in the first chapter of this manuscript, we focus in the two following ones on families of quantum codes based on 4-dimensional hyperbolic objects. We highlight the interplay between their geometric side, given by manifolds, and their combinatorial sides, given by abstract polytopes. We use both sides to analyze the corresponding quantum codes. In the fourth and last chapter we analyze a family of quantum codes based on spherical objects of arbitrary dimension. To have more flexibility in the design of quantum codes, we use combinatorial objects that

realize this spherical geometry: hypercube complexes. This setting allows us to introduce a new link between classical and quantum error correction where classical codes are used to introduce homology in hypercube complexes.

Résumé

La mémoire quantique est constituée de matériaux présentant des effets quantiques comme la superposition. C'est cette possibilité de superposition qui distingue l'élément élémentaire de mémoire quantique, le qubit, de son analogue classique, le bit. Contrairement à un bit classique, un qubit peut être dans un état différent de l'état 0 et de l'état 1. Il existe aujourd'hui un grand nombre de propositions d'implémentations physiques de mémoire quantique envisagées. On peut citer par exemple les qubits supraconducteurs, les qubits topologiques, les qubits implémentés par un spin nucléaire ou électronique et les qubits implémentés par un degré de liberté photonique. Une difficulté majeure de la réalisation physique de mémoire quantique est la nécessité d'isoler le système utilisé de son environnement. En effet l'interaction d'un système quantique avec son environnement entraîne un phénomène appelé décohérence qui se traduit par des erreurs sur l'état du système quantique. Dit autrement, à cause de la décohérence, il est possible que les qubits ne soient pas dans l'état dans lequel il est prévu qu'ils soient. Lorsque ces erreurs s'accumulent le résultat d'un calcul quantique a de grandes chances de ne pas être le résultat attendu.

La correction d'erreur quantique est un ensemble de techniques permettant de protéger l'information quantique de ces erreurs. Elle consiste à réaliser un compromis entre le nombre de qubits et leur qualité. Plus précisément un code correcteur d'erreur permet à partir de N qubits physiques bruités de simuler un nombre plus petit K de qubits logiques, c'est-à-dire virtuels, moins bruités.

La tâche semble a priori irréalisable. En effet le théorème de non clonage quantique affirme qu'il est impossible de copier l'information quantique. C'est pourquoi les techniques de correction d'erreur classique fondée sur la redondance ne fonctionnent pas dans le cas quantique. Une solution à cette énigme a été trouvée en 1995 par le mathématicien Peter Shor [Sho95]. Il parvient à répartir l'information quantique que pourrait porter un seul qubit (dit qubit logique) sur un ensemble de 9 qubits physiques. De plus si l'un des 9 qubits physiques subit une erreur, l'information logique n'est pas perdue. Cela fonctionne quelque soit l'erreur sur le qubit physique. C'est le premier code correcteur d'erreur quantique: le code de Shor. Depuis de nombreux autres codes quantiques ont été trouvés. La famille de codes la plus connue est sans doute celle découverte par le physicien Alexei Kitaev: le code torique.

Le code torique consiste à agencer un nombre N de qubits physiques selon une grille carrée munie de conditions aux limites périodiques (c'est-à-dire selon un tore) et de les mesurer par ensembles de 4 qubits. On peut alors montrer que l'information quantique encodée dans les lacets non contractiles du tore n'est pas détruite par ces mesures. Comme un tore possède une base de l'ensemble de ses lacets non contractiles constituée de 2 éléments, cette information quantique correspond à 2 qubits logiques. De plus cette information quantique n'est pas perdue tant qu'au plus $\sqrt{N}/4$ qubits physiques subissent une erreur. La supériorité du code torique par rapport au code de Shor est donc importante en termes de nombre de qubits physiques pour lesquels on peut tolérer des erreurs. De plus le

code torique possède la propriété essentielle de nécessiter un nombre constant de qubits impliqués dans chaque mesure (4 en l'occurrence), indépendamment du nombre de qubits physiques. Le code torique appartient donc à la famille des codes quantiques LDPC (Low Density Parity Check).

Cette construction peut être généralisée à des formes géométriques (variétés) autres qu'un tore. La famille des codes hyperboliques 2D proposée en 2002 par Michael Freedman, David Meyer et Feng Luo permet d'encoder un nombre de qubits logiques proportionnel au nombre de qubits physiques. En revanche elle fait moins bien que le code torique en termes de nombres de qubits pouvant avoir subi une erreur. En effet elle n'en tolère qu'un nombre proportionnel au logarithme du nombre de qubits physiques.

En 2013 Nicolas Delfosse a montré qu'il existait des limites à cette méthode consistant à agencer des qubits physiques selon une surface. On note désormais N le nombre de qubits physiques, K le nombre de qubits logiques et D le double du nombre de qubits physiques sur lesquels on peut tolérer des erreurs sans perdre l'information logique d'un code quantique. Alors tout code quantique défini par une surface voit ses paramètres contraints par l'inégalité $K(\frac{D}{\log K})^2 \leq CN$ pour une constante universelle C . De plus le code torique sature cette inégalité avec K constant, certaines familles de codes hyperboliques la sature avec K proportionnel à N et les codes semi-hyperboliques définis dans [BVC⁺17] la sature en interpolant entre ces deux cas extrêmes.

Toutefois la borne de Delfosse ne concerne que les codes quantiques définis à partir d'une surface, c'est-à-dire une variété de dimension 2. En 2014, Larry Guth et Alexander Lubotzky montrent qu'une famille de code définie à partir de variétés hyperboliques de dimension 4 dépasse la borne $K(\frac{D}{\log K})^2 \leq CN$.

Dans cette thèse, nous sommes partis de la construction de Guth et Lubotzky et en avons donné une version plus explicite et plus régulière. Pour définir un pavage régulier de l'espace hyperbolique de dimension 4, nous utilisons le groupe de symétrie de symbole de Schläfli $\{4, 3, 3, 5\}$. Nous en donnons la représentation matricielle correspondant au modèle de l'hyperboloïde et à un hypercube centré sur l'origine et dont les faces sont orthogonales aux quatre axes de coordonnées. Ce groupe arithmétique admet une infinité de sous-groupes de congruence. Si l'on considère le quotient de l'espace hyperbolique de dimension 4 par chacun de ces groupes, on obtient une variété close (avec singularités pour les groupes les plus grands). La relation d'inclusion des groupes correspond une relation de revêtement des variétés (de façon contravariante). À part éventuellement pour les premiers sous-groupes, on obtient les relations d'incidence du pavage $\{4, 3, 3, 5\}$ en considérant les cosets correspondants à quatre des cinq générateurs de Coxeter du groupe. Cette construction permet d'obtenir une famille de codes quantiques encodant un nombre de qubits logiques proportionnel au nombre de qubits physiques et dont la distance minimale croît au moins comme $N^{0.1}$. Bien que ces paramètres soient également ceux de la construction de Guth et Lubotzky, la régularité de cette construction permet de construire explicitement des exemples de taille raisonnable et d'envisager des algorithmes de décodage qui exploitent cette régularité.

Dans un second chapitre nous considérons une famille de codes quantiques hyperboliques 4D de symbole de Schläfli $\{5, 3, 3, 5\}$. Après avoir énoncé une façon de prendre le quotient des groupes correspondant en conservant la structure locale du groupe, nous construisons les matrices de parité correspondant à des codes quantiques ayant 144, 720, 9792, 18 000 et 90 000 qubits physiques. Ces codes ont respectivement 72, 144, 2 200, 3624 et 18 024 qubits logiques. Nous appliquons un algorithme de Belief Propagation au décodage de ces

codes et analysons les résultats numériquement.

Dans un troisième et dernier chapitre nous définissons une nouvelle famille de codes quantiques à partir de cubes de dimension arbitrairement grande. En prenant le quotient d'un cube de dimension n par un code classique de paramètres $[n, k, d]$ et en identifiant les qubits physiques avec les faces de dimension p du polytope quotient ainsi défini, on obtient un code quantique ayant les paramètres suivants: $N = \binom{n}{p} 2^{n-p-k}$, $K = \binom{p+k-1}{k-1}$ et $D = \min\left(\binom{d}{p}, 2^{n-p-k}\right)$. De cette famille multi-paramètre, il est possible d'extraire une sous-famille telle que K et D croissent comme une puissance de N . Cette famille de codes quantiques a l'originalité de considérer des quotients par des codes classiques. En cela elle s'éloigne de la topologie et appartient plutôt à la famille des codes homologiques. L'étude de la testabilité locale de cette famille de code est étudiée dans [LLZ19].

Chapter 1

Quantum computing concepts

In this chapter, quantum computing concepts are introduced. §1.1 to §1.7 strive to introduce technical terminology in an intuitive manner. §1.8 gives an original exposition of the orthogonal decomposition of $\mathbb{C}^{2^n}$ defined by a CSS code. It highlights the link between the homology of the discrete space and the decomposition of the complex physical space in subspaces indexed by syndrome values. §1.9 continues to highlight the relationship between the discrete space and the complex space of a CSS code by linking their duality theories.

1.1 Quantum memory: qubits

The elementary memory element of a quantum computer is a qubit, *i.e.* a two level quantum system. When observed (we will also use the terminology “measured”), qubits can only take two values, usually denoted 0 and 1 or rather $|0\rangle$ and $|1\rangle$. However the set of values a qubit can take before being observed is uncountably infinite: it can be any linear combination of these two values up to normalization and a global phase. Topologically it corresponds to the set of points of a 2 dimensional sphere. Algebraically it is the following set:

$$\{\alpha|0\rangle + \beta|1\rangle \mid (\alpha, \beta) \in \mathbb{C}^2 \setminus \{(0, 0)\}\} / \mathbb{C}$$

where the quotient by $\mathbb{C}$ corresponds to the following equivalence relation $\sim$:

$$\forall \lambda \in \mathbb{C} \setminus \{0\}, \alpha|0\rangle + \beta|1\rangle \sim \lambda\alpha|0\rangle + \lambda\beta|1\rangle.$$

This distinction between the set of observable values and the intrinsic set of values is characteristic of quantum mechanics. Indeed for a bit, the elementary memory element of a classical computer, both the set of observable values and the intrinsic set of values are $\{0, 1\}$. A qubit has the same set of observable values as a bit but a much larger intrinsic set of values.

A legitimate question concerns the definition of the intrinsic set of values. Indeed since we cannot observe such values, in which sense do they even exist? Before answering this question we want to describe how the state of a qubit relates to the observed values. The relationship is probabilistic: if a qubit in the state $\alpha|0\rangle + \beta|1\rangle$ is observed, *i.e.* measured, in the $(|0\rangle, |1\rangle)$ basis, the observed outcome is 0 with probability $\frac{|\alpha|^2}{|\alpha|^2 + |\beta|^2}$ and 1 with probability $\frac{|\beta|^2}{|\alpha|^2 + |\beta|^2}$. Moreover if 0 is observed, the state of the qubit after the measurement is $|0\rangle$ and if 1 is observed, the so-called post-measurement state is $|1\rangle$. We will see later on that quantum error correction makes crucial use of this side effect of a measurement.

If a qubit were merely a probabilistic bit, it would be possible to describe its intrinsic state with a single parameter $p \in [0, 1]$. The fact that the set of possible states of a qubit is a 2-dimensional sphere and not the $[0, 1]$ line segment gives a hint that a qubit is more than a probabilistic bit.

In general, an n -qubit state measured in the so-called canonical basis $(|x\rangle)_{x \in \{0,1\}^n}$ yields a probability distribution on 2^n values. Indeed any quantum state $|\psi\rangle$ on n qubits can be written as follows:

$$|\psi\rangle = \sum_{x \in \{0,1\}^n} \sqrt{p_x} e^{i\theta_x} |x\rangle$$

where $(p_x)_{x \in \{0,1\}^n}$ is a probability distribution over 2^n values. When measured in the canonical basis, this quantum state gives the output x with probability p_x . Moreover for every x such that $p_x \neq 0$, the quantum state is also described by a phase θ_x . Even though the probability distribution corresponding to a measurement in the canonical basis is independent of these phases θ_x , the probability distribution corresponding to first processing $|\psi\rangle$ and then measuring it depends on $(\theta_x)_{x \in \{0,1\}^n}$ in general.

Now that we know how quantum memory is defined mathematically, we can describe how quantum information, *i.e.* the intrinsic state of n qubits, is processed. We will now refer to this intrinsic state as the quantum state. We saw that an n -qubit quantum state measured in the canonical basis defines a probability distribution over the 2^n elements of the canonical basis. In this manuscript, we will not mention mixed states. Therefore the terminology “quantum state” refers in this manuscript to a pure state. Readers not familiar with pure and mixed states can simply refer to the aforementioned mathematical definition of a quantum state.

1.2 Quantum processing: Unitary matrices

Any invertible matrix $M \in \text{GL}_{2^n}(\mathbb{C})$ sends any legitimate (at least one amplitude is not zero) n -qubit state to another legitimate n -qubit state. Since qubit states are defined up to global multiplication by a non-zero complex scalar, this description is redundant. Indeed there exist distinct matrices of $\text{GL}_{2^n}(\mathbb{C})$ that have the same action on the set of quantum states. Without loss of generality, we can restrict attention to normalized qubit states, *i.e.* such that the sum of the squared modules of their amplitudes is 1. By doing so, we restrict quantum state transformations to unitary matrices. However this description is still redundant because a normalized qubit state is defined up to a global phase. Therefore the set of transformations of n -qubit states is the set of unitaries $\text{U}_{2^n}(\mathbb{C})$ up to a global phase.

This set is uncountably infinite. This is problematic in order to design a programmable quantum computer because at first sight each unitary matrix corresponds to some tailored hardware. The situation is actually more practical once we have realized that any unitary matrix can be approximated to arbitrary precision by a finite set of unitaries called a universal quantum gate set. A quantum processor is able to implement any unitary of this universal quantum gate set and thus to approximate any general unitary to arbitrary precision. In chapter 4.5 of [NC02], it is proven that the set $\{\text{CNOT}, \text{H}, \text{T}\}$ of three unitary matrices, or gates, is a universal gate set. Remarkably, CNOT is a two-qubit gate and H and T are one-qubit gates. When we say that a k -qubit gate G acts on an n -qubit state, we actually mean that the unitary matrix $G \otimes I_2 \otimes \dots \otimes I_2$ is applied to the n -qubit state. I_2 is the identity matrix on $\mathbb{C}^2$. For completeness we give the matrix expressions of the gate CNOT in the basis $(|00\rangle, |01\rangle, |10\rangle, |11\rangle)$, and the gates H and T in the basis $(|0\rangle, |1\rangle)$:

$$\text{CNOT} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}, \quad \text{H} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}, \quad \text{T} = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\pi/4} \end{pmatrix}.$$

The advantage of a quantum computer over a classical computer arguably doesn't lie in its use of quantum memory. Indeed we have seen above that an n -qubit state is nothing more than the possibility to sample once from a probability distribution over 2^n values. It rather stems from my point of view from the ability to process this memory efficiently with quantum gates: even a 1-qubit gate can modify the outcome probability of any of the 2^n possible outcomes of an n -qubit state. The tensor product structure of an n -qubit state ($\mathbb{C}^{2^n} = (\mathbb{C}^2)^{\otimes n}$) plays a central role once we have restricted attention to a universal gate set made of one-qubit and two-qubit gates.

1.3 Quantum measurements: Hermitian operators

In this section we only consider orthogonal measurements since they are sufficient to introduce quantum error correction.

An orthogonal measurement is defined by the decomposition of $\mathbb{C}^{2^n}$ into a direct sum of orthogonal subspaces. To each subspace corresponds an outcome of the measurement and after the measurement the quantum state is orthogonally projected onto the subspace corresponding to the outcome. Both aspects of the measurement (the output and the post-measurement projection) are used in quantum error correction.

It is possible to compactly encode the data defining a quantum measurement in a Hermitian operator acting on $\mathbb{C}^{2^n}$. Indeed Hermitian operators are diagonalizable in an orthonormal basis. Therefore the eigenspaces of a Hermitian operator yield an orthogonal decomposition of $\mathbb{C}^{2^n}$. It is useful to label the outcome of the measurement by the eigenvalue of the Hermitian operator. This way the post-measurement state following the outcome λ is the projection of the pre-measurement state onto the eigenspace E_λ .

1.4 Pauli operators

Pauli operators are central to most constructions of quantum error correcting codes. There are four one-qubit Pauli operators:

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

They act on the qubit $\alpha|0\rangle + \beta|1\rangle$ like they act on the vector $(\alpha, \beta)^T \in \mathbb{C}^2$.

n -qubit Pauli operators are defined as n -fold tensor products of the above operators. Note that Pauli operators are both unitary and Hermitian. Therefore they can define both a measurement and a gate.

1.5 Quantum Error Correcting Codes

Quantum error correction seems at first sight to be a daunting task because of the no-cloning theorem which colloquially states that quantum information cannot be copied. Therefore the redundant schemes used in classical error correction seem useless in a quantum setting.

The approach used in quantum error correction consists in spreading the quantum information that K qubits could hold on N physical qubits with $N > K$. This way the quantum information is not copied but merely distributed in a larger quantum system. If any quantum operation on any $\lfloor \frac{D-1}{2} \rfloor$ physical qubits does not alter the logical quantum information (the one which could be held by only K qubits), we say that the quantum error correcting code has minimum distance at least D . The minimum distance of the code is the maximal such D .

In this manuscript we will focus on CSS codes which we now define formally. A CSS code on N physical qubits is defined by two sets of measurements respectively called X -measurements and Z -measurements. Any of these measurements is defined by a list of qubits. The corresponding X -measurement (respectively Z -measurement) is the Hermitian operator defined as the tensor product of the Pauli matrix X (respectively Z) for the qubits that are in the list and the identity I_2 for the qubits that are not in the list.

Since $X^2 = Y^2 = Z^2 = I$, the eigenvalues of any n -qubit Pauli operator belong to $\{-1, 1\}$. Therefore the measurement defined by an n -qubit Pauli operator gives a binary output and projects the quantum state onto one of two orthogonal subspaces of $\mathbb{C}^{2^n}$. A quantum error correcting code is a set of compatible measurements. By compatible we mean that any subspace appearing in the orthogonal decomposition of the measurement A is a stable subspace for the operator corresponding to the measurement B. This way, measurements A and B together define a decomposition of $\mathbb{C}^{2^n}$ which is finer than both the decomposition defined by B and the decomposition defined by A. A set of compatible measurement $(A)_i$ can thus be understood as a single finer measurement called a joint measurement. Indeed it defines a single finer decomposition of $\mathbb{C}^{2^n}$ into orthogonal subspaces, *i.e.* a single measurement. Observe that a set of measurements defined by Hermitian operators defines a joint measurement if and only if the Hermitian operators pairwise commute.

In the case of CSS codes, the commutation of two measurement operators is straightforward to assess: two X -measurements always commute, two Z -measurements always commute and an X -measurement A and a Z -measurement B commute if and only if there is an even number of qubits on which A acts with an X and B acts with a Z . Indeed at the level of a single qubit the following anticommutation relation holds: $XZ = -ZX$. It is common to represent the set of X -measurements as a binary matrix H_X . H_X has n columns, where n is the number of physical qubits. Each row of H_X represents an X -measurement with the convention that the qubits on which the operator acts like the X Pauli matrix correspond to a 1 and the qubits on which the operator acts like the identity I_2 correspond to a 0. For instance the operator $I \otimes X \otimes X \otimes I$ acting on 4 qubits corresponds to the row 0110 of H_X . The matrix H_Z is defined similarly. Note that the commutation condition can be written compactly as:

$$H_X H_Z^T = 0.$$

Borrowing terminology from classical error correction, the matrices H_X and H_Z are sometimes referred to as parity-check matrices. The measurement corresponding to one of their rows can be called a check. And the list of all measurement outputs constitutes the syndrome. Note that the measurement outputs are labeled by eigenvalues of Pauli operators and therefore belong to $\{-1, 1\}$. A check is said to be not satisfied if its measurement

output is -1 and satisfied if its measurement output is 1 . To use the terminology of classical error-correction, measurement outputs are commonly relabeled under the transformation $x \mapsto \frac{x+1}{2}$ and thus belong to $\{1, 0\}$ (with this convention, 1 means that the check is not satisfied and 0 means that it is satisfied).

1.6 Decoding and logical errors

The purpose of error correcting codes is to bring back a lightly corrupted quantum state to its original value. Once some noise has corrupted the quantum state encoded in a quantum code, the joint measurement defined by the quantum code has two effects: it projects the quantum state to one of the orthogonal subspaces and it yields the syndrome which labels this subspace. We make the assumption that the original quantum state belongs to the subspace corresponding to the all zero syndrome, the one corresponding to every check being satisfied. This subspace is called the codespace. If the measurement syndrome is not zero, we know that we are out of the codespace and look for a transformation that brings the quantum state back to the codespace. Let us introduce some notations. $|\psi\rangle$ is the original quantum state. $|\psi\rangle_{noisy}$ is the pre-measurement corrupted quantum state. $|\psi\rangle_{projected}$ is the post-measurement quantum state. We know that $|\psi\rangle$ belongs to the codespace and that $|\psi\rangle_{projected}$ belongs to one of the orthogonal subspace defining the quantum code. Let C be the Hermitian operator corresponding to a check. We have:

$$\begin{aligned} C|\psi\rangle &= |\psi\rangle \\ C|\psi\rangle_{projected} &= \pm |\psi\rangle_{projected}. \end{aligned}$$

We are looking for a unitary matrix R called a recovery operation such that $R|\psi\rangle_{projected}$ belongs to the codespace. Let us denote by $\mathcal{C}_-$ the set of unsatisfied checks and by $\mathcal{C}_+$ the set of satisfied checks. It is straightforward to verify that $R|\psi\rangle_{projected}$ belongs to the codespace if and only if the following commutation relations hold:

$$\begin{aligned} \forall C \in \mathcal{C}_-, \quad CR &= -RC \\ \forall C \in \mathcal{C}_+, \quad CR &= RC \end{aligned}$$

For a CSS code it is always possible to find such a recovery operation in the form of a Pauli operator. Note that here we use a Pauli operator as a gate and not as a measurement like previously. Finding the most probable (or more generally a probable) recovery Pauli operator for a given syndrome is called decoding the quantum code.

We now make more precise the term “most probable recovery Pauli operator”. Without loss of generality, we can write the post-measurement state $|\psi\rangle_{projected} = N_{Pauli} N_{codespace} |\psi\rangle$ where $N_{codespace}$ is a unitary that globally stabilizes the codespace and N_{Pauli} is a Pauli operator that sends the codespace to the subspace corresponding to the observed syndrome. The noise $N_{codespace}$ cannot be corrected since it cannot even be detected by the checks. The noise N_{Pauli} however is detected by the syndrome and we try to invert it with the recovery operation. Under some locality assumptions about the quantum noise, we can assume that $N_{codespace}$ is neglectable and that N_{Pauli} has weight w with probability p^w for some $p \in]0, 1[$. The weight of an n -qubit Pauli operator is the number of non identity terms in the n -fold tensor product. For instance $N_{Pauli} = I \otimes I \otimes X \otimes Z \otimes I \otimes Y \otimes I$ has weight 3. With a CSS code, it is common practice to decompose N_{Pauli} into the product of an X term and a Z term and to make the slightly different assumption that each of this term has weight w_X (respectively w_Z) with probability p^{w_X} (respectively p^{w_Z}) for some $p \in]0, 1[$. This slightly different assumption has no physical justification, it is done for simplicity of the analysis. Keeping the above example, we have $N_{Pauli} = iN_X N_Z$ with $N_X = I \otimes I \otimes X \otimes I \otimes I \otimes X \otimes I$ and $N_Z = I \otimes I \otimes I \otimes Z \otimes I \otimes Z \otimes I$. We used the fact that

$Y = iXZ$. The Z -checks give information about N_X and the X -checks give information about N_Z .

It is easier to use classical error correction notations at this point. Let n_X be the binary vector corresponding to N_X : n_X has length n , the number of physical qubits, a 1 at qubit indices where N_X has an X and a 0 at qubit indices where N_X has an I . The syndrome s_Z corresponding to the Z -checks is:

$$s_Z = H_Z n_X.$$

Similarly, n_Z being the binary vector corresponding to N_Z , the syndrome s_X corresponding to the X -checks is:

$$s_X = H_X n_Z.$$

Finding a probable recovery Pauli operator boils down to two classical decoding problems: we know s_Z and H_Z and try to find the smallest weight e_X satisfying

$$s_Z = H_Z e_X. \quad (1.1)$$

Similarly we know s_X and H_X and try to find the smallest weight e_Z satisfying

$$s_X = H_X e_Z. \quad (1.2)$$

However there is a little more flexibility in the quantum decoding problem than in the classical decoding problem. Assume we have found e_X satisfying 1.1. It corresponds to the X part of the recovery operation

$$R_X = \bigotimes_{i=1}^n X^{(e_X)_i}.$$

Similarly e_Z satisfies eq. 1.2 and corresponds to

$$R_Z = \bigotimes_{i=1}^n Z^{(e_Z)_i}.$$

We know that $|\psi\rangle_{\text{corrected}} = R_X R_Z |\psi\rangle_{\text{projected}}$ belongs to the codespace. Let e_C be the binary vector corresponding to the Z -check C . Note that the operator $\bigotimes_{i=1}^n X^{(e_C)_i}$ is both the Hermitian operator defining the check C and a unitary defining a valid recovery operator. Since the codespace is point-wise fixed by any Hermitian operator corresponding to a check, $\bigotimes_{i=1}^n X^{(e_C)_i} |\psi\rangle_{\text{corrected}} = |\psi\rangle_{\text{corrected}}$. In other words,

$$\bigotimes_{i=1}^n X^{(e_C)_i} R_X R_Z |\psi\rangle_{\text{projected}} = R_X R_Z |\psi\rangle_{\text{projected}}.$$

Therefore the recovery operation

$$\tilde{R}_X = \bigotimes_{i=1}^n X^{(e_C)_i} R_X = \bigotimes_{i=1}^n X^{(e_X \oplus e_C)_i}$$

yields the same quantum state as $R_X = \bigotimes_{i=1}^n X^{(e_X)_i}$. Since this is true for any Z -check C , the quantum decoding problem consists in finding e_X satisfying eq. 1.1 and such that $n_X \oplus e_X \in \text{Im}(H_Z^T)$. This gives more flexibility than imposing $e_X = n_X$ like in a classical decoding problem. This extra flexibility makes the quantum decoding problem somewhat different from its classical counterpart. Of course we are also looking for a solution e_Z of

eq. 1.2 equal to n_Z up to an element of $\text{Im}(H_X^T)$.

Because of this peculiarity, the quantum decoding problem has its own terminology which we now introduce. Given a syndrome s_Z , a parity-check matrix H_Z and a so-called generator matrix H_X (H_Z and H_X are alternatively called check matrices and generator matrices depending on the type of errors we focus on), any binary vector e_X satisfying $s_Z = H_Z e_X$ is called a valid error. Any two valid errors e_X and $e_{\tilde{X}}$ satisfying $e_X \oplus e_{\tilde{X}} \in \text{Im}(H_X^T)$ are called equivalent errors. Any valid error equivalent to n_X is called a successful decoding error. Any valid error not equivalent to n_X is called a logical error.

In other words, e_X is a valid error if and only if $e_X \oplus n_X \in \text{Ker}(H_Z)$. If $e_X \oplus n_X \in \text{Im}(H_X^T)$, e_X is a successful decoding error. Otherwise, *i.e.* if it is valid but not successful, it is a logical error. Note that the CSS condition $H_Z H_X^T = 0$ can be restated as $\text{Im}(H_X^T) \subset \text{Ker}(H_Z)$. Therefore the above terminology is consistent (any successful decoding error is a valid error).

We now characterize the minimum distance D of a CSS quantum error correcting code in terms of the above terminology:

$$\begin{aligned} D_X &= \min_{l_X \in \text{Ker}(H_Z) \setminus \text{Im}(H_X^T)} w(l_X) \\ D_Z &= \min_{l_Z \in \text{Ker}(H_X) \setminus \text{Im}(H_Z^T)} w(l_Z) \\ D &= \min(D_X, D_Z) \end{aligned}$$

where $w(l_X)$, respectively $w(l_Z)$, denote the weight of the logical binary vector l_X , respectively l_Z .

To summarize, a decoding algorithm for a CSS quantum code takes as inputs a syndrome s_Z , the parity-check matrix H_Z and the generator matrix H_X . It outputs a small weight error e_X such that $s_Z = H_Z e_X$. It optionally takes into account the fact that e_X must have a small weight only up to an element of $\text{Im}(H_X^T)$. If the decoding algorithm does not manage to output such an error e_X , we say that the decoding did not terminate. If the decoding terminates, we need to know the noise n_X from which s_Z was computed to determine whether the decoding was successful. If $e_X \oplus n_X \in \text{Im}(H_X^T)$, the decoding was successful. Otherwise the decoding suffered from a logical error. The same work must but done with the Z -noise binary vector n_Z .

1.7 Homological quantum codes

An elegant way to define a CSS quantum code is to define the matrices H_X and H_Z as the incidence matrices of a polytope. More precisely an n -dimensional polytope is defined by $(n+1)$ sets of i -faces for $i \in \{0, \dots, n\}$ and incidences for any pair made of an i -face and a j -face for $i \neq j$. A homological quantum code can be defined by choosing $i \in \{1, \dots, n-1\}$, identifying qubits with i -faces, X -checks with $(i-1)$ -faces and Z -checks with $(i+1)$ -faces. H_X is given by the $(i-1)$ -faces - i -faces incidence matrix and H_Z by the $(i+1)$ -faces - i -faces incidence matrix. For a non degenerate polytope, the essential CSS property $H_X H_Z^T = 0$ is automatically satisfied.

Moreover if the polytope is defined from a tessellation of a closed manifold, there are some deep connections between the $[[N, K, D]]$ parameters of the homological quantum code and the characteristics of the manifold: N is proportional to the volume of the manifold, K is equal to the rank of the i^{th} -homology group of the manifold and D is proportional to the i^{th} homological systole of the manifold. Rigorous definitions and statements will be

given in Chapter 2.

In Chapter 2 we construct such a homological quantum code family from a 4-dimensional (4D) hyperbolic manifold. In Chapter 3 we give small instances of 4D hyperbolic codes and study their decoding under a Belief Propagation algorithm. In Chapter 4 we construct a homological quantum code family from quotients of the n -dimensional hypercube by classical error correcting codes.

1.8 The orthogonal decomposition of $\mathbb{C}^{2^n}$ defined by a CSS code

We saw in §1.5 that a CSS code on n physical qubits is defined by two orthogonal to each other binary parity check matrices H_X of size $l_X \times n$ and H_Z of size $l_Z \times n$. The complex vector space $(\mathbb{C}^2)^{\otimes n} \simeq \mathbb{C}^{(2^n)}$ associated to the n physical qubits is called the physical space. Let r_X , respectively r_Z , be the rank of H_X (respectively H_Z). A subspace of the physical space defined by H_X and H_Z and of dimension 2^k where $k = n - r_X - r_Z$ is called the quantum codespace $\mathcal{C}_Q$. We say that the quantum code defined by (H_X, H_Z) has 2^k logical qubits. We will see in this section how the quantum codespace $\mathcal{C}_Q$ is defined by H_X and H_Z and why it has dimension 2^k . Instead of following the general approach for stabilizer codes exposed in [NC02], we will emphasise the homological viewpoint on CSS codes and give mathematical details.

The orthogonality of the rows of H_X with the rows of H_Z can be expressed equivalently in different languages. We give three such languages below:

In parity-check matrices language:

$$\mathbb{F}_2^{l_X} \xleftarrow{H_X} \mathbb{F}_2^n \xrightarrow{H_Z} \mathbb{F}_2^{l_Z}$$

with $H_Z H_X^T = 0$.

H_Z^T denotes the transpose of H_Z . $\mathbb{F}_2$ is the finite field with 2 elements.

In chain complex language:

$$\mathbb{F}_2^{n_{p-1}} \xleftarrow{\partial_p} \mathbb{F}_2^{n_p} \xleftarrow{\partial_{p+1}} \mathbb{F}_2^{n_{p+1}}$$

with $\partial_{p+1} \circ \partial_p = 0$.

In cochain complex language:

$$\mathbb{F}_2^{n_{p-1}} \xrightarrow{\delta_p} \mathbb{F}_2^{n_p} \xrightarrow{\delta_{p+1}} \mathbb{F}_2^{n_{p+1}}$$

with $\delta_{p+1} \circ \delta_p = 0$.

We will use these three languages interchangeably. It is easy to translate from one to another using the following identifications:

$$\begin{aligned}
H_X &= \partial_p = \delta_p^T. \\
H_Z &= \partial_{p+1}^T = \delta_{p+1}. \\
l_X &= n_{p-1}. \\
n &= n_p. \\
l_Z &= n_{p+1}.
\end{aligned}$$

Let us now explain how we can use the two orthogonal to each other parity-check matrices H_X and H_Z to construct a 2^k dimensional subspace of $(\mathbb{C}^2)^{\otimes n}$. First we will define the representations ρ_X , respectively ρ_Z , from $(\mathbb{F}_2)^n$ to X-Pauli operators (respectively Z-Pauli operators).

$$\begin{aligned}
\rho_X : (\mathbb{F}_2)^n &\rightarrow \text{GL}_{2^n}(\mathbb{C}) \\
x_1 \dots x_n &\mapsto \bigotimes_{i=1}^n X^{x_i}.
\end{aligned}$$

$$\begin{aligned}
\rho_Z : (\mathbb{F}_2)^n &\rightarrow \text{GL}_{2^n}(\mathbb{C}) \\
z_1 \dots z_n &\mapsto \bigotimes_{i=1}^n Z^{z_i}.
\end{aligned}$$

Recall that $X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ and $Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$.

Observe that $XZ = -ZX$. Therefore $\rho_X(x)\rho_Z(z) = (-1)^{\langle z, x \rangle} \rho_Z(z)\rho_X(x)$. $\langle z, x \rangle$ is the canonical scalar product ¹ on $(\mathbb{F}_2)^n$:

$$\begin{aligned}
\langle \cdot, \cdot \rangle : (\mathbb{F}_2)^n \times (\mathbb{F}_2)^n &\rightarrow \mathbb{F}_2 \\
(z, x) &\mapsto \sum_{i=1}^n x_i z_i.
\end{aligned}$$

We consider four subspaces of $(\mathbb{F}_2)^n$ defined by H_X and H_Z :

- $B_p = \text{Im}(H_Z^T)$ is the space of boundaries.
- $B^p = \text{Im}(H_X^T)$ is the space of coboundaries.
- $Z_p = \ker(H_X)$ is the space of cycles.
- $Z^p = \ker(H_Z)$ is the space of cocycles.

Observe that these four spaces satisfy the following relations:

- $(B_p)^\perp = Z^p$.
- $(B^p)^\perp = Z_p$.
- $B_p \subset Z_p$.
- $B^p \subset Z^p$.

¹This is the scalar product for which δ_p is the adjoint of ∂_p .

These relations imply that if z is a boundary and x is a coboundary, $\langle z, x \rangle = 0$ and therefore $\rho_X(x)\rho_Z(z) = \rho_Z(z)\rho_X(x)$.

The codespace $\mathcal{C}_Q$ is defined as the common 1-eigenspace of the commuting family of $\mathbb{C}$ -endomorphisms $\{\rho_Z(bd)\rho_X(cobd) \mid bd \in B_p, cobd \in B^p\}$.

To prove that $\dim(\mathcal{C}_Q) = 2^k$, we will use the following Lemma:

Lemma 1. *Let F be a subspace of $(\mathbb{F}_2)^n$ such that $\dim(F) = r$. The linear application f_F defined as*

$$\begin{aligned} f_F : (\mathbb{F}_2)^n &\rightarrow (\mathbb{F}_2)^{2^r} \\ z &\mapsto (\langle v, z \rangle \mid v \in F). \end{aligned}$$

satisfies $\text{rk}(f_F) = \dim(F) = r$.

Proof. $\ker(f_F) = F^\perp$.

Therefore $\text{rk}(f_F) = n - \dim(F^\perp) = \dim(F)$. $\square$

Corollary 2. *Let F, G be subspaces of $(\mathbb{F}_2)^n$ such that $\dim(F) = r$ and $\dim(G) = s$. The linear application $f_{(F,G)}$ defined as*

$$\begin{aligned} f_{(F,G)} : (\mathbb{F}_2)^n \times (\mathbb{F}_2)^n &\rightarrow (\mathbb{F}_2)^{2^{r+s}} \\ (z, x) &\mapsto (\langle v, z \rangle + \langle w, x \rangle \mid (v, w) \in F \times G). \end{aligned}$$

satisfies $\text{rk}(f_{(F,G)}) = \dim(F) + \dim(G) = r + s$.

Proof. $F \times G$ is a subspace of $(\mathbb{F}_2)^n \times (\mathbb{F}_2)^n$ of dimension $\dim(F) + \dim(G)$ and the standard inner product on $(\mathbb{F}_2)^n \times (\mathbb{F}_2)^n$ is $((v, w), (z, x)) \mapsto \langle v, z \rangle + \langle w, x \rangle$. $\square$

We consider the following set of commuting Pauli operators:

$$S_{(B^p, Z_p)} = \{\rho_X(cobd)\rho_Z(cyc) \mid cobd \in B^p, cyc \in Z_p\}.$$

Since B^p and Z_p are orthogonal, any two operators of $S_{(B^p, Z_p)}$ commute. As Hermitian commuting operators, they are simultaneously diagonalisable. We define the following vector of $\mathbb{C}^{2^n}$:

$$|\psi\rangle = \sum_{cobd \in B^p} \rho_X(cobd) |0\dots 0\rangle.$$

It is straightforward to verify that $|\psi\rangle$ is a joint eigenvector of $S_{(B^p, Z_p)}$ with joint eigenvalue $(1, \dots, 1)$.

Observe that $\rho_Z(z)\rho_X(x)|\psi\rangle$ is a joint eigenvector of $S_{(B^p, Z_p)}$ with joint eigenvalue $\{(-1)^{(f_{(B^p, Z_p)}(z,x))_i} \mid i \in \{1, \dots, 2^{r+s}\}\}$. Indeed,

$$\begin{aligned} \forall v \in B^p, w \in Z_p, \rho_X(v)\rho_Z(w)(\rho_Z(z)\rho_X(x)|\psi\rangle) \\ = (-1)^{\langle v, z \rangle + \langle w, x \rangle} \rho_Z(z)\rho_X(x)\rho_X(v)\rho_Z(w)|\psi\rangle \\ = (-1)^{\langle v, z \rangle + \langle w, x \rangle} \rho_Z(z)\rho_X(x)|\psi\rangle. \end{aligned}$$

Since $\text{rk}(f_{(B^p, Z_p)}) = r_X + (n - r_X) = n$, $|\text{Im}(f_{(B^p, Z_p)}(z, x))| = 2^n = \dim(\mathbb{C}^{2^n})$. Therefore $S_{(B^p, Z_p)}$ has a joint basis of diagonalisation made of vectors in $\{\rho_Z(z)\rho_X(x)|\psi\rangle \mid z, x \in (\mathbb{F}_2)^n\}$ and every joint eigenspace has dimension 1.

Let us now consider the subset of $S_{(B^p, Z_p)}$ defining the codespace $\mathcal{C}_Q$:

$$S_{(B^p, B_p)} = \{\rho_X(cobd)\rho_Z(bd) \mid cobd \in B^p, bd \in B_p\}.$$

$\rho_Z(z)\rho_X(x)|\psi\rangle$ is a joint eigenvector of $S_{(B^p, B_p)}$ with joint eigenvalue $f_{(B^p, B_p)}(z, x)$. Since $(B^p \times B_p) \subset (B^p \times Z_p)$, we can define a map from $\text{Im}(f_{(B^p, Z_p)})$ to $\text{Im}(f_{(B^p, B_p)})$ by restriction. This map is of course surjective. To prove that each element of $\text{Im}(f_{(B^p, B_p)})$ has the same number of preimages under this restriction map, observe that $f_{(B^p, B_p)}$ and $f_{(B^p, Z_p)}$ are linear applications with respective kernels:

$$\ker f_{(B^p, B_p)} = Z_p \times Z^p,$$

$$\ker f_{(B^p, Z_p)} = Z_p \times B^p.$$

Therefore each element of $\text{Im}(f_{(B^p, B_p)})$ is attained by $|\ker f_{(B^p, B_p)}|/|\ker f_{(B^p, Z_p)}| = |Z^p|/|B^p| = 2^k$ elements of $\text{Im}(f_{(B^p, Z_p)})$ by the restriction map.

This proves that every joint eigenspace of $S_{(B^p, B_p)}$ is the direct sum of 2^k joint eigenspaces of $S_{(B^p, Z_p)}$. Since every joint eigenspace of $S_{(B^p, Z_p)}$ has dimension 1, every joint eigenspace of $S_{(B^p, B_p)}$ has dimension 2^k . In particular the codespace $\mathcal{C}_Q$, which is by definition the joint eigenspace of $S_{(B^p, B_p)}$ with joint eigenvalue $(1, \dots, 1)$, has dimension 2^k .

This exposition of the orthogonal decomposition defined by a CSS code is arguably more intuitive than the general exposition given in [NC02] for stabilizer codes. By considering the joint eigenspaces of every operator of $S_{(B^p, Z_p)}$ and of every operator of $S_{(B^p, B_p)}$, we don't have to choose arbitrary preferred representatives. Thus the relationship between the discrete spaces B^p , B_p , Z^p and Z_p and orthogonal decompositions of $\mathbb{C}^{2^n}$ is, in our opinion, clearer. In the next section we will dive deeper in the relationship between $\mathbb{F}^n$ and $\mathbb{C}^{2^n}$ by linking their respective dualities.

1.9 $\mathbb{F}^n$ and $\mathbb{C}^{2^n}$ dualities

Given a vector space E , we denote by E^* the space of linear forms on E .

On $(\mathbb{F}_2)^n$, linear forms and vectors can be identified by the canonical binary scalar product: $\langle x, y \rangle_{\mathbb{F}_2} = \sum_{i=1}^n x_i y_i \in \mathbb{F}_2$. It yields the following vector space isomorphism:

$$\begin{aligned} \varphi : (\mathbb{F}_2)^n &\rightarrow ((\mathbb{F}_2)^n)^* \\ v &\mapsto (x \mapsto \langle v, x \rangle_{\mathbb{F}_2}) \end{aligned} \quad .$$

In §1.8 we took advantage of the above identification to represent cochains as elements of $(\mathbb{F}_2)^n$. However strictly speaking cochains are elements of $((\mathbb{F}_2)^n)^*$. This is how we consider them in this section.

Moreover we consider in this section two versions of the physical space $\mathbb{C}^{2^n}$: $\mathbb{C}[(\mathbb{F}_2)^n]$ on which chains will act and $\mathbb{C}[(\mathbb{F}_2)^n]^*$ on which cochains will act.

Recall that given a finite group G , its group algebra $\mathbb{C}[G]$ is the complex vector space with a basis indexed by the elements of G . We denote by $\{b_v \mid v \in (\mathbb{F}_2)^n\}$ the canonical basis of $\mathbb{C}[(\mathbb{F}_2)^n]$ and by $\{b_{\varphi(w)} \mid \varphi(w) \in ((\mathbb{F}_2)^n)^*\}$ the canonical basis of $\mathbb{C}[(\mathbb{F}_2)^n]^*$. In the field of quantum information, b_x is usually written $|x\rangle$ but we will rather use the notation b_x here.

We define a Hermitian bilinear form on $\mathbb{C}[(\mathbb{F}_2)^n]^* \times \mathbb{C}[(\mathbb{F}_2)^n]$ by its value on basis elements (extended by bilinearity):

$$\begin{aligned}
h : \mathbb{C}[(\mathbb{F}_2)^n]^* \times \mathbb{C}[(\mathbb{F}_2)^n] &\rightarrow \mathbb{C} \\
(b_{\varphi(x)}, b_y) &\mapsto (-1)^{\varphi(x)[y]} \\
&= (-1)^{\langle x, y \rangle_{\mathbb{F}_2}} \quad .
\end{aligned}$$

The matrix of h in the two canonical bases of $\mathbb{C}[(\mathbb{F}_2)^n]^*$ and $\mathbb{C}[(\mathbb{F}_2)^n]$ is $H^{\otimes n}$ where $H = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$ is the Hadamard matrix up to a normalization factor. Therefore h is a Hermitian nondegenerate form.

With this Hermitian nondegenerate form, we can define an isomorphism $\tilde{h}$ between $\mathbb{C}[(\mathbb{F}_2)^n]^*$ and $(\mathbb{C}[(\mathbb{F}_2)^n])^*$ in the same way we used $\langle _, _ \rangle_{\mathbb{F}_2}$ to define an isomorphism between $(\mathbb{F}_2)^n$ and $((\mathbb{F}_2)^n)^*$:

$$\begin{aligned}
\tilde{h} : \mathbb{C}[(\mathbb{F}_2)^n]^* &\rightarrow (\mathbb{C}[(\mathbb{F}_2)^n])^* \\
b_{\varphi(x)} &\mapsto (b_y \mapsto h(b_{\varphi(x)}, b_y)) \quad .
\end{aligned}$$

Moreover we can use the canonical (positive-definite hence nondegenerate) Hermitian form on $\mathbb{C}[(\mathbb{F}_2)^n] \times \mathbb{C}[(\mathbb{F}_2)^n]$:

$$\begin{aligned}
\beta : \mathbb{C}[(\mathbb{F}_2)^n] \times \mathbb{C}[(\mathbb{F}_2)^n] &\rightarrow \mathbb{C} \\
(b_x, b_y) &\mapsto \delta_{x,y}
\end{aligned}$$

to define in the same way an isomorphism $\tilde{\beta}$ between $\mathbb{C}[(\mathbb{F}_2)^n]$ and $(\mathbb{C}[(\mathbb{F}_2)^n])^*$:

$$\begin{aligned}
\tilde{\beta} : \mathbb{C}[(\mathbb{F}_2)^n] &\rightarrow (\mathbb{C}[(\mathbb{F}_2)^n])^* \\
b_x &\mapsto (b_y \mapsto \beta(b_x, b_y)) \quad .
\end{aligned}$$

Finally, since $\mathbb{C}[G]$ is also the space of functions from G to $\mathbb{C}$, the following commutative diagram defines the isomorphism $F(\varphi)$ between $\mathbb{C}[(\mathbb{F}_2)^n]^*$ and $\mathbb{C}[(\mathbb{F}_2)^n]$:

$$\begin{array}{ccc}
(\mathbb{F}_2)^n & \xrightarrow{\varphi} & ((\mathbb{F}_2)^n)^* \\
& \searrow F(\varphi)[f] & \downarrow f \\
& & \mathbb{C} \\
& \text{= } f \circ \varphi &
\end{array}$$

The following diagram of vector space isomorphisms summarises the situation:

$$\begin{array}{ccc}
\mathbb{C}[(\mathbb{F}_2)^n] & \xrightarrow{\tilde{\beta}} & (\mathbb{C}[(\mathbb{F}_2)^n])^* \\
\uparrow F(\varphi) & \nearrow \tilde{h} & \\
\mathbb{C}[(\mathbb{F}_2)^n]^* & &
\end{array}$$

This diagram is not commutative as the following proposition shows:

Propositon 3. *Going around the above diagram once counterclockwise amounts to performing the Hadamard transform $H^{\otimes n}$.*

Proof. We first apply $F(\varphi)^{-1}$:

$$\forall x \in (\mathbb{F}_2)^n, \quad (F(\varphi))^{-1}(b_x) = b_{\varphi(x)} \quad .$$

We then apply $\tilde{h}$:

$$\begin{aligned} \forall x \in (\mathbb{F}_2)^n, \forall y \in (\mathbb{F}_2)^n, \quad \tilde{h}((F(\varphi))^{-1}(b_x)) [b_y] &= h(b_{\varphi(x)}, b_y) \\ &= (-1)^{\langle x, y \rangle_{\mathbb{F}_2}} \\ &= \sum_{z \in (\mathbb{F}_2)^n} (-1)^{\langle x, z \rangle_{\mathbb{F}_2}} \delta_{y, z} \\ &= \sum_{z \in (\mathbb{F}_2)^n} (-1)^{\langle x, z \rangle_{\mathbb{F}_2}} \tilde{\beta}(b_z) [b_y] \quad . \end{aligned}$$

Therefore,

$$\forall x \in (\mathbb{F}_2)^n, \quad \tilde{h}((F(\varphi))^{-1}(b_x)) = \sum_{z \in (\mathbb{F}_2)^n} (-1)^{\langle x, z \rangle_{\mathbb{F}_2}} \tilde{\beta}(b_z) \quad .$$

We finally apply $\tilde{\beta}^{-1}$:

$$\forall x \in (\mathbb{F}_2)^n, \quad \tilde{\beta}^{-1}(\tilde{h}((F(\varphi))^{-1}(b_x))) = \sum_{z \in (\mathbb{F}_2)^n} (-1)^{\langle x, z \rangle_{\mathbb{F}_2}} b_z \quad .$$

□

To show a symmetric behaviour between chains and cochains, we choose to identify $\mathbb{C}[(\mathbb{F}_2)^n]^*$ and $\mathbb{C}[(\mathbb{F}_2)^n]$ through $\tilde{\beta}$ and $\tilde{h}$:

$$T \stackrel{\text{def}}{=} \tilde{\beta}^{-1} \circ \tilde{h} : \mathbb{C}[(\mathbb{F}_2)^n]^* \rightarrow \mathbb{C}[(\mathbb{F}_2)^n].$$

In §1.8 we defined the action ρ_Z of chains on $\mathbb{C}[(\mathbb{F}_2)^n]$ as Z operators. It means that a chain z acts on an element b_v of the basis of $\mathbb{C}[(\mathbb{F}_2)^n]$ as follows:

$$z \cdot b_v = \rho_Z(z)[b_v] = (-1)^{\langle z, v \rangle_{\mathbb{F}_2}} b_v.$$

Propositon 4. *If we define the action ρ_Z^* of cochains on $\mathbb{C}[(\mathbb{F}_2)^n]^*$ as Z operators, it translates under the identification T of $\mathbb{C}[(\mathbb{F}_2)^n]^*$ and $\mathbb{C}[(\mathbb{F}_2)^n]$ into an action ρ_X of cochains on $\mathbb{C}[(\mathbb{F}_2)^n]$ as X operators:*

$$\forall x, w \in (\mathbb{F}_2)^n, \quad T(\rho_Z^*(x)[b_{\varphi(w)}]) = \rho_X(x)[T(b_{\varphi(w)})].$$

In other words, T intertwines the representations ρ_Z^ and ρ_X .*

Proof. Prop. 3 shows that $T = H^{\otimes n} \circ (F(\varphi)^{-1})$.

We first compute the right hand side:

$$\begin{aligned} \forall w \in (\mathbb{F}_2)^n, \quad T(b_{\varphi(w)}) &= H^{\otimes n}(b_w) \\ &= \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle w, a \rangle_{\mathbb{F}_2}} b_a. \\ \forall x, w \in (\mathbb{F}_2)^n, \quad \rho_X(x)[T(b_{\varphi(w)})] &= \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle w, a \rangle_{\mathbb{F}_2}} b_{x+a}. \end{aligned}$$

We now compute the left hand side:

$$\begin{aligned}
\forall x, w \in (\mathbb{F}_2)^n, \quad \rho_Z^*(x)b_{\varphi(w)} &= (-1)^{\langle x, w \rangle_{\mathbb{F}_2}} b_{\varphi(w)}. \\
\forall x, w \in (\mathbb{F}_2)^n, \quad T(\rho_Z^*(x)b_{\varphi(w)}) &= (-1)^{\langle x, w \rangle_{\mathbb{F}_2}} H^{\otimes n}(b_w) \\
&= (-1)^{\langle x, w \rangle_{\mathbb{F}_2}} \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle w, a \rangle_{\mathbb{F}_2}} b_a \\
&= \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle x, w \rangle_{\mathbb{F}_2}} (-1)^{\langle w, a \rangle_{\mathbb{F}_2}} b_a \\
&= \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle w, x+a \rangle_{\mathbb{F}_2}} b_a \\
&= \sum_{a \in (\mathbb{F}_2)^n} (-1)^{\langle w, a \rangle_{\mathbb{F}_2}} b_{x+a}.
\end{aligned}$$

□

Therefore acting as a Z operator on $\mathbb{C}[(\mathbb{F}_2)^n]^*$ is the same under the identification given by the isomorphism T as acting as an X operator on $\mathbb{C}[(\mathbb{F}_2)^n]$. With this viewpoint, the symmetry between chains and cochains is highlighted since they both act as Z operators on distinct but identified complex vector spaces. Moreover the duality between chains and cochains is reflected in the duality between the complex vector spaces they act on.

Note that it is unfortunate that the established practise makes cochains act as X operators and chains act as Z operators. Imagine we use the converse definition: chains act as X and cochains as Z . Then in the reformulation of this paragraph chains would act as X on $\mathbb{C}[(\mathbb{F}_2)^n]$ and cochains as X on $\mathbb{C}[(\mathbb{F}_2)^n]^*$. Since acting as X is the natural action of a group on its group algebra by translation of the basis elements:

$$\begin{aligned}
(\mathbb{F}_2)^n &\hookrightarrow \mathbb{C}[(\mathbb{F}_2)^n] \rightarrow \mathbb{C}[(\mathbb{F}_2)^n] \\
x \cdot b_y &\mapsto b_{x+y} = \rho_X(x)[b_y]
\end{aligned}$$

the construction would have been even more natural with this converse definition.

Chapter 2

Golden codes

In this chapter we expose a homological code construction which stems from the tessellation of a 4 dimensional hyperbolic manifold. It is based on this article: [LL19].

A major advantage of homological codes with a fixed and compact local structure is that they are naturally of the low-density parity-check (LDPC) type, meaning that generators of the stabilizer group act nontrivially on a constant number of qubits and that each qubit is acted upon by a constant number of generators. This is of course especially interesting for potential implementations, but also at a more mathematical level since classical LDPC codes play a central role in classical coding theory. A second advantage of homological codes is that they can lead to simple and efficient decoding algorithms which directly exploit the local structure of the code on the manifold [DKLP02, Har04, DZ17, DN17].

The parameters $[[N, K, D]]$ of homological codes can be derived from the properties of the underlying manifold: the number N of physical qubits of the code is given by the number of i -faces in the tessellation, the number K of logical qubits is given by the rank of the i^{th} homology group, and the minimum distance, that is the minimum weight of a nontrivial Pauli error, is related to the i^{th} homological systole of the manifold, that is the minimal number of i -faces forming a homologically nontrivial i -cycle. Exploiting this connection with manifolds exhibiting systolic freedom, Freedman, Meyer and Luo [FML02] were able to construct the quantum LDPC codes with the best minimum distance presently known, achieving $d = \Theta(n^{1/2} \log^{1/4} n)$ ¹.

An important question is to understand what parameters $[[N, K, D]]$ can be achieved with quantum LDPC codes. The toric code and the code of Ref. [FML02] display a large minimum distance but only encode a constant number of qubits, $k = O(1)$. If the manifold is Euclidean, strong constraints apply on the code parameters: namely the parameters have to satisfy $KD^2 \leq cN$ for some constant c [BPT10]. For tessellations of 2-dimensional manifolds, Delfosse showed that $KD^2 \leq c(\log K)^2 N$ [Del13]. In particular, these results show that one cannot get a good minimum distance for codes with constant rate built from 2-dimensional manifolds.

In many cases, it is natural to consider constant-rate codes where $K = \Theta(N)$: such codes for instance allow one to obtain quantum fault-tolerant computation with constant space overhead [Got13, FGL18a]. For a long time, it was believed that constant-rate homological codes could not have a large minimum distance, that is growing polynomially with their length [Zém09]. A recent breakthrough was the work of Guth and Lubotzky [GL14] who

¹Ref. [FML02] mentions $d = \Theta(n^{1/2} \log^{1/2} n)$ but it seems that there is a typo in their numerical application.

gave a construction of homological codes in hyperbolic 4-space that combine a constant rate with a minimum distance $D = N^\alpha$, for some constant $\alpha > 0$. It was later shown by Murillo that the construction could be adapted to yield $\alpha \in [0.2, 0.3]$. Quickly after this result, Hastings proposed a decoding algorithm for such codes [Has16]. Unfortunately, the analysis of Hastings' decoding algorithm is only valid when local brute-force decoding is performed at a scale that may not be computationally practical. In fact, it is difficult to precisely analyze the performance of Hastings' decoder because the local structure of the codes of [GL14] is not completely explicit.

In this work, we give a variant of the construction of Guth and Lubotzky which admits a simple explicit local structure: it is based on a regular tessellation of the 4-dimensional hyperbolic space by means of hypercubes. We then exploit this local structure to design an efficient decoding algorithm which tries to locally shorten cycles. In Section 2.1, we give an overview of our approach compared to that of Guth and Lubotzky. In Section 2.2, we explain how to obtain a regular tessellation of hyperbolic 4-space with hypercubes. In Section 2.3, we detail how to quotient the space to get a compact manifold, which then yields the quantum code. We finally describe our local decoder and analyze its performances in Section 2.4.

2.1 A variant of Guth and Lubotzky's construction based on a Regular Tessellation of Hyperbolic Space

The family of manifolds considered in [GL14] is a family of 4-dimensional hyperbolic coverings. The tessellations can be obtained by pulling back the natural tessellation of the base space. Each covering equipped with its natural tessellation gives rise to a quantum error correcting code. Unfortunately the fundamental polytope of this natural tessellation is not regular. In particular, it is nontrivial to obtain the local structure of the tessellation, and therefore an explicit description of the code generators. While this did not prevent Hastings from designing a decoding algorithm for these codes [Has16], simulating its performance for the codes of [GL14] appears quite impractical. (Note, however, that Hastings' decoder was recently implemented for the 4-dimensional toric code, in Euclidean space [BDMT16].)

It is useful to see the 4-dimensional homological quantum error correcting codes that Guth and Lubotzky and we construct as generalisations of the 2-dimensional toric code. Let us therefore give the arithmetic manifold viewpoint on the toric code. We consider the ordinary tessellation of the Euclidean plane by unit squares such that vertices have integer coordinates. The translation group of Euclidean plane is $\mathbb{R} \times \mathbb{R}$. We denote by Γ_{toric} the subgroup $\mathbb{Z} \times \mathbb{Z}$ of this translation group. Elements of Γ_{toric} stabilize the ordinary tessellation of Euclidean plane. Let $I = p\mathbb{Z}$ be an ideal of $\mathbb{Z}$, with p a positive integer and define $\Gamma(I)_{\text{toric}}$ to be $I \times I$. The quotient $\mathcal{M}_{\text{toric}}(I)$ of the Euclidean plane by $\Gamma(I)_{\text{toric}}$ is a torus, that naturally inherits the tessellation by unit squares from the Euclidean plane. The constructions of [GL14] and of the present work are generalisations of the 2-dimensional Euclidean toric code in a 4-dimensional hyperbolic setting. To help the reader draw analogies with the toric code, we introduced in this paragraph notations similar to the notations used in the sequel.

We now summarise the construction of Guth and Lubotzky and explain the advantages of our approach. In Ref. [GL14], the construction is based on tessellations of the hyperbolic 4-space $\mathbb{H}^4$. To each code corresponds a manifold equipped with a tessellation. The base space $\mathcal{M}$ is constructed by considering the action of a cocompact discrete group of

isometries Γ on hyperbolic 4-space: $\mathcal{M} = \Gamma \backslash \mathbb{H}^4$. To each finite index subgroup $\Gamma(I)$ of Γ corresponds a covering $\mathcal{M}(I)$ of $\mathcal{M}$ given by $\mathcal{M}(I) = \Gamma(I) \backslash \mathbb{H}^4$. It is natural to tessellate $\mathcal{M}$ with a single 4-face and to tessellate $\mathcal{M}(I)$ with a number of 4-faces equal to the index of $\Gamma(I)$ in Γ . Each 4-face is isometric to the first one. Unfortunately the 4-face is not regular in [GL14], which makes the local description of the quantum code rather complicated. To obtain a similar construction with a regular 4-face, we reverse the process: we start with a convenient regular 4-face and then build a corresponding discrete group of isometries Γ .

For its symmetries and because it tessellates the hyperbolic 4-space, we choose the 4-dimensional hypercube as our targeted regular 4-face. We embed it in hyperbolic 4-space and scale it according to the $\{4, 3, 3, 5\}$ tessellation of hyperbolic 4-space (see Section 2.2.2.1 for a definition of Schafli symbols). The group Γ is generated by the direct isometries of hyperbolic 4-space sending opposing faces of the hypercube to each other with no rotation. The tessellating 4-face we obtain is a hypercube by construction.

The tricky part of the construction is to define finite index subgroups of our discrete group of isometries Γ in a way similar to [GL14]. Indeed, arithmeticity of subgroups $\Gamma(I)$ plays a central role in lower bounding the minimum distance of the corresponding error correcting codes. To achieve this goal, we rely on arithmetic structures defined over the number field $\mathbb{Q}(\sqrt{5})$. Replacing $\mathbb{Q}$ by this number field, $\mathbb{Z}$ by $\mathbb{Z}[\phi]$ - the ring of integers of $\mathbb{Q}(\sqrt{5})$ - and ideals $p\mathbb{Z}$ by ideals of $\mathbb{Z}[\phi]$ where ϕ is the *golden ratio* (giving its name to our construction), it is possible to define principal congruence subgroups $\Gamma(I)$ such that the corresponding family of error correcting codes satisfies the same asymptotic estimates as in [GL14]. We therefore obtain a family of codes with a regular local structure, a non-vanishing rate and a minimum distance lower bounded by $n^{0.1}$, where n is the number of physical qubits.

We take advantage of the regular local structure to design an efficient decoding algorithm. This algorithm is highly local. For X errors, it decreases the syndrome at the scale of a single 4-face. In particular, our algorithm is more local and explicit than Hastings' decoder [Has16]. We prove that syndromes associated with errors of weight below the injectivity radius of the manifold always contain a pattern that can be locally shortened so as to decrease the weight of the syndrome. In other words, the algorithm simply consists in examining the syndrome in the neighborhood of an error and acting on qubits to decrease the syndrome weight. We show that arbitrary errors of size $O(\log N)$ are corrected by this algorithm, which in turn implies that random errors will be corrected with high probability if the error rate is small enough. These results are similar to those of Hastings' decoder, but with the advantage of an entirely explicit algorithm with precise bounds on its performances.

2.2 Hyperbolic 4-space and its Regular Tessellation by Hypercubes

In this section, we first introduce the minimal background on hyperbolic 4-space and regular tessellations. We then focus on the tessellation of hyperbolic 4-space by 4-dimensional hypercubes on which our quantum code construction is based.

2.2.1 Hyperbolic space

We use the hyperboloid model to describe 4-dimensional hyperbolic space. As a set, 4-dimensional hyperbolic space is identified with

$$\mathbb{H}^4 = \{(x_0, x_1, x_2, x_3, x_4) \in \mathbb{R}^5 / -x_0^2 + x_1^2 + x_2^2 + x_3^2 + x_4^2 = -1, \quad x_0 > 0\}.$$

It is endowed with a Riemannian metric such as to make it a space of constant negative sectional curvature. Its orientation-preserving isometry group is $SO^o(1, 4)$, the identity component of the special indefinite orthogonal group. The four coordinates x_1, x_2, x_3 and x_4 are sufficient to parametrise $\mathbb{H}^4$. Indeed x_0 can be retrieved from the condition $x_0^2 = 1 + x_1^2 + x_2^2 + x_3^2 + x_4^2$. Therefore in the sequel we will ignore the coordinate x_0 and refer to x_1 as the *first* coordinate and not as the second.

The reader is referred to [Rat06] for a comprehensive introduction to hyperbolic geometry. To give some intuition about hyperbolic space we will merely compare the perimeter growth of a hyperbolic circle of radius r with its Euclidean counterpart. In hyperbolic space, such a circle has perimeter $2\pi \sinh(r)$. The growth is exponentially faster than its Euclidean counterpart $2\pi r$. In spherical space on the other hand, the perimeter of a circle of radius r is only $2\pi \sin(r)$, for $r < \pi$. Informally speaking, there is more room in the angular direction in hyperbolic space than in Euclidean space just like there is less room in the angular direction in spherical space than in Euclidean space. One can make this statement more precise by considering regular tessellations and their combinatorial properties.

2.2.2 Regular polytopes and tessellations

The geometric point of view on tessellations is probably the most intuitive. By geometric, we mean that vertices, edges and higher dimensional faces of the tessellation are subsets of a geometric space such as for example the hyperbolic plane or the Euclidean 3-space. However a tessellation also entails purely combinatorial data, namely the incidences between its i -faces and its $(i+1)$ -faces. We will refer to this combinatorial data as the abstract polytope attached to a tessellation. For a comprehensive exposition of this so-called abstract point of view we refer to the book of McMullen and Schulte [MS02] (Chapter 2 for the abstract point of view and Chapter 5 for its interplay with geometric realisations). The abstract point of view is especially relevant to quantum error correction since the combinatorial data is sufficient to define a quantum error correcting code. We will only mention here that an abstract polytope is called regular if its automorphism group is transitive on the set of its flags. Moreover the realisation of a abstract regular polytope as a tessellation of a geometric space is called regular if its automorphism group can be represented as an isometry group of the geometric space.

Interestingly the combinatorial data of a abstract regular polytope (its incidences) determines in which geometric space it can be embedded. We can thus talk about spherical, Euclidean and hyperbolic abstract regular polytopes:

Definition 5. *An abstract regular polytope is called spherical (respectively Euclidean, respectively hyperbolic) if it can be realised with regular faces in a spherical (respectively Euclidean, respectively hyperbolic) manifold.*

Informally speaking, if a polytope is locally too small to fit in Euclidean space, it curves inwards and yields a spherical tessellation. If it is too big, it yields a hyperbolic tessellation. In the Euclidean case, the faces of the tessellation can be scaled by multiplying all lengths by a given positive real λ . In the spherical and hyperbolic cases, however, the volumes

of faces are imposed by the combinatorics of the tessellation: tessellations far from being Euclidean lead to faces with a large volume.

2.2.2.1 Combinatorial point of view on tessellations: Schläfli symbols

Results of this section come from Ref. [Cox54]. Since the realisation of an abstract regular polytope is essentially unique (up to a scaling factor if it is euclidean) we will often not distinguish a regular tessellation of a geometric space from its abstract regular polytope: the combinatorial data attached to it. Therefore we can describe regular tessellations *via* their *Schläfli symbols*, which are defined recursively for $p, q, r, \dots$ positive integers:

- $\{p\}$ refers to a regular p -sided polygon.
- $\{p, q\}$ refers to a regular tessellation by regular p -sided polygons such that each vertex is incident to q regular p -sided polygons.

One obtains a tessellation of the Euclidean plane if $(p - 2)(q - 2) = 4$, or of the hyperbolic plane if $(p - 2)(q - 2) > 4$. Finally if $(p - 2)(q - 2) < 4$, then $\{p, q\}$ can represent either a tessellation of the two-dimensional sphere or a 3-dimensional polyhedron.

There are five regular 3-dimensional polyhedrons called the *Platonic solids*: the regular *icosahedron* $\{3, 5\}$; the regular *octahedron* $\{3, 4\}$; the regular *tetrahedron* $\{3, 3\}$; the *cube* $\{4, 3\}$ and the regular *dodecahedron*, $\{5, 3\}$.

- If $\{p, q\}$ and $\{r, q\}$ are 3-dimensional polyhedrons², then $\{p, q, r\}$ refers to a regular tessellation by $\{p, q\}$ -polyhedrons such that each edge of the tessellation is incident to r $\{p, q\}$ -polyhedrons. Note that the terminology *honeycomb* is sometimes used instead of tessellation to insist on 3-dimensionality. The terminology *mosaic* can be encountered as well. We will use tessellation in the sequel regardless of the dimension. Similarly as before, the nature of the tessellation depends on the relation between the integers p, q, r . If $\cos(\frac{\pi}{q}) = \sin(\frac{\pi}{p}) \sin(\frac{\pi}{r})$, one obtains a tessellation of the Euclidean 3-dimensional space. If $\cos(\frac{\pi}{q}) > \sin(\frac{\pi}{p}) \sin(\frac{\pi}{r})$, one gets a tessellation of the hyperbolic 3-dimensional space. Finally, if $\cos(\frac{\pi}{q}) < \sin(\frac{\pi}{p}) \sin(\frac{\pi}{r})$, it can represent either a tessellation of the spherical 3-dimensional space or a 4-dimensional polytope.

There are six regular 4-dimensional polytopes: $\{3, 3, 5\}$ is the *600-cell*, $\{3, 3, 4\}$ is the *4-orthoplex*, $\{3, 4, 3\}$ is the *24-cell*, $\{3, 3, 3\}$ is the regular *4-simplex*, $\{4, 3, 3\}$ is the 4-dimensional *hypercube*, and $\{5, 3, 3\}$ is the *120-cell*.

- If $\{p, q, r\}$ and $\{s, r, q\}$ are 4-dimensional polytopes, $\{p, q, r, s\}$ refers to a regular tessellation by $\{p, q, r\}$ -polytopes such that each 2-face of the tessellation is incident to s $\{p, q, r\}$ -polytopes.

If $\frac{\cos^2(\frac{\pi}{q})}{\sin^2(\frac{\pi}{p})} + \frac{\cos^2(\frac{\pi}{r})}{\sin^2(\frac{\pi}{s})} = 1$, it is a tessellation of the 4-dimensional Euclidean space. If $\frac{\cos^2(\frac{\pi}{q})}{\sin^2(\frac{\pi}{p})} + \frac{\cos^2(\frac{\pi}{r})}{\sin^2(\frac{\pi}{s})} > 1$, it is a tessellation of hyperbolic 4-space.

There are five regular tessellations of hyperbolic 4-space: $\{3, 3, 3, 5\}$, $\{4, 3, 3, 5\}$, $\{5, 3, 3, 5\}$, $\{5, 3, 3, 4\}$ and $\{5, 3, 3, 3\}$.

Given a tessellation or a polytope described by Schläfli symbol $\{p_1, \dots, p_n\}$, the tessellation or polytope described by $\{p_n, \dots, p_1\}$ is called the *dual* tessellation or polytope. It is the tessellation obtained by mapping every i -face to an $(n - i)$ -face. Note that duality doesn't change the hyperbolic, Euclidean or spherical type of a tessellation.

²The condition that $\{r, q\}$ is also a 3-dimensional polyhedron is necessary, for instance, to ensure that the dual tessellation $\{r, q, p\}$ is well-defined.

2.2.2.2 The $\{4,3,3,5\}$ regular tessellation of hyperbolic 4-space

In this work we will focus on the $\{4,3,3,5\}$ regular tessellation of hyperbolic 4-space. The 4-faces $\{4,3,3\}$ of this tessellation are 4-dimensional hypercubes, which are especially convenient. In particular, one can exploit the fact that the symmetries of the hypercube are compatible with its description in coordinates in the hyperboloid model to find a nice description of a discrete subgroup of $SO^o(1,4)$ corresponding to the $\{4,3,3,5\}$ regular tessellation. The other regular tessellations of hyperbolic 4-space could lead to similar constructions and would yield quantum codes with similar asymptotic properties. Since their symmetries are less compatible with coordinates in the hyperboloid model, they would require more work to make computations explicit and we will not consider them in this work.

2.2.3 Isometry group of the tessellation

We consider a regular hypercube centered at the origin of the hyperboloid model and such that each of its eight 3-faces is orthogonal to a coordinate axis. We denote this hypercube by T in the sequel (a 4-dimensional hypercube is also called a tesseract). Since the hypercube T is regular, there exist direct isometries of the hyperbolic 4-space $\mathbb{H}^4$ sending any one of them onto the opposite one. These isometries can be thought of as the hyperbolic equivalent of Euclidean translations. Requiring that these direct isometries act trivially on three coordinates defines them uniquely. For example the direct isometries that send the 3-faces orthogonal to the first coordinate axis onto each other are given by the following matrices:

$$g_1 = \begin{pmatrix} \cosh t & \sinh t & 0 & 0 & 0 \\ \sinh t & \cosh t & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \quad (g_1)^{-1} = \begin{pmatrix} \cosh t & -\sinh t & 0 & 0 & 0 \\ -\sinh t & \cosh t & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}.$$

The pair of direct isometries sending the 3-faces orthogonal to the second coordinate axis onto each other is given by these two matrices:

$$g_2 = \begin{pmatrix} \cosh t & 0 & \sinh t & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ \sinh t & 0 & \cosh t & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \quad (g_2)^{-1} = \begin{pmatrix} \cosh t & 0 & -\sinh t & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ -\sinh t & 0 & \cosh t & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}.$$

The two remaining pairs of direct isometries $g_3, (g_3)^{-1}, g_4$ and $(g_4)^{-1}$ are obtained from the above matrices by permuting two coordinates. Recall that the coordinate x_0 is redundant with the four others. Therefore the zeroth coordinate should not be permuted.

The angle between two adjacent 3-faces of the hypercube T depends on the volume of T or equivalently on the parameter t : as we will show, the greater t , the greater the volume of T and the smaller the angle between two adjacent 3-faces. We will compute the value of t such that this angle is $2\pi/5$. Indeed in the $\{4,3,3,5\}$ regular tessellation of hyperbolic 4-space, five hypercubes meet along each 2-face, which means that the *dihedral angle* between two 3-faces of the same hypercube must be $2\pi/5$. Note that the dihedral angle between 3-faces is sometimes called dichoral angle to insist on higher dimension. In the sequel we will use the terminology dihedral angle regardless of dimension. To compute dihedral angles in the hyperboloid model we need some definitions.

Definition 6 (Ratcliffe [Rat06] §3.1). *The Lorentzian inner product denoted $\circ$ is the bilinear map defined on $\mathbb{R}^5 \times \mathbb{R}^5$ by:*

$$u \circ v = -u_0v_0 + u_1v_1 + u_2v_2 + u_3v_3 + u_4v_4.$$

Two vectors u, v are Lorentz orthogonal if $u \circ v = 0$.

Definition 7 (Ratcliffe [Rat06] §3.1). *The Lorentzian norm of a vector u is the complex number denoted $\|u\|$ satisfying $u \circ u = \|u\|^2$ and such that $\|u\|$ is either positive imaginary, 0 or positive.*

Note that if $\|u\|$ is positive imaginary, $\|u\|$ denotes its modulus.

Definition 8 (Ratcliffe [Rat06] §3.2). *The space-like angle η between two space-like vectors u and v is defined by: $u \circ v = \|u\| \cdot \|v\| \cos(\eta)$ and $0 \leq \eta \leq \pi$.*

Definition 9 (Ratcliffe [Rat06] §6.4). *Let S and T be two adjacent sides of a convex polytope P . Let u , respectively v , be a vector that is Lorentz orthogonal to S , respectively T , and directed away from P . Let η the space-like angle between u and v . Then the dihedral angle $\theta(S, T)$ between S and T is defined by:*

$$\theta(S, T) = \pi - \eta(u, v).$$

With this definition the dihedral angle is invariant under global Lorentz transformations. Indeed Lorentz orthogonality and space-like angles are Lorentz invariant. This property is necessary since Lorentz transformations are the isometries of the hyperbolic metric.

To fully justify the definition it remains to show that we obtain the expected dihedral angle when the two sides intersect at the origin $(1, 0, 0)^T$ of the hyperboloid. For simplicity we assume that S and T are two lines intersecting at the origin of the hyperbolic plane. In the hyperboloid model we can assume that $S = \{(\cosh(x), \sinh(x), 0) \mid x \in \mathbb{R}\}$ and $T = \{(\cosh(x), \cos(\theta(S, T)) \sinh(x), \sin(\theta(S, T)) \sinh(x)) \mid x \in \mathbb{R}\}$. Then $u = (0, 0, -1)$ is Lorentz orthogonal to S and $v = (0, -\sin(\theta(S, T)), \cos(\theta(S, T)))$ is Lorentz orthogonal to T . Vectors u and v are directed away from P (P is defined consistently with $\theta(S, T)$). A computation yields $u \circ v = -\cos(\theta(S, T))$. Since $\|u\| = \|v\| = 1$ we obtain that $\cos(\eta(u, v)) = -\cos(\theta(S, T))$. Since by definition both $\theta(S, T)$ and $\eta(u, v)$ are in $[0, \pi]$, this gives $\theta(S, T) = \pi - \eta(u, v)$.

We can now come back to the hypercube T centered at the origin of hyperbolic space. Let C_1 , respectively C_2 , be the 3-face of T orthogonal in hyperbolic 4-space to the first, respectively second, axis and such that its first, respectively second, coordinate in the hyperboloid model is positive. Recall that x_0 is referred to as the zeroth coordinate. Thus the first coordinate is x_1 and the second is x_2 . Points of C_1 have coordinates of the form $\lambda(\cosh(t/2), \sinh(t/2), a, b, c)^T$ for some $a, b, c \in \mathbb{R}$ and a normalising constant λ . Similarly points of C_2 have coordinates of the form $\lambda(\cosh(t/2), a, \sinh(t/2), b, c)^T$. It is straightforward to verify that $N_1 = (\sinh(t/2), \cosh(t/2), 0, 0, 0)^T$ is Lorentz orthogonal to C_1 and $N_2 = (\sinh(t/2), 0, \cosh(t/2), 0, 0)^T$ is Lorentz orthogonal to C_2 . We have:

$$\begin{aligned} \eta(N_1, N_2) &= \arccos\left(\frac{N_1 \circ N_2}{\|N_1\| \|N_2\|}\right) = \arccos(-\sinh^2(t/2)), \\ \theta(C_1, C_2) &= \pi - \eta(N_1, N_2), \\ \theta(C_1, C_2) &= \pi - \arccos(-\sinh^2(t/2)). \end{aligned}$$

As announced, the dihedral angle between two adjacent 3-faces of the hypercube T decreases with parameter t , or equivalently when the volume of T increases.

Since we want to build a $\{4,3,3,5\}$ tessellation, five hypercubes have to be incident to each 2-face of the hypercube. This imposes $\theta(C_1, C_2) = 2\pi/5$ and leads to $t = 2\operatorname{arsinh}(\sqrt{\cos(2\pi/5)})$. We eventually obtain

$$\cosh(t) = \frac{1 + \sqrt{5}}{2} \quad \text{and} \quad \sinh(t) = \sqrt{\frac{1 + \sqrt{5}}{2}},$$

the golden ratio ϕ and its square root.

We denote by Γ the discrete subgroup of $SO^o(1, 4)$ generated by the four direct isometries g_1, g_2, g_3 and g_4 sending a 3-face of the hypercube onto the opposite 3-face. Note that there are eight such direct isometries but they are pairwise inverse of each other.

2.2.4 Coxeter group approach

The problem with the group Γ defined above is that its fundamental domain is not the hypercube. Indeed there are elements of Γ that fix the hypercube globally but not pointwise. They correspond to symmetries of the hypercube. Therefore the fundamental domain of Γ is only a fraction of the hypercube.

However we would like to work with $\{4, 3, 3, 5\}$ tessellations. This way the tessellation is regular with a local structure that is easy to describe. To achieve this goal we will define a double extension of Γ which is a representation of the $\{4, 3, 3, 5\}$ Coxeter group. The standard technique of considering cosets of this Coxeter group will then yield the tessellation by hypercubes (see [MS02]). We can define $\Gamma^{\{4,3,3,5\}} = \langle r_0, r_1, r_2, r_3, r_4 \rangle$, with its five generators given by:

$$\begin{aligned} r_0 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_1 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_2 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \\ r_3 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 \end{pmatrix}, & r_4 &= \begin{pmatrix} \phi & 0 & 0 & 0 & -\sqrt{\phi} \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ \sqrt{\phi} & 0 & 0 & 0 & -\phi \end{pmatrix}. \end{aligned}$$

It is straightforward to verify that these five generators satisfy the relations defining the $\{4, 3, 3, 5\}$ string Coxeter group:

$$r_0^2 = r_1^2 = r_2^2 = r_3^2 = r_4^2 = (r_0 r_1)^4 = (r_1 r_2)^3 = (r_2 r_3)^3 = (r_3 r_4)^5 = \operatorname{id}.$$

To obtain a $\{4, 3, 3, 5\}$ tessellation of $\mathbb{H}^4$, we can follow [MS02] and identify S_i -cosets in $\Gamma^{\{4,3,3,5\}}$ with i -faces of the $\{4, 3, 3, 5\}$ tessellation. For $i \in \{0, \dots, 4\}$, the group S_i is the subgroup of $\Gamma^{\{4,3,3,5\}}$ generated by the four generators $(r_j)_{j \neq i}$ (for instance $S_1 \stackrel{\text{def}}{=} \langle r_0, r_2, r_3, r_4 \rangle$). By definition, an i -face F_a and a j -face F_b are incident if the

corresponding cosets $g_a S_i$ and $g_b S_j$ have a non-empty intersection.

With these definitions we only have a combinatorial description of the $\{4, 3, 3, 5\}$ tessellation. To obtain the geometrical version of the tessellation, observe that each element of $\Gamma^{\{4,3,3,5\}}$ corresponds to a simplex of $\mathbb{H}^4$: the identity of $\Gamma^{\{4,3,3,5\}}$ corresponds to a fundamental domain $\mathcal{S}$ of $\mathbb{H}^4$ (which happens to be a simplex in this case) and any $g \in \Gamma^{\{4,3,3,5\}}$ corresponds to $g\mathcal{S}$. Now an i -face of the tessellation corresponds to a coset of S_i in $\Gamma^{\{4,3,3,5\}}$. This coset can be considered as a set of elements of $\Gamma^{\{4,3,3,5\}}$ or, in other words, as a set of closed simplices of $\mathbb{H}^4$. If $i = 4$ the geometrical 4-face is defined as the union of these closed simplices. If $i \in \{0, \dots, 3\}$ the geometrical i -face is defined as the intersection of the 4-faces incident to this i -face.

This $\{4, 3, 3, 5\}$ tessellation of hyperbolic 4-space has an infinite number of i -faces for every $i \in \{0, \dots, 4\}$. To build a code with a finite number of qubits, we need a tessellation with a finite number of 2-faces. We use in the sequel number theoretical tools to construct quotients of the hyperbolic 4-space equipped with a $\{4, 3, 3, 5\}$ tessellation.

2.3 Compact Manifolds equipped with a $\{4, 3, 3, 5\}$ Tessellation

We want to define a quantum code by identifying physical qubits with 2-faces of a tessellation. To obtain a code with a finite number of physical qubits, we will consider tessellations of compact manifolds. We will therefore consider the $\{4, 3, 3, 5\}$ tessellation of compact manifolds obtained as quotients of hyperbolic 4-space. These manifolds are called *arithmetic* because they are quotients of $\mathbb{H}^4$ by arithmetic subgroups of $\Gamma^{\{4,3,3,5\}}$. We first review the definitions of a number field and its ring of integers. We then use these tools to associate an arithmetic subgroup $\Gamma^{\{4,3,3,5\}}(I)$ to every ideal I of the ring of integers $\mathbb{Z}[\phi]$.

2.3.1 Number fields and rings of integers

Definition 10. A number field K is a finite degree field extension of the field of rational numbers $\mathbb{Q}$.

Definition 11. A complex number is an algebraic number if it is a root of a non-zero polynomial over $\mathbb{Q}$.

Theorem 12 (e.g. Marcus, [Mar77] Appendix 2). Every number field has the form $\mathbb{Q}(\alpha)$ for some algebraic number $\alpha \in \mathbb{C}$. If α is a root of an irreducible polynomial over $\mathbb{Q}$ having degree n , then

$$\mathbb{Q}(\alpha) = \{a_0 + a_1\alpha + \dots + a_{n-1}\alpha^{n-1} \mid \forall i \in \{0, \dots, n-1\}, a_i \in \mathbb{Q}\}.$$

Since $\sqrt{5}$ is a root of $X^2 - 5$, which is irreducible over $\mathbb{Q}$, we have

$$\mathbb{Q}(\sqrt{5}) = \{a_0 + a_1\sqrt{5} \mid a_0, a_1 \in \mathbb{Q}\}.$$

Definition 13. A complex number is an algebraic integer if it is a root of a monic (leading coefficient equal to 1) polynomial with coefficients in $\mathbb{Z}$.

Definition 14. *The ring of integers of a number field K is the subset of its algebraic integers. It is denoted O_K .*

Propositon 15 (e.g. Marcus, [Mar77] p. 15). *Let $m \in \mathbb{Z}$ satisfy $m \equiv 1 \pmod{4}$ and let K be the quadratic number field $\mathbb{Q}(\sqrt{m})$. Then,*

$$O_K = \left\{ \frac{a + b\sqrt{m}}{2} \mid a, b \in \mathbb{Z} \right\}.$$

Applying this characterization to the case $K = \mathbb{Q}(\sqrt{5})$ yields:

$$\begin{aligned} O_{\mathbb{Q}(\sqrt{5})} &= \left\{ \frac{a + b\sqrt{m}}{2} \mid a, b \in \mathbb{Z}, \quad a \equiv b \pmod{2} \right\} \\ O_{\mathbb{Q}(\sqrt{5})} &= \mathbb{Z}[\phi]. \end{aligned}$$

where ϕ is the golden ratio $\frac{1+\sqrt{5}}{2}$.

2.3.2 Arithmetic subgroups $\Gamma^{\{4,3,3,5\}}(I)$

Since ϕ and its square root are algebraic numbers, $\mathbb{Q}(\sqrt{\phi})$ is a number field. Its ring of integers is $\mathbb{Z}[\sqrt{\phi}]$, and therefore every matrix of $\Gamma^{\{4,3,3,5\}}$ has its coefficients in the ring $\mathbb{Z}[\sqrt{\phi}]$.

Definition 16. *A number field is totally real if all its embeddings in $\mathbb{C}$ are embeddings in $\mathbb{R}$.*

In order to prove the same asymptotic behaviour of the code parameters n , k and d as in Refs [GL14] and [Mur16], we need to work with a totally real number field. Note that the totally real number field condition is not explicit in [GL14] but it is implicitly used to show that their arithmetic group Γ is discrete. However $\mathbb{Q}(\sqrt{\phi})$ is not a totally real number field. Indeed $\sqrt{\phi}$ has minimal polynomial $X^4 - X^2 + 1$ which factorises as $(X - \sqrt{\phi})(X + \sqrt{\phi}) \left(X - i\sqrt{\frac{\sqrt{5}-1}{2}} \right) \left(X + i\sqrt{\frac{\sqrt{5}-1}{2}} \right)$ and thus $\mathbb{Q}(\sqrt{\phi})$ admits the two non-real embeddings determined by $\sqrt{\phi} \mapsto i\sqrt{\frac{\sqrt{5}-1}{2}}$ and by $\sqrt{\phi} \mapsto -i\sqrt{\frac{\sqrt{5}-1}{2}}$. We therefore conjugate matrices of $\Gamma^{\{4,3,3,5\}}$ in such a way that all their entries now belong to a totally real number field. Since matrix multiplication is defined through addition and multiplication of their entries, it is sufficient to ensure that the four matrices generating $\Gamma^{\{4,3,3,5\}}$ have their entries in a totally real number field.

Observe that $\begin{pmatrix} 1/\sqrt{\phi} & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} \phi & -\sqrt{\phi} \\ \sqrt{\phi} & -\phi \end{pmatrix} \begin{pmatrix} \sqrt{\phi} & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} \phi & -1 \\ \phi & -\phi \end{pmatrix}$ and $\begin{pmatrix} \phi & 1 \\ \phi & \phi \end{pmatrix}^{-1} = \begin{pmatrix} \phi & -1 \\ \phi & -\phi \end{pmatrix}$. Therefore, defining $P = \text{diag}(\sqrt{\phi}, 1, 1, 1, 1)$, the group $\tilde{\Gamma}^{\{4,3,3,5\}}$ defined as $P^{-1}\Gamma^{\{4,3,3,5\}}P$ has all its matrices with entries in the number field $K = \mathbb{Q}(\phi) = \mathbb{Q}(\sqrt{5})$, and even in its ring of integers $\mathbb{Z}[\phi]$. Note that since $\Gamma^{\{4,3,3,5\}}$ is a subgroup of $O(1, 4)$, matrices g in $\tilde{\Gamma}^{\{4,3,3,5\}}$ satisfy the equation $g^T \tilde{J} g = \tilde{J}$ where $\tilde{J} = \text{diag}(-\phi, 1, 1, 1, 1)$.

Now the minimal polynomial of $\sqrt{5}$ is $X^2 - 5$ which factorises as $(X - \sqrt{5})(X + \sqrt{5})$. Hence the two embeddings of $\mathbb{Q}(\sqrt{5})$ in $\mathbb{C}$ are the identity and the embedding determined by $\sqrt{5} \mapsto -\sqrt{5}$. $\mathbb{Q}(\sqrt{5})$ is thus a totally real number field.

Definition 17. Let I be an ideal of a ring A . Let G be a matrix group with coefficients in A . The principal congruence subgroup of level I of G is the kernel of the reduction modulo I morphism. It is denoted $G(I) = \ker \pi_I$ with

$$\begin{aligned} \pi_I : M_n(A) &\rightarrow M_n(A/I) \\ (a_{i,j}) &\mapsto (a_{i,j} + I). \end{aligned}$$

It is natural to consider ideals of the ring A because we want the quotient A/I to be a ring in order for $M_n(A/I)$ to be defined. Hence to each ideal I of $\mathbb{Z}[\phi]$ corresponds a normal subgroup $\Gamma^{\{4,3,3,5\}}(I)$ of $\Gamma^{\{4,3,3,5\}}$. We denote by $\mathcal{M}(I)$ the quotient of $\mathbb{H}^4$ by $\Gamma^{\{4,3,3,5\}}(I)$. By definition $\mathcal{M}(I) = \Gamma^{\{4,3,3,5\}}(I) \backslash \mathbb{H}^4$ is the set of orbits of $\mathbb{H}^4$ under the action of $\Gamma^{\{4,3,3,5\}}(I)$. Note that we use the notation $\Gamma^{\{4,3,3,5\}}(I) \backslash \mathbb{H}^4$ and not $\mathbb{H}^4 / \Gamma^{\{4,3,3,5\}}(I)$ because $\Gamma^{\{4,3,3,5\}}(I)$ acts on $\mathbb{H}^4$ on the left. $\mathcal{M}(I)$ naturally inherits the hyperbolic structure of $\mathbb{H}^4$.

For completeness, we will now detail how $\mathcal{M}(I)$ inherits the $\{4, 3, 3, 5\}$ tessellation of $\mathbb{H}^4$. By definition of $\Gamma(I)$ the following short sequence is exact:

$$1 \rightarrow \Gamma^{\{4,3,3,5\}}(I) \rightarrow \Gamma^{\{4,3,3,5\}} \rightarrow \pi_I(\Gamma^{\{4,3,3,5\}}) \rightarrow 1.$$

Therefore, by the first isomorphism theorem, the quotient group $\Gamma^{\{4,3,3,5\}} / \Gamma^{\{4,3,3,5\}}(I)$ is isomorphic to $\pi_I(\Gamma^{\{4,3,3,5\}})$. This quotient group acts on $\mathcal{M}(I)$ in the following manner: for any $g \cdot \Gamma^{\{4,3,3,5\}}(I) \in \Gamma^{\{4,3,3,5\}} / \Gamma^{\{4,3,3,5\}}(I)$ and $\Gamma^{\{4,3,3,5\}}(I) \cdot x$ in $\mathcal{M}(I)$,

$$(g \cdot \Gamma^{\{4,3,3,5\}}(I)) \cdot (\Gamma^{\{4,3,3,5\}}(I) \cdot x) = \Gamma^{\{4,3,3,5\}}(I) \cdot (g \cdot x).$$

Since $\Gamma^{\{4,3,3,5\}}(I)$ is normal in $\Gamma^{\{4,3,3,5\}}$, this is well defined and it is a group action. We will see in subsection 2.3.4 that for ideals with sufficiently large norms (see Definition 20), $\Gamma^{\{4,3,3,5\}}(I)$ acts freely (without fixed points) on $\mathbb{H}^4$. Therefore for such ideals $\mathcal{M}(I)$ is a manifold.

Moreover for $i \in \{0, \dots, 4\}$, i -faces of the $\{4, 3, 3, 5\}$ tessellation of $\mathbb{H}^4$ have a diameter upper bounded by some constant c depending on the local structure. Again, for ideals with sufficiently large norms, $\Gamma^{\{4,3,3,5\}}(I)$ acts on $\mathbb{H}^4$ in a way such that no pair of points $x, y \in \mathbb{H}^4$ satisfying $d(x, y) \leq c$ belong to the same orbit. For such ideals I , the $\{4, 3, 3, 5\}$ local structure is preserved by $\Gamma^{\{4,3,3,5\}}(I)$. We can retrieve it by considering cosets of $\pi_I(\Gamma^{\{4,3,3,5\}})$.

More precisely the group $\tilde{\Gamma}^{\{4,3,3,5\}}(I)$ is generated by $r_{0,\bar{J}}, r_{1,\bar{J}}, r_{2,\bar{J}}, r_{3,\bar{J}}$ and $r_{4,\bar{J}}$:

$$\begin{aligned} r_{0,\bar{J}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_{1,\bar{J}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_{2,\bar{J}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \\ r_{3,\bar{J}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 \end{pmatrix}, & r_{4,\bar{J}} &= \begin{pmatrix} \phi & 0 & 0 & 0 & -1 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ \phi & 0 & 0 & 0 & -\phi \end{pmatrix}. \end{aligned}$$

Therefore for any ideal I of $\mathbb{Z}[\phi]$, the group $\pi_I(\tilde{\Gamma}^{\{4,3,3,5\}})$ is generated by $r_{0,I}, r_{1,I}, r_{2,I}, r_{3,I}$ and $r_{4,I}$:

$$\begin{aligned}
r_{0,I} &= \begin{pmatrix} 1+I & 0 & 0 & 0 & 0 \\ 0 & -1+I & 0 & 0 & 0 \\ 0 & 0 & 1+I & 0 & 0 \\ 0 & 0 & 0 & 1+I & 0 \\ 0 & 0 & 0 & 0 & 1+I \end{pmatrix}, & r_{1,I} &= \begin{pmatrix} 1+I & 0 & 0 & 0 & 0 \\ 0 & 0 & 1+I & 0 & 0 \\ 0 & 1+I & 0 & 0 & 0 \\ 0 & 0 & 0 & 1+I & 0 \\ 0 & 0 & 0 & 0 & 1+I \end{pmatrix}, \\
r_{2,I} &= \begin{pmatrix} 1+I & 0 & 0 & 0 & 0 \\ 0 & 1+I & 0 & 0 & 0 \\ 0 & 0 & 0 & 1+I & 0 \\ 0 & 0 & 1+I & 0 & 0 \\ 0 & 0 & 0 & 0 & 1+I \end{pmatrix}, & r_{3,I} &= \begin{pmatrix} 1+I & 0 & 0 & 0 & 0 \\ 0 & 1+I & 0 & 0 & 0 \\ 0 & 0 & 1+I & 0 & 0 \\ 0 & 0 & 0 & 0 & 1+I \\ 0 & 0 & 0 & 1+I & 0 \end{pmatrix}, \\
r_{4,I} &= \begin{pmatrix} \phi+I & 0 & 0 & 0 & -1+I \\ 0 & 1+I & 0 & 0 & 0 \\ 0 & 0 & 1+I & 0 & 0 \\ 0 & 0 & 0 & 1+I & 0 \\ \phi+I & 0 & 0 & 0 & -\phi+I \end{pmatrix}.
\end{aligned}$$

For ideals I whose norm is large enough, we can define i -faces of $\mathcal{M}(I)$ with the same coset method we used for $\mathbb{H}^4$: i -faces correspond to cosets of $\pi_I(\tilde{\Gamma}^{\{4,3,3,5\}})$ by its subgroup $S_{i,I}$ generated by $(r_{j,I})_{j \neq i}$. Incident faces correspond to cosets whose intersection is not empty. We will use this method in Subsection 2.3.5 to construct explicit quantum codes.

The results stated in the sequel of this paper are valid for ideals I whose norm is large enough to have a $\{4, 3, 3, 5\}$ tessellation of $\mathcal{M}(I)$.

Definition 18. Let H be a subgroup of a group G . The index of H in G , denoted $[G : H]$, is the cardinal of the quotient G/H .

Lemma 19. The number of 2-faces of the $\{4, 3, 3, 5\}$ tessellation of $\mathcal{M}(I)$ is $n(I) = [\Gamma^{\{4,3,3,5\}} : \Gamma^{\{4,3,3,5\}}(I)]/80$.

Proof. $\mathcal{M}(I)$ admits a tessellation by $[\Gamma^{\{4,3,3,5\}} : \Gamma^{\{4,3,3,5\}}(I)]$ simplices isometric to a fundamental domain of the action of $\Gamma^{\{4,3,3,5\}}$ on $\mathbb{H}^4$. A 2-face of the $\{4, 3, 3, 5\}$ tessellation of $\mathcal{M}(I)$ corresponds to a coset of $S_{2,I}$ in $\Gamma^{\{4,3,3,5\}}(I)$. The result follows from the value of the cardinal of $S_{2,I}$:

$$|S_{2,I}| = |S_2| = |\langle r_0, r_1 \rangle| \times |\langle r_3, r_4 \rangle| = 8 \times 10 = 80.$$

□

Definition 20. The norm $N(I)$ of an ideal I of a ring A is the cardinal of the quotient A/I .

It is shown in Ref. [Mur16] that $[\Gamma^{\{4,3,3,5\}} : \Gamma^{\{4,3,3,5\}}(I)] \leq 4N(I)^{\dim(O(1,4))} = 4N(I)^{10}$. This provides an upper bound on the size of the quantum code associated with an ideal I :

$$n(I) \leq N(I)^{10}/20. \quad (2.1)$$

Note that the ring $\mathbb{Z}[\phi]$ admits a family of ideals whose norms are unbounded. Indeed the norm of the ideal of $\mathbb{Z}[\phi]$ generated by m is m^2 . This translates into a family of quantum codes with an unbounded number of physical qubits. Moreover there are other ideals in

$\mathbb{Z}[\phi]$. For example, the ideal generated by $\sqrt{5}$ has norm 5.

We will now paraphrase the correspondence exposed in Ref. [GL14] between a family of coverings and a family of quantum codes. From each 4-dimensional manifold equipped with a $\{4, 3, 3, 5\}$ tessellation $\mathcal{M}(I)$, a code is constructed: qubits are identified with 2-faces of $\mathcal{M}(I)$, X -type stabilizers are identified with 1-faces (edges) of $\mathcal{M}(I)$ and Z -type stabilizers are identified with 3-faces of $\mathcal{M}(I)$. Each X -type, respectively Z -type, stabilizer acts by an X Pauli matrix, respectively a Z Pauli matrix, on every qubit it is incident to. The codespace is the common $(+1)$ -eigenspace of the set of stabilizers. The length n of the code, *i.e.* its number of physical qubits, is the number of 2-faces of the tessellation. It is proportional to the volume of $\mathcal{M}(I)$. The dimension k of the code, *i.e.* its number of logical qubits, is the second Betti number of $\mathcal{M}(I)$, *i.e.* the rank of its second homology group. The minimum distance d of the code is the minimal number of 2-faces forming a homologically nontrivial 2-cycle in $\mathcal{M}(I)$. It is lower bounded by a quantity proportional to the least area of a homologically nontrivial surface of $\mathcal{M}(I)$. These proportionality coefficients do not depend on the ideal I . With this correspondence, the asymptotic behaviour of n , k and d is understood in terms of the family of manifolds $(\mathcal{M}(I))_{I \in \mathbb{Z}[\phi]}$ independently of the $\{4, 3, 3, 5\}$ tessellation.

To each ideal I of the ring of integers $\mathbb{Z}[\phi]$ corresponds a manifold $\mathcal{M}(I)$ equipped with a $\{4, 3, 3, 5\}$ tessellation and a quantum error correcting code $\mathcal{C}(I)$.

2.3.3 Lower bound on the rate of the quantum codes

Quantum codes based on regular tessellations of hyperbolic spaces have a non-vanishing rate. It is well-known for tessellations of the hyperbolic plane and it is also true for tessellations of the hyperbolic 4-space. The argument is given in Ref. [GL14] (Theorem 7 and Corollary 9) and we can sketch it here and make it more quantitative than in Ref. [GL14]:

As a consequence of Gauss-Bonnet-Chern's theorem [CHE96], the Euler characteristic $\chi(I)$ of the closed oriented hyperbolic 4-manifolds $\mathcal{M}(I)$ satisfies $\chi(I) = c \text{vol}(\mathcal{M}(I))$. It is possible to generalise the definition of the Euler characteristic (see *e.g.* [Mar15]) to orbifolds (roughly speaking, manifolds that can have singularities) in such a way that this definition still holds. We can illustrate this by computing the Euler characteristic of the hypercube T of the $\{4, 3, 3, 5\}$ tessellation of hyperbolic 4-space. For each i in $\{0, \dots, 4\}$ we have to divide the number of i -faces of T by the number of hypercubes an i -face would be incident to in the $\{4, 3, 3, 5\}$ tessellation of hyperbolic 4-space. We obtain:

$$\begin{aligned} \chi(T) &= \frac{1}{1} - \frac{8}{2} + \frac{24}{5} - \frac{32}{20} + \frac{16}{600} \\ &= \frac{17}{75}. \end{aligned}$$

Gauss-Bonnet-Chern theorem also yields $\chi(T) = c \text{vol}(T)$.

$$\begin{aligned}
n(I) &= \#(2\text{-faces}) \\
&= \frac{24}{5} \#(4\text{-faces}) \\
&= \frac{24}{5} \frac{\text{vol}(\mathcal{M}(I))}{\text{vol}(T)} \\
&= \frac{24}{5} \frac{\chi(I)}{\chi(T)} \\
&= \frac{360}{17} \chi(I).
\end{aligned}$$

Moreover, by definition of the Euler characteristic, $\chi(I) = \sum_{i=0}^4 (-1)^i \dim H_i(\mathcal{M}(I), \mathbb{Z}_2)$, where $H_i(\mathcal{M}(I), \mathbb{Z}_2)$ is the i^{th} homology group of $\mathcal{M}(I)$ with coefficients in $\mathbb{Z}_2$.

Since $\mathcal{M}(I)$ is a connected 4-manifold, $\dim H_0(\mathcal{M}(I), \mathbb{Z}_2) = \dim H_4(\mathcal{M}(I), \mathbb{Z}_2) = 1$.

Since physical qubits are identified with 2-faces of the tessellation, the number of logical qubits $k(I)$ of the quantum code corresponding to $\mathcal{M}(I)$ is $\dim H_2(\mathcal{M}(I), \mathbb{Z}_2)$.

$$\begin{aligned}
k(I) &= \chi(I) + \dim H_1(\mathcal{M}(I), \mathbb{Z}_2) + \dim H_3(\mathcal{M}(I), \mathbb{Z}_2) - 2 \\
&\geq \chi(I) - 2 \\
&\geq \frac{17}{360} n(I) - 2
\end{aligned}$$

This proves that the asymptotic rate of this family of quantum codes is greater than or equal to $\frac{17}{360} \approx 0.0472$.

Note that with the $\{5, 3, 3, 5\}$ tessellation, the lower bound on the asymptotic rate is $\frac{5}{720} \times 26 \approx 0.18$. The other regular tessellations of hyperbolic 4-space yield lower rates.

2.3.4 Lower bound on the minimum distances of the quantum codes

Following Ref. [GL14] we could prove that the minimum distance d asymptotically satisfies $n^\epsilon \leq d \leq n^{0.3}$ for an $\epsilon > 0$. But we will rather follow Ref. [Mur16] and Ref. [Mur17] and derive a tighter lower bound for the minimum distance: $d = \Omega(n^{0.1})$. We will also mention a variant of the construction yielding the even better $d = \Omega(n^{0.2})$. These two lower bounds rely on algebraic arguments.

The first lower bound on the minimum distance is obtained by lower-bounding the trace of matrices of $\tilde{\Gamma}^{\{4,3,3,5\}}(I)$. Indeed this lower bound on the trace of a matrix g then yields a lower bound on the distance between a point $x \in \mathbb{H}^4$ and its image $g \cdot x$. Finally, through Anderson's theorem (Ref. [GL14] th. 17) the size of the smallest homologically nontrivial 2-cycle is exponentially controlled by the size of the smallest homologically nontrivial 1-cycle.

We will start by deriving the lower bound on the trace of a matrix g of $\tilde{\Gamma}^{\{4,3,3,5\}}(I)$. g satisfies the matrix equation $g^T \tilde{J} g = \tilde{J}$ where $\tilde{J} = \text{diag}(-\phi, 1, 1, 1, 1)$. This translates into 10 quadratic equations on the entries of g . We will only need the five equations coming from entries on the diagonal:

$$-\phi g_{0,0}^2 + g_{1,0}^2 + g_{2,0}^2 + g_{3,0}^2 + g_{4,0}^2 = -\phi, \quad (2.2)$$

$$-\phi g_{0,j}^2 + g_{1,j}^2 + g_{2,j}^2 + g_{3,j}^2 + g_{4,j}^2 = 1 \quad \text{for } j \in \{1, \dots, 4\}. \quad (2.3)$$

Denoting by σ the nontrivial embedding of $\mathbb{Q}(\sqrt{5})$ in $\mathbb{C}$ that sends $\sqrt{5}$ to $-\sqrt{5}$ and applying it to these equations yields:

$$\begin{aligned} -\sigma(\phi) \sigma(g_{0,0})^2 + \sigma(g_{1,0})^2 + \sigma(g_{2,0})^2 + \sigma(g_{3,0})^2 + \sigma(g_{4,0})^2 &= -\sigma(\phi), \\ -\sigma(\phi) \sigma(g_{0,j})^2 + \sigma(g_{1,j})^2 + \sigma(g_{2,j})^2 + \sigma(g_{3,j})^2 + \sigma(g_{4,j})^2 &= 1 \quad \text{for } j \in \{1, \dots, 4\}. \end{aligned}$$

Observing that $-\sigma(\phi)$ is positive, we obtain from Eq.(2.2) that $|\sigma(g_{0,0})| \leq 1$. Similarly Eq.(2.3) yields $|\sigma(g_{j,j})| \leq 1$ for $j \in \{1, \dots, 4\}$. Defining for $j \in \{0, \dots, 4\}$, $y_j := g_{j,j} - 1$, we have $|\sigma(y_j)| \leq 2$ for $j \in \{0, \dots, 4\}$. Moreover we can rewrite Eq.(2.2) and Eq.(2.3):

$$-2\phi y_0 - \phi y_0^2 + g_{1,0}^2 + g_{2,0}^2 + g_{3,0}^2 + g_{4,0}^2 = 0, \quad (2.4)$$

$$2y_j + y_j^2 - \phi g_{0,j}^2 + \sum_{i \neq j} g_{i,j}^2 = 0 \quad \text{for } j \in \{1, \dots, 4\}. \quad (2.5)$$

From Eq.(2.4) and Eq.(2.5) we obtain that $2\phi y_0$ and $2y_j$ for $j \in \{1, \dots, 4\}$ belong to I^2 .

Definition 21. *The norm $N(x)$ of an element x of a number field is the product of its conjugates. For a quadratic field with non trivial embedding σ , $N(x) = x\sigma(x)$.*

Propositon 22 (e.g. [Mur17]). *The absolute value of the norm of an element of an ideal is greater than or equal to the norm of this ideal.*

By multiplication and summation we know that $2\phi(y_0 + y_1 + \dots + y_4)$ belongs to I^2 . Hence,

$$\begin{aligned} |N(2\phi(y_0 + y_1 + \dots + y_4))| &\geq N(I)^2 \\ |N(y_0 + y_1 + \dots + y_4)| &\geq \frac{N(I)^2}{N(2\phi)} \\ &\geq \frac{N(I)^2}{4}. \end{aligned}$$

Therefore,

$$\begin{aligned} |y_0 + y_1 + \dots + y_4| &= \frac{|N(y_0 + y_1 + \dots + y_4)|}{|\sigma(y_0 + y_1 + \dots + y_4)|} \\ &\geq \frac{\frac{N(I)^2}{4}}{|\sigma(y_0)| + |\sigma(y_1)| + \dots + |\sigma(y_4)|} \\ &\geq \frac{N(I)^2}{40}. \end{aligned}$$

Since $\text{tr}(g) = y_0 + y_1 + \dots + y_4 + 5$, we obtain $|\text{tr}(g)| \geq \frac{N(I)^2}{40} - 5$.

Injecting Eq. 2.1, we can lower-bound the absolute value of the trace of a matrix $g \in \Gamma^{\{4,3,3,5\}}(I)$ by the number of physical qubits $n(I)$:

$$|\text{tr}(g)| \geq \frac{1}{2 \times 20^{0.8}} n(I)^{0.2} - 5.$$

We will now define the displacement function of a matrix $g \in \Gamma^{\{4,3,3,5\}}(I)$ acting on $\mathbb{H}^4$ and lower-bound it by $|\text{tr}(g)|$.

Definition 23. *The displacement function ρ of a matrix M acting on a space X is the infimum over $x \in X$ of the distance between x and Mx :*

$$\rho_M = \inf_{x \in X} d(x, Mx).$$

In our case, since the quotient manifold $\mathcal{M}(I)$ is closed and compact, the 1-systole is nothing but the infimum over $g \in \Gamma^{\{4,3,3,5\}}(I)$ of the displacement function of g :

$$1\text{-syst}(\mathcal{M}(I)) = \inf_{g \in \Gamma^{\{4,3,3,5\}}(I)} \rho_g.$$

Observing that both the trace and the displacement function are invariant by conjugation it is easy to prove (see [Mur17] Proposition 6.1.1 p.64):

$$|\text{tr}(g)| \leq 2 \cosh(\rho_g) + 3.$$

Proposition 6.1.1 in [Mur17] only concerns direct isometries of $\mathbb{H}^4$ but its proof generalizes straightforwardly to general isometries of $\mathbb{H}^4$.

Thus,

$$\begin{aligned} 1\text{-syst}(\mathcal{M}(I)) &\geq \rho_g \\ &\geq \ln(|\text{tr}(g)| - 4) \\ &\geq \ln\left(\frac{1}{2 \times 20^{0.8}} n(I)^{0.2} - 9\right) \\ &\geq \ln\left(\frac{n(I)^{0.2} - 18 \times 20^{0.8}}{2 \times 20^{0.8}}\right). \end{aligned}$$

Definition 24. *The injectivity radius of a hyperbolic manifold is the supremum of the radii r such that the restriction of the covering projection $\mathbb{H}^4 \rightarrow \mathcal{M}$ to any ball of radius r is injective.*

The injectivity radius $R(I)$ of the closed compact manifold $\mathcal{M}(I)$ is half its 1-systole:

$$\begin{aligned} R(I) &\geq \frac{1\text{-syst}(\mathcal{M}(I))}{2}, \\ &\geq \ln\left(\left(\frac{n(I)^{0.2} - 18 \times 20^{0.8}}{2 \times 20^{0.8}}\right)^{0.5}\right). \end{aligned}$$

A specific case of Anderson's theorem yields:

Theorem 25 ([GL14] th. 17). *Let M be a closed manifold with a hyperbolic metric. Let Z be a homologically non-trivial 2-cycle with coefficients in $\mathbb{Z}_2$. Let R be the injectivity radius of M . Then the volume of Z is greater than or equal to the volume of a disk of radius R in the hyperbolic plane:*

$$\text{vol}(Z) \geq 2\pi(\cosh(R) - 1).$$

Since every 2-face of the $\{4,3,3,5\}$ tessellation has the same volume v , for every 2-chain C of $\mathcal{M}(I)$ with its tessellation, $\text{vol}(C) = \text{wt}(C) \times v$ where $\text{wt}(C)$ is the number of faces of the chain C . To have a fully explicit result, we will compute the value of v , which is also the area of a square in the regular $\{4,5\}$ tessellation of the hyperbolic plane. We can compute its value thanks to the (2-dimensional) Gauss-Bonnet theorem:

$$\begin{aligned}
v &= -2\pi \chi(\{4, 5\}\text{-square}) \\
&= -2\pi \left(\frac{1}{1} - \frac{4}{2} + \frac{4}{5} \right) \\
&= \frac{2}{5}\pi.
\end{aligned}$$

Applying Theorem 25 to $\mathcal{M}(I)$ yields:

$$D(I) \geq \frac{\pi}{v}(\exp(R(I)) - 2),$$

where $D(I)$ is the minimal distance of the quantum code corresponding to $\mathcal{M}(I)$. This gives the bound on the minimal distance:

$$D(I) \geq \frac{5}{2} \left(\frac{N(I)^{0.2} - 18 \times 20^{0.8}}{2 \times 20^{0.8}} \right)^{0.5} - 5.$$

Asymptotically, $D(I)$ is greater than or equal to $\frac{5}{2^{1.5} \times 20^{0.4}} N(I)^{0.1} \approx 0.53 N(I)^{0.1}$.

Similarly to Ref. [Mur16], we can consider the spin group $\text{Spin}(1, 4)$, which is a double covering of $SO^o(1, 4)$. Defining principal congruence subgroups at the level of the spin group $\text{Spin}(1, 4)$, Murillo shows that the minimum distance d of the corresponding codes satisfies $D = \Omega(n^{0.2})$ [Mur16]. We note that the arithmetic manifolds defined at the level of the spin group are not strictly speaking the same as the ones defined at the level of the indefinite orthogonal group. Indeed the arithmetic subgroups of $\Gamma^{\{4, 3, 3, 5\}}$ by which the hyperbolic 4-space is quotiented are different. To derive the lower bound $N^{0.2}$ on the minimum distance, the whole construction has to be done at the level of the spin group. Doing so does not alter the rate of the family of codes nor its $\{4, 3, 3, 5\}$ local structure. Therefore it does not modify the local decoders designed in Section 2.4. However since using the spin group makes the exposition less intuitive and does not improve the main result qualitatively, we will not state it in the main theorem:

Theorem 26. *There exists a family of homological quantum error correcting codes $[[N, K, D]]$ defined from hyperbolic 4-manifolds equipped with $\{4, 3, 3, 5\}$ tessellations. This family has non-vanishing rate $\frac{K}{N}$ which is asymptotically lower bounded by $17/360$. The minimum distance D of its codes grows at least like $N^{0.1}$.*

2.3.5 Estimates of the number of physical qubits

The family of codes used to state Theorem 26 has the drawback of being sparse. We show now that the smallest value of N corresponding to a proper ideal of $\mathbb{Z}[\phi]$ is 234 000. However there are normal subgroups of $\Gamma^{\{4, 3, 3, 5\}}$ which are not constructed from an ideal of $\mathbb{Z}[\phi]$. Finding such normal subgroups with small index in $\Gamma^{\{4, 3, 3, 5\}}$ would lead to quantum codes with a more reasonable, *i.e.* small enough to be practical, number of physical qubits. Even though the control over the minimum distance is lost when considering non arithmetic normal subgroups, the rate of the family of codes and the local decoders are valid for any normal subgroup. Moreover it could be interesting to use the technique of Ref. [BVC⁺17] to interpolate between arithmetic hyperbolic 4-dimensional codes and *e.g.* Euclidean 4-dimensional codes. This can be done by refining the hyperbolic tessellation by a Euclidean tessellation of the hypercubes.

Since $SO^o(1, 4)$ has dimension 10, the number of hypercubes in the manifold equipped with a $\{4, 3, 3, 5\}$ tessellation $\mathcal{M}(I)$ is proportional to $N(I)^{10}$. Therefore the number of qubits of the quantum error correcting code is also proportional to $N(I)^{10}$. In $\mathbb{Z}[\phi]$, the smallest proper ideals we have found have norm 4, 5, 9, 11. The ideal whose norm is 4 is $2\mathbb{Z}[\phi]$. But we have to ignore this ideal because $r_{0, 2\mathbb{Z}[\phi]}$ is the identity. The ideal whose norm is 5 is $\sqrt{5}\mathbb{Z}[\phi]$. It corresponds to a number of qubits of the order of $5^{10} \approx 10^7$.

More precisely, we can compute an upper bound on the cardinal of $\pi_I(\Gamma^{\{4, 3, 3, 5\}})$ for the ideal $I = \sqrt{5}\mathbb{Z}[\phi]$. Since $\mathbb{Z}[\phi]/(\sqrt{5}\mathbb{Z}[\phi])$ is the finite field $\mathbb{F}_5$, $\pi_{\sqrt{5}\mathbb{Z}[\phi]}(\Gamma^{\{4, 3, 3, 5\}})$ is isomorphic to a subgroup of the orthogonal group with dimension 5 and entries in $\mathbb{F}_5$. Using the result of [Wil09], Sec. 3.7.2 p. 72, we obtain $|\pi_{\sqrt{5}\mathbb{Z}[\phi]}(\Gamma^{\{4, 3, 3, 5\}})| \leq 18\,720\,000$.

Moreover, since $\phi + \sqrt{5}\mathbb{Z}[\phi]$ in $\mathbb{Z}[\phi]/(\sqrt{5}\mathbb{Z}[\phi])$ corresponds to 3 in $\mathbb{F}_5$, we have the following generating set for $\pi_{\sqrt{5}\mathbb{Z}[\phi]}(\Gamma^{\{4, 3, 3, 5\}})$:

$$\begin{aligned} r_{0, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 4 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_{1, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, & r_{2, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \\ r_{3, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 \end{pmatrix}, & r_{4, \mathbb{F}_5} &= \begin{pmatrix} 3 & 0 & 0 & 0 & 4 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 3 & 0 & 0 & 0 & 2 \end{pmatrix}. \end{aligned}$$

Using the software GAP for computational discrete algebra and this set of generators, we find that $|\pi_{\sqrt{5}\mathbb{Z}[\phi]}(\Gamma^{\{4, 3, 3, 5\}})| = 18\,720\,000$ (which implies that $\pi_{\sqrt{5}\mathbb{Z}[\phi]}(\Gamma^{\{4, 3, 3, 5\}})$ is the whole orthogonal group with dimension 5 and entries in $\mathbb{F}_5$). Since $|S_{2, (\sqrt{5}\mathbb{Z}[\phi])}| = 80$, the number $N(\sqrt{5}\mathbb{Z}[\phi])$ of physical qubits of the corresponding code is $18\,720\,000/80 = 234\,000$. Using a computational discrete algebra software like GAP and the coset method, we can compute the parity check matrices of this code.

2.4 Local Decoders

In this section, we design efficient decoding algorithms for the family of codes constructed in the previous section. These decoders are tailored for the whole hyperbolic 4-space equipped with a $\{4, 3, 3, 5\}$ tessellation. Of course we want to apply these decoders to codes with a finite number of physical qubits, *i.e.* to the hyperbolic manifolds $\mathcal{M}(I)$ equipped with a $\{4, 3, 3, 5\}$ tessellation.

In this work we consider arbitrary errors of weight logarithmic in the number of physical qubits. Indeed the injectivity radii of the arithmetic hyperbolic manifolds associated with the golden code family scale logarithmically with their volumes. In terms of decoding, this implies that decoding a number of errors logarithmic in the number of physical qubits is strictly equivalent in the arithmetic hyperbolic manifolds and in the whole hyperbolic 4-space equipped with a $\{4, 3, 3, 5\}$ tessellation. In other words our decoder provably succeeds for any error pattern of weight logarithmic in the number of physical qubits. Second, the same decoder will succeed with high probability to correct random error patterns of weight linear in the number of physical qubits, for instance if the qubits are affected independently by depolarizing noise.

The advantage of our decoders over the generic hyperbolic 4-dimensional decoder by Hastings [Has16] is their high locality. Indeed Hastings' decoder is local at the level of a ball of radius R_{dec} where R_{dec} is constant but unknown. Since in hyperbolic 4-space

the number of 2-faces in a ball of radius R_{dec} grows like $e^{3R_{dec}}$, even small values of R_{dec} can lead to an unpractical degree of locality. For instance the authors of [BDMT16] use the value $R_{dec} = 1.5$ to implement a version of Hastings' decoder in a 4-dimensional toric code setting. With such a small value of R_{dec} the analysis of the performance of Hastings' decoder probably does not apply. The analysis of our decoders, on the other hand, is valid at a level of locality that is computationally practical.

Since the codes we consider are CSS, it is possible to decode X -type and Z -type errors independently, and this is what our algorithm does. Because correcting these two types of errors on a qubit is sufficient to correct an arbitrary single-qubit error, we can state our decoding theorem as follows.

Theorem 27. *There exists a constant C such that for any error E corrupting less than $C \log N$ physical qubits, the decoding algorithm returns a set of qubits E' such that E and E' differ by a sum of stabilizers.*

Since stabilizers act trivially on the codespace, Theorem 27 implies that any codestate corrupted on at most $C \log N$ physical qubits is perfectly recovered by the active error correction procedure.

Moreover, standard results in percolation theory show that for a random error model where each qubit is affected independently and identically with a depolarizing node, then below some constant noise threshold, the error will affect qubits that belong to small connected components of the tessellation of size $O(\log N)$. This is because the tessellation has constant degree. Using the same ideas as in [FGL18b], the decoding algorithm will correct the error with high probability.

Theorem 28. *There exists a constant $p_0 > 0$ such that if each qubit is independently and identically affected by an X or a Z error with probability $p < p_0$, then the decoding algorithm corrects the error with high probability.*

2.4.1 Decoding Z -errors

As mentioned, the algorithm successively decodes Z -errors then X -errors. It succeeds if it recovers the right error patterns, up to some element of the stabilizer group. We first consider Z -errors. A Z -decoder takes as input a syndrome on X -type stabilizers and outputs a set of Z -errors consistent with this syndrome. For golden codes, X -type stabilizers are defined by edges in the $\{4, 3, 3, 5\}$ tessellation. The error pattern is by definition the set of 2-faces corresponding to qubits having a Z -error. The syndrome is the boundary of the error pattern. Since every boundary is a cycle, the syndrome consists of several loops of edges.

Definition 29. *A path of edges from vertex v_1 to vertex v_2 is minimal if no other path of edges from vertex v_1 to vertex v_2 is shorter.*

The Z -decoder follows from following lemma:

Lemma 30. *In the $\{4, 3, 3, 5\}$ tessellation, every loop of edges has at least one subpath of length at most 8 which is not minimal.*

Lemma 30 is proven in the appendix.

With Lemma 30 at hand, it is now easy to design a local decoder:

- From every edge of the syndrome, explore every path of edges in the syndrome of length at most 8.
- If such a path is not minimal, flip qubits to decrease its length.
- Iterate, until no non-minimal path of length at most 8 can be found.

While the complexity of the Z -decoder appears at first sight to be quadratic in the size of the syndrome, it can be made linear if one only explores in the $(i + 1)^{\text{th}}$ step paths that were not already explored during the i -th round of the algorithm. Indeed flipping a qubit only affects a constant number of paths of length at most 8. Moreover, as long as the error weight is below the injectivity radius of the manifold, or if the error consists of many such small connected components, then the syndrome weight is proportional to the error weight. This fact comes from the hyperbolicity of the tessellation. In other words, the decoding algorithm has a complexity linear in the error weight.

2.4.2 Decoding X -errors

We now turn our attention to decoding X -errors. An X -decoder takes as input a syndrome on Z -type stabilizers and outputs a set of X -errors consistent with this syndrome. For golden codes, Z -type stabilizers are defined by polyhedrons (3-faces) in the $\{4, 3, 3, 5\}$ tessellation. It is more convenient for us to work with edges than with polyhedrons. We therefore consider the $\{5, 3, 3, 4\}$ dual tessellation. With this point of view, Z -type stabilizers are defined by edges in the $\{5, 3, 3, 4\}$ dual tessellation.

The X -decoder follows from the following lemma:

Lemma 31. *In the $\{5, 3, 3, 4\}$ tessellation, every loop of edges has at least one subpath incident to a single 4-face and which is not minimal.*

Lemma 31 is proven in the appendix.

With Lemma 31 at hand, it is now easy to design a local decoder:

- From every edge of the syndrome, explore every path of edges in the syndrome incident to a single 4-face.
- If such a path is not minimal, flip qubits to decrease its length.
- Iterate, until no non-minimal path incident to a single 4-face can be found.

The complexity of this X -decoder is linear in the size of the error for the same reason as the Z -decoder.

2.4.3 Proof of the Z -decoder Lemma

Before proving Lemma 30, we first establish a 2-dimensional version of it. Even though this 2-dimensional version is irrelevant to decoding homological quantum codes, it allows us to illustrate the main ideas with figures and may help the reader understand the key role of hyperbolicity in Lemma 30.

Lemma 32. *In the $\{4, 5\}$ tessellation of hyperbolic plane, every loop of edges has at least one subpath of length at most 4 which is not minimal.*

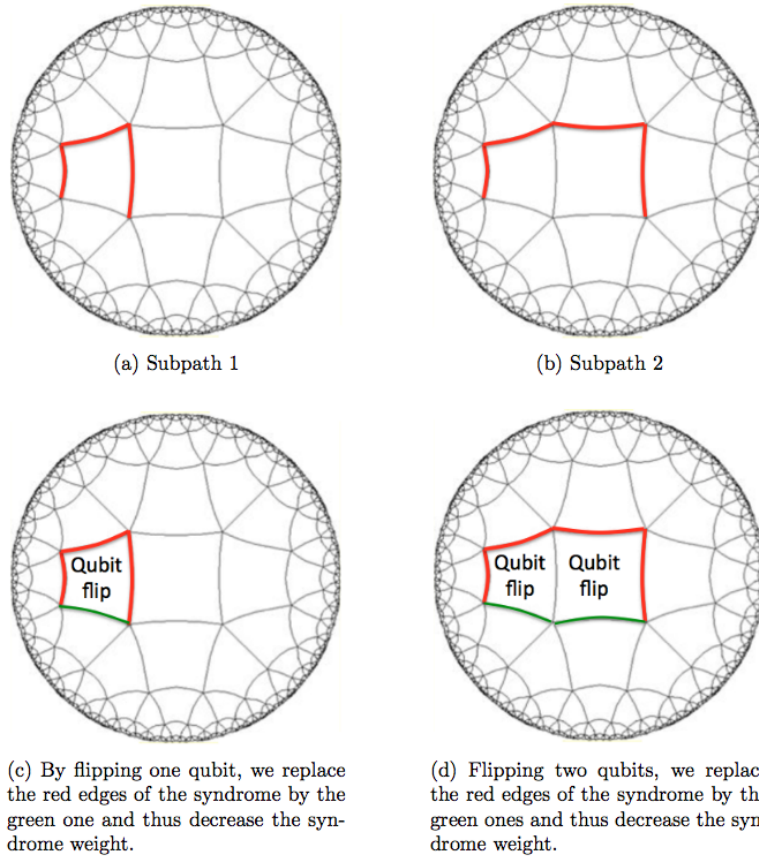


Figure 2.1: In the $\{4, 5\}$ tessellation of hyperbolic plane, every loop of edges contains one of the two subpaths in red. These two subpaths are not minimal: they can be replaced by the shorter green ones by flipping one or two qubits. (source for image: [tes])

Proof. Equivalently, in the $\{4, 5\}$ tessellation of hyperbolic plane every loop of edges admits at least one of the two subpaths depicted on Fig. 2.1.

It is sufficient to prove it on a single loop of edges of the $\{4, 5\}$ tessellation of hyperbolic plane. We choose an arbitrary orientation on this loop. An edge e is written $e = \{v_1, v_2\}$ if it is oriented from v_1 to v_2 . To each edge $e = \{v_1, v_2\}$ we assign a cone C_e defined as the set of points of hyperbolic plane closer to e than to any other edge incident to v_2 . The cone C_e divides the hyperbolic plane in two regions: the outside of the cone and the inside of the cone.

We suppose by contradiction that there exists a loop L of edges in the $\{4, 5\}$ tessellation of hyperbolic plane such that every subpath of L of length at most 4 is minimal. Figure 2.2 shows by an exhaustive search that for any edge e , there exists f in $L \setminus \{e\}$ such that C_f contains C_e . By immediate induction, it is then possible to construct a sequence $(e_i)_{i \in \mathbb{N}}$ of edges in L such that $j > i$ implies that C_{e_j} contains C_{e_i} . This contradicts the fact that L is a loop. $\square$

We can now prove Lemma 30.

of Lemma 30. : It follows the exact same line. To each edge $e = \{v_1, v_2\}$, we assign a cone C_e defined as the set of points of hyperbolic 4-space closer to e than to any other edge incident to v_2 . Since the number of length 8 paths on the edge graph of the $\{4, 3, 3, 5\}$ tessellation is too high to check every case manually, we used a computer program to find that every minimal path of length 8 contains a subpath such that the cone assigned to its last edge contains the cone assigned to its first edge. Therefore in order to form a loop, at least one subpath of length at most 8 has to not be minimal. The decoder consists in flipping qubits in order to shorten this subpath. $\square$

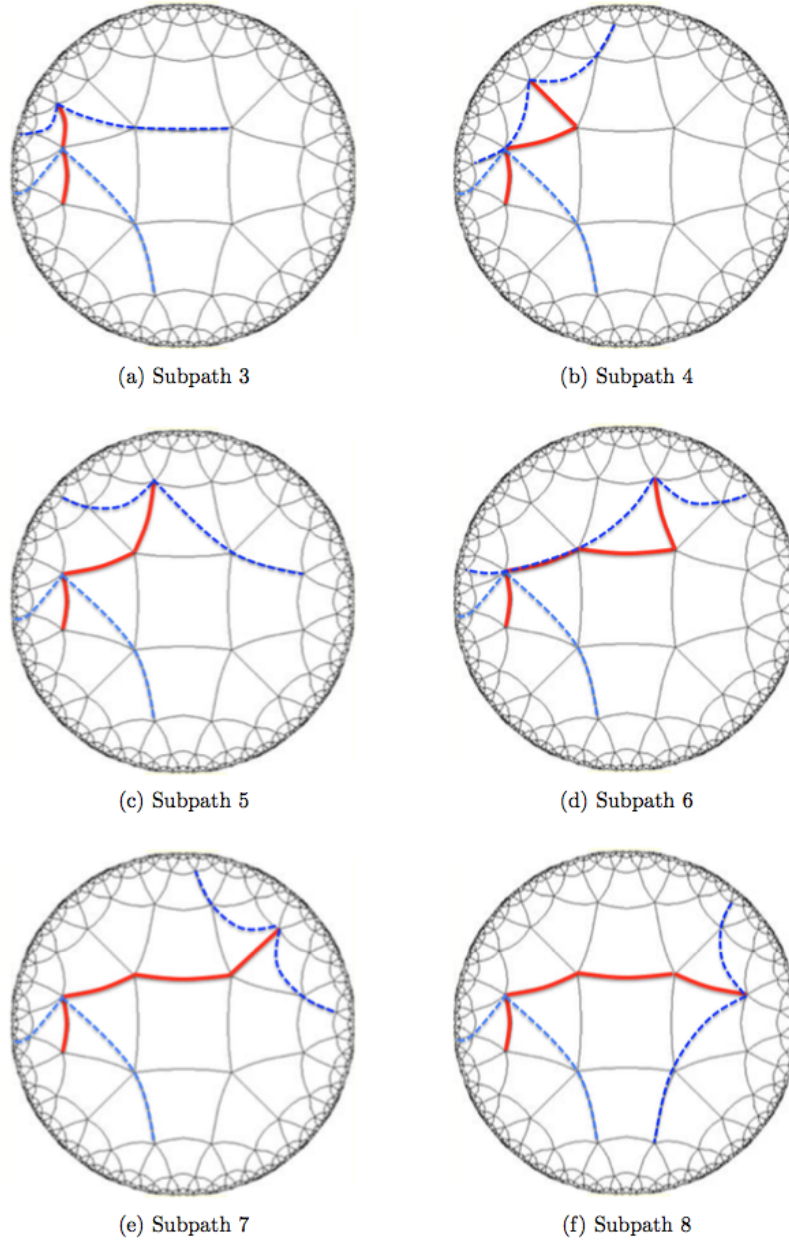


Figure 2.2: The dark blue cone assigned to the last edge of the path contains the light blue cone assigned to the first edge of the path. Every minimal path of length 4 contains one of these six subpaths (or a subpath symmetric to it). Therefore if every length 4 subpath is minimal, it is impossible to form a loop. (source for image: [\[tes\]](#))

2.4.4 Proof of the X -decoder Lemma

Before proving Lemma 31 we prove a 2-dimensional version of it. Even though the 2-dimensional version is irrelevant to decoding homological quantum codes, it allows us to illustrate the main ideas with figures and may help the reader understand the key role of hyperbolicity in Lemma 31.

Lemma 33. *In the $\{5,4\}$ tessellation of hyperbolic plane, every loop of edges admits at least one subpath incident to a single pentagon and which is not minimal.*

Equivalently, in the $\{5,4\}$ tessellation of hyperbolic plane every loop of edges has at least one subpath consisting of three edges incident to the same pentagon. After flipping the qubit corresponding to this pentagon, this subpath of length 3 (or 4) is replaced by a subpath of length 2 (or 1) and thus the syndrome weight is reduced.

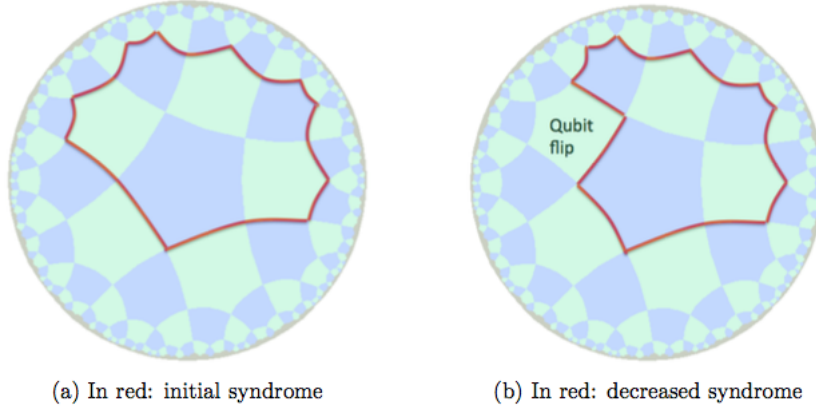


Figure 2.3: Every loop of edges in the $\{5,4\}$ tessellation of hyperbolic plane contains a subpath of three edges incident to the same pentagon. Flipping the qubit corresponding to this pentagon reduces the syndrome weight. (source for image: [tes])

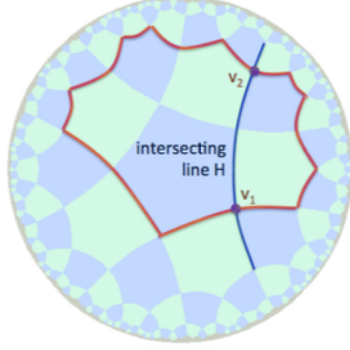
Proof. We consider a loop L of edges in the $\{5,4\}$ tessellation.

As shown in Fig. 2.4(a), there exists a geodesic line H in the $\{5,4\}$ tessellation which intersects the loop L at two of its vertices v_1 and v_2 . v_1 and v_2 define a partition of L into two subpaths. We denote these two subpaths by L_g and L_r . Without loss of generality, assume that the geodesic line H is extremal with respect to L_g in the sense that every edge in L_g is incident to a pentagon incident to an edge of H . This is illustrated in Fig 2.4(b). Without loss of generality, assume that the edge of L_g incident to v_1 does not belong to the extremal geodesic line H .

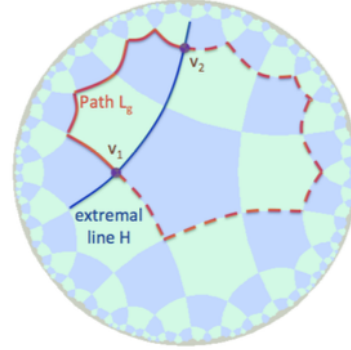
If there exists a pentagon P such that every edge in L_g is incident to P , then L_g is not minimal because the path in the geodesic line H going from v_1 to v_2 consists of a single edge. It is thus shorter than L_g and Lemma 33 is proven in this case.

If such a pentagon P doesn't exist, we denote by w the last vertex of L_g such that every vertex between v_1 and w in L_g is incident to a single pentagon (see Fig. 2.4(c)). We consider the subpath S of L_g going from v_1 to w . S is incident to a single pentagon. It has length 3. We denote by x the vertex of H at edge-distance 1 from w . The path S'

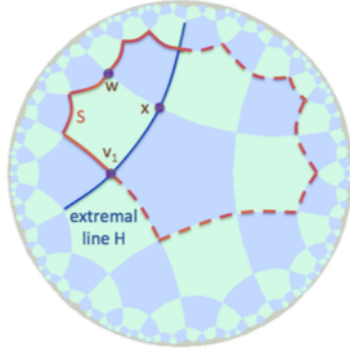
consisting of the edge $\{v_1, x\}$ and the edge $\{x, w\}$ has length 2 (see Fig. 2.4(d)). It is shorter than S . $\square$



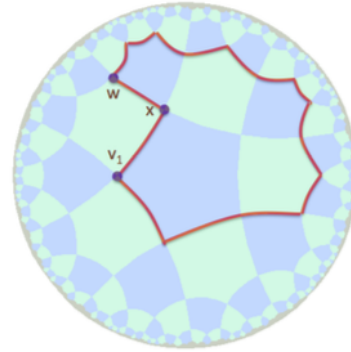
(a) The geodesic line H intersects the loop L at vertices v_1 and v_2 .



(b) The geodesic line H is extremal: every edge of L_g is incident to a pentagon incident to H .



(c) Vertex w is the last vertex of L_g incident to the green pentagon. We define S as the subpath of L_g from v_1 to w and x as the vertex of H at distance 1 from w .



(d) The path going from v_1 to v_2 through x is shorter than S . Path S is incident to a single pentagon.

Figure 2.4: Every loop L in the $\{5, 4\}$ tessellation has a subpath S consisting of three edges incident to the same pentagon. S is not minimal since it can be replaced by a path of length 2. A similar property holds for loops in the $\{5, 3, 3, 4\}$ tessellation. (source for image: [tes])

We are now ready to prove Lemma 31.

of Lemma 31. : The proof is very similar to the proof of Lemma 33. We consider a loop L of edges in the $\{5, 3, 3, 4\}$ tessellation.

As shown in Fig. 2.4(a), there exists a geodesic hyperplane H in the $\{5, 3, 3, 4\}$ tessellation which intersects the loop L at two of its vertices v_1 and v_2 . Vertices v_1 and v_2 define a partition of L into two subpaths. We denote these two subpaths by L_g and L_r . Without loss of generality, assume that the geodesic hyperplane H is extremal with respect to L_g in the sense that every edge in L_g is incident to a 4-face incident to an edge of H . This is illustrated in Fig. 2.4(b). Without loss of generality, assume that the edge of L_g incident

to v_1 does not belong to the extremal geodesic hyperplane H .

If there exists a 4-face P such that every edge in L_g is incident to P , then L_g is not minimal. Indeed the path in the geodesic hyperplane H going from v_1 to v_2 is shorter than L_g and Lemma 31 is proven in this case.

If such a 4-face P doesn't exist, we denote by w the last vertex of L_g such that every vertex between v_1 and w in L_g is incident to a single 4-face (see Fig. 2.4(c)). We consider the subpath S of L_g going from v_1 to w . S is incident to a single 4-face. We denote by x the vertex of H at edge-distance 1 from w . We define S' as one of the shortest path in H going from v_1 to x concatenated with the single edge path going from x to w (see Fig. 2.4(d)). An exhaustive search on the 1-skeleton of a 120-cell shows that S' is always shorter than S . $\square$

2.5 Conclusion of the chapter

In this work, we have presented a variant of the quantum LDPC code family due to Guth and Lubotzky. Like theirs, our family is also obtained by considering tessellations of hyperbolic 4-space, but the crucial new feature of our construction is that the tessellation is regular. We then exploit this regularity to design an efficient and explicit decoding algorithm that provably corrects arbitrary errors of weight $O(\log N)$ and decodes with high probability random independent and identically distributed errors provided the error rate is below some constant threshold.

We note that both the dimension 4 and hyperbolicity present advantages for decoding. Placing the qubits on 2-faces yields syndromes which are cycles of edges (or of coedges) and a decoder should simply try to shorten such cycles, which can be done efficiently by means of a local algorithm as we demonstrated. This algorithm is also more efficient in hyperbolic space since the syndrome weight increases linearly with the error weight (for small errors). This is arguably simpler than pairing vertices as required in 2-dimensional codes. Another advantage of 1-dimensional syndromes is that they contain redundant information, which should be helpful when considering more realistic scenarios where syndrome measurements are not assumed to be ideal.

Chapter 3

The decoding problem for Hyperbolic 4D codes

In this chapter we investigate numerically the decoding of hyperbolic 4D codes. For the numerical simulations to terminate in a reasonable time, we need small instances of such hyperbolic 4D codes. We saw however in Chapter 2 that it was a real challenge to construct arithmetic hyperbolic 4D codes of reasonably small sizes: the smallest we found with the $\{4,3,3,5\}$ local structure has 234 000 physical qubits.

In this chapter we address this issue by focusing on the $\{5,3,3,5\}$ hyperbolic 4D local structure and by considering a novel quotienting technique. We manage therefore to define $\{5,3,3,5\}$ quantum codes with 720, 9792, 18 000 and 90 000 physical qubits. By searching numerically for normal subgroups of the $\{5,3,3,5\}$ Coxeter group, we also define non arithmetic codes with the same $\{5,3,3,5\}$ local structure that have 144, 18432 and 19 584 physical qubits.

We finally study the performance of a Belief Propagation decoder on the codes with 144, 720, 9792, 18 000 and 19584 qubits. The code with 90 000 qubits was not considered because it has too many physical qubits. The one with 18432 qubits also not because it is too close in size to the code with 18 000 qubits.

3.1 Arithmetic construction of $\{5,3,3,5\}$ quantum codes

In Chapter 2 we insisted on the geometric properties of hyperbolic 4D codes. In this chapter we will focus on its algebraic properties. We therefore give more background on Coxeter groups.

3.1.1 Canonical Representation of Coxeter groups

This section follows chapter 3 of Ref. [MS02].

Definition 34. A Coxeter group of rank n is a finitely presented group defined by n generators $R_0, \dots, R_{n-1}$ and $n(n+1)/2$ relations:

$$\forall i \in \{0, \dots, n-1\}, \forall j \in \{i, \dots, n-1\}, (R_i R_j)^{a_{i,j}} = id.$$

$\forall i \in \{0, \dots, n-1\}, a_{i,i} = 1$. The other $a_{i,j}$ parametrise Coxeter groups.¹

The parameters $a_{i,j}$ are often given through the corresponding so-called Coxeter matrix:

$$M = (-2 \cos(\pi/a_{i,j}))_{0 \leq i,j \leq n-1}$$

¹For simplicity, we don't allow $a_{i,j} = \infty$ here.

where $a_{j,i} = a_{i,j}$ by definition.

To this Coxeter matrix we can associate the symmetric bilinear form $g(_, _)$ on $\mathbb{R}^n$:

$$\forall i, j \in \{0, \dots, n-1\}, g(e_i, e_j) = m_{i,j} = -2 \cos(\pi/a_{i,j})$$

The canonical representation of a Coxeter group is defined by mapping its generators $(R_i)_{i \in \{0, \dots, n-1\}}$ to the $\mathbb{R}^n$ endomorphisms $(r_i)_{i \in \{0, \dots, n-1\}}$.

$$\begin{aligned} \forall i, j \in \{0, \dots, n-1\}, r_i(e_j) &\stackrel{\text{def}}{=} e_j - 2 \frac{g(e_i, e_j)}{g(e_i, e_i)} e_i \\ &= e_j - g(e_i, e_j) e_i. \end{aligned}$$

We note $e_i^{\perp_g} = \ker(x \mapsto g(e_i, x))$.

r_i is the reflection in the hyperplane $e_i^{\perp_g}$ along e_i (i.e. mapping e_i to $-e_i$).

Let us verify that the relations defining the Coxeter group are satisfied by this representation.

$$\begin{aligned} \forall i, j \in \{0, \dots, n-1\}, r_i^2(e_j) &= r_i(e_j) - g(e_i, e_j) r_i(e_i) \\ &= e_j - g(e_i, e_j) e_i - g(e_i, e_j) (e_i - g(e_i, e_i) e_i) \\ &= e_j - 2g(e_i, e_j) e_i + 2g(e_i, e_j) e_i \\ &= e_j \end{aligned}$$

Therefore $\forall i \in \{0, \dots, n-1\}, r_i^2 = \text{id}$.

Since $\dim e_i^{\perp_g} = \dim e_j^{\perp_g} = n-1$ and $e_i^{\perp_g} \neq e_j^{\perp_g}$,² $\dim(e_i^{\perp_g} \cap e_j^{\perp_g}) = n-2$.

Since e_i and e_j don't belong to $e_i^{\perp_g} \cap e_j^{\perp_g}$, we can complete a basis of $e_i^{\perp_g} \cap e_j^{\perp_g}$ with e_i and e_j to form a basis $\mathcal{B}$ of $\mathbb{R}^n$. Let us write r_i and r_j in $\mathcal{B}$:

$$\begin{aligned} r_{i,\mathcal{B}} &= \begin{pmatrix} I_{n-2} & 0 & 0 \\ 0 & -1 & g(e_i, e_j) \\ 0 & 0 & 1 \end{pmatrix} & r_{j,\mathcal{B}} &= \begin{pmatrix} I_{n-2} & 0 & 0 \\ 0 & 1 & 0 \\ 0 & g(e_i, e_j) & -1 \end{pmatrix} \\ r_{i,\mathcal{B}} r_{j,\mathcal{B}} &= \begin{pmatrix} I_{n-2} & 0 & 0 \\ 0 & g(e_i, e_j)^2 - 1 & -g(e_i, e_j) \\ 0 & g(e_i, e_j) & -1 \end{pmatrix} \end{aligned}$$

We can focus on the bottom-right two by two submatrix A and compute its trace and determinant:

$$\begin{aligned} \det(A) &= 1. \\ \text{tr}(A) &= g(e_i, e_j)^2 - 2 \\ \text{tr}(A) &= 4 \cos^2(\pi/a_{i,j}) - 2 \\ \text{tr}(A) &= 2 \cos(2\pi/a_{i,j}). \end{aligned}$$

Therefore A has two distinct eigenvalues $\lambda_+ = \exp(i2\pi/a_{i,j})$ and $\lambda_- = \exp(-i2\pi/a_{i,j})$, is similar over $\mathbb{C}$ to the diagonal matrix with entries λ_+ and λ_- and hence satisfies $A^{a_{i,j}} = I_2$.

We have thus proven that $(r_{i,\mathcal{B}} r_{j,\mathcal{B}})^{a_{i,j}} = I_n$.

We refer to [MS02] theorem 3A10 for a proof that this representation is faithful (equivalently the reflections $(r_i)_{i \in \{0, \dots, n-1\}}$ don't satisfy other relations than the ones satisfied by the generators $(R_i)_{i \in \{0, \dots, n-1\}}$ of the Coxeter group).

²true because we have ruled out $a_{i,j} = \infty$

3.1.2 Representation of the $\{5,3,3,5\}$ string Coxeter group

A string Coxeter group is a Coxeter group satisfying $a_{i,j} = 2$ for all $i, j \in \{0, \dots, n-1\}$ such that $|i-j| > 2$. We call such a group the $\{a_{0,1}, \dots, a_{n-2,n-1}\}$ string Coxeter group or simply the $\{a_{0,1}, \dots, a_{n-2,n-1}\}$ Coxeter group.

We start with the Coxeter matrix of the $\{5,3,3,5\}$ string Coxeter group:

$$g = \begin{pmatrix} 2 & -\phi & 0 & 0 & 0 \\ -\phi & 2 & -1 & 0 & 0 \\ 0 & -1 & 2 & -1 & 0 \\ 0 & 0 & -1 & 2 & -\phi \\ 0 & 0 & 0 & -\phi & 2 \end{pmatrix}$$

Indeed:

$$\begin{aligned} -2 \cos(\pi/1) &= 2 & -2 \cos(\pi/2) &= 0 \\ -2 \cos(\pi/3) &= -1 & -2 \cos(\pi/5) &= -\phi \end{aligned}$$

where $\phi = \frac{1+\sqrt{5}}{2}$ is the golden ratio.

From this Coxeter matrix, we define the canonical representation $r_0, \dots, r_4$. In matrix form, in the canonical basis, we obtain:

$$\begin{aligned} r_{0,\{5,3,3,5\}} &= \begin{pmatrix} -1 & \phi & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{1,\{5,3,3,5\}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ \phi & -1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\ r_{2,\{5,3,3,5\}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & -1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{3,\{5,3,3,5\}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & -1 & \phi \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\ r_{4,\{5,3,3,5\}} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & \phi & -1 \end{pmatrix} \end{aligned}$$

We denote $\langle r_{i,\{5,3,3,5\}} \rangle_{i \in \{0, \dots, 4\}}$ by Γ .

Even though we won't actually need it in this chapter, we want to show that Γ is isomorphic to a group of isometries of $\mathbb{H}^4$. This way we know that the asymptotic results of Chapter 2 apply to arithmetic subgroups of Γ .

Note first that elements of Γ preserve the symmetric bilinear form g :

$$\begin{aligned}
\forall i, j, k \in \{0, \dots, 4\}, g(r_i(e_j), r_i(e_k)) &= g(e_j - g(e_i, e_j)e_i, e_k - g(e_i, e_k)e_i) \\
g(r_i(e_j), r_i(e_k)) &= g(e_j, e_k) - g(e_i, e_j)g(e_i, e_k) - g(e_i, e_k)g(e_j, e_i) \\
&\quad + g(e_i, e_j)g(e_i, e_k)g(e_i, e_i) \\
g(r_i(e_j), r_i(e_k)) &= g(e_j, e_k) - 2g(e_i, e_j)g(e_i, e_k) + 2g(e_i, e_j)g(e_i, e_k) \\
g(r_i(e_j), r_i(e_k)) &= g(e_j, e_k).
\end{aligned}$$

The following Lemma shows that the signature of a quadratic form determines its isometry group up to isometry:

Lemma 35. *Let Q_1 and Q_2 be two real symmetric matrices having the same signature. Let $\mathcal{G}_1$ (respectively $\mathcal{G}_2$) be the group of matrices preserving Q_1 (respectively Q_2):*

$$\mathcal{G}_i = \{M \in M_n(\mathbb{R}) \mid M^T Q_i M = Q_i\}$$

Then $\mathcal{G}_1$ and $\mathcal{G}_2$ are isomorphic.

Proof. Without loss of generality, we can assume that $Q_2 = \text{diag}(\epsilon_1, \dots, \epsilon_n)$ with $\epsilon_i \in \{-1, 1\}$. Since Q_1 is real symmetric, there is an orthogonal matrix O such that $Q_1 = O^T D O$ where $D = \text{diag}(d_1, \dots, d_n)$ is a diagonal matrix. Since Q_1 and Q_2 have the same signature we can assume that for all i , $d_i = \epsilon_i a_i^2$, with $a_i \in \mathbb{R}$. Defining $D_a = \text{diag}(a_1, \dots, a_n)$, we have:

$$Q_1 = O^T D_a Q_2 D_a O.$$

Since $D_a = D_a^T$, a simple computation finishes the proof:

$$\begin{aligned}
M^T Q_1 M &= Q_1 \\
\iff M^T O^T D_a Q_2 D_a O M &= O^T D_a Q_2 D_a O \\
\iff D_a^{-1} O M^T O^T D_a Q_2 D_a O M O^T D_a^{-1} &= Q_2 \\
\iff \tilde{M}^T Q_2 \tilde{M} &= Q_2 \quad \text{with } \tilde{M} \stackrel{\text{def}}{=} D_a O M O^T D_a^{-1}.
\end{aligned}$$

The isomorphism is thus given by $M \mapsto D_a O M O^T D_a^{-1}$. □

To prove that Γ is isomorphic to a group of isometries, it remains to show that the quadratic form defined by g has the same signature as the Lorentzian quadratic form (denoted J in chapter 2), namely $\{-1, 1, 1, 1, 1\}$ (also denoted $(1, 4)$). Let us compute the characteristic polynomial of g :

$$\begin{aligned}
\chi_g &= \det \begin{pmatrix} X-2 & \phi & 0 & 0 & 0 \\ \phi & X-2 & 1 & 0 & 0 \\ 0 & 1 & X-2 & 1 & 0 \\ 0 & 0 & 1 & X-2 & \phi \\ 0 & 0 & 0 & \phi & X-2 \end{pmatrix} \\
\chi_g &= (X-2) \det \begin{pmatrix} X-2 & 1 & 0 & 0 \\ 1 & X-2 & 1 & 0 \\ 0 & 1 & X-2 & \phi \\ 0 & 0 & \phi & X-2 \end{pmatrix} - \phi \det \begin{pmatrix} \phi & 0 & 0 & 0 \\ 1 & X-2 & 1 & 0 \\ 0 & 1 & X-2 & \phi \\ 0 & 0 & \phi & X-2 \end{pmatrix} \\
\chi_g &= (X-2)^2 \det \begin{pmatrix} X-2 & 1 & 0 \\ 1 & X-2 & \phi \\ 0 & \phi & X-2 \end{pmatrix} - (X-2) \det \begin{pmatrix} X-2 & \phi \\ \phi & X-2 \end{pmatrix} \\
&\quad - \phi^2 \det \begin{pmatrix} X-2 & 1 & 0 \\ 1 & X-2 & \phi \\ 0 & \phi & X-2 \end{pmatrix} \\
\chi_g &= ((X-2)^2 - \phi^2)((X-2)((X-2)^2 - \phi^2) - (X-2)) - (X-2)((X-2)^2 - \phi^2) \\
\chi_g &= ((X-2)^2 - \phi^2)(X-2)((X-2)^2 - \phi^2 - 2) \\
\chi_g &= (X-2-\phi)(X-2+\phi)(X-2)(X-2-\sqrt{2+\phi^2})(X-2+\sqrt{2+\phi^2})
\end{aligned}$$

Therefore the spectrum of g is $Sp_g = \{2 - \sqrt{2+\phi^2}, 2 - \phi, 2, 2 + \phi, 2 + \sqrt{2+\phi^2}\}$ and each eigenspace has dimension 1. We have proven that g has signature $(1, 4)$ and thus that Γ is isomorphic to a group of isometries of hyperbolic 4-space.

3.1.3 Abstract regular polytopes and string C-groups

This paragraph follows closely sections 2A and 2E of [MS02].

An abstract polytope of rank n is a partially ordered set with properties (P1), ..., (P4) below. The elements of $\mathcal{P}$ are called its faces. Two faces F and G are said to be incident if $F \leq G$ or $G \leq F$. Totally ordered subsets of $\mathcal{P}$ are called chains of $\mathcal{P}$. Maximal chains are called flags of $\mathcal{P}$.

(P1) $\mathcal{P}$ contains a least face and a greatest face; they are denoted by F_{-1} and F_n respectively.

(P2) Each flag of $\mathcal{P}$ has length $n+1$ (that is, contains exactly $n+2$ faces including F_{-1} and F_n).

$\mathcal{P}$ is said to be connected if for any two faces F and G of $\mathcal{P}$ there exists $k \in \mathbb{N}$ and a sequence $(H_t)_{t \in \{0, \dots, k\}}$ of faces of $\mathcal{P}$ such that $H_0 = F$, $H_k = G$ and for all $t \in \{0, \dots, k-1\}$, H_t is incident to H_{t+1} .

For any two incident faces F and G of $\mathcal{P}$ with $F \leq G$, we call

$$G/F \stackrel{\text{def}}{=} \{H \in \mathcal{P} \mid F \leq H \leq G\}$$

a section of $\mathcal{P}$. $\mathcal{P}$ is said to be strongly connected if each of its sections (including itself) is connected.

(P3) $\mathcal{P}$ is strongly connected.

(P4) For all $i \in \{0, \dots, n-1\}$, if F and G are incident faces of $\mathcal{P}$, of ranks $i-1$ and $i+1$ respectively, then there are exactly two i -faces H of $\mathcal{P}$ such that $F < H < G$.

Note that in terms of quantum error correcting codes, (P4) implies that parity-check matrices defined as $(i, i-1)$ and $(i, i+1)$ face incidences are orthogonal.

The automorphism group of an abstract polytope is the permutation group of its faces preserving their ranks and their incidences.

An abstract polytope is called regular if its automorphism group is transitive on its set of flags. Note that this is a stronger condition than being transitive on the set of faces of rank i for all $i \in \{0, \dots, n-1\}$.

Abstract regular polytopes are in one to one correspondence with string C-groups. After defining string C-groups, we will show how we can construct a string C-group from an abstract regular polytope and conversely.

Let G be a group generated by involutions $(\rho_i)_{i \in \{0, \dots, n-1\}}$. G is called a C-group if it has the intersection property with respect to these generators, namely if for all $I, J \subseteq \{0, \dots, n-1\}$,

$$\langle \rho_i \mid i \in I \rangle \cap \langle \rho_j \mid j \in J \rangle = \langle \rho_k \mid k \in I \cap J \rangle.$$

A C-group is called a string C-group if its generators satisfy the relations:

$$(\rho_i \rho_j)^2 = \text{id} \quad \text{for all } i, j \in \{0, \dots, n-1\} \text{ such that } |i - j| \geq 2.$$

Given an abstract regular polytope $\mathcal{P}$ of rank n , we consider its automorphism group $\Gamma_{\mathcal{P}} = \text{Aut}(\mathcal{P})$. Γ has a natural action on the set of flags of $\mathcal{P}$ and since $\mathcal{P}$ is regular, we know that this action is transitive.

Let us choose a base flag $\mathcal{F}$ of $\mathcal{P}$. For any $i \in \{0, \dots, n-1\}$, there exists a unique flag such that for all $j \in \{0, \dots, n-1\}$, its j -faces are the same as $\mathcal{F}$'s j -faces except for $i = j$. We denote this flag by $\mathcal{F}_i$. We denote by ρ_i the element of Γ sending $\mathcal{F}$ to $\mathcal{F}_i$. ρ_i is unique because of the diamond-shape property (P4). For the same reason, $\rho_i^2 = \text{id}$ for all $i \in \{0, \dots, n-1\}$. Γ is generated by $(\rho_i)_{i \in \{0, \dots, n-1\}}$ because of the strong-connectedness property (P3). $(\rho_i \rho_j)^2 = \text{id}$ for all $i, j \in \{0, \dots, n-1\}$ such that $|i - j| \geq 2$ because the new face F_i only depends on F_{i-1} and F_{i+1} and $|i - j| \geq 2$ implies that $j \neq i+1$, $j \neq i-1$, $i \neq j+1$ and $i \neq j-1$. For more detailed proofs of the above statements and of the fact that the automorphism group Γ is a string C-group we refer to [MS02] section 2B.

From a string C-groups of rank n (*i.e.* with n generators) and its preferred generators $(\rho_i)_{i \in \{0, \dots, n-1\}}$, we consider the subgroups S_i generated by the $(n-1)$ generators $(\rho_j)_{j \neq i}$. We consider the abstract polytope $\mathcal{P}_{\Gamma}$ whose set of i -faces is the set of S_i left cosets in G : $\{gS_i \mid g \in \Gamma\}$. By definition, two faces $F_a = g_a S_i$ and $F_b = g_b S_j$ are incident if the corresponding cosets have a non empty intersection: $g_a S_i \cap g_b S_j \neq \emptyset$.

We refer to [MS02] section 2E for a proof that the abstract polytope defined from a string C-group is regular and that the string C-group is its automorphism group.

Therefore the two maps defined above (from abstract regular polytopes to string C-groups and conversely) are each other's inverse and we can think of abstract regular polytopes and string C-groups as the same object. We will use the notations $\mathcal{P}_{\Gamma}$ and $\Gamma_{\mathcal{P}}$ in the sequel to go back and forth between these equivalent objects.

3.1.4 Normal quotients of Coxeter groups are C-groups

We refer to [Bou07] chapter 4 §1.8 theorem 2 for a proof that Coxeter groups satisfy the intersection property. Coxeter groups are therefore C-groups and the name C-groups is now justified.

We will now see that the quotient of a Coxeter group by one of its normal subgroup N is a C-group. Let $I, J \subset \{0, \dots, n-1\}$. S_I and S_J are the corresponding subgroups of the Coxeter group Γ .

$\tilde{S}_I$ and $\tilde{S}_J$ are the corresponding subgroups of the quotient group $N \backslash \Gamma$:

$$\begin{aligned}\tilde{S}_I &= \langle [r_i]_N \mid i \in I \rangle \\ \tilde{S}_I &= N \backslash \langle r_i \mid i \in I \rangle.\end{aligned}$$

where $[r_i]_N$ denotes the equivalence class of r_i in $N \backslash \Gamma$. Note that quotienting on the left or on the right by a normal subgroup doesn't make any difference. We use the less common notation $N \backslash \Gamma$ here in order to be consistent with the next section, where we will consider quotients of Γ by some of its not necessarily normal subgroup H .

$$\begin{aligned}\tilde{S}_I \cap \tilde{S}_J &= (N \backslash \langle r_i \mid i \in I \rangle) \cap (N \backslash \langle r_j \mid j \in J \rangle) \\ \tilde{S}_I \cap \tilde{S}_J &= N \backslash (\langle r_i \mid i \in I \rangle \cap \langle r_j \mid j \in J \rangle) \\ \tilde{S}_I \cap \tilde{S}_J &= \langle [r_k]_N \mid k \in I \cap J \rangle.\end{aligned}$$

This shows that the intersection property is invariant under quotienting by N . Therefore quotients by normal subgroups correspond to abstract regular polytopes.

However the type of the C-group may not be the one of the Coxeter group it comes from. For instance if a generator is sent to the identity. Or if the product of two generators has a smaller order in the quotient than in the Coxeter group. In the next paragraph we give a condition sufficient to avoid such pathological cases.

3.1.5 The non-local subgroup condition

Before we state the non-local subgroup condition, we want to generalize the discussion to quotients of Γ by one of its subgroup H , which is not necessarily normal. Even though the quotient $H \backslash \Gamma$ does not in general have a group structure when H is not normal, it is still possible to define the abstract polytope associated to $H \backslash \Gamma$ as follows:

The orbit of an i -face $F_a = g_a S_i$ of $\mathcal{P}_\Gamma$ under the left action of H is $\{h g_a S_i \mid h \in H\}$. By definition this orbit is a face of the quotient abstract polytope $\mathcal{P}_{H \backslash \Gamma} \stackrel{\text{def}}{=} H \backslash (\mathcal{P}_\Gamma)$. We denote it by $H F_a$. In terms of elements of Γ , it corresponds to the double coset $H g_a S_i$. We use the usual incidence definition: $H F_a$ and $H F_b$ are incident if $H g_a S_i \cap H g_b S_j \neq \emptyset$. Let us stress that when H is not a normal subgroup, $H \backslash \Gamma$ is not a group and $\mathcal{P}_{H \backslash \Gamma}$ is not a regular abstract polytope but a mere abstract polytope. However we are about to show that the non-local subgroup condition is sufficient to define a CSS quantum code from $\mathcal{P}_{H \backslash \Gamma}$.

We want to find a condition under which quotienting on the right by S_i , $i \in \{0, \dots, 4\}$ "doesn't interact" with quotienting on the left by H :

Definition 36. A subgroup H of a C-group G satisfies the non-local subgroup condition if:

$$\forall i, j \in \{0, \dots, 4\}, \forall g \in \Gamma, g H g^{-1} \cap S_i S_j = \{id\}. \quad (3.1)$$

The term non-local refers to the subgroup H : since the subgroups S_i , $i \in \{0, \dots, 4\}$ are "local" (with respect for instance to the distance in the Caley graph $(\Gamma, (r_i)_{i \in \{0, \dots, 4\}})$, the subgroup H has to be non-local in order to not interact with the S_i , $i \in \{0, \dots, 4\}$.

The non-local subgroup condition is sufficient to prove the following lifting Lemma:

Lemma 37 (lifting of S_i cosets). *Let H be a subgroup of a string C-group Γ of rank n satisfying the non-local subgroup condition 3.1. For $i, j \in \{0, \dots, n\}$, let $H F_i$ be an i -face of $H \backslash \mathcal{P}_\Gamma$ and let $H F_j$ be a j -face of $H \backslash \mathcal{P}_\Gamma$ incident to $H F_i$. For any i -face K_i of $\mathcal{P}$ such that $H K_i = H F_i$, there exists a unique face K_j of $\mathcal{P}$ such that $H K_j = H F_j$ and K_j is incident to K_i .*

Proof. There exists $g_i \in \Gamma$ such that $K_i = g_i S_i$. Then $HK_i = HF_i = Hg_i S_i$. There exists $g_j \in \Gamma$ such that $HF_j = Hg_j S_j$. There exist $h_i, h_j \in H$, $s_i \in S_i$ and $s_j \in S_j$ such that $h_i g_i s_i = h_j g_j s_j$. Therefore

$$g_i s_i = h_i^{-1} h_j g_j s_j \quad (3.2)$$

We can define $K_j = h_i^{-1} h_j g_j S_j$. Clearly $HK_j = HF_j$ and K_j is incident to K_i .

To prove unicity, suppose that a face L_j of $\mathcal{P}$ satisfies $HL_j = HF_j$ and L_j is incident to K_i . There exists $g \in \Gamma$ such that $L_j = g S_j$.

Since $Hg S_j = Hg_j S_j$, there exist $h \in H$ and $s'_j \in S_j$ such that

$$g = h g_j s'_j \quad (3.3)$$

Since L_j is incident to K_i , there exists $s'_i \in S_i$ and $s'_j \in S_j$ such that $g s'_j = g_i s'_i$.

Injecting 3.2, we obtain $g s'_j = h_i^{-1} h_j g_j s_j s_i^{-1} s'_i$.

Injecting 3.3, we have $h g_j s'_j s'_j = h_i^{-1} h_j g_j s_j s_i^{-1} s'_i$.

We can rewrite this as $s_j^{-1} g_j^{-1} h_j^{-1} h_i h g_j s_j = s_i^{-1} s'_i (s'_j)^{-1} (s''_j)^{-1} s_j$.

Defining $\bar{g} = g_j s_j$, $\bar{h} = h_j^{-1} h_i h$, $\bar{s}_i = s_i^{-1} s'_i$ and $\bar{s}_j = (s'_j)^{-1} (s''_j)^{-1} s_j$, we have

$$\bar{g}^{-1} \bar{h} \bar{g} = \bar{s}_i \bar{s}_j.$$

Using the non-local subgroup condition 3.1, it implies that $\bar{g}^{-1} \bar{h} \bar{g} = \text{id}$ and therefore that $\bar{h} = \text{id}$.

We have proven that $h = h_i^{-1} h_j$, which means that $L_j = K_j$. $\square$

Thus under the non-local subgroup condition 3.1 we can use the lifting Lemma 37 to prove that the orthogonality of parity-check matrices is preserved by such quotients:

Let HF_{i-1} be an $(i-1)$ -face and HF_{i+1} be an $(i+1)$ -face of the quotient polytope. Let $\{\tilde{F}_{i_1}, \dots, \tilde{F}_{i_m}\}$ be the set of i -faces incident to both HF_{i-1} and HF_{i+1} . Under the non-local subgroup condition 3.1, Lemma 37 shows that there exist faces of the covering polytope K_{i-1} covering HF_{i-1} , K_{i+1} covering HF_{i+1} and $\forall k \in \{1, \dots, n\}, K_{i_k}$ covering HF_{i_k} such that $\{K_{i_1}, \dots, K_{i_m}\}$ is the set of i -faces incident to both K_{i-1} and K_{i+1} . Note that we say that a face K of $\mathcal{P}_\Gamma$ covers a face HF of $\mathcal{P}_{H \setminus \Gamma}$ if and only if $HK = HF$.

Parity check matrices are orthogonal if and only if m is even for every pair of an $(i-1)$ -face and an $(i+1)$ -face. Orthogonality of parity-check matrices is thus clearly preserved by quotient by subgroups satisfying the non-local condition.

We have generalized in this section the discussion to quotients by non normal subgroups because such quotients are useful to find small quantum codes with a given local structure. In this chapter we even consider repeated quotients: first by a normal subgroup N of Γ and second by a not necessarily normal subgroup H of Γ/N . To find normal subgroups of Γ we use the arithmetic method exposed in Chapter 2 and which we now summarize.

3.1.6 Reducing Γ modulo a prime ideal I of $\mathbb{Z}[\phi]$.

Recall the definitions given in Chapter 2:

$$\begin{aligned} \pi_I : M_n(B) &\rightarrow M_n(B/I) \\ (b_{i,j}) &\mapsto (b_{i,j} + I). \end{aligned}$$

$$\Gamma(I) = \ker(\pi_I : \Gamma \rightarrow M_n(\mathbb{Z}[\phi]/I))$$

The following short sequence $1 \rightarrow \Gamma(I) \rightarrow \Gamma \rightarrow \pi_I(\Gamma) \rightarrow 1$ is therefore exact and by the first isomorphism theorem,

$$\pi_I(\Gamma) \simeq \Gamma/\Gamma(I).$$

In this case, the subgroup $N = \Gamma(I)$ by which we are quotienting is the kernel of the group homomorphism π_I and is thus normal. For a normal subgroup, we can rephrase the non-local subgroup condition 3.1 as:

$$\forall i, j \in \{0, \dots, 4\}, \ker(\pi_I) \cap S_i S_j = \{\text{id}\} \quad (3.4)$$

In other words, the non-local subgroup condition is satisfied if and only if for all $i, j \in \{0, \dots, 4\}$, π_I restricted to $S_i S_j$ is a bijection. Since S_i are finite groups, this is a condition that we can check numerically, or even manually in most cases. We now detail the quotients corresponding to the two smallest norms ideals of $\mathbb{Z}[\phi]$: $2\mathbb{Z}[\phi]$ and $\sqrt{5}\mathbb{Z}[\phi]$. We also verify that the next smallest norm ideal would yield a quantum codes with too many qubits for numerical simulations.

3.1.6.1 Reducing modulo $I = 2\mathbb{Z}[\phi]$ gives matrices over $\mathbb{F}_4$ and a code with 9 792 qubits.

We have to understand how $\mathbb{Z}[\phi]$ projects onto $\mathbb{F}_4$ when quotienting by $2\mathbb{Z}[\phi]$. Let us write the four elements of $\mathbb{F}_4$: 0, 1, ω and $\omega + 1$. ω is a root of the irreducible degree 2 polynomial $X^2 + X + 1$. The multiplication table is as follows:

*	0	1	ω	$\omega+1$
0	0	0	0	0
1	0	1	ω	$\omega+1$
ω	0	ω	$\omega+1$	1
$\omega+1$	0	$\omega+1$	1	ω

-1 and 1 project on 1 in $\mathbb{F}_4$.

We want to determine on which element of $\mathbb{F}_4$ ϕ projects. The equation $\phi^2 - \phi - 1 = 0$ in $\mathbb{Z}[\phi]$ becomes $\phi^2 + \phi + 1 = 0$ in $\mathbb{F}_4$. Therefore ϕ projects on either ω or $\omega + 1$. Since exchanging ω and $\omega + 1$ is an isomorphism of $\mathbb{F}_4$, both projections are correct and we can choose arbitrarily. We choose ω to be the image of ϕ in $\mathbb{F}_4$ in the sequel.

Taking $(r_{i,\{5,3,3,5\}})_{i \in \{0, \dots, 4\}}$ modulo $2\mathbb{Z}[\phi]$, we obtain matrices with entries in $\mathbb{F}_4$:

$$\begin{aligned}
r_{0,\mathbb{F}_4} &= \begin{pmatrix} 1 & \omega & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{1,\mathbb{F}_4} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ \omega & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\
r_{2,\mathbb{F}_4} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{3,\mathbb{F}_4} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 & \omega \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\
r_{4,\mathbb{F}_4} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & \omega & 1 \end{pmatrix}
\end{aligned}$$

The group $G_{\mathbb{F}_4} = \langle r_{i,\mathbb{F}_4} \rangle_{i \in \{0, \dots, 4\}} = \pi_I(\Gamma)$ for $I = 2\mathbb{Z}[\phi]$ has order 979 200. This was verified with the computer algebra system GAP. Identifying qubits with 2-faces, it corresponds to a quantum code with $979\,200 / 100 = 9\,792$ physical qubits. Indeed $|S_2| = 100$. We checked numerically that the non-local subgroup condition 3.4 holds. It is thus a valid quantum code. We found numerically that this code has 2 220 logical qubits.

3.1.6.2 Reducing modulo $I = \sqrt{5}\mathbb{Z}[\phi]$ gives matrices over $\mathbb{F}_5$ and a code with 90 000 qubits.

In $\mathbb{F}_5$, only 3 satisfies $x^2 - x - 1 = 0$. Therefore ϕ modulo $\sqrt{5}\mathbb{Z}[\phi]$ is 3 in $\mathbb{F}_5$.

$$\begin{aligned}
r_{0,\mathbb{F}_5} &= \begin{pmatrix} 4 & 3 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{1,\mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 3 & 4 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\
r_{2,\mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & 4 & 1 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{3,\mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 4 & 3 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \\
r_{4,\mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 3 & 4 \end{pmatrix}
\end{aligned}$$

The group $G_{\mathbb{F}_5} = \langle r_{i,\mathbb{F}_5} \rangle_{i \in \{0, \dots, 4\}} = \pi_I(\Gamma)$ for $I = \sqrt{5}\mathbb{Z}[\phi]$ has order 9 000 000. It corresponds to a quantum code with $9\,000\,000 / 100 = 90\,000$ physical qubits. We checked numerically that the non-local subgroup condition 3.4 holds. It is thus a valid quantum code. We found numerically that this code has 18 024 logical qubits.

3.1.6.3 Reducing modulo $I = 3\mathbb{Z}[\phi]$ gives matrices over $\mathbb{F}_9$ and a code with 34 432 128 qubits

This number of physical qubits is too large for numerical simulations. $2\mathbb{Z}[\phi]$, $\sqrt{5}\mathbb{Z}[\phi]$ and $3\mathbb{Z}[\phi]$ are the three ideals of $\mathbb{Z}[\phi]$ with the smallest norms. Therefore other ideals of $\mathbb{Z}[\phi]$ would lead to quantum codes with even larger numbers of physical qubits and we don't consider them. However we show in the next section that we can define other quantum codes with the $\{5,3,3,5\}$ local structure and a small number of physical qubits by quotienting quotient groups.

3.1.7 Subgroups of $G_{\mathbb{F}_5}$ satisfying the non-local condition.

It is possible to quotient a code which was already defined as a quotient. This is what we do with the code corresponding to the group $G_{\mathbb{F}_5}$, *i.e.* to the ideal $\sqrt{5}\mathbb{Z}[\phi]$.

Following Ref. [AS⁺92] we define a new basis $\mathcal{B} = (b_i)_{i \in \{0, \dots, 4\}}$ from the canonical basis $\mathcal{B}_c = (e_i)_{i \in \{0, \dots, 4\}}$ by:

$$\begin{aligned} b_0 &= 2e_0 + 4e_1 + e_3 + 4e_4 \\ b_1 &= 4e_0 + e_1 + 4e_3 + 2e_4 \\ b_2 &= 3e_0 + 2e_1 + 2e_3 + 3e_4 \\ b_3 &= e_2 \\ b_4 &= 4e_0 + e_1 + 4e_3 + e_4 \end{aligned}$$

We can change basis and give $(r_{i, \mathcal{B}, \mathbb{F}_5})_{i \in \{0, \dots, 4\}}$ in the new basis $\mathcal{B}$:

$$\begin{aligned} r_{0, \mathcal{B}, \mathbb{F}_5} &= \begin{pmatrix} 4 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 3 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{1, \mathcal{B}, \mathbb{F}_5} &= \begin{pmatrix} 4 & 0 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 \\ 2 & 0 & 1 & 4 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 2 & 0 & 0 & 4 & 1 \end{pmatrix} \\ r_{2, \mathcal{B}, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 4 & 4 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} & r_{3, \mathcal{B}, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 4 & 0 & 1 & 0 \\ 0 & 2 & 1 & 4 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 3 & 0 & 1 & 1 \end{pmatrix} \\ r_{4, \mathcal{B}, \mathbb{F}_5} &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 4 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 2 & 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

Working in this new basis, we consider two subgroups of $G_{\mathbb{F}_5} = \langle (r_{i, \mathcal{B}, \mathbb{F}_5})_{i \in \{0, \dots, 4\}} \rangle$ defined in Ref. [AS⁺92]:

$$\begin{aligned}
h_1 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 3 & 0 & 0 & 0 & 1 \end{pmatrix} & h_2 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 3 & 0 & 0 & 2 & 1 \end{pmatrix} \\
h_3 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 2 & 0 & 0 & 1 \end{pmatrix} & h_4 &= \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 2 & 0 & 3 & 1 \end{pmatrix}
\end{aligned}$$

$$H_{125} = \langle h_1, h_2, h_3, h_4 \rangle.$$

$$h = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 3 & 1 \end{pmatrix}$$

$$H_5 = \langle h \rangle.$$

We checked numerically that H_5 and H_{125} both satisfy the non-local subgroup condition 3.1. Therefore we obtain a valid quantum code when we quotient the code with 90 000 qubits constructed from $G_{\mathbb{F}_5}$ by these subgroups. Since H_5 has order 5 in $G_{\mathbb{F}_5}$, it yields a code with 18 000 qubits. We found numerically that this code has 3 624 logical qubits. Since H_{125} has order 125 in $G_{\mathbb{F}_5}$, it yields a code with 720 qubits. We found numerically that this code has 184 logical qubits.

We restricted our attention to the two subgroups H_5 and H_{125} because they had already been studied in Ref. [AS⁺92]. However there are possibly many other subgroups of $G_{\mathbb{F}_5}$ satisfying the non-local subgroup condition 3.1. Each of them would yield another quantum code.

3.1.8 A note on normal subgroups and regular codes

Let $\mathcal{P}$ be a regular abstract polytope and Γ be its automorphism group. For a general subgroup H of $\Gamma = \text{Aut}(\mathcal{P})$, we have: $\text{Aut}(H \backslash \mathcal{P}) \simeq H \backslash N_{\text{Aut}(\mathcal{P})}(H)$ [MS02]. $N_{\text{Aut}(\mathcal{P})}(H)$ is the normalizer of H in $\text{Aut}(\mathcal{P})$, i.e. the largest subgroup of $\text{Aut}(\mathcal{P})$ such that H is normal in this group.

If H is normal in $\text{Aut}(\mathcal{P})$, $N_{\text{Aut}(\mathcal{P})}(H) = \text{Aut}(\mathcal{P})$ and therefore $\text{Aut}(H \backslash \mathcal{P}) \simeq H \backslash \text{Aut}(\mathcal{P})$.

If H is not normal in $\text{Aut}(\mathcal{P})$, $N_{\text{Aut}(\mathcal{P})}(H)$ is a proper subgroup of $\text{Aut}(\mathcal{P})$ and therefore $H \backslash \mathcal{P}$ is not a regular abstract polytope ($\text{Aut}(H \backslash \mathcal{P})$ doesn't act flag-transitively on $H \backslash \mathcal{P}$).

We can define quantum codes from subgroups H which are not normal in $\text{Aut}(\mathcal{P})$ but satisfy the non-locality condition. The corresponding polytope is not as regular as when we quotient by a normal subgroup but the local structure of the covering polytope is nonetheless preserved.

However it is not clear to us how to explicitly compute the faces and incidences of $H \setminus \mathcal{P}$ without explicitly computing the faces and incidences of $\mathcal{P}$ before. If $\mathcal{P}$ is finite we can compute it with a computer algebra system. But the case where $\mathcal{P}$ is infinite whereas $H \setminus \mathcal{P}$ is finite is problematic: it is impossible to explicitly compute $\mathcal{P}$ and therefore we don't know how to explicitly compute $H \setminus \mathcal{P}$. We also don't know how to compute $H \setminus \Gamma$ in the infinite case. Indeed $H \setminus \Gamma$ is not a group in general and therefore we can't use the images of the generators of Γ like we do when we quotient by a normal subgroup.

The codes with 720 and 18 000 qubits defined above are quotients of the finite polytope $\mathcal{P}$ corresponding to the code with 90 000 qubits. Note that a qubit of the codes with 720 or 18 000 qubits thus corresponds to a triple quotient of the $\{5, 3, 3, 5\}$ Coxeter group Γ : it is an element of $H \setminus (N \setminus \Gamma) / S_2$, where N is the normal congruence subgroup $\Gamma(\sqrt{5}\mathbb{Z}[\phi])$.

3.2 Numerical search for normal subgroups of the Coxeter group

It is also possible to find finite codes with the $\{5, 3, 3, 5\}$ local structure by adding an extra relation to the Coxeter group. By searching numerically through such extra relations, my coauthor Nickolas Breuckmann defines 3 codes with the same $\{5, 3, 3, 5\}$ local structure and with respectively 144, 18 432 and 19 584 physical qubits. Since they have the same local structure, we use arithmetic and numerically searched codes together for the decoding numerical simulations. We found numerically that these three codes have respectively 72, 4 232 and 4 324 logical qubits.

3.3 Decoding with Belief Propagation

The decoding algorithm we consider belongs to the Belief Propagation family. We used this decoding algorithm because it is easy to parallelize and can therefore run very fast on quantum hardware. In the next paragraph we define the Belief Propagation algorithm (BP) used and show that it would correspond to maximum likelihood if the Tanner graph of the quantum code were a tree. Since the Tanner graph of a $\{5, 3, 3, 5\}$ quantum code is not a tree, BP gives a heuristic decoding algorithm in this setting.

3.3.1 Belief propagation on a tree

$X^{(j)}$ are random variables corresponding to variable nodes. They are independently and identically distributed (*iid*) Bernoulli distributions with parameter $p \in [0, 1]$.

$Y^{(k)}$ are random variables corresponding to check nodes. $Y^{(k)} = \bigoplus_{j \text{ neighbour of } k} X^{(j)}$. The values of $Y^{(k)}$ are observed and we denote them by $y_{\text{obs}}^{(k)}$.

We want to compute marginals of the random variables $X^{(j)}$ conditioned on the observations $y_{\text{obs}}^{(k)}$. If the Tanner graph is a tree, it is possible to iteratively compute these marginals exactly.

We assume in this paragraph that the Tanner graph is a tree and that a root has been chosen. We will use the notation $k > j$ to denote that k is a descendant of j . For each variable node j , we define the following function whose domain is $\{0, 1\}$:

$$p^{(j)}(x) = \Pr(X^{(j)} = x \mid \{Y^{(k)} = y_{\text{obs}}^{(k)}\}_{k > j}).$$

For each check node k , denoting by j its parent variable node, we define the following function whose domain is $\{0, 1\}$:

$$q^{(k)}(x) = \Pr(Y^{(k)} = y_{\text{obs}}^{(k)} \mid X^{(j)} = x, \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k}).$$

To compute $p^{(j)}(x)$ from $(q^{(k)}(x))_k$ children of j , we need the following variation of Bayes' formula:

$$\Pr(A \mid B, C) \Pr(B \mid C) = \Pr(B \mid A, C) \Pr(A \mid C)$$

Indeed the left hand side of the above equation equals $\frac{\Pr(A \cap B \cap C)}{\Pr(B \cap C)} \frac{\Pr(B \cap C)}{\Pr(C)} = \frac{\Pr(A \cap B \cap C)}{\Pr(C)}$ which is symmetric in (A, B) and therefore equals the right hand side.

We apply this formula with

$$\begin{aligned} A &= (X^{(j)} = x) \\ B &= (\{Y^{(k)} = y_{\text{obs}}^{(k)}\}_{k \text{ children of } j}) \\ C &= (\{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>j, m \text{ not a child of } j}) \end{aligned}$$

and define the normalisation constant $Z = \Pr(B \mid C)$. We know that $\Pr(A) = p$.

$$\begin{aligned} p^{(j)}(x) &= \Pr(A \mid B, C) \\ p^{(j)}(x) &= \frac{1}{Z} \Pr(B \mid A, C) \Pr(A \mid C) \\ p^{(j)}(x) &= \frac{1}{Z} \Pr(B \mid A, C) \Pr(A) \quad \text{since } A \text{ and } C \text{ are independent.} \\ p^{(j)}(x) &= \frac{1}{Z} p \prod_{k \text{ children of } j} \Pr(Y^{(k)} = y_{\text{obs}}^{(k)} \mid A, C) \\ p^{(j)}(x) &= \frac{1}{Z} p \prod_{k \text{ children of } j} \Pr(Y^{(k)} = y_{\text{obs}}^{(k)} \mid X^{(j)} = x, \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>j, m \text{ not a child of } j}) \\ p^{(j)}(x) &= \frac{1}{Z} p \prod_{k \text{ children of } j} \Pr(Y^{(k)} = y_{\text{obs}}^{(k)} \mid X^{(j)} = x, \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k}) \\ &\quad \text{since } Y \text{ random variables at distance more than 2 in the tree are independent.} \\ p^{(j)}(x) &= \frac{1}{Z} p \prod_{k \text{ children of } j} q^{(k)}(x). \end{aligned} \tag{3.5}$$

Since $p^{(j)}(0) + p^{(j)}(1) = 1$, we obtain that

$$Z = p \prod_{k \text{ children of } j} q^{(k)}(x) + (1 - p) \prod_{k \text{ children of } j} q^{(k)}(1 - x).$$

We now compute $q^{(k)}(x)$ from $(p^{(l)}(x))_l$ children of k :

$$\begin{aligned}
q^{(k)}(x) &= \Pr(Y^{(k)} = y_{\text{obs}}^{(k)} \mid X^{(j)} = x, \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k}) \\
q^{(k)}(x) &= \Pr\left(\bigoplus_{l \text{ children of } k} X^{(l)} \oplus X^{(j)} = y_{\text{obs}}^{(k)} \mid X^{(j)} = x, \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k}\right) \\
q^{(k)}(x) &= \Pr\left(\bigoplus_{l \text{ children of } k} X^{(l)} = y_{\text{obs}}^{(k)} \oplus x \mid \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k}\right) \\
1 - 2q^{(k)}(x) &= (-1)^{y_{\text{obs}}^{(k)} + x + 1} \prod_{l \text{ children of } k} (1 - 2\Pr(X^{(l)} = 1 \mid \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>k})) \\
1 - 2q^{(k)}(x) &= (-1)^{y_{\text{obs}}^{(k)} + x + 1} \prod_{l \text{ children of } k} (1 - 2\Pr(X^{(l)} = 1 \mid \{Y^{(m)} = y_{\text{obs}}^{(m)}\}_{m>l})) \\
1 - 2q^{(k)}(x) &= (-1)^{y_{\text{obs}}^{(k)} + x + 1} \prod_{l \text{ children of } k} (1 - 2p^{(l)}(1)) \tag{3.6}
\end{aligned}$$

We could use eqs. (3.5) and (3.6) to define the iterative Belief Propagation algorithm. However for numerical stability reasons we will follow [RU08] and use logarithmic ratios:

$$\begin{aligned}
lp^{(j)} &= \log \frac{p^{(j)}(0)}{p^{(j)}(1)} \\
lq^{(k)} &= \log \frac{q^{(k)}(0)}{q^{(k)}(1)}
\end{aligned}$$

Under this transformation, eq. (3.5) translates into:

$$lp^{(j)} = \log\left(\frac{1-p}{p}\right) + \sum_{k \text{ children of } j} lq^{(k)}. \tag{3.7}$$

Observing that $q^{(k)}(1) = \frac{1}{\exp(lq^{(k)})+1}$, we obtain:

$$\begin{aligned}
1 - 2q^{(k)}(1) &= \frac{\exp(lq^{(k)}) - 1}{\exp(lq^{(k)}) + 1} \\
1 - 2q^{(k)}(1) &= \frac{\exp(lq^{(k)}/2) - \exp(-lq^{(k)}/2)}{\exp(lq^{(k)}/2) + \exp(-lq^{(k)}/2)} \\
1 - 2q^{(k)}(1) &= \tanh \frac{lq^{(k)}}{2}.
\end{aligned}$$

Similarly $1 - 2p^{(j)}(1) = \tanh \frac{lp^{(j)}}{2}$ and therefore eq. (3.6) translates into:

$$\begin{aligned}
\tanh \frac{lq^{(k)}}{2} &= (-1)^{y_{\text{obs}}^{(k)}} \prod_{l \text{ children of } k} \tanh \frac{lp^{(l)}}{2} \\
lq^{(k)} &= \frac{(-1)^{y_{\text{obs}}^{(k)}}}{2} \tanh^{-1} \left(\prod_{l \text{ children of } k} \tanh \frac{lp^{(l)}}{2} \right) \tag{3.8}
\end{aligned}$$

The Belief Propagation algorithm we use is defined from eqs. (3.7) and (3.8): the check node k sends the message $lq^{(k)}$ to its parent node. The variable node j sends the message $lp^{(j)}$ to its parent node.

The first message is sent by the leaves of the tree, which we assume are variable nodes. It is initialized to $\log\left(\frac{1-p}{p}\right)$.

The last message is received by the root of the tree, which we assume is a variable node. The value $\frac{1}{\exp(lp^{(\text{root})})+1}$ gives the probability that the random variable corresponding to the root is 1 conditioned on the observation of all the check variables.

3.3.1.1 The loopy Belief Propagation heuristic

When the Tanner graph is not a tree we can still use the Belief Propagation algorithm as it was described above. However it does not compute exact probabilities any more: it is a heuristic whose performance we investigate numerically.

In the case of a $\{5, 3, 3, 5\}$ quantum code, the Tanner graph $\mathcal{T}$ on which we apply the belief propagation algorithm has the union of the qubits and the X -checks as its vertex set and has an edge between a qubit and an X -check if they correspond to incident faces. Z -errors are corrected on this graph. The graph whose vertex set is the union of the qubits and the Z -checks is used to correct X -errors. Since these two graph are isomorphic we won't mention the second one any more.

The graph $\mathcal{T}$ is not a tree. We use the letter $\mathcal{T}$ to refer to a Tanner graph. However we will use BP on $\mathcal{T}$ as if it were a tree. In this setting BP doesn't compute exact marginals but can still be used as a heuristic. Let us give the length of the smallest cycles in $\mathcal{T}$. These cycles are the first suspects when BP doesn't perform optimally. In a $\{5, 3, 3, 5\}$ quantum code, qubits have weight 5 and checks have weight 12. Moreover the qubits and X -checks incident to a given Z -generator correspond respectively to the pentagons and the edges of a dodecahedron. Therefore there are cycles in $\mathcal{T}$ of weight 6: 3 qubits and 3 X checks. Since no pair of qubits is incident to the same pair of X -checks, these are the shortest cycles in a $\{5, 3, 3, 5\}$ quantum code. It means that the Belief Propagation algorithm follows a cycle after 3 rounds.

The fact that the Tanner graph has a large number of very short cycles is unavoidable with quantum LDPC codes. This is why Belief Propagation algorithms work much better with classical LDPC codes than with quantum LDPC codes. However the situation is intuitively more favorable in a 4-dimensional hyperbolic setting than in a 2-dimensional setting or in a Euclidean setting where the probability that a random path in the Tanner graph intersects itself would be higher.

3.4 Numerical Results

In this section we evaluate the performances of $\{5, 3, 3, 5\}$ quantum codes with the BP decoder against random noise. We consider the phenomenological noise model where X and Z errors are distributed independently and identically as Bernoulli variables with parameter $p \in [0, 1]$. We correct X and Z errors independently and we will only discuss Z errors since $\{5, 3, 3, 5\}$ quantum codes have isomorphic X and Z sides.

We first consider noiseless measurements. We apply the Belief Propagation algorithm in parallel. A round of message-passing consists in each variable node sending a message to each of its neighbor check node and each check node sending a message to each of its neighbor variable node. After each round r of Belief Propagation we compute w_r , the weight of the syndrome if we were to flip the qubits whose belief to have an error is higher than 0.5. We stop as soon as $w_r \geq w_{r-1}$ or when $w_r = 0$. If we stopped because $w_r = 0$ and there is no logical error, we say that the decoding succeeded. Figure 3.1 shows the statistical frequency of unsuccessful decoding as a function of the physical error rate. It gives numerical evidence for a noiseless threshold above 5 % physical error rate.

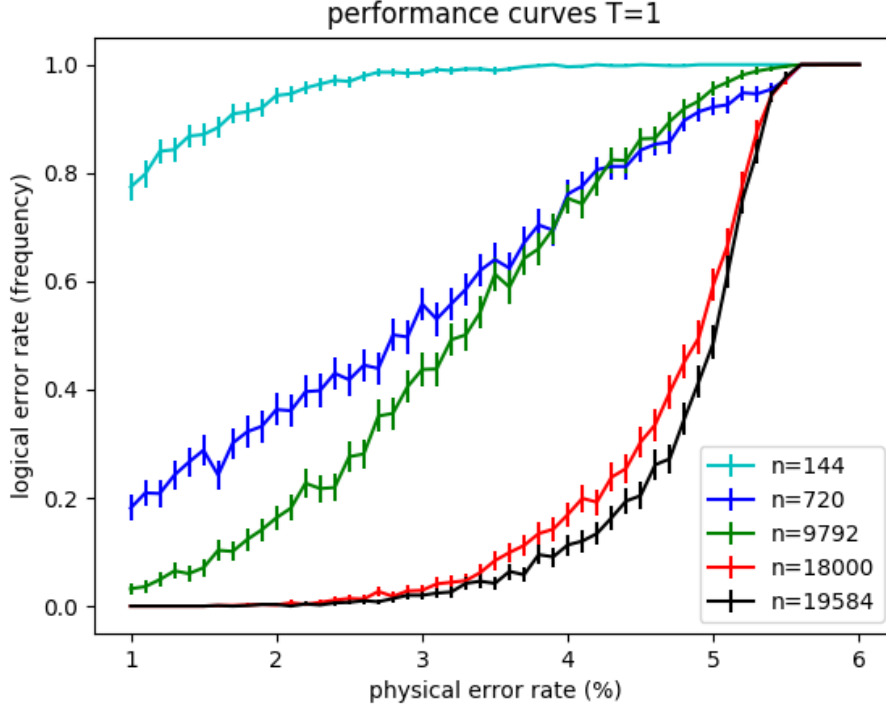


Figure 3.1: Performance curves for noiseless measurements. Vertical error bars correspond to the approximate 95 % confidence interval given by $p = \hat{p} \pm 1.96 \sqrt{\frac{\hat{p}(1-\hat{p})}{n_{\text{trials}}}}$. At a given physical error rate, $\hat{p}$ is the mean value for unsuccessful decoding and n_{trials} is the number of independent simulations. The decoding is said unsuccessful if it terminates with a non zero syndrome or with a logical error. Here $n_{\text{trials}} = 1000$ for each physical error rate and each quantum code. For the codes with more than 10 000 qubits, the transition between the successful and the unsuccessful decoding is quite sharp and gives numerical evidence for a noiseless threshold above 5 %.

We also consider noisy measurements. More precisely we consider T rounds of error correction. At each round $t \in \{1, \dots, T\}$, each qubit independently undergoes a Z error e_t^{noise} with probability p , where p is the physical error rate. If $t \neq 1$, this error e_t is added to $e_{t-1}^{\text{res.}}$, the residual error at round $t-1$. The noiseless syndrome is computed:

$$s_t^{\text{noiseless}} = H(e_{t-1}^{\text{res.}} \oplus e_t^{\text{noise}}).$$

For $t \in \{1, \dots, T-1\}$, each check node independently undergoes an error with probability q . This defines a syndrome noise s_t^{noise} . In our numerical simulations we only consider the case $p = q$ for simplicity. The noisy syndrome is given to the BP decoder:

$$s_t^{\text{noisy}} = s_t^{\text{noiseless}} \oplus s_t^{\text{noise}}.$$

The BP decoder outputs an inferred error:

$$e_t^{\text{inf.}} = \text{BP}_{\text{dec.}}(s_t^{\text{noisy}}).$$

The residual error is updated:

$$e_t^{\text{res.}} = e_{t-1}^{\text{res.}} \oplus e_t^{\text{noise}} \oplus e_t^{\text{inf.}}$$

For the last round, $t = T$, we assume perfect measurements and therefore have $s_T^{noise} = 0$. If the weight of the syndrome after the BP correction of this last round is zero and the residual error $e_T^{res.}$ is not a logical error, we say that the decoding succeeded. Figure 3.2 shows the statistical frequency of unsuccessful decoding as a function of the physical error rate for $T \in \{2, \dots, 5\}$. Note that the noiseless measurement scenario corresponds to $T = 1$.

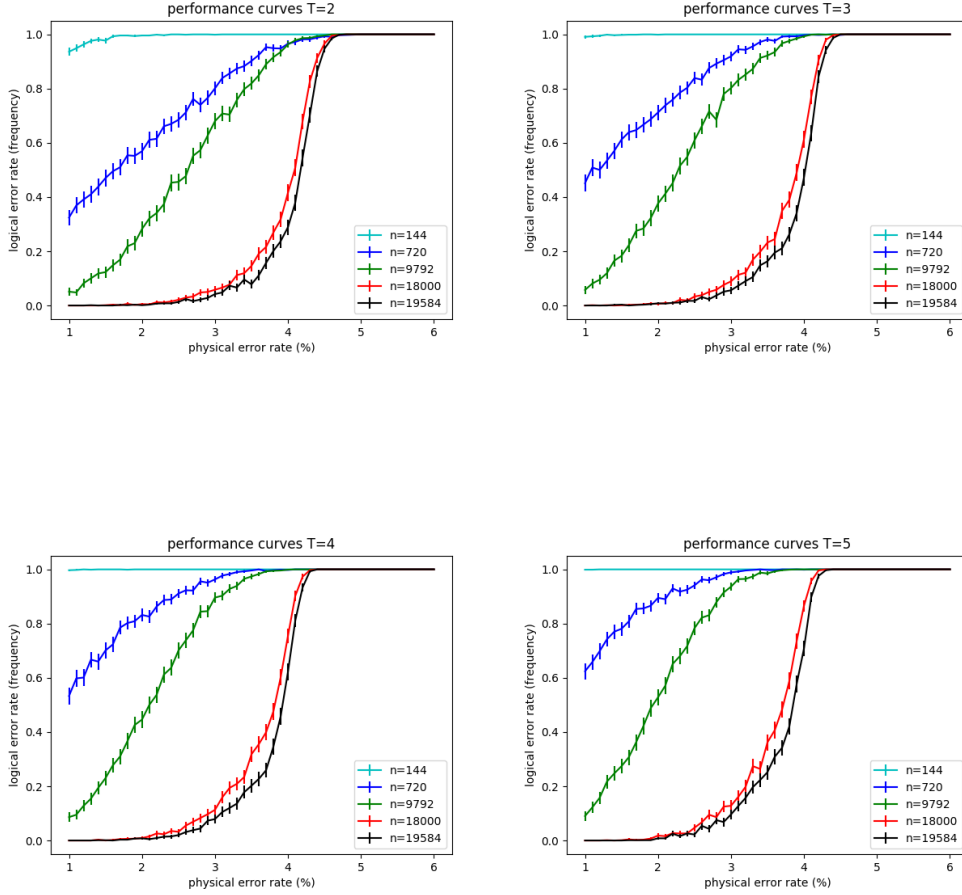


Figure 3.2: Performance curves for noisy measurements with T rounds and single-shot error correction at each round. $T \in \{2, \dots, 5\}$. The case $T = 1$ corresponds to noiseless measurement and has been given above. At each round, new errors occur on the physical qubits and the syndrome. At each round, the BP decoder flips qubits until the syndrome doesn't decrease any more. The measurement of the last round is assumed to be noiseless. The decoding is successful if the syndrome is zero and there is no logical error after this last round. The quantum codes with more than 10 000 qubits give numerical evidence for a threshold that decreases with T and is around 4 % at $T = 5$.

We also sort the performance curves by quantum code on Figure 3.3 to highlight the influence of the number of rounds T on a given quantum code.

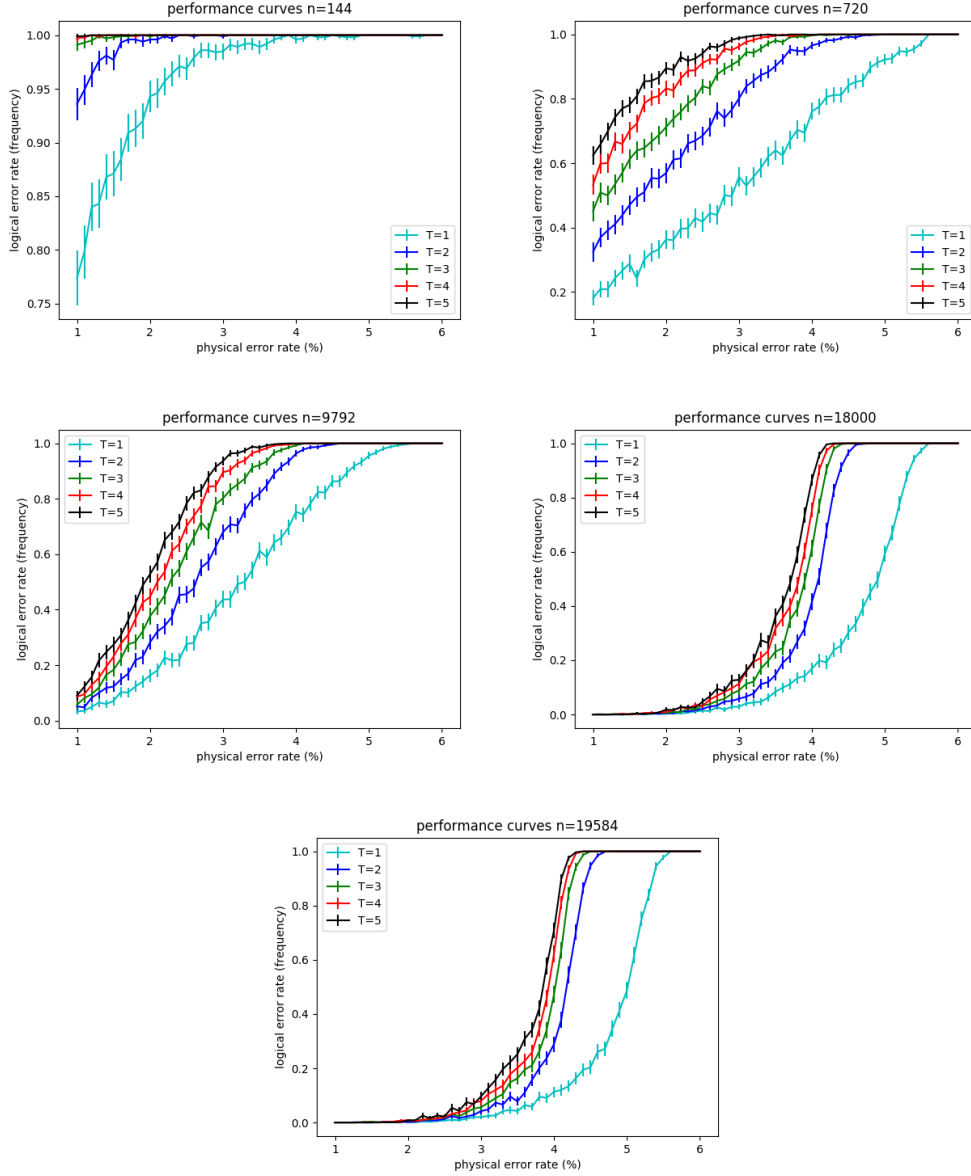


Figure 3.3: Performance curves sorted by quantum codes. As expected, the probability of successful decoding decreases with the number of error correcting rounds T .

A threshold around 4% at $T = 5$, *i.e.* after 5 rounds of error correction, is high for a family of quantum codes with a rate around 20%. It gives numerical evidence that the 4-dimensional hyperbolic codes have very code performance in practise. It would be interesting to see if this performance improves significantly with a custom decoding algorithm instead of the generic Belief Propagation algorithm used in this chapter.

Chapter 4

Cube-quotient codes

This chapter introduces a new family of quantum codes. Like the codes of the two previous chapter, they are homological codes, *i.e.* defined from the incidences of a polytope. However the polytopes in this chapter are not 4-dimensional: they are hypercubes of arbitrary dimensions. In the sequel we drop the prefix hyper and use to terminology cube to refer to an n -cube, regardless of the dimension n .

The intuition behind considering homological codes defined from a cube is that codes coming from spherical geometry objects may have large minimum distances. Indeed in dimension 2, hyperbolic codes have logarithmic minimum distances, Euclidean codes like the toric code have minimum distances scaling like $\sqrt{n}$ and a code defined from a projective plane also has $\sqrt{n}$ minimum distance but with a larger constant than its Euclidean counterpart. However in dimension 2, we don't know how to define an infinite family of quantum codes without losing the benefits of positive curvature. In this chapter we keep spherical objects by considering cubes and forget the fixed dimension requirement by allowing the dimension of the cubes to be arbitrarily large. This approach was also considered in [Has16]. However a major downside of this construction is that the quantum code defined from an n -cube by identifying qubits with p -faces for any $p \in \{1, \dots, n-2\}$ encodes zero logical qubit. We solve this problem by considering quotients of cubes.

We actually construct a multi-parameter family of quantum codes: the first integer parameter is the dimension of the cube and the second integer parameter is the dimension of the faces that correspond to qubits. The last parameter corresponds to the way we quotient the cube. For some asymptotic results, we extract from this multi-parameter family a one-parameter family of quantum codes optimizing the properties considered.

We first review a result from homological algebra which we use repeatedly in this chapter. We then focus on the hemicube: the quotient of a cube by the antipodal map. We find the number of logical qubits and the cycle and cocycle minimum distances of these codes. Interestingly the same codes were already defined with a completely different approach in [Aud13]. The case of the hemicube is a stepping stone for the general case: we give the number of logical qubits and the cycle and cocycle minimum distances of the general cube-quotient code. Finally, as a side remark, we explain how the non-local subgroup condition defined in chapter 3 applies to the cube-quotient codes.

4.1 The Long Exact Sequence

Possible references for this section are [Wei95] §1.3 p. 10 and [Rot08] Theorem 6.10 p. 333.

Definition 38. A chain complex $\mathcal{C}_\bullet$ is a sequence of objects $(C_n)_{n \in \mathbb{Z}}$ and of morphisms $(d_n : C_n \rightarrow C_{n-1})_{n \in \mathbb{Z}}$ called differentials such that the composition of any two successive differentials is zero: $d_{n-1}d_n = 0$.

Elements of $\ker(d_n)$ are called cycles and elements of $\text{Im}(d_{n+1})$ are called boundaries. Every boundary is a cycle but the converse is not necessarily true. Homology groups give information about this defect.

Definition 39. The n th-homology groups $H_n(\mathcal{C}_\bullet)$ of a chain complex $\mathcal{C}_\bullet$ is defined as the quotient $H_n = \ker(d_n) / \text{Im}(d_{n+1})$.

Definition 40. A chain map f from the chain complex $\mathcal{C}_\bullet$ to the chain complex $\mathcal{D}_\bullet$ is a sequence of morphisms $f_n : C_n \rightarrow D_n$ such that for all $n \in \mathbb{Z}$, $f_{n-1}d_n = d_n f_n$. By abuse of notation we used the same symbol d_n to refer to the distinct differentials $d_n : C_n \rightarrow C_{n-1}$ and $d_n : D_n \rightarrow D_{n-1}$.

A chain map sequence $\mathcal{A}_\bullet \xrightarrow{f} \mathcal{B}_\bullet \xrightarrow{g} \mathcal{C}_\bullet$ is called exact if for all $n \in \mathbb{Z}$, the sequence $A_n \xrightarrow{f_n} B_n \xrightarrow{g_n} C_n$ is exact.

Theorem 41 (Long Exact Sequence). A short exact sequence $0 \rightarrow \mathcal{A}_\bullet \rightarrow \mathcal{B}_\bullet \rightarrow \mathcal{C}_\bullet \rightarrow 0$ of chain complexes induces the following long exact sequence of homology groups:

$$\dots \rightarrow H_n(\mathcal{A}_\bullet) \rightarrow H_n(\mathcal{B}_\bullet) \rightarrow H_n(\mathcal{C}_\bullet) \rightarrow H_{n-1}(\mathcal{A}_\bullet) \rightarrow H_{n-1}(\mathcal{B}_\bullet) \rightarrow \dots$$

We refer to [Wei95] or [Rot08] for a proof. However we will make more explicit the induced morphism $H_n(\mathcal{A}_\bullet) \rightarrow H_n(\mathcal{B}_\bullet)$ (or equivalently $H_n(\mathcal{B}_\bullet) \rightarrow H_n(\mathcal{C}_\bullet)$) and the connecting morphism $H_n(\mathcal{C}_\bullet) \rightarrow H_{n-1}(\mathcal{A}_\bullet)$.

The homology group morphism $H_n(\mathcal{A}_\bullet) \rightarrow H_n(\mathcal{B}_\bullet)$ is induced by $f_n : A_n \rightarrow B_n$. It is well defined because f_n takes cycles to cycles and boundaries to boundaries. To avoid confusions we will sometimes denote the chain group morphism by $f_{\text{chain},n}$ and the homology group morphism by $f_{\text{hom},n}$.

The connecting morphism $H_n(\mathcal{C}_\bullet) \rightarrow H_{n-1}(\mathcal{A}_\bullet)$ takes more work to construct.

Let $[c_n]$ be a class in $H_n(\mathcal{C}_\bullet)$ represented by the element c_n of C_n . There exists $b_n \in B_n$ such that $g_n(b_n) = c_n$. Now, $g_{n-1}(d_n(b_n)) = d_n(g_n(b_n)) = d_n(c_n) = 0$. Therefore there exists $a_{n-1} \in A_{n-1}$ such that $f_{n-1}(a_{n-1}) = d_n(b_n)$. The connecting morphism is defined by sending $[c_n]$ to $[a_{n-1}]$.

We leave it to the reader to prove that a_{n-1} is a cycle, that its class $[a_{n-1}]$ in $H_{n-1}(\mathcal{A}_\bullet)$ doesn't depend on the representative c_{n-1} chosen for $[c_{n-1}]$ and that the connecting map actually is a morphism. To avoid confusions we will sometimes denote the chain group differential by $d_{\text{chain},n}$ and the connecting homology group morphism by $d_{\text{hom},n}$.

4.2 Number of logical qubits in a hemicube code

The n -hemicube is the quotient of the cube by the antipodal map. Since it is topologically the projective n -space, the quantum code obtained from identifying qubits with

p -faces of the n -hemicube has as many logical qubits as the rank of the p^{th} homology group of the projective n -space with coefficients in $\mathbb{F}_2$. Hence it has one logical qubit for $p \in \{1, \dots, n-2\}$. In this section, we give a more algebraic proof of this result using the long exact sequence of §4.1 because it generalises better to other quotients of the n -cube.

Quotienting by the antipodal map is the same as quotienting by the repetition code $C_r = \{0\dots 0, 1\dots 1\}$. Therefore we denote the hemicube by Q_n/C_r . We define now a short exact sequence of chain maps involving the n -hemicube and the n -cube and derive the associated long exact sequence of homology groups.

The short exact sequence of chain complexes we consider is

$$0 \rightarrow \mathcal{C}_\bullet(Q_n/C_r) \xrightarrow{i} \mathcal{C}_\bullet(Q_n) \xrightarrow{\pi} \mathcal{C}_\bullet(Q_n/C_r) \rightarrow 0.$$

$\mathcal{C}_\bullet(Q_n)$ and $\mathcal{C}_\bullet(Q_n/C_r)$ denote the chain complexes of the n -cube and the n -hemicube.

The projection π is the linear extension at the level of chains of the projection given by quotienting by the repetition code. A face f_p of the n -cube is sent to its equivalence class $[f_p]$ in the hemicube: $\forall f_p \in F_p(Q_n), \pi_p(f_p) = [f_p]$.

The injective map i is the linear extension of the application sending a p -face $[f_p]$ of the hemicube to the sum of the two faces of the cube belonging to this equivalence class: $\forall [f_p] \in F_p(Q_n/C_r), i_p([f_p]) = f_p + (f_p \oplus 1\dots 1)$. $f_p \oplus 1\dots 1$ is the translate of f_p by the non-zero codeword of the repetition code: $1\dots 1$.

It is not hard to verify that this defines a short exact sequence for each p and that π_p and i_p commute with the differential ∂ .

We thus have a short exact sequence of chain maps and we can write the associated long exact sequence of homology groups given by Thm. 41:

$$\dots \rightarrow H_p(Q_n/C_r) \rightarrow H_p(Q_n) \rightarrow H_p(Q_n/C_r) \rightarrow H_{p-1}(Q_n/C_r) \rightarrow H_{p-1}(Q_n) \rightarrow \dots$$

Since $H_p(Q_n) = 0$ for all $p \in \{1, \dots, n-1\}$, we obtain:

$$\forall p \in \{1, \dots, n-1\}, 0 \rightarrow H_p(Q_n/C_r) \rightarrow H_{p-1}(Q_n/C_r) \rightarrow 0.$$

$H_0(Q_n/C_r)$ has dimension 1 since the hemicube is path-connected. Therefore we obtain by immediate induction that $H_p(Q_n/C_r)$ has dimension 1 for all $p \in \{0, \dots, n-1\}$.

Therefore if we construct a quantum code by identifying qubits with p -faces for any $p \in \{1, \dots, n-2\}$, this quantum code has 1 logical qubit.

In the two following sections, we will give the cycle and cocycle minimum distances for the hemicube code. Even though these minimum distances are specific cases of the results of §4.9 and §4.10, we treat them separately as a stepping stone for the general case.

4.3 Cycle minimum distance for the hemicube code

We will construct homologically non trivial p -cycles in the n -hemicube recursively. We denote such a cycle by $\begin{bmatrix} n \\ p \end{bmatrix}_r \in \mathcal{C}_p(Q_n/C_r)$. We construct $\begin{bmatrix} n \\ p \end{bmatrix}_r$ as the image under the projection π_p of a chain $\begin{bmatrix} n \\ p \end{bmatrix} \in \mathcal{C}_p(Q_n)$.

We initialise the construction at $p = 0$ for any $n \geq 1$ by defining $\begin{bmatrix} n \\ 0 \end{bmatrix}$ as the chain made of a single 0-face: $0\dots 0$. $\begin{bmatrix} n \\ 0 \end{bmatrix} := 0\dots 0 \in \mathcal{C}_0(Q_n)$.

We define $\begin{bmatrix} n \\ p \end{bmatrix} := \frac{(-1)^{p+1}}{2} \begin{bmatrix} n-1 \\ p \end{bmatrix} + * \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \in \mathcal{C}_p(Q_n)$ with the convention that $\begin{bmatrix} m \\ p \end{bmatrix} = \emptyset$ if $p \leq -1$ or $p \geq m+1$.

Note that $\frac{(-1)^{p+1}}{2}$ is 1 when p is even and 0 when p is odd. By $\frac{(-1)^{p+1}}{2} \begin{bmatrix} n-1 \\ p \end{bmatrix}$, we mean the concatenation of either a 1 or a 0 to every face of $\begin{bmatrix} n-1 \\ p \end{bmatrix}$ depending on the parity of p .

For a chain $chain \in \mathcal{C}_p(Q_n)$, we denote its translate by $1\dots 1$ by $chain \oplus 1\dots 1$. A more general and precise definition will be given in §4.7.

To show that $\begin{bmatrix} n \\ p \end{bmatrix}_r$ is a p -cycle in the n -hemicube, we need the following lemma about chains of the n -cube:

Lemma 42. $\partial \begin{bmatrix} n \\ p \end{bmatrix} = \begin{bmatrix} n \\ p-1 \end{bmatrix} + (\begin{bmatrix} n \\ p-1 \end{bmatrix} \oplus 1\dots 1)$

Proof. The proof, by induction, follows from the recursive definition of $\begin{bmatrix} n \\ p \end{bmatrix}$ and a calculation:

$$\begin{aligned} \partial \begin{bmatrix} n \\ p \end{bmatrix} &= \frac{(-1)^p + 1}{2} \partial \begin{bmatrix} n-1 \\ p \end{bmatrix} + * \partial \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} + 0 \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} + 1 \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \\ &= \frac{(-1)^p + 1}{2} \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} + \frac{(-1)^p + 1}{2} (\begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \oplus 1\dots 1) \\ &\quad + * \begin{bmatrix} n-1 \\ p-2 \end{bmatrix} + * (\begin{bmatrix} n-1 \\ p-2 \end{bmatrix} \oplus 1\dots 1) + 0 \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} + 1 \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \\ &= \frac{(-1)^p + 1}{2} (\begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \oplus 1\dots 1) + * \begin{bmatrix} n-1 \\ p-2 \end{bmatrix} + * (\begin{bmatrix} n-1 \\ p-2 \end{bmatrix} \oplus 1\dots 1) + \frac{(-1)^{p-1} + 1}{2} \begin{bmatrix} n-1 \\ p-1 \end{bmatrix} \\ &= \begin{bmatrix} n \\ p-1 \end{bmatrix} + (\begin{bmatrix} n \\ p-1 \end{bmatrix} \oplus 1\dots 1). \end{aligned}$$

To go from line 2 to line 3, observe that $0 + 1 + \frac{(-1)^p + 1}{2} = \frac{(-1)^{p-1} + 1}{2}$ is true since it can be written $0 + 1 + 0 = 1$ if p is even and $0 + 1 + 1 = 0$ if p is odd. $\square$

From Lemma 42 we infer that in the n -hemicube, since $chain = chain \oplus 1\dots 1$, $\begin{bmatrix} n \\ p \end{bmatrix}_r$ is a p -cycle. Prop. 43 shows that $\begin{bmatrix} n \\ p \end{bmatrix}_r$ has minimum weight among homologically non-trivial p -cycles:

Proposition 43. *For every $n \geq 3$, for every $p \in \{1, \dots, n-2\}$, the quantum hemicube code has cycle minimum distance $\binom{n}{p}$.*

Proof. Observe by immediate induction that the $\binom{n}{p}$ possible positions of the stars are each attained by exactly one p -face of the $\begin{bmatrix} n \\ p \end{bmatrix}_r$ p -cycle.

Consider now the partition of the p -faces of a p -cycle of the hemicube by the position of its stars. Adding a boundary to a p -chain doesn't change the parity of the number of p -faces of the p -chain in any set of this partition. Therefore any p -cycle in the same homology class as $\begin{bmatrix} n \\ p \end{bmatrix}_r$ has at least one p -face in any set of this partition. Since there is only one non-zero homology class in the hemicube, the cycle minimum distance of the hemicube code is $\binom{n}{p}$. $\square$

We defined the chain $\begin{bmatrix} n \\ p \end{bmatrix}$ recursively in this section. However it is possible to define it by a closed formula: $\begin{bmatrix} n \\ p \end{bmatrix}$ is the sum over every possible position of p stars in a word of length n of the p -face f_p defined by $f_p(i) = s_i \bmod 2$ where s_i is the number of stars on coordinates smaller than i . In §4.7 we will define this closed formula more precisely and generalise it to quotients by other classical codes than the repetition code. We will call the resulting p -cycles product cycles.

4.4 Cocycle minimum distance for the hemicube code

Proposition 44. *For every $n \geq 3$, for every $p \in \{1, \dots, n-2\}$, the quantum hemicube code has cocycle minimum distance 2^{n-p-1} .*

Proof. Recall that the p -cycle $\begin{bmatrix} n \\ p \end{bmatrix}_r$ has exactly one p -face in each of the $\binom{n}{p}$ possible positions of p stars among n coordinates.

For every $y \in (\mathbb{F}_2)^n / C_r$, consider the translate $\begin{bmatrix} n \\ p \end{bmatrix}_r \oplus y$. It is a p -cycle of the n -hemicube because ∂ commutes with $\oplus y$. Moreover it also has exactly one p -face in each of the $\binom{n}{p}$ possible positions of p stars among n coordinates.

Let $\text{cocyc}_{p,r}$ be a cohomologically non trivial p -cocycle of the n -hemicube. Since the hemicube code has only one logical qubit, $\text{cocyc}_{p,r}$ is not orthogonal to $\begin{bmatrix} n \\ p \end{bmatrix}_r \oplus y$ for every $y \in (\mathbb{F}_2)^n / C_r$. There are 2^{n-1} cycles $\begin{bmatrix} n \\ p \end{bmatrix}_r \oplus y$ and every p -face of the n -hemicube belongs to 2^p of them. Therefore $\text{cocyc}_{p,r}$ has weight at least 2^{n-p-1} .

Moreover the p -cocycle of the n -hemicube defined as the sum of all the p -faces whose p stars are on the p first coordinates has weight 2^{n-p-1} . It is cohomologically non trivial since it is not orthogonal to the homologically non trivial cycle $\begin{bmatrix} n \\ p \end{bmatrix}_r$. $\square$

We can now state the following theorem:

Theorem 45. *Identifying qubits with p -faces of the n -hemicube yields a quantum code with 2^{n-1} physical qubits, 1 logical qubit and a minimum distance equal to $\min(2^{n-p-1}, \binom{n}{p})$, i.e. a $[[2^{n-1}, 1, \min(2^{n-p-1}, \binom{n}{p})]]$ quantum code.*

In §4.6 we will generalise the definition of such p -cocycles to quotients by other classical codes than the repetition code and call them canonical cocycles.

4.5 Number of logical qubits in a cube-quotient code

By cube-quotient code we mean that we quotient the cube Q_n by a classical code $C = [n, k, d]$ with $d \geq p + 2$. (Qubits correspond to p -faces of the quotient.) We will denote the quotient polytope by Q_n/C .

For all $p \in \mathbb{N}$, we could construct the following short exact sequence of p -chains:

$$0 \rightarrow \bigoplus_{j=1}^{2^k-1} \mathcal{C}_p^j(Q_n/C) \xrightarrow{i_p} \mathcal{C}_p(Q_n) \xrightarrow{\pi_p} \mathcal{C}_p(Q_n/C) \rightarrow 0.$$

However it doesn't give a short exact sequence of chain maps because we don't know how we could construct a differential on the direct sum module such that we have the desired commuting relation: $\delta \circ i_p = i_{p-1} \circ \delta$.

Instead we will proceed step by step by considering a sequence of classical codes $(C_j = [n, j, d_j])_{j \in \{1, \dots, k\}}$ such that $C_k = C$ and for all $j \in \{1, \dots, k-1\}$, C_{j+1} contains C_j .

For all $p \in \mathbb{N}$, $j \in \{0, \dots, k-1\}$, we can construct the following short exact sequence of p -chains:

$$0 \rightarrow \mathcal{C}_p(Q_n/C_{k+1}) \xrightarrow{i_{chain,p}} \mathcal{C}_p(Q_n/C_k) \xrightarrow{\pi_{chain,p}} \mathcal{C}_p(Q_n/C_{k+1}) \rightarrow 0.$$

This time it does define a short exact sequence of chain complexes since $i_{chain,p}$ and $\pi_{chain,p}$ commute with the boundary operator ∂_{chain} :

$$0 \rightarrow \mathcal{C}_\bullet(Q_n/C_{k+1}) \xrightarrow{i_{complex}} \mathcal{C}_\bullet(Q_n/C_k) \xrightarrow{\pi_{complex}} \mathcal{C}_\bullet(Q_n/C_{k+1}) \rightarrow 0.$$

The associated long exact sequence of homology groups is:

$$\begin{aligned} \dots & \xrightarrow{\partial_{hom}} H_p(Q_n/C_{k+1}) \xrightarrow{i_{hom,p}} H_p(Q_n/C_k) \xrightarrow{\pi_{hom,p}} H_p(Q_n/C_{k+1}) \xrightarrow{\partial_{hom}} \\ & H_{p-1}(Q_n/C_{k+1}) \xrightarrow{i_{hom,p-1}} H_{p-1}(Q_n/C_k) \xrightarrow{\pi_{hom,p-1}} H_{p-1}(Q_n/C_{k+1}) \xrightarrow{\partial_{hom}} \dots \end{aligned}$$

Let us now derive the similar long exact sequence in cohomology. Since chains and cochains are canonically identified (they both can be considered as subsets of the set of faces) and since i_p and π_p commute with the coboundary operator δ , we have the following short exact sequence of cochain complexes:

$$0 \rightarrow \mathcal{C}^\bullet(Q_n/C_{k+1}) \xrightarrow{i_{cocomplex}} \mathcal{C}^\bullet(Q_n/C_k) \xrightarrow{\pi_{cocomplex}} \mathcal{C}^\bullet(Q_n/C_{k+1}) \rightarrow 0.$$

The associated long exact sequence of cohomology groups is:

$$\begin{aligned} \dots & \xrightarrow{\delta_{cohom}} H^p(Q_n/C_{k+1}) \xrightarrow{i_{cohom,p}} H^p(Q_n/C_k) \xrightarrow{\pi_{cohom,p}} H^p(Q_n/C_{k+1}) \xrightarrow{\delta} \\ & H^{p+1}(Q_n/C_{k+1}) \xrightarrow{i_{cohom,p+1}} H^{p+1}(Q_n/C_k) \xrightarrow{\pi_{cohom,p+1}} H^{p+1}(Q_n/C_{k+1}) \xrightarrow{\delta_{cohom}} \dots \end{aligned}$$

This long exact sequence of cohomology groups is actually easier to manipulate than its homology counterpart. In the next section we will see what the existence of a preferred basis of the cohomology groups $H^p(Q_n/C_k)$ implies for the long exact sequence.

4.6 Cohomology basis and short exact sequence in cohomology in a cube-quotient code

Definition 46. The p -direction of a p -face $f_p \in F_p(Q_n/C_k)$ is the subset of coordinates where f_p has a star: $\{i \in \{1, \dots, n\} \mid f_p(i) = *\}$.

Definition 47. For a p -direction $I \subset \{1, \dots, n\}$ with $|I| = p$, the corresponding canonical cocycle is the sum of all the p -faces in $F_p(Q_n/C_k)$ having this p -direction I .

It is straightforward to verify that it is indeed a cocycle. We denote it by $\text{cocyc}_{I,p,k}$.

Theorem 48. The cohomology group $H^p(Q_n/C_k)$ has a basis such that each basis element is represented by a canonical cocycle.

Proof. Let $(c_1, \dots, c_{k-1})$ be a basis of C_{k-1} completed in a basis $(c_1, \dots, c_k)$ of C_k . We consider a fixed $j \in \text{Supp}(c_k)$. We can assume without loss of generality that $\forall i \in \{1, \dots, k-1\}, j \notin \text{Supp}(c_i)$ (just add c_k to c_i if needed).

By the induction hypothesis, the cohomology group $H^p(Q_n/C_{k-1})$ has a basis such that each basis element is represented by a canonical cocycle. Since π_{cochain} applied to a canonical cocycle gives the empty cochain $\emptyset$, π_{cohom} is zero on $H^p(Q_n/C_{k-1})$. For the same reason π_{cohom} is zero on $H^{p-1}(Q_n/C_{k-1})$ and the long exact sequence in cohomology breaks into the following short exact sequence:

$$0 \rightarrow H^{p-1}(Q_n/C_k) \xrightarrow{\delta_{\text{cohom}}} H^p(Q_n/C_k) \xrightarrow{i_{\text{cohom}}} H^p(Q_n/C_{k-1}) \rightarrow 0.$$

We will use the above short exact sequence and apply the induction hypothesis to $H^p(Q_n/C_{k-1})$ and $H^{p-1}(Q_n/C_k)$:

Let $I \subset \{1, \dots, n\}$ with $|I| = p$ be a p -direction such that $[\text{cocyc}_{I,p,k-1}]$ is an element of the basis of $H^p(Q_n/C_{k-1})$. Since $i_{\text{cochain}}(\text{cocyc}_{I,p,k}) = \text{cocyc}_{I,p,k-1}$, $i_{\text{cohom}}([\text{cocyc}_{I,p,k}]) = [\text{cocyc}_{I,p,k-1}]$. Therefore the basis of cohomology classes of $H^p(Q_n/C_{k-1})$ represented by canonical cocycles has a free family of preimages by i_{cohom} represented by canonical cocycles of $H^p(Q_n/C_k)$.

Let $I \subset \{1, \dots, n\}$ with $|I| = p-1$ be a $(p-1)$ -direction such that $[\text{cocyc}_{I,p-1,k}]$ is an element of the basis of $H^{p-1}(Q_n/C_k)$. $j \notin I$ because $\forall x \in C_{k-1}, j \notin \text{Supp}(x)$. Also because $\forall x \in C_{k-1}, j \notin \text{Supp}(x)$, it makes sense to say that the j^{th} coordinate of a p -face of $\text{cocyc}_{I,p-1,k}$ is zero or one. Keeping only the faces of $\text{cocyc}_{I,p-1,k}$ whose j^{th} coordinate is zero gives a preimage of $\text{cocyc}_{I,p-1,k}$ by π_{cochain} . Applying δ_{cochain} to this preimage gives $i_{\text{cochain}}(\text{cocyc}_{I \cup \{j\},p,k})$. Since δ_{cohom} corresponds to $i_{\text{cochain}}^{-1} \circ \delta_{\text{cochain}} \circ \pi_{\text{cochain}}^{-1}$ at the level of cochains, we obtain that $\delta_{\text{cohom}}([\text{cocyc}_{I,p-1,k}]) = [\text{cocyc}_{I \cup \{j\},p,k}]$. Therefore the basis of cohomology classes of $H^{p-1}(Q_n/C_k)$ represented by canonical cocycles is sent by δ_{cohom} to a free family of cohomologically classes represented by canonical cocycles of $H^p(Q_n/C_k)$.

The exactness of the short exact sequence implies that the concatenation of these two free families forms a basis of $H^p(Q_n/C_k)$. $\square$

As side products, we obtain that for every p, k , $\pi_{\text{cochain},p,k} = 0$ and that the long exact sequences in cohomology break into pieces of small exact sequences:

$$0 \leftarrow H^p(Q_n/C_{k-1}) \xleftarrow{i_{\text{cohom}}} H^p(Q_n/C_k) \xleftarrow{\delta_{\text{cohom}}} H^{p-1}(Q_n/C_k) \leftarrow 0.$$

for every p, k . We wrote the above short exact sequence in cohomology from right to left to prepare its adjunction property with its homology counterpart.

4.6.1 Adjunction and short exact sequence in homology in a cube-quotient code

The following "quasi-equations" depicted with $\approx$ summarise how the connecting homology and cohomology morphisms are constructed from applications at the level of chains and cochains:

$$\partial_{hom} \approx i_{chain}^{-1} \circ \partial_{chain} \circ \pi_{chain}^{-1}, \quad (4.1)$$

$$\delta_{cohom} \approx i_{cochain}^{-1} \circ \delta_{cochain} \circ \pi_{cochain}^{-1}. \quad (4.2)$$

On the right hand side of $\approx$ are (co)chain morphisms and preimages of chain morphisms. On the left hand side of $\approx$ are (co)homology group morphisms. $\approx$ means that if we consider a (co)chain representing a (co)homology class, any preimage or image by the right hand side (co)chain morphisms yields a representative of the image by the left hand side (co)homology morphism. This is true by construction of the connecting (co)homology morphisms ∂_{hom} and δ_{cohom} .

Lemma 49. π_{chain} and $i_{cochain}$ are adjoint. (Since chains and cochains are canonically identified, $\pi_{cochain}$ and i_{chain} are also adjoint.)

Proof. By linearity it is sufficient to prove it at the level of faces.

Let $f_{p,k-1}$ be a $(p, k-1)$ -face and $f_{p,k}$ be a (p, k) -face.

$$\begin{aligned} & \langle \pi_{chain}(f_{p,k-1}), f_{p,k} \rangle = 1 \\ \Leftrightarrow & \pi_{chain}(f_{p,k-1}) = f_{p,k} \\ \Leftrightarrow & f_{p,k-1} + (f_{p,k-1} \oplus c_k) = i_{cochain}(f_{p,k}) \\ \Leftrightarrow & \langle f_{p,k-1}, i_{cochain}(f_{p,k}) \rangle = 1 \end{aligned}$$

□

We already know that ∂_{chain} and $\delta_{cochain}$ are adjoint.

We also know that the bilinear form $\langle, \rangle$ is well defined at the level of homology and cohomology groups.

Using eqs. (4.1) and (4.2), we see that the connecting morphisms ∂_{hom} and δ_{cohom} are adjoint at the level of homology and cohomology groups.

It is straightforward to prove that π_{hom} and i_{cohom} are adjoint because they correspond to π_{chain} and $i_{cochain}$ on representatives. Similarly π_{cohom} and i_{hom} are adjoint.

In section §4.6 we have proved that π_{cohom} is zero. Thus its adjoint i_{hom} is also zero and the long exact sequences in homology break into pieces of short exact sequences:

$$0 \rightarrow H_p(Q_n/C_{k-1}) \xrightarrow{\pi_{hom}} H_p(Q_n/C_k) \xrightarrow{\partial_{hom}} H_{p-1}(Q_n/C_k) \rightarrow 0.$$

To summarise we have obtained for every p, k two short exact sequences adjoint to each other, one in homology and one in cohomology:

$$\begin{aligned} 0 \rightarrow H_p(Q_n/C_{k-1}) & \xrightarrow{\pi_{hom}} H_p(Q_n/C_k) \xrightarrow{\partial_{hom}} H_{p-1}(Q_n/C_k) \rightarrow 0, \\ 0 \leftarrow H^p(Q_n/C_{k-1}) & \xleftarrow{i_{cohom}} H^p(Q_n/C_k) \xleftarrow{\delta_{cohom}} H^{p-1}(Q_n/C_k) \leftarrow 0. \end{aligned}$$

4.7 Product cycles in a cube-quotient code

Before we define product cycles, we need to define translations at the level of coordinates, faces and chains in Q_n and in Q_n/C_k .

For every p -face $f = f(1)...f(n)$ of Q_n we define $f \oplus y$, its translation by $y = y(1)...y(n) \in (\mathbb{F}_2)^n$, by

$$\forall i \in \{1, \dots, n\}, (f \oplus y)(i) = f(i) \oplus y(i).$$

At the level of coordinates, we define $* \oplus 1 = * \oplus 0 = *$. For example for $f = 001 * 1*$ and $y = 010111$, we have $f \oplus y = 011 * 0*$.

For every p -chain $c = \sum f \in \mathcal{C}_p(Q_n)$, we define $c \oplus y$ its translation by $y = y(1)...y(n) \in (\mathbb{F}_2)^n$ by

$$c \oplus y = \sum (f \oplus y).$$

Since the translation $\oplus y$ is compatible with quotienting by C_k , we use the same definitions in Q_n/C_k .

Recall that the p -direction $I \subset \{1, \dots, n\}$ of a p -face $f_p \in F_p$ is the subset of coordinates where f_p has a star: $I = \{i \in \{1, \dots, n\} \mid f_p(i) = *\}$.

In Q_n , there are 2^{n-p} p -faces having a given p -direction I . One of them is called the standard p -face with p -direction I and denoted by f_I , which we now define: for every $i \in \{1, \dots, n\}$, we define s_i as the cardinal of $I \cap \{1, \dots, i\}$. We define $f_I(i)$, the i^{th} coordinate of f_I to be s_i modulo 2.

For example with $n = 8$ and $p = 2$, the standard 2-face for the 2-direction $\{3, 7\} = _ * _ * _ * _$ is $00 * 111 * 0$ in Q_n .

In Q_n/C_k the standard p -face $f_{I,k}$ is the image under $\Pi_k = \pi_k \circ \dots \circ \pi_1$ of f_I .

For $x_1, \dots, x_k \in (\mathbb{F}_2)^n$ and $p_1, \dots, p_k \in \mathbb{N}$ such that $p_1 + \dots + p_k = p$, we define a product chain $chain\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}$ in $\mathcal{C}_p(Q_n)$ as follows:

A k -tuple $(I_1, \dots, I_k)$ of subsets of $\{1, \dots, n\}$ is adapted to $\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}$ if it satisfies the following conditions:

- $\forall i, j \in \{1, \dots, k\}, I_i \cap I_j = \emptyset$.
- $\forall i \in \{1, \dots, k\}, I_i \subset \text{Supp}(x_i)$.
- $\forall i \in \{1, \dots, k\}, |I_i| = p_i$.

$chain\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k} \in \mathcal{C}_p(Q_n)$ is the sum of the standard p -faces $f_{I_1 \cup \dots \cup I_k}$ over every k -tuple $(I_1, \dots, I_k)$ satisfying the above conditions. Note that the sum is over k -tuples $(I_1, \dots, I_k)$ and not over p -directions $I_1 \cup \dots \cup I_k$. It means that if a p -direction I admits an even number of adapted partitions $(I_1, \dots, I_k)$, it actually doesn't belong to $chain\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}$.

$chain_k\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k} \in \mathcal{C}_p(Q_n/C_k)$ is the image of $chain\left(\begin{smallmatrix} x_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}$ under Π_k .

Lemma 50.

$$\begin{aligned} & \partial(\text{chain}(\binom{x_i}{p_i}_{1 \leq i \leq k})) \\ &= \sum_{j=1}^k \text{chain}(\binom{x_i}{p_i - \delta_{i,j}}_{1 \leq i \leq k}) + (\text{chain}(\binom{x_i}{p_i - \delta_{i,j}}_{1 \leq i \leq k}) \oplus x_j) \end{aligned}$$

Proof. Taking the boundary of a p -chain amounts to replacing each star of each of its p -faces by either a 0 or a 1. Let $I_1 \subset \text{Supp}(x_1), \dots, I_k \subset \text{Supp}(x_k)$ satisfy $|I_1| = p_1, \dots, |I_k| = p_k$ and $I_i \cap I_j = \emptyset$ for every $i, j \in \{1, \dots, k\}$. $f_{(I_1, \dots, I_k)}$ is the corresponding p -face of $\text{chain}(\binom{x_i}{p_i}_{1 \leq i \leq k})$.

Choosing a star from $f_{(I_1, \dots, I_k)}$ amounts to choosing $j \in \{1, \dots, k\}$ and a star in I_j . It therefore yields k intervals $(I_i)'$ defined by $(I_i)' = I_i$ if $i \neq j$ and $(I_j)' = I_j \setminus \{i_*\}$ where i_* is the coordinate of the chosen star.

Replacing $\{i_*\}$ by a zero or a one gives two translates of $f_{(I_1, \dots, I_j \setminus \{i_*\}, \dots, I_k)}$. To each $\tilde{I}_j$ such that $(I_1, \dots, \tilde{I}_j, \dots, I_k)$ is adapted to $\binom{x_i}{p_i}_{1 \leq i \leq k}$ and such that $\tilde{I}_j = (I_j \setminus \{i_*\}) \cup \{i_{x_j}\}$ for $i_{x_j} \in \text{Supp}(x_j)$ correspond two other translates of $f_{(I_1, \dots, I_j \setminus \{i_*\}, \dots, I_k)}$. When summed, some of these translates cancel pairwise and we are left with $f_{(I_1, \dots, I_j \setminus \{i_*\}, \dots, I_k)} + (f_{(I_1, \dots, I_j \setminus \{i_*\}, \dots, I_k)} \oplus x_j)$.

Summing over every possible $(I_1, \dots, I_j \setminus \{i_*\}, \dots, I_k)$ finishes the proof. $\square$

Corollary 51. For $c_1, \dots, c_k \in C_k$, $\text{chain}_k(\binom{c_i}{p_i}_{1 \leq i \leq k})$ is a p -cycle of $\mathcal{C}_p(Q_n/C_k)$. We denote it by $\text{cyc}_{p,k}(\binom{c_i}{p_i}_{1 \leq i \leq k})$.

Proof.

$$\forall j, \Pi_k(\text{chain}(\binom{c_i}{p_i - \delta_{i,j}}_{1 \leq i \leq k})) = \Pi_k((\text{chain}(\binom{c_i}{p_i - \delta_{i,j}}_{1 \leq i \leq k}) \oplus c_j))$$

$\square$

Corollary 52.

$$\partial_{\text{hom},p,k}[\text{cyc}_{p,k}(\binom{c_i}{p_i}_{1 \leq i \leq k})] = [\text{cyc}_{p-1,k}(\binom{c_i}{p_i - \delta_{i,k}}_{1 \leq i \leq k})]$$

Proof.

$$\begin{aligned}
& \partial(\pi_k^{-1}(\text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}))) \\
& \partial(\pi_k^{-1}(\Pi_k(\text{chain}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k})))) \\
& \partial(\Pi_{k-1}(\text{chain}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}))) \\
& = \Pi_{k-1}(\text{chain}(\left(\begin{smallmatrix} c_i \\ p_i - \delta_{i,j} \end{smallmatrix}\right)_{1 \leq i \leq k})) + \Pi_{k-1}(\text{chain}(\left(\begin{smallmatrix} c_i \\ p_i - \delta_{i,j} \end{smallmatrix}\right)_{1 \leq i \leq k}) \oplus c_k) \\
& = i_k(\Pi_{k-1}(\text{chain}(\left(\begin{smallmatrix} c_i \\ p_i - \delta_{i,j} \end{smallmatrix}\right)_{1 \leq i \leq k}))) \\
& = i_k(\text{cyc}_{p-1,k}(\left(\begin{smallmatrix} c_i \\ p_i - \delta_{i,k} \end{smallmatrix}\right)_{1 \leq i \leq k}))
\end{aligned}$$

Recalling that $\partial_{\text{hom},p,k}$ corresponds to $i_{\text{chain},p,k}^{-1} \circ \partial_{\text{chain},p,k} \circ \pi_{\text{chain},p,k}^{-1}$ finishes the proof. $\square$

4.8 Homology basis in a cube-quotient code

We are now ready to prove Thm. 53 by induction on $(p+k)$:

Theorem 53. *Let $c_1, \dots, c_k$ form a basis of C_k . $H_p(Q_n/C_k)$ has a basis indexed by k -tuples $(p_1, \dots, p_k)$ satisfying $p_1 + \dots + p_k = p$ and such that each basis element is the homology class represented by the product cycle $\text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k})$.*

Proof. We will use the following short exact sequence and apply the induction hypothesis to $H_p(Q_n/C_{k-1})$ and $H_{p-1}(Q_n/C_k)$:

$$0 \rightarrow H_p(Q_n/C_{k-1}) \xrightarrow{\pi_{\text{hom}}} H_p(Q_n/C_k) \xrightarrow{\partial_{\text{hom}}} H_{p-1}(Q_n/C_k) \rightarrow 0.$$

Since $\pi_{\text{chain},k}(\text{cyc}_{p,k-1}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k-1})) = \text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k})$ with $p_k = 0$, the basis of homology classes of $H_p(Q_n/C_{k-1})$ represented by product cycles is sent by π_{hom} to a free family of homology classes represented by the product cycles of $H_p(Q_n/C_k)$ satisfying $p_k = 0$.

Since $\partial_{\text{hom}}([\text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k})] = [\text{cyc}_{p-1,k}(\left(\begin{smallmatrix} c_i \\ p_i - \delta_{i,k} \end{smallmatrix}\right)_{1 \leq i \leq k})]$, the basis of homology classes of $H_{p-1}(Q_n/C_k)$ represented by product cycles has a free family of preimages by ∂_{hom} represented by the product cycles of $H_p(Q_n/C_k)$ satisfying $p_k \neq 0$.

The exactness of the short sequence implies that the concatenation of these two free families forms a basis of $H_p(Q_n/C_k)$. $\square$

4.9 Cocycle minimum distance in a cube-quotient code

Lemma 54. *For any product cycle $\text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k})$, for any $y \in (\mathbb{F}_2)^n$, the translate $\text{cyc}_{p,k}(\left(\begin{smallmatrix} c_i \\ p_i \end{smallmatrix}\right)_{1 \leq i \leq k}) \oplus y$ is a cycle which belongs to the same homology class as*

$$cyc_{p,k}\left(\begin{pmatrix} c_i \\ p_i \end{pmatrix}_{1 \leq i \leq k}\right).$$

Proof. The translate is a cycle since ∂_{chain} and $\oplus y$ commute.

To prove that translation doesn't alter the homology class we will show that

$cyc_{p,k}\left(\begin{pmatrix} c_i \\ p_i \end{pmatrix}_{1 \leq i \leq k}\right) + (cyc_{p,k}\left(\begin{pmatrix} c_i \\ p_i \end{pmatrix}_{1 \leq i \leq k}\right)) \oplus y$ is a boundary. Equivalently we will show that it is orthogonal to every cohomology class in $H^p(Q_n/C_k)$. It is sufficient to consider the canonical cocycles representing a basis of $H^p(Q_n/C_k)$. Observing that $cyc_{p,k}\left(\begin{pmatrix} c_i \\ p_i \end{pmatrix}_{1 \leq i \leq k}\right) + (cyc_{p,k}\left(\begin{pmatrix} c_i \\ p_i \end{pmatrix}_{1 \leq i \leq k}\right)) \oplus y$ has exactly 0 or 2 p -faces per p -direction finishes the proof. $\square$

Therefore each homology class of the product cycles basis of $H_p(Q_n/C_k)$ is represented by 2^{n-k} different cycles corresponding to the 2^{n-k} different translations $y \in \mathbb{F}_2/C_k$. Each p -face belongs to exactly 0 or 2^p of the 2^{n-k} different cycles. This observation leads to the following proposition:

Propositon 55. *The cocycle minimum distance $D_{p,k}^{(cohom)}$, i.e. the minimum weight of a cohomologically non trivial p -cocycle in $\mathcal{C}^p(Q_n/C_k)$ verifies:*

$$D_{p,k}^{(cohom)} = 2^{n-p-k}.$$

Proof. Let $cocyc_{p,k}$ be a cohomologically non trivial (p,k) -cocycle. $cocyc_{p,k}$ is not orthogonal to at least one product cycle representing an element of the basis of $H_p(Q_n/C_k)$. Therefore $cocyc_{p,k}$ is not orthogonal to any of the 2^{n-k} different cycles obtained by translating this product cycle. Since each p -face of $cocyc_{p,k}$ belongs to at most 2^p translated product cycles, $cocyc_{p,k}$ contains at least 2^{n-k-p} p -faces. Moreover the value 2^{n-p-k} is attained by canonical cocycles. $\square$

4.10 Cycle minimum distance in a cube-quotient code

Propositon 56. *The cycle minimum distance $D_{p,k}^{(hom)}$, i.e. the minimum weight of a homologically non trivial p -cycle in $\mathcal{C}_p(Q_n/C_k)$ verifies:*

$$D_{p,k}^{(hom)} = \binom{d}{p}.$$

Proof. We prove the following proposition by induction on $p + j$:

A homologically non trivial (p,k) -cycle is not orthogonal to at least $\binom{d}{p}$ canonical (p,k) -cocycles.

Since canonical cocycles are disjoint, the value of the cycle minimum distance follows immediately.

The initialisation is straightforward.

Let $cyc_{p,k}$ be a cycle representing a nontrivial homology class $h_{p,k} \in H_p(Q_n/C_k)$.

first case: $\partial_{hom}(h_{p,k}) = 0$ in $H_{p-1}(Q_n/C_k)$ for at least one decomposition $C_k = C_{k-1} \cup (C_{k-1} \oplus C_k)$.

Then there exists a non trivial homology class $h_{p,k-1} \in H_p(Q_n/C_{k-1})$ such that $h_{p,k} = \pi_{hom}(h_{p,k-1})$. Let $cyc_{p,k-1}$ be a $(p,k-1)$ -cycle representing $h_{p,k-1}$.

By the induction hypothesis there are $\binom{d}{p}$ canonical $(p, k-1)$ -cocycles not orthogonal to $cyc_{p,k-1}$. Let $cocyc_{p,k-1}$ be such a cocycle. Since π_{chain} and $i_{cochain}$ are adjoint:

$$\begin{aligned} & \langle i_{cochain}^{-1}(cocyc_{p,k-1}), cyc_{p,k} \rangle \\ &= \langle cocyc_{p,k-1}, \pi_{chain}^{-1}(cyc_{p,k}) \rangle \\ &= \langle cocyc_{p,k-1}, cyc_{p,k-1} \rangle \\ &= 1. \end{aligned}$$

Therefore applying $i_{cochain}^{-1}$ to the $\binom{d}{p}$ canonical $(p, k-1)$ -cocycles not orthogonal to $cyc_{p,k-1}$ yields $\binom{d}{p}$ canonical (p, k) -cocycles not orthogonal to $cyc_{p,k}$. The induction hypothesis is proven in this case.

second case: $\partial_{hom}(h_{p,k}) \neq 0$ in $H_{p-1}(Q_n/C_k)$ for every decomposition $C_k = C_{k-1} \cup (C_{k-1} \oplus c_k)$.

By definition of ∂_{hom} , any preimage $i_{chain}^{-1} \circ \partial_{chain} \circ \pi_{chain}^{-1}(cyc_{p,k})$ represents $\partial_{hom}(h_{p,k})$.

By the induction hypothesis there exists $\binom{d}{p-1}$ distinct canonical $(p-1, k)$ -cocycles orthogonal to $i_{chain}^{-1} \circ \partial_{chain} \circ \pi_{chain}^{-1}(cyc_{p,k})$. Let $cocyc_{p-1,k}$ be such a cocycle. Any preimage $i_{cochain}^{-1} \circ \delta_{cochain} \circ \pi_{cochain}^{-1}(cocyc_{p-1,k})$ is a (p, k) -cocycle orthogonal to $cyc_{p,k}$:

$$\begin{aligned} &= \langle i_{cochain}^{-1} \circ \delta_{cochain} \circ \pi_{cochain}^{-1}(cocyc_{p-1,k}), cyc_{p,k} \rangle \\ &= \langle \delta_{cochain} \circ \pi_{cochain}^{-1}(cocyc_{p-1,k}), \pi_{chain}^{-1}(cyc_{p,k}) \rangle \\ &= \langle \pi_{cochain}^{-1}(cocyc_{p-1,k}), \partial_{chain} \circ \pi_{chain}^{-1}(cyc_{p,k}) \rangle \\ &= \langle cocyc_{p-1,k}, i_{chain}^{-1} \circ \partial_{chain} \circ \pi_{chain}^{-1}(cyc_{p,k}) \rangle \\ &= 1. \end{aligned}$$

Let us count the number of canonical (p, k) -cocycles $i_{cochain}^{-1} \circ \delta_{cochain} \circ \pi_{cochain}^{-1}(cocyc_{p-1,k})$ that we can construct from the $\binom{d}{p-1}$ distinct canonical $(p-1, k)$ -cocycles $cocyc_{p-1,k}$.

Since $i_{cochain}$ is a bijection, $i_{cochain}^{-1}$ is uniquely defined. But $\pi_{cochain}^{-1}(cocyc_{p-1,k})$ can be any preimage of $cocyc_{p-1,k}$ by $\pi_{cochain}$. We use the same technique as in the construction of the cohomology basis represented by canonical cocycles.

The k^{th} element of the basis of the classical code c_k has weight at least d . Let I be the $p-1$ -direction of the canonical cocycle $cocyc_{p-1,k}$. At least $(d-p+1)$ coordinates are in $\text{Supp}(c_k) \setminus I$. Denoting by j one of these $(d-p+1)$ coordinates, the $(p-1)$ -cochain obtained by only keeping the $(p-1)$ -faces of $cocyc_{p-1,k}$ having a 0 at coordinate j is a preimage of $cocyc_{p-1,k}$ by the $\pi_{cochain}$ associated to a decomposition $C_k = C_{k-1} \cup (C_{k-1} \oplus c_k)$ such that $\forall x \in C_{k-1}, x_j = 0$. Applying $\delta_{cochain}$ to this cochain amounts to replacing this 0 at coordinate j of every $(p-1)$ -face by a $*$ and yields $cocyc_{I \cup \{j\}, p, k-1}$. Applying $i_{cochain}^{-1}$ gives $cocyc_{I \cup \{j\}, p, k}$.

With this procedure each canonical (p, k) -cocycle $cocyc_{I \cup \{j\}, p, k}$ has been counted at most p times. We have therefore constructed at least $\frac{d-p+1}{p} \binom{d}{p-1} = \binom{d}{p}$ distinct canonical (p, k) -cocycle orthogonal to $cyc_{p,k}$. The induction hypothesis is proven in this case too.

Moreover the value $\binom{d}{p}$ is attained by the product cycles $cyc_{p,k}(\binom{c_i}{p_i}_{1 \leq i \leq k})$ such that $p_1 = p, p_{i \neq 1} = 0$ and c_1 has weight d . $\square$

4.11 The non-local subgroup condition on cube-quotient codes

It is possible to express cube-quotient codes in the language of Chapter 3. Indeed an n -cube is a regular abstract polytope in the sense of Chapter 3. It corresponds to the $\{4, 3, \dots, 3\}$ string Coxeter group. For an n -cube, there are $(n - 2)$ integers 3 in the Schläfli symbol $\{4, 3, \dots, 3\}$. This notation encodes a presentation of the symmetry group of the n -cube. This group is called the hyperoctahedral group and we denote it by Γ_n .

$$\begin{aligned}\Gamma_n &= \langle r_0, \dots, r_{n-1} \rangle \\ \forall i \in \{0, \dots, n-1\}, r_i^2 &= \text{id} \\ (r_0 r_1)^4 &= \text{id} \\ \forall i \in \{0, \dots, n-2\}, (r_i r_{i+1})^3 &= \text{id} \\ \forall i \in \{0, \dots, n-3\}, \forall j \in \{i+2, \dots, n-1\}, (r_i r_j)^2 &= \text{id}.\end{aligned}$$

The hyperoctahedral group can also be written as a semidirect product of $(\mathbb{F}_2)^n$ and Sym_n . Each $\mathbb{F}_2$ factor corresponds to the translation of a coordinate (0 and 1 labels are interchanged). The symmetric group Sym_n corresponds to permutations of the n coordinates. Since permuting coordinates and translating them are non commutative operations, the group law is a semidirect product and not a direct product:

$$\begin{aligned}\Gamma &\simeq \text{Sym}_n \ltimes (\mathbb{F}_2)^n \text{ with this product law:} \\ (\sigma, x)(\tau, y) &= (\sigma\tau, x \oplus \sigma(y))\end{aligned}$$

Under this isomorphism, Γ_n acts on the vertices of the hypercube Q^n like it acts on $(\mathbb{F}_2)^n$: $(\tau, y) \cdot z = y \oplus \tau(z)$. It is straightforward to verify that $(\sigma, x)(\tau, y) \cdot z = x \oplus \sigma(y) \oplus \sigma\tau(z) = (\sigma\tau, x \oplus \sigma(y)) \cdot z$. The action of Γ_n on higher dimensional faces of Q^n is induced by its action on vertices.

Interestingly, quotienting by a classical code is the same as quotienting by a subgroup of Γ_n which is actually a subgroup of $(\mathbb{F}_2)^n$: only the translation symmetries are considered and the permutation symmetries are not used for quotienting. Usually, these subgroups are not normal subgroups of Γ_n since a classical code is not invariant under arbitrary permutations of the coordinates. However we saw in Chapter 3 that the non-local subgroup condition is sufficient to define a quantum code from a quotient of a regular abstract polytope. Let us verify that this condition is satisfied by the subgroup corresponding to a classical code with minimum distance at least $p + 2$.

We consider the subgroups S_p and S_{p+1} corresponding respectively to qubits and Z -checks.

$$S_p = \langle r_0, \dots, r_{p-1}, r_{p+1}, \dots, r_{n-1} \rangle.$$

Since we are working with the $\text{Sym}_n \ltimes (\mathbb{F}_2)^n$ version of Γ_n , we have to compute the images of the r_i 's under this isomorphism. From the cases $n = 2$ and $n = 3$ we see that r_0 corresponds to the pure translation by 10...0, that r_1 corresponds to the pure permutation (12) of the two first coordinates and that r_2 corresponds to the pure permutation (23) of the coordinates 2 and three. Based on these observations, it is reasonable to guess that for all $i \in \{1, \dots, n-1\}$, r_i corresponds to the pure permutation $(i \ i+1)$. It is straightforward to verify a posteriori that this defines the correct isomorphism. Indeed every relationship of the Coxeter group generators is satisfied by their images in $S_n \ltimes (\mathbb{F}_2)^n$ and the two

groups have the same cardinal (namely $2^n n!$). The following identifications summarize the isomorphism:

$$\begin{aligned} r_0 &\leftrightarrow (\text{id}, 10\dots 0) \\ r_1 &\leftrightarrow ((12), 0\dots 0) \\ &\dots \\ r_{n-1} &\leftrightarrow ((n-1 \quad n), 0\dots 0). \end{aligned}$$

Observe that

$$\begin{aligned} S_i &= \langle r_0, \dots, r_{i-1}, r_{i+1}, \dots, r_{n-1} \rangle \\ S_i &= \langle r_0, \dots, r_{i-1} \rangle \times \langle r_{i+1}, \dots, r_{n-1} \rangle. \\ S_i &\simeq \Gamma_i \times \text{Sym}_{n-i} \end{aligned}$$

since r_i and r_j commute when $|i - j| \geq 2$. We recognize the direct product of the hyperoctahedral group on the i first coordinates and the symmetric group on the $(n - i)$ last coordinates.

Given a pure translation subgroup H corresponding to a classical code with minimum distance $p + 2$, we want to show that for all $g \in \Gamma_n$, $gHg^{-1} \cap S_p S_{p+1} = gHg^{-1} \cap S_{p+1} S_p = \{\text{id}\}$. We first focus on gHg^{-1} . Let $g = (\sigma, y)$ be an element of $\text{Sym}_n \ltimes (\mathbb{F}_2)^n$. Its inverse is $g^{-1} = (\sigma^{-1}, \sigma^{-1}(y))$. Let $h = (\text{id}, x)$ be an element of the subgroup H .

$$\begin{aligned} ghg^{-1} &= (\sigma, y)(\text{id}, x)(\sigma^{-1}, \sigma^{-1}(y)) \\ ghg^{-1} &= (\sigma, y)(\sigma^{-1}, x \oplus \sigma^{-1}(y)) \\ ghg^{-1} &= (\text{id}, y \oplus \sigma(x) \oplus \sigma(\sigma^{-1}(y))) \\ ghg^{-1} &= (\text{id}, \sigma(x)). \end{aligned}$$

This computation shows that ghg^{-1} is a pure translation by a word $\sigma(x)$ that has the same Hamming weight as x . Therefore for $h \in H \setminus \{\text{id}\}$, ghg^{-1} is a pure translation by a word $\sigma(x)$ of weight at least $p + 2$.

Suppose by contradiction that there is a non trivial element in $ghg^{-1} = s_p s_{p+1}$. Then the permutation parts of s_p and s_{p+1} must be each other's inverse. Since S_p permutes the p first coordinates and the $(n - p)$ last coordinates independently and a similar statement holds for S_{p+1} , this common permutation must permute the p first coordinates and the $(n - p - 1)$ last coordinates independently (and fix the p^{th} coordinate). Let us denote such a permutation by τ and the translation parts of s_p and s_{p+1} by y and z respectively.

$$\begin{aligned} s_p s_{p+1} &= (\tau, y)(\tau^{-1}, z) \\ s_p s_{p+1} &= (\text{id}, y \oplus \tau(z)) \end{aligned}$$

Moreover we know that y and z are non zero only on the p and $p + 1$ first coordinates respectively. Therefore the same is true of $\tau(z)$ and $y \oplus \tau(z)$. It implies that $y \oplus \tau(z)$ has weight at most $p + 1$ and therefore cannot equal $\sigma(x)$. We have proven that quotienting by a classical code whose minimum distance is at least $p + 2$ corresponds to quotienting by a subgroup H such that for all $g \in \Gamma_n$, $gHg^{-1} \cap S_p S_{p+1} = \{\text{id}\}$. The same reasoning shows that the corresponding subgroup H also satisfies $gHg^{-1} \cap S_{p+1} S_p = \{\text{id}\}$, $gHg^{-1} \cap S_p S_{p-1} = \{\text{id}\}$ and $gHg^{-1} \cap S_{p-1} S_p = \{\text{id}\}$.

Conclusion

In this PhD thesis, we have constructed and analyzed two families of quantum error correcting codes.

First the family of 4-dimensional hyperbolic codes. Following Guth and Lubotzky [GL14], we have given a regular version of the 4-dimensional hyperbolic code family. This regularity has many practical advantages. It makes possible a finer analysis of local decoders for 4-dimensional hyperbolic codes. It allows us to explicitly compute parity-check matrices for small instances of this family of codes. It gives the possibility to further quotient such codes without modifying their local 4-dimensional hyperbolic structure. All these improvements made the 4-dimensional hyperbolic code family amenable to numerical simulations. We were able to show numerically a high threshold for the 4-dimensional family of code with a Belief Propagation decoder with both noiseless and noisy measurements.

We also tried, without success so far, to give a purely combinatorial proof of the minimum distance of 4-dimensional hyperbolic codes. Indeed the proof today is the one given in [GL14]. It relies on the representation of groups as isometry groups in a hyperbolic space and Anderson's theorem (see [GL14] Theorem 17 p. 15). We believe that the use of Coxeter groups could make possible a combinatorial analog of Anderson's theorem. Such a proof could open the door to designing new quantum codes, which share some properties with 4-dimensional quantum codes but which may be less geometric in their construction. Unfortunately we failed to find such a proof, which means that it could be material for future work.

Second the family of cube-quotient codes. Investigated with the hope of finding good minimum distance properties, this family actually yields an interesting trade-off between its number of logical qubits and its minimum distance. It is indeed possible to extract a one-parameter family such that both scale like a power of the number of physical qubits. Even though these $[[N, K, D]]$ parameters are not as good as the state of the art given by hypergraph product codes [TZ14], this family is remarkable for its originality. Indeed taking the quotient of the cube by a classical error correcting code corresponds to taking the quotient of the symmetry group of the cube by a translation subgroup. Such a subgroup is not normal in general. The geometric connection is therefore lost by this type of quotients and novel proof techniques had to be introduced: the long exact sequence Lemma from homological algebra played a central role in our proofs.

It would be interesting to study decoding algorithms on cube-quotient codes. We tried the small set flip decoder ([LTZ15]) on the hemicube code corresponding to the 8-cube with qubits on 2-faces and saw a good decoding capacity. However we did not include this work in this manuscript because we think it would be more interesting to study systematically a less generic decoding algorithm. Our attempts to design such a quotient-cube tailored decoding algorithm were unfortunately not conclusive and could be material for future work.

It is interesting to note that hypergraph product codes ([TZ14]), 4-dimensional hyperbolic

codes and quotient-cube codes don't attract as much attention from the quantum computing community as could be expected since they exhibit much better $[[N, K, D]]$ parameters than the toric and surface codes. The reason for the success of the toric and surface codes is that they can be implemented isometrically on a (Euclidean) plane. This property is very appealing to experimentalists. However a quantum code is always defined on a finite number of physical qubits. Therefore any quantum code can be engineered on a (Euclidean) plane if non local measurements are allowed. I would not be surprised if hypergraph product codes, 4-dimensional codes and quotient-cube codes would gain popularity once the toric and surface codes have been realised experimentally and the fact that they encode only a constant number of logical qubits becomes the limiting factor.

Bibliography

- [AS⁺92] Uwe Abresch, Viktor Schroeder, et al. Graph manifolds, ends of negatively curved spaces and the hyperbolic 120-cell space. *Journal of Differential Geometry*, 35(2):299–336, 1992.
- [Aud13] Benjamin Audoux. An application of khovanov homology to quantum codes. *arXiv preprint arXiv:1307.4677*, 2013.
- [BDMT16] Nikolas P Breuckmann, Kasper Duivenvoorden, Dominik Michels, and Barbara M Terhal. Local decoders for the 2d and 4d toric code. *arXiv preprint arXiv:1609.00510*, 2016.
- [Bou07] N. Bourbaki. *Groupes et algèbres de Lie: Chapitres 4, 5 et 6*. Bourbaki, Nicolas. Springer Berlin Heidelberg, 2007.
- [BPT10] Sergey Bravyi, David Poulin, and Barbara Terhal. Tradeoffs for reliable quantum information storage in 2d systems. *Physical review letters*, 104(5):050503, 2010.
- [BVC⁺17] Nikolas P Breuckmann, Christophe Vuillot, Earl Campbell, Anirudh Krishna, and Barbara M Terhal. Hyperbolic and semi-hyperbolic surface codes for quantum storage. *arXiv preprint arXiv:1703.00590*, 2017.
- [CHE96] BY SHIING-SHEN CHERN. On the curvatura integra in a riemannian manifold. In *A Mathematician And His Mathematical Work: Selected Papers of SS Chern*, pages 121–131. 1996.
- [Cox54] Harold Stephen Macdonald Coxeter. Regular honeycombs in hyperbolic space. In *Proceedings of the International Congress of Mathematicians*, volume 3, pages 155–169, 1954.
- [Del13] Nicolas Delfosse. Tradeoffs for reliable quantum information storage in surface codes and color codes. In *Information Theory Proceedings (ISIT), 2013 IEEE International Symposium on*, pages 917–921. IEEE, 2013.
- [DKLP02] Eric Dennis, Alexei Kitaev, Andrew Landahl, and John Preskill. Topological quantum memory. *Journal of Mathematical Physics*, 43(9):4452–4505, 2002.
- [DN17] Nicolas Delfosse and Naomi H Nickerson. Almost-linear time decoding algorithm for topological codes. *arXiv preprint arXiv:1709.06218*, 2017.
- [DZ17] Nicolas Delfosse and Gilles Zémor. Linear-time maximum likelihood decoding of surface codes over the quantum erasure channel. *arXiv preprint arXiv:1703.01517*, 2017.

- [FGL18a] Omar Fawzi, Antoine Grospellier, and Anthony Leverrier. Constant overhead quantum fault-tolerance with quantum expander codes. In *2018 IEEE 59th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 743–754. IEEE, 2018.
- [FGL18b] Omar Fawzi, Antoine Grospellier, and Anthony Leverrier. Efficient decoding of random errors for quantum expander codes. In *Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing*, pages 521–534. ACM, 2018.
- [FML02] Michael H Freedman, David A Meyer, and Feng Luo. Z2-systolic freedom and quantum codes. *Mathematics of quantum computation, Chapman & Hall/CRC*, pages 287–320, 2002.
- [Gal62] Robert Gallager. Low-density parity-check codes. *IRE Transactions on information theory*, 8(1):21–28, 1962.
- [GL14] Larry Guth and Alexander Lubotzky. Quantum error correcting codes and 4-dimensional arithmetic hyperbolic manifolds. *Journal of Mathematical Physics*, 55(8):082202, 2014.
- [Got13] Daniel Gottesman. Fault-tolerant quantum computation with constant overhead. *arXiv preprint arXiv:1310.2984*, 2013.
- [Har04] James William Harrington. *Analysis of quantum error-correcting codes: symplectic lattice codes and toric codes*. PhD thesis, California Institute of Technology, 2004.
- [Has16] MB Hastings. Quantum codes from high-dimensional manifolds. *arXiv preprint arXiv:1608.05089*, 2016.
- [Kit03] A Yu Kitaev. Fault-tolerant quantum computation by anyons. *Annals of Physics*, 303(1):2–30, 2003.
- [LL19] Vivien Londe and Anthony Leverrier. Golden codes: quantum ldpc codes from regular tessellations of hyperbolic 4-manifolds. *Quantum Information and Computation*, 19(5, 6):361–391, May 2019.
- [LLZ19] Anthony Leverrier, Vivien Londe, and Gilles Zémor. Towards local testability for quantum coding. *preprint*, 2019.
- [LTZ15] Anthony Leverrier, Jean-Pierre Tillich, and Gilles Zémor. Quantum expander codes. In *2015 IEEE 56th Annual Symposium on Foundations of Computer Science*, pages 810–824. IEEE, 2015.
- [Mar77] Daniel A Marcus. *Number fields*, volume 8. Springer, 1977.
- [Mar15] Bruno Martelli. Hyperbolic four-manifolds. *arXiv preprint arXiv:1512.03661*, 2015.
- [MS02] Peter McMullen and Egon Schulte. *Abstract regular polytopes*, volume 92. Cambridge University Press, 2002.
- [Mur16] Plinio GP Murillo. Systole of congruence coverings of arithmetic hyperbolic manifolds. *arXiv preprint arXiv:1610.03870*, 2016.
- [Mur17] Plinio GP Murillo. *On Arithmetic Manifolds with Large Systole*. PhD thesis, PhD thesis, IMPA, 2017.

- [NC02] Michael A Nielsen and Isaac Chuang. Quantum computation and quantum information, 2002.
- [Rat06] John Ratcliffe. *Foundations of hyperbolic manifolds*, volume 149. Springer Science & Business Media, 2006.
- [Rot08] Joseph J Rotman. *An introduction to homological algebra*. Springer Science & Business Media, 2008.
- [RU08] Tom Richardson and Ruediger Urbanke. *Modern coding theory*. Cambridge university press, 2008.
- [Sho95] Peter W Shor. Scheme for reducing decoherence in quantum computer memory. *Physical review A*, 52(4):R2493, 1995.
- [tes] <https://mathcs.clarku.edu/~djoyce/poincare/poincare.html>. Accessed: 2010-09-30.
- [TZ14] Jean-Pierre Tillich and Gilles Zémor. Quantum ldpc codes with positive rate and minimum distance proportional to the square root of the blocklength. *IEEE Transactions on Information Theory*, 60(2):1193–1202, 2014.
- [Wei95] Charles A Weibel. *An introduction to homological algebra*. Number 38. Cambridge university press, 1995.
- [Wil09] Robert Wilson. *The finite simple groups*, volume 251. Springer Science & Business Media, 2009.
- [Zém09] Gilles Zémor. On cayley graphs, surface codes, and the limits of homological coding for quantum error correction. In *International Conference on Coding and Cryptology*, pages 259–273. Springer, 2009.

Titre : Codes correcteurs d'erreur quantique topologiques au-delà de la dimension 2

Résumé : La mémoire quantique est constituée de matériaux présentant des effets quantiques comme la superposition. C'est cette possibilité de superposition qui distingue l'élément élémentaire de mémoire quantique, le qubit, de son analogue classique, le bit. Contrairement à un bit classique, un qubit peut être dans un état différent de l'état 0 et de l'état 1. Une difficulté majeure de la réalisation physique de mémoire quantique est la nécessité d'isoler le système utilisé de son environnement. En effet l'interaction d'un système quantique avec son environnement entraîne un phénomène appelé décohérence qui se traduit par des erreurs sur l'état du système quantique. Dit autrement, à cause de la décohérence, il est possible que les qubits ne soient pas dans l'état dans lequel il est prévu qu'ils soient. Lorsque ces erreurs s'accumulent le résultat d'un calcul quantique a de grandes chances de ne pas être le résultat attendu.

La correction d'erreur quantique est un ensemble de techniques permettant de protéger l'information quantique de ces erreurs. Elle consiste à réaliser un compromis entre le nombre de qubits et leur qualité. Plus précisément un code correcteur d'erreur permet à partir de N qubits physiques bruités de simuler un nombre plus petit K de qubits logiques, c'est-à-dire virtuels, moins bruités. La famille de codes la plus connue est sans doute celle découverte par le physicien Alexei Kitaev: le code torique. Cette construction peut être généralisée à des formes géométriques (variétés) autres qu'un tore. En 2014, Larry Guth et Alexander Lubotzky proposent une famille de code définie à partir de variétés hyperboliques de dimension 4 et montrent que cette famille fournit un compromis intéressant entre le nombre K de qubits logiques et le nombre d'erreurs qu'elle permet de corriger.

Dans cette thèse, nous sommes partis de la construction de Guth et Lubotzky et en avons donné une version plus explicite et plus régulière. Pour définir un pavage régulier de l'espace hyperbolique de dimension 4, nous utilisons le groupe de symétrie de symbole de Schläfli $\{4, 3, 3, 5\}$. Nous en donnons la représentation matricielle correspondant au modèle de l'hyperboloïde et à un hypercube centré sur l'origine et dont les faces sont orthogonales aux quatre axes de coordonnées. Cette construction permet d'obtenir une famille de codes quantiques encodant un nombre de qubits logiques proportionnel au nombre de qubits physiques et dont la distance minimale croît au moins comme $N^{0.1}$. Bien que ces paramètres soient également ceux de la construction de Guth et Lubotzky, la régularité de cette construction permet de construire explicitement des exemples de taille raisonnable et d'envisager des algorithmes de décodage qui exploitent cette régularité.

Dans un second chapitre nous considérons une famille de codes quantiques hyperboliques 4D de symbole de Schläfli $\{5, 3, 3, 5\}$. Après avoir énoncé une façon de prendre le quotient des groupes correspondant en conservant la structure locale du groupe, nous construisons les matrices de parité correspondant à des codes quantiques ayant 144, 720, 9792, 18 000 et 90 000 qubits physiques. Nous appliquons un algorithme de Belief Propagation au décodage de ces codes et analysons les résultats numériquement.

Dans un troisième et dernier chapitre nous définissons une nouvelle famille de codes quantiques à partir de cubes de dimension arbitrairement grande. En prenant le quotient d'un cube de dimension n par un code classique de paramètres $[n, k, d]$ et en identifiant les qubits physiques avec les faces de dimension p du polytope quotient ainsi défini, on obtient un code quantique. Cette famille de codes quantiques a l'originalité de considérer des quotients par des codes classiques. En cela elle s'éloigne de la topologie et appartient plutôt à la famille des codes homologiques.

Mots clés : codes correcteurs, quantique, topologie, homologie

Title : Topological Quantum Error-Correcting Codes beyond dimension 2

Abstract : Error correction is the set of techniques used in order to store, process and transmit information reliably in a noisy context. The classical theory of error correction is based on encoding classical information redundantly. A major endeavor of the theory is to find optimal trade-offs between redundancy, which we try to minimize, and noise tolerance, which we try to maximize.

The quantum theory of error correction cannot directly imitate the redundant schemes of the classical theory because it has to cope with the no-cloning theorem: quantum information cannot be copied. Quantum error correction is nonetheless possible by spreading the information on more quantum memory elements than would be necessary. In quantum information theory, dilution of the information replaces redundancy since copying is forbidden by the laws of quantum mechanics.

Besides this conceptual difference, quantum error correction inherits a lot from its classical counterpart. In this PhD thesis, we are concerned with a class of quantum error correcting codes whose classical counterpart was defined in 1961 by Gallager [Gal62]. At that time, quantum information was not even a research domain yet. This class is the family of low density parity check (LDPC) codes. Informally, a code is said to be LDPC if the constraints imposed to ensure redundancy in the classical setting or dilution in the quantum setting are local.

More precisely, this PhD thesis focuses on a subset of the LDPC quantum error correcting codes: the homological quantum error correcting codes. These codes take their name from the mathematical field of homology, whose objects of study are sequences of linear maps such that the kernel of a map contains the image of its left neighbour. Originally introduced to study the topology of geometric shapes, homology theory now encompasses more algebraic branches as well, where the focus is more abstract and combinatorial. The same is true of homological codes: they were introduced in 1997 by Kitaev [Kit03] with a quantum code that has the shape of a torus. They now form a vast family of quantum LDPC codes, some more inspired from

geometry than others. Homological quantum codes were designed from spherical, Euclidean and hyperbolic geometries, from 2-dimensional, 3-dimensional and 4-dimensional objects, from objects with increasing and unbounded dimension and from hypergraph or homological products.

After introducing some general quantum information concepts in the first chapter of this manuscript, we focus in the two following ones on families of quantum codes based on 4-dimensional hyperbolic objects. We highlight the interplay between their geometric side, given by manifolds, and their combinatorial side, given by abstract polytopes. We use both sides to analyze the corresponding quantum codes. In the fourth and last chapter we analyze a family of quantum codes based on spherical objects of arbitrary dimension. To have more flexibility in the design of quantum codes, we use combinatorial objects that realize this spherical geometry: hypercube complexes. This setting allows us to introduce a new link between classical and quantum error correction where classical codes are used to introduce homology in hypercube complexes.

Keywords : error-correcting codes, quantum, homology, topology

Unité de recherche

Institut de Mathématiques de Bordeaux

UMR 5251

351, cours de la Liberation
F-33405 Talence Cedex, France