

Propriedades Algébricas e Geométricas dos Códigos de Bloco Quânticos

Tese apresentada à Faculdade de Engenharia Elétrica e de Computação da Universidade Estadual de Campinas como parte dos requisitos exigidos para a obtenção do título de Mestre em Engenharia Elétrica.

por

Wanessa Carla Gazzoni

Orientador: **Prof. Dr. Reginaldo Palazzo Júnior** FEEC/UNICAMP

Banca Examinadora

Prof. Dr. Reginaldo Palazzo Júnior (Presidente)
Prof. Dr. Celso de Almeida
Prof. Dr. Marcelo Firer
Prof^a. Dr^a. Sueli Irene Rodrigues Costa

Propriedades Algébricas e Geométricas dos Códigos de Bloco Quânticos

Este exemplar corresponde à redação final da dissertação devidamente corrigida e defendida por Wanessa Carla Gazzoni e aprovada pela banca examinadora.

Campinas, 2 de abril de 2004.

Prof. Dr. Reginaldo Palazzo Júnior (Orientador)

Banca Examinadora:

Prof. Dr. Reginaldo Palazzo Júnior (FEEC/UNICAMP)

Prof. Dr. Celso de Almeida (FEEC/UNICAMP)

Prof. Dr. Marcelo Firer (IMECC/UNICAMP)

Prof^a. Dr^a. Sueli Irene Rodrigues Costa (IMECC/UNICAMP)

Tese apresentada na Faculdade de Engenharia Elétrica e de Computação da Universidade Estadual de Campinas como parte dos requisitos exigidos para a obtenção do título de Mestre em Engenharia Elétrica.

Resumo

Este trabalho tem como objetivo apresentar a idéia geral da teoria matemática envolvida no processo de codificação de dados utilizando estados quânticos. Para isso, apresentaremos conceitos e propriedades da Teoria Quântica, que caracterizam o novo ambiente para a construção de códigos. Definiremos o grupo de operadores de erros que podem atuar no processo de transmissão através de um canal quântico e, baseados neste grupo, estudaremos a estrutura dos códigos corretores de erros, visando encontrar condições para a eficiência destes. Entre os códigos corretores de erros quânticos e suas propriedades, destacaremos a classe de códigos estabilizadores, cuja estrutura tem correspondência com a formalização dos códigos clássicos gerados sobre $GF(4)$.

Palavras-Chave: codificação por estados quânticos, códigos corretores de erros quânticos, grupo de operadores de erros quânticos, códigos estabilizadores.

Abstract

This research aims at presenting the general idea of the mathematical concepts and procedures for data encoding by use of quantum states. In this direction, basic concepts and properties from Quantum Mechanics are presented with the objective of code constructions. The quantum error groups acting on general quantum states is defined. Based on the properties of this group the algebraic structure of the error correcting codes is specified with the purpose of establishing the conditions under which the code achieves its maximum efficiency. Among the classes of quantum error correcting codes we consider the class of stabilizer codes for its robustness, rich algebraic structures, and its correspondence with classical codes over $GF(4)$.

Key-words: encoding of quantum states, quantum error correcting codes, quantum error groups, stabilizer codes.

Agradecimentos

A Deus, que na sua infinita bondade me reservou uma família maravilhosa: Osmar, Wilma e Rafael.

Ao Prof. Dr. Reginaldo Palazzo Jr., cuja convivência é para mim grande honra: obrigada pela atenção, pelas boas conversas e, principalmente, pela paciência.

A meus amigos de graduação do IMECC, em especial a Gabriel Lima e Daniel Miranda. Aos também amigos Jayme Vaz, Marcelo Firer, Alberto Saa e Benjamin Bordin, meus agradecimentos pelo apoio e pelo que pude aprender com cada um nas pequenas turbulências da vida. Pelo mesmo motivo devo recordar-me de Luciana Diógenes.

A Luiz Henrique Tizei pelo constante carinho e atenção.

Aos amigos do DT Antônio Carlos Aido, Rodrigo Cavalcante e Luis Milla-Leon pela colaboração e a Leandro Aguiar pelo auxílio na fase das correções.

Aos funcionários da FEEC e do IMECC que colaboram, sempre que possível, para nossa melhor formação e vivência acadêmica, com tudo o que isto pode incluir.

A fundação CAPES que financiou meus estudos de Mestrado.

A MEUS PAIS

Conteúdo

Resumo e Abstract	i
Agradecimentos	ii
Dedicatória	iii
Conteúdo	iv
Lista de Figuras	vi
Lista de Símbolos	vii
1 Introdução	1
1.1 Regras de um novo jogo	2
1.2 Do Bit para o Qubit	4
1.2.1 O produto tensorial	6
1.3 Emaranhados	10
1.4 Algoritmos Quânticos	11
1.5 A Construção de Computadores Quânticos	12
1.6 A Internet Quântica	13
2 O Grupo de Erros	15
2.1 Matriz Densidade	15
2.2 O Grupo de Pauli	18
2.3 Grupos de Clifford	22
2.4 Estrutura Matemática do Grupo de Erros	23
2.4.1 Considerações sobre $\bar{E}$	27
2.5 A Idéia dos Códigos	31
3 Códigos Corretores de Erros Quânticos	34
3.1 Analogias na Teoria Quântica de Correção	35
3.2 Modelo para Erros de Fase	40

3.3	Código de Repetição de Shor com Nove Qubits	42
3.4	Crítérios Gerais para Correção de Erros	44
3.5	Exemplo de um Código Corretor de Erros Quânticos	48
3.5.1	Formalismo dos códigos clássicos lineares	48
3.5.2	Códigos CSS	50
3.6	Regras para a Construção de Códigos	50
3.6.1	Limitantes para a correção quântica	51
4	Códigos Estabilizadores	54
4.1	O Formalismo Estabilizador	55
4.2	Alguns Exemplos de Códigos Estabilizadores	60
4.3	Códigos Clássicos Construídos a partir de Códigos Estabilizadores	68
4.4	Códigos sobre $GF(4)$	71
	Conclusão	75
	Trabalhos Futuros	78
	Bibliografia	79

Lista de Figuras

1.1	Esfera de Bloch e a representação de um estado $ \psi\rangle$	5
4.1	Eficiência de um código estabilizador.	55

Lista de Símbolos

$ \cdot\rangle$	vetor coluna	3
$\langle\cdot $	$ \cdot\rangle^t$	3
$\otimes$	produto tensorial	3
$ \cdot\rangle$	$ \cdot\rangle \otimes \cdot\rangle$	3
$\langle\cdot \cdot\rangle$	produto interno usual	4
$*$	operador de complexo conjugado	4
$\ \cdot\ $	norma vetorial usual	5
$\neg$	complementar binário	6
$\dagger$	operador de complexo conjugado transposto	6
$ \cdot\rangle\langle\cdot $	produto interno transposto	7
$O(2^n)$	grupo de matrizes ortogonais $2^n \otimes 2^n$ sobre $\mathbb{Z}_2^n$	21
$O(2n)$	grupo de matrizes ortogonais $2n \otimes 2n$ sobre $\mathbb{Z}_2^n$	21
$U(2^n)$	grupo de matrizes unitárias $2^n \otimes 2^n$ sobre $\mathbb{C}_2^n$	21
L/G_n	grupo de classes laterais de L quocientado pelo grupo G_n	21
$o(E)$	número de elementos do grupo E	22
$(a b)$	vetor de comprimento $2n$ e entradas em $\{0,1\}$	27

Capítulo 1

Introdução

Neste trabalho, nosso objetivo foi aliar conceitos da Teoria Quântica, propriedades da Teoria de Grupos e condições para a correção de erros no contexto quântico seguidas de exemplos dos códigos corretores quânticos mais conhecidos. Entre estes está a classe dos códigos estabilizadores, cujo formalismo será apresentado. As analogias com a teoria clássica, construídas via códigos sobre $GF(4)$, também foram apontadas de forma a tornar possível a compreensão desta nova pesquisa, que é a codificação de informações utilizando estados quânticos, e as possibilidades para a sua expansão.

Neste primeiro capítulo, nos preocupamos em apontar a evolução da questão em estudo, principalmente na área da implementação. Para isto, definimos a unidade básica da informação quântica e o caracter não local dos estados quânticos (emaranhamento), que torna possível a realização de certas tarefas no processamento da informação.

No Capítulo 2, estudamos o grupo de erros e a redução do mesmo ao grupo de Pauli baseados em conceitos da Teoria de Grupos. Sobre a redução construída, definimos a operação fundamental e propriedades como a ortogonalidade, o que nos possibilitou a introdução dos fundamentos das transformações do grupo de erros sobre autoespaços dos códigos.

No Capítulo 3, apresentamos as condições necessárias e suficientes para um código corrigir um determinado conjunto de erros quânticos. Em seguida, definimos as propriedades dos códigos e dos canais quânticos. Apresentamos os limitantes para a correção quântica traçando analogias com o caso clássico. Mostramos também como exemplos dois códigos quânticos bastante discutidos na literatura: o *código de repetição de nove qubits de Shor* e o código do tipo CSS com parâmetros $(7, 1, 3)$.

Por fim, no Capítulo 4 trataremos da formalização matemática da estrutura dos códigos estabilizadores. Para tal propósito, fizemos uso das propriedades algébricas envolvidas com

subgrupos do grupo de erros. Apresentamos então exemplos de códigos estabilizadores e discutimos a completa identificação entre o caso clássico e o quântico, no que diz respeito à construção de códigos. Exibindo os conceitos da codificação sobre $GF(4)$, temos os elementos fundamentais para gerar esta identificação.

1.1 Regras de um novo jogo

A história da tecnologia está envolvida com a descoberta de novos caminhos de se utilizar as leis da Natureza, explorando-se recursos físicos como os materiais, as forças e as fontes de energia [7].

A Mecânica Quântica é um novo caminho. Na escala da percepção humana, as leis da Física Clássica são boas aproximações para os fenômenos naturais. Entretanto, na escala atômica, elas falham e a teoria clássica perde seu espaço para a teoria quântica. A tecnologia envolve-se com esse fato quando se percebe que os computadores ficam a cada dia menores e mais rápidos. Por exemplo, portas lógicas, utilizadas no processo de codificação, são alocadas em chips de silício com menos de 1 micron. Este é, portanto, o cenário da Mecânica Quântica.

O conhecimento que se tem sobre a teoria e a fundamentação da Mecânica Quântica pode ser comparado ao *conhecimento de um jogador novato de xadrez que sabe todas as regras do jogo, mas que faz jogadas absurdas sobre o tabuleiro por não dominar os princípios heurísticos da teoria*, [20]. Desde os tempos de Einstein conhece-se os postulados da Mecânica Quântica, mas ainda falta o avanço na área das aplicações.

Algumas partes desta grande empreitada rumo à inovação são bastante interessantes. Por exemplo, no caso da compressão de dados, pergunta-se: qual é o número mínimo de bits necessários para armazenar a informação produzida por uma fonte? No caso clássico, a pergunta foi respondida em 1948 por Shannon que desenvolveu a Teoria da Informação, aplicada até nossos dias para a compreensão de fenômenos em jogos de azar e mercado de ações. No caso quântico, o que temos com respeito a esta questão, e muitas outras, é um paradigma proposto por Schumacher em 2001, que apontou o caminho a ser seguido para a procura das respostas, [20]:

1. Quais recursos físicos estão disponíveis na Mecânica Quântica?

2. Que tarefas de processamento da informação pode-se utilizar no problema estudado?
3. Quais são os critérios de sucesso?

A compreensão do paradigma de Schumacher no estudo de um sistema e as aplicações que dela virão parecem ser objetivo para pesquisas no século 21. Entre as aplicações, com certeza a que mais nos motiva é a criação de um computador quântico.

É importante ressaltar que para a teoria de codificação quântica a existência de códigos corretores de erros é fundamental. O fato é que os sistemas quânticos possuem grande fragilidade, o que torna qualquer transmissão duvidosa mediante à quantidade de erros que formam um contínuo de possibilidades, conforme veremos nas seções seguintes. Além da fragilidade, existe ainda a dificuldade de se obter medidas confiáveis a respeito dos dados enviados: os resultados dependem do processo escolhido para tal. Se tal parâmetro não estiver adequado à situação, o estado é danificado e com isso perde-se toda a informação contida nele.

Um exemplo usual para a excessiva instabilidade de um sistema quântico é o experimento conhecido como gato de Schrödinger. Suponha que Fibby, uma gata, seja trancada em uma gaiola com um vidro de cianeto que está ligado a um complicado aparelho que contém uma pequena amostra de uma substância radioativa, um contador de Geiger e um martelo. (Ressalta-se o fato de que o aparelho deve estar “protegido” contra a intervenção direta por parte do gato.)

A quantidade de material radioativo deve ser suficientemente pequena para que no intervalo de uma hora a probabilidade de que um dos átomos da amostra se desintegre seja de 50 %. Se um dos átomos se desintegrar, o contador de Geiger será disparado, ativando um relé e liberando o martelo. O martelo, por sua vez, quebrará o vidro, liberando o veneno, que matará Fibby.

Depois que Fibby é trancada na gaiola, ninguém sabe em que momento um dos átomos da amostra se desintegrará. A questão é a seguinte: Fibby estará viva ou morta depois de uma hora? A resposta correta é a seguinte: é impossível saber, a não ser abrindo a gaiola e examinando o interior. As equações quânticas que descrevem a desintegração radioativa não podem prever o que o observador encontrará quando abrir a gaiola, apenas a probabilidade de Fibby estar viva ou morta. De acordo com a descrição quântica, Fibby se encontra em uma espécie de limbo, meio morta e meio viva, até alguém abrir a gaiola. Em outras palavras,

toda vez que alguém realiza uma observação, novos “universos” se abrem, um para cada uma das diferentes possibilidades [25]. Este fenômeno é conhecido como *decoerência*.

Um outro exemplo, menos usual porém mais cotidiano, é o caso de se ter nas mãos um balde contendo H_2O . Não se sabe se a água está na forma líquida ou na forma de cubos de gelo e não é permitido olhar para dentro do recipiente. Assim, decide-se fazer um experimento para descobrir o estado da água. Coloca-se no balde uma pinça e retira-se um cubo de gelo. Quando a resposta já era quase definitiva, mergulha-se um pano no balde e ele sai molhado...

A explicação que parece mais lógica é a de que o balde contém cubos de gelo flutuando em água. Entretanto, optando por executar mais um experimento, vira-se o balde de cabeça para baixo, descobrindo que algumas vezes caem apenas cubos de gelo, enquanto em outras cai apenas água. Este último experimento garante uma resposta mais completa em relação aos anteriores. Porém, colapsa o sistema utilizado [25].

1.2 Do Bit para o Qubit

O que faz um computador quântico tão diferente de um computador clássico? Começamos analisando a unidade básica da informação clássica: o bit. Do ponto de vista físico, um bit é um sistema de dois estados: este sistema pode ser determinado por um dos dois estados distintos representados pelos valores lógicos 0 ou 1. Por exemplo, num computador digital, a amplitude de um pulso com 5 volts na entrada de um registro pode representar um bit de informação, denotado por 1, enquanto que a ausência deste pulso denota 0. Um bit clássico de informação pode ser codificado usando duas polarizações diferentes da luz ou dois estados diferentes de um átomo, [7].

A Mecânica Quântica nos diz que se um bit quântico pode existir em um dos dois estados distintos ele também pode existir na superposição destes estados. Este fato não tem análogo clássico: aqui o qubit pode representar os valores 0 e 1 simultaneamente. Este é um fenômeno quântico conhecido como *interferência de uma partícula*. Baseado neste fato, é que definiu-se o *bit quântico* ou simplesmente o *qubit*.

Um exemplo usual de qubit são os dois estados de um elétron orbitando em um átomo. Um deles é o estado fundamental, denotado por $|0\rangle$ e o outro o estado excitado, representado por $|1\rangle$. O símbolo $|\cdot\rangle$ é chamado *ket* na notação de Dirac, típica da Mecânica Quântica. Um *ket* representa um vetor do tipo coluna. O seu transposto, portanto um vetor linha, é

chamado *bra*, representado por $\langle . |$.

A partir desta idéia, podemos definir qubits como a combinação linear dos estados $|0\rangle$ e $|1\rangle$ na forma

$$|\psi\rangle = \alpha |0\rangle + \beta |1\rangle, \quad (1.1)$$

onde α e β são números complexos.

Na representação usual da Física, utiliza-se a esfera de Bloch, apresentada na Figura 1.1. Nesta esfera, o estado $|0\rangle$ está associado ao pólo norte e $|1\rangle$ ao pólo sul. Qualquer superposição na forma da eq. (1.1) de amplitudes α e β , constitui um qubit e está contido na casca da esfera.

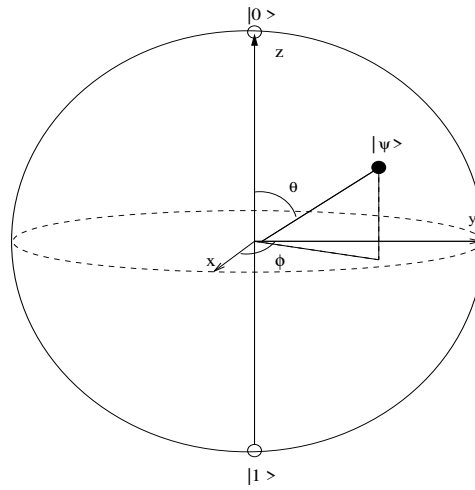


Figura 1.1: Esfera de Bloch e a representação de um estado $|\psi\rangle$.

Sob o ponto de vista de geradores do espaço de representação dos qubits, consideramos o fato que a casca da esfera é uma superfície e portanto uma variedade de ordem 2. Isto implica que a base para o espaço necessita de dois geradores, que correspondem a $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ e $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$.

Estes dois vetores são identificados, na descrição adotada, com os estados $|0\rangle$ e $|1\rangle$, respectivamente.

Esta interpretação na esfera de Bloch permite-nos supor a existência de uma quantidade contínua de possibilidades de estados para o qubit. Este fato dá origem a muitas das propriedades extraordinárias da informação quântica.

Uma comparação interessante é a seguinte: que quantidade de informação clássica pode ser armazenada em um qubit de forma que ela possa ser extraída e utilizada. *A priori*, pode-

mos deduzir que essa quantidade é infinita. Para especificarmos um estado quântico temos que definir sua latitude e longitude correspondente à representação na superfície da esfera de Bloch. Estes números codificam uma cadeia longa de bits. Por exemplo, 011101101... poderia ser codificado como sendo um estado com latitude 01 grau 11 minutos e 01101... segundos. Utilizamos então infinitos bits 0 e 1 nesta descrição e concluímos que o qubit carrega uma quantidade de informação clássica infinita [20] !

Embora pareça correto, o raciocínio não o é. Podemos codificar uma quantidade infinita de informação clássica em um único qubit, mas não há como extrair essa informação. A mais simples tentativa de “ler” o estado do qubit, que seria uma medida usual direta, resultaria em 0 ou 1. A mensuração compararia a probabilidade associada ao qubit de estar no estado 0 e à de estar no estado 1. A maior delas fornece o estado final do qubit. Entretanto, qualquer medida adotada “apaga” todas as informações contidas no qubit, com exceção daquela que de fato revela.

Assim, os princípios da Mecânica Quântica nos impedem de extrair mais de um único bit de informação de cada qubit, independente da medida ou da codificação. Este resultado surpreendente foi provado em 1973 por A. S. Holevo, depois de uma conjectura formulada por J. P. Gordon em 1964. *É como se o qubit contivesse informações ocultas que podemos manipular, mas não podemos acessar diretamente* [20].

Observamos que o resultado discutido por Holevo e Gordon respeitam o paradigma de Schumacher para a Teoria de Informação. O problema a ser resolvido é quantos qubits (recurso físico) seriam necessários para armazenar uma dada quantidade de informação clássica (a tarefa), de modo a permitir que a informação possa ser recuperada confiavelmente (o critério de sucesso).

Os pesquisadores criaram também um conceito matemático conhecido como *Holevo chi*, representado por χ , que têm sido usado para simplificar a análise de fenômenos mais complexos; a utilização de χ para o estudo de uma determinada ação num sistema resulta numa simplificação analítica semelhante às simplificações viabilizadas pela inserção do conceito de entropia de Shannon, [20].

1.2.1 O produto tensorial

Para apresentarmos conceitos sobre sistemas quânticos voltados à realização de tarefas de processamento computacional é necessário definirmos uma base que contenha múltiplos

qubits, para que a caracterização seja de fato eficaz. A operação usual para definirmos esta base é o produto tensorial entre qubits sobre o espaço vetorial complexo, onde as estruturas da Mecânica Quântica são descritas. Comentaremos brevemente a estrutura destes espaços e em seguida apresentaremos as propriedades do produto tensorial.

Um espaço vetorial complexo V é um espaço de Hilbert se existir um produto interno, escrito na forma $\langle \phi | \psi \rangle$, definido pelas seguintes regras [18]: considerando $a, b \in \mathbb{C}$ e $|\phi\rangle, |\psi\rangle, |u\rangle, |v\rangle \in V$, então

1. $\langle \psi | \phi \rangle = \langle \phi | \psi \rangle^*$, onde $*$ é o complexo conjugado,
2. $\langle \phi | (a|u\rangle + b|v\rangle) \rangle = a\langle \phi | u \rangle + b\langle \phi | v \rangle$,
3. $\langle \phi | \phi \rangle > 0$ se $|\phi\rangle \neq 0$.

A norma de um vetor $|\phi\rangle$ é dada por

$$\| |\phi\rangle \| = \sqrt{\langle \phi | \phi \rangle}.$$

A notação $\langle \phi |$ é usada para o vetor dual do vetor $|\phi\rangle$. O dual é um operador linear do espaço vetorial V para os números complexos, definido por

$$\langle \phi | (|v\rangle) = \langle \phi | v \rangle, \quad \forall |v\rangle \in V.$$

Se $|\phi\rangle = a|0\rangle + b|1\rangle$ e $|\psi\rangle = c|0\rangle + d|1\rangle$, então as matrizes que representam o produto interno e o transposto deste são, respectivamente [18]:

$$\begin{aligned} \langle \phi | \psi \rangle &= \begin{bmatrix} a^* & b^* \end{bmatrix} \begin{bmatrix} c \\ d \end{bmatrix} = a^*c + b^*d, \\ |\phi\rangle\langle\psi| &= \begin{bmatrix} a \\ b \end{bmatrix} \begin{bmatrix} c^* & d^* \end{bmatrix} = \begin{bmatrix} ac^* & ad^* \\ bc^* & bd^* \end{bmatrix}. \end{aligned}$$

Apresentada a caracterização dos espaços vetoriais complexos, passamos a descrever as propriedades do produto tensorial.

Suponha que V e W sejam espaços vetoriais de dimensões m e n , respectivamente. O produto tensorial $V \otimes W$ é um espaço vetorial mn -dimensional. Os elementos de $V \otimes W$ são combinações lineares dos produtos tensoriais $|v\rangle \otimes |w\rangle$, satisfazendo as seguintes propriedades [18]: para $z \in \mathbb{C}$, $|v\rangle, |v_1\rangle, |v_2\rangle \in V$ e $|w\rangle, |w_1\rangle, |w_2\rangle \in W$, então

1. $z(|v\rangle \otimes |w\rangle) = (z|v\rangle) \otimes |w\rangle = |v\rangle \otimes (z|w\rangle),$
2. $(|v_1\rangle + |v_2\rangle) \otimes |w\rangle = (|v_1\rangle \otimes |w\rangle) + (|v_2\rangle \otimes |w\rangle),$
3. $|v\rangle \otimes (|w_1\rangle + |w_2\rangle) = (|v\rangle \otimes |w_1\rangle) + (|v\rangle \otimes |w_2\rangle).$

Utilizamos as notações $|v\rangle |w\rangle$, $|v, w\rangle$ ou $|vw\rangle$ para o produto tensorial $|v\rangle \otimes |w\rangle$. Note que o produto tensorial é não comutativo e por isso a notação deve preservar a ordem.

Dados dois operadores lineares A e B definidos sobre os espaços vetoriais V e W , respectivamente o operador linear $A \otimes B$ em $V \otimes W$ é descrito por

$$(A \otimes B)(|v\rangle \otimes |w\rangle) = A|v\rangle \otimes B|w\rangle, \quad (1.2)$$

onde $|v\rangle \in V$ e $|w\rangle \in W$. A matriz que representa $A \otimes B$ é dada por

$$A \otimes B = \begin{bmatrix} A_{11}B & . & . & . & A_{1m}B \\ . & . & . & . & . \\ . & . & . & . & . \\ . & . & . & . & . \\ A_{m1}B & . & . & . & A_{mm}B \end{bmatrix},$$

onde A e B são matrizes de ordem m e n , respectivamente. Então, a matriz $A \otimes B$ tem ordem mn . Por exemplo, dadas

$$A = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, B = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix},$$

o produto tensorial $A \otimes B$ é

$$A \otimes B = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \end{bmatrix}.$$

As operações definidas podem ser utilizadas também para matrizes não quadradas, assim como para produto tensorial de dois vetores. Por exemplo, se tomamos os estados $|0\rangle$ e $|1\rangle$, o resultado do produto tensorial $|0\rangle \otimes |1\rangle$ é

$$|0\rangle \otimes |1\rangle = |01\rangle = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix},$$

um vetor que está no espaço vetorial 4-dimensional.

O estado geral $|\mu\rangle$ para 2 qubits é a superposição dos estados $|00\rangle, |01\rangle, |10\rangle$ e $|11\rangle$,

$$|\mu\rangle = \alpha |00\rangle + \beta |01\rangle + \gamma |10\rangle + \delta |11\rangle,$$

com a restrição

$$|\alpha|^2 + |\beta|^2 + |\gamma|^2 + |\delta|^2 = 1.$$

Para o caso geral, o estado $|\nu\rangle$ com N qubits é uma superposição de 2^N estados $|0\rangle, |1\rangle, \dots, |2^N - 1\rangle$,

$$|\nu\rangle = \sum_{i=0}^{2^N-1} \alpha_i |i\rangle,$$

com amplitudes α_i sujeitas a

$$\sum_{i=0}^{2^N-1} |\alpha_i|^2 = 1.$$

A base ortonormal $\{|0\rangle, \dots, |2^N - 1\rangle\}$ é chamada *base computacional*. O estado com N qubits é um elemento do espaço vetorial complexo 2^N dimensional. Quando o número de qubits aumenta linearmente, a dimensão do espaço vetorial associado cresce exponencialmente. Assim, a medida de um estado genérico, $|\psi\rangle$, resulta em $|i_0\rangle$ com probabilidade $|\alpha_{i_0}|^2$, onde $0 \leq i_0 < 2^N$ [18].

A descrição do caso geral nos permite comentar a respeito da larga vantagem, em questão de tempo de operação, da Computação Quântica sobre a Clássica para a realização de tarefas. Isso porque um único operador pode checar a informação contida nas N entradas do *ket*. Supondo que a tarefa seja uma busca em lista não ordenada, um algoritmo quântico pode encontrar a resposta no tempo de um único acesso de memória.

Portanto, a atuação de um computador quântico é similar à de computadores clássicos ligados em paralelo. Este conceito é conhecido como *paralelismo quântico*, introduzido por Deutsch em 1985.

Para exemplificar de forma mais quantitativa os resultados impressionantes que nos possibilita o paralelismo de Deutsch, considere o caso de mil processadores elementares e suponha que cada um seja capaz de representar um bit de informação (0 ou 1). Nesse caso, o número de combinações possíveis de 0 e 1 que se pode obter é igual a 2 elevado à milésima potência, o que corresponde a 1 seguido por 301 zeros. Para contar todas essas combinações, à razão de um trilhão de combinações por segundo, seria necessário um tempo igual a 10 bilhões de vezes a idade do Universo. É evidente que nenhum computador comum seria capaz

de explorar todas as combinações possíveis de um número tão grande de processadores; em outras palavras, haveria muitos problemas que um computador comum com mil processadores em paralelo simplesmente não teria tempo suficiente para resolver. Entretanto, se os microprocessadores fossem quânticos, cada um capaz de representar um qubit de informação, todas as combinações possíveis de zeros e uns seriam representadas simultaneamente. Isto equivale a dizer que um computador com mil processadores quânticos em paralelo seria 2 elevado à milésima potência mais rápido que um computador com um único processador quântico [25].

1.3 Emaranhados

Em 1935, Einstein, Podolsky e Rosen publicaram um trabalho que, partindo de preceitos teóricos da Mecânica Quântica, mostrou ser possível duas partículas se emaranharem, de modo que qualquer tentativa de determinar certa característica em uma seria refletida imediatamente sobre as configurações da outra, não importando a distância entre elas.

Esta *fantasmagórica ação à distância* é o fenômeno de *partículas emaranhadas* que comprovadamente existe na Natureza e é considerado hoje um recurso físico quantizável (como a energia) que permite a execução de tarefas como as de processamento da informação, [2].

Schrödinger ficou tão impressionado com o emaranhamento que, em um estudo pioneiro e inspirador publicado em 1935, referiu-se a este recurso físico como o *traço característico da Mecânica Quântica, aquele que nos obriga a abandonar inteiramente as linhas do pensamento clássico*. Talvez o que mais tenha chamado a atenção de Schrödinger é que membros de um conjunto de objetos emaranhados não têm seus estados quânticos individuais. Somente o grupo como um todo tem seu estado bem definido. Qualquer interação com o objeto emaranhado afeta simultaneamente tudo o que com ele está enlaçado, [20].

Devido a sua estranheza, por muito tempo o emaranhamento foi considerado uma curiosidade e ignorado pelos físicos. Isso mudou na década de 1960 quando Bell previu que os estados quânticos emaranhados permitiram realizar testes experimentais que distinguiam de uma vez por todas a Mecânica Quântica da Mecânica Clássica, mesmo que pudessemos alterar as previsões clássicas para embutir os resultados quânticos, [20]. Em 1964, Bell provou que se existisse outra teoria ainda mais completa que a Mecânica Quântica, ela necessariamente teria a mesma característica não-local, dada pelo emaranhamento. Este carácter não-local foi experimentalmente provado também em 1997 por N. Gisin, [1], que utilizou

pares de fótons viajando em cabos de fibra óptica distantes 30 quilômetros.

A pergunta que se passou a fazer é se o emaranhamento valeria também para um grupo de partículas. O matemático N. Linden, [1], em um estudo teórico, obteve a resposta: sim.

Na década de 1990, as pesquisas com o emaranhamento de objetos concentraram-se no contexto do processamento da informação. Em 1991, Ekert, [9], demonstrou como utilizar o recurso físico para distribuir chaves criptográficas imunes à espionagem. Em 1992, Bennett e Wiesner demonstraram que o emaranhamento pode ajudar no envio de informações clássicas de um local para outro, [20], num processo denominado *codificação superdensa*, no qual dois bits clássicos são transferidos através de um único qubit.

Em 1997, cientistas na Áustria usaram o emaranhamento para teletransportar um fóton instantaneamente de um ponto a outro em uma sala. Havia três partículas no experimento. Já em 2001, o número de partículas envolvidas num experimento parecido subiu para a casa dos trilhões. Nesta experiência, Julsgaard e Bolzik, [2], dispararam um laser sobre duas amostras de césio, fazendo com que todas as partículas ficassem emaranhadas por metade de um milésimo de segundo. O fato é que as amostras não interagiram mutuamente e isso significa que o emaranhamento pode ser obtido a distâncias consideráveis, como já se havia comentado teoricamente. Esse foi o primeiro passo para o princípio da Comunicação Quântica.

1.4 Algoritmos Quânticos

Em termos de algoritmos quânticos falaremos brevemente dos dois que são os mais comentados na literatura e, em seguida, apresentaremos as implementações tecnológicas onde esses algoritmos e tantos outros são simulados.

Lov Grover, [14], propôs um algoritmo de busca determinada em uma lista não ordenada de N itens com complexidade $\sqrt{N}$ passos.

Considere, por exemplo, uma pesquisa de um número telefônico específico em um diretório contendo um milhão de entradas, estocadas na memória de um computador em ordem alfabética dos nomes correspondentes aos números. Um algoritmo clássico faria em média 500.000 acessos de memória. Um computador quântico, com o algoritmo de Grover, examina todas as entradas simultaneamente, exibindo a resposta no tempo de um único acesso. O problema é quando o computador exibe de fato a resposta. Para tal é necessário realizar uma medida para a obtenção da informação. Todavia, se esta medida não for escolhida de forma adequada, o processo altera o resultado da busca, tornando possível a exibição de uma

resposta errada, [7].

O algoritmo de Grover foi utilizado para atacar o sistema criptográfico clássico *DES* (*Data Encryption Standard*) que apresenta $2^{56} = 7 \times 10^{16}$ caminhos possíveis. Se um computador clássico, operando com um algoritmo clássico puder checar 1 milhão de caminhos por segundo, levará 100 anos para descobrir o caminho correto, enquanto pelo algoritmo de Grover obtém-se tal resultado em menos de 4 minutos. Por este motivo, o algoritmo de Grover tem vasta aplicação em Criptografia, [7].

O outro algoritmo que vamos comentar é o de Shor, [24], proposto em 1994. O algoritmo fatora números inteiros com centenas de dígitos eficientemente em um computador quântico. A vantagem deste algoritmo sobre os clássicos é impressionante: para decompor um número inteiro com 500 dígitos em fatores primos, o melhor algoritmo clássico necessitaria de 5×10^{24} passos ou cerca de 150 mil anos na velocidade de um terahertz. Empregando o emaranhamento dos estados quânticos, o algoritmo proposto por Shor necessitaria de 5×10^{10} passos, ou menos de um segundo na velocidade de um terahertz. Analisando a questão temporal, temos que este algoritmo quântico opera em tempo polinomial enquanto os algoritmos clássicos operam em tempo exponencial ¹, [20].

O algoritmo de Shor, aplicado a um computador quântico, ameaça os esquemas utilizados em nossos dias para proteger informações eletrônicas, como o sistema *RSA* (*Rivest, Shamir, Adleman*) que é freqüentemente utilizado na segurança das informações bancárias. O sistema *RSA* baseia-se na impossibilidade de computadores clássicos fatorarem números inteiros com centenas de dígitos em pouco tempo, [20]. Porém, para que a ameaça se torne completa falta ainda um grande passo: a construção de computadores quânticos em escala capaz de atender pelo menos as grandes corporações. Por enquanto, não se tem notícias que possibilitem acreditarmos que estamos às vésperas de tal situação, mas o processo científico anda a passos rápidos, o que nos dá a certeza que isso ocorrerá num futuro não imediato, mas não muito distante.

1.5 A Construção de Computadores Quânticos

O potencial do fenômeno quântico na computação foi abordado primeiramente por Richard Feynman na *Primeira Conferência sobre Física da Computação*, realizada em 1981, [10]. Ele observou que parecia ser impossível simular, de maneira eficiente, a evolução de

¹Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer.

um sistema quântico num computador clássico. Entretanto, Feynman encarou tal obstáculo como possibilidade e, nos anos seguintes a 1981, avançou muito na pesquisa *do impossível*, obtendo resultados que foram somados ao trabalho de David Deutsch, que em 1985 publicou um artigo teórico no qual descrevia um computador quântico universal, [7].

Sabe-se que empecilhos não faltam para a implementação de uma máquina quântica. Além da dificuldade técnica de se trabalhar em escalas atômicas, um importante obstáculo é a *decoerência*. Denominamos *decoerência* às ações do meio que afetam as interações que geram as superposições quânticas às quais associamos as informações.

Entretanto, em 2000 a IBM projetou uma máquina que opera com cinco qubits. Tal máquina resolve um determinado problema em 1600 operações enquanto uma máquina clássica utilizaria 1,1 trilhão de operações para resolvê-lo. Em seguida, o Laboratório Nacional de Los Alamos, Novo México, projetou uma máquina com sete qubits, [2].

Como quanto maior o número de qubits envolvidos no processo, maior o rendimento das máquinas, concluímos que caminhamos para o sucesso da implementação dos computadores quânticos práticos.

Por outro lado, uma equipe de químicos e pesquisadores fundou uma empresa para desenvolver circuitos de computação muito potentes baseados em trilhões de blocos individuais, cada um do tamanho de uma molécula, nos quais dados podem ser armazenados e posteriormente extraídos. Os circuitos eletrônicos são minúsculos e criados por processos químicos com o objetivo de serem utilizados em computadores muito rápidos. O principal problema desta iniciativa, fora o alto custo financeiro, é o alto índice de defeito na produção dos componentes, muito maior que todos os que até hoje foram verificados.

Também em 2000, a IBM utilizou apenas sete átomos para fatorar o número 15 num computador quântico. Numa operação similar, os computadores clássicos mais modernos operam com chaves com bilhões de átomos em cada uma.

1.6 A Internet Quântica

Cirac e Zoller, em 1997, na Áustria formularam o primeiro projeto para a internet quântica. Em março de 2000, Lloyd, Shahrian e Hemmer melhoraram o projeto anterior trazendo-o para mais perto da realidade. A idéia foi atacar o problema da comunicação quântica com a criação de um par de fótons emaranhados enviados via fibra óptica, [2].

Os átomos, através de uma armadilha de *laser*, são super resfriados, consequência da

absorção de fótons. Verificando as absorções simultâneas seria possível saber quando os átomos absorveram um par emaranhado. Quando isso acontecer, os átomos em si se tornam emaranhados e quem recebeu a mensagem codificada agora compartilha um par de partículas emaranhadas, [2].

A esta idéia somou-se a contribuição de Chuang e Gottesman: a proposta de softwares que executam cálculos quânticos, projetam o conteúdo das mensagens e mantêm os qubits livres de erros. Também, Shor e Steane, em 1995, demonstraram que os erros podem ser corrigidos por meio da execução de uma série de cálculos sobre os dados da informação.

Todas essas contribuições importantes apontam para a inovação tecnológica que John Preskill classifica como a *fluorescente indústria do software quântico*, [2].

Capítulo 2

O Grupo de Erros

A obra de Hermann Weyl *Gruppentheorie und Quantenmechanik* de 1931, cuja tradução *The Theory of Groups and Quantum Mechanics* foi publicada em 1950, [28], foi a primeira a mostrar a importância da Teoria de Grupos para a formalização e compreensão de fenômenos muito complexos na Mecânica Quântica. Em seguida, John von Neumann em 1932 apresentou a fundamentação matemática da Física Quântica, [22]. A partir daí, a teoria de grupos e a teoria da Mecânica Quântica nunca mais se separaram.

Neste capítulo, faremos uma breve análise matemática sobre o grupo de erros. Na Seção 2.1, definiremos a *matriz densidade*, relacionada a probabilidade de se determinar o estado de um qubit. Na Seção 2.2, trataremos dos operadores do grupo de Pauli, que é o grupo que representa as transformações em Mecânica Quântica. Na Seção 2.3, veremos as principais propriedades dos Grupos de Clifford e como elas nos auxiliam na definição do grupo de erros e de suas operações, que foi o objetivo da Seção 2.4. Todos os conceitos vistos serão aplicados na Seção 2.5, onde apresentamos o exemplo de um código gerado adequadamente para uma situação particular da correção de erros quânticos.

2.1 Matriz Densidade

Todas as transformações que ocorrem em Mecânica Quântica são unitárias. Isto pode ser justificado pelo segundo Postulado da Mecânica Quântica, [21].

Postulado: A evolução de um sistema quântico fechado é descrita por uma transformação unitária. Isto é, o estado $|\psi\rangle$ de um sistema no tempo t_1 é levado para o estado $|\psi'\rangle$ do sistema no tempo t_2 por um operador U que depende somente dos tempos t_1 e t_2 , de forma que podemos escrever

$$|\psi'\rangle \equiv U |\psi\rangle,$$

para U satisfazendo $U^\dagger U = UU^\dagger = I$ ¹.

Esta restrição ou propriedade decorre da forma como a Teoria Quântica foi fundamentada sobre a Teoria de Probabilidades. Em geral, não se pode afirmar com certeza o estado quântico de uma partícula em um sistema sem a realização de uma medida. Porém, sabemos que a soma das probabilidades relacionadas aos possíveis estados é 1. Esta propriedade pode ser descrita pela *matriz densidade* P associada a um estado quântico $|\psi\rangle$ definida da seguinte forma:

$$\text{Considere } |\psi\rangle = \alpha |0\rangle + \beta |1\rangle \text{ para } |0\rangle = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \text{ e } |1\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

A matriz densidade associada é dada por:

$$P = |\psi\rangle\langle\psi| = \begin{bmatrix} |\alpha|^2 & \alpha\beta^* \\ \alpha^*\beta & |\beta|^2 \end{bmatrix},$$

onde $|\alpha|^2$ representa a probabilidade de $|\psi\rangle$ estar no estado $|0\rangle$ e $|\beta|^2$ a probabilidade de estar no estado $|1\rangle$. Sendo $|0\rangle$ e $|1\rangle$ os únicos resultados possíveis para a medida, temos, pela Teoria de Probabilidades, que $|\alpha|^2 + |\beta|^2 = 1$.

A ação que a matriz P representa é o que chamamos *projeção*, que resultará no *processo de medida quântica*. Este processo de medida altera a quantidade de informação associada ao qubit caso não seja realizada de forma adequada. Esta característica decorre do *Teorema da Não-Clonagem* dos estados quânticos, que apresentamos em seguida, [11].

Teorema 2.1.1 (Não-Clonagem) *Não há operação quântica que leve um estado $|\psi\rangle$ para o estado $|\psi\rangle \otimes |\psi\rangle$, qualquer que seja o estado $|\psi\rangle$.*

Demonstração: Consequência direta da linearidade da Mecânica Quântica. Suponha que exista uma operação que identifique $|\psi\rangle$ e $|\psi\rangle \otimes |\psi\rangle$ e que $|\psi\rangle$ e $|\phi\rangle$ sejam estados distintos. Pela definição da suposta operação:

$$|\psi\rangle \longrightarrow |\psi\rangle \otimes |\psi\rangle$$

$$|\phi\rangle \longrightarrow |\phi\rangle \otimes |\phi\rangle,$$

o que nos permite escrever

¹O símbolo $\dagger$ indica o complexo conjugado transposto.

$$|\psi\rangle + |\phi\rangle \longrightarrow (|\psi\rangle + |\phi\rangle) \otimes (|\psi\rangle + |\phi\rangle) = |\psi\psi\rangle + |\psi\phi\rangle + |\phi\psi\rangle + |\phi\phi\rangle. \quad (2.1)$$

Porém, pela linearidade temos que

$$|\psi\rangle + |\phi\rangle \longrightarrow (|\psi\rangle \otimes |\psi\rangle) + (|\phi\rangle \otimes |\phi\rangle) = |\psi\psi\rangle + |\phi\phi\rangle. \quad (2.2)$$

Como as equações (2.1) e (2.2) são distintas e representam a mesma quantidade temos uma contradição que surgiu exatamente ao supor que existe uma operação que identifique $|\psi\rangle$ e $|\psi\rangle \otimes |\psi\rangle$. ■

Por este resultado, concluímos que o estado

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

que é um qubit múltiplo, não pode ser escrito como o produto tensorial de 2 qubits. De fato, para $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, o produto $|\psi\rangle \otimes |\psi\rangle$ resultaria em

$$|\psi\rangle \otimes |\psi\rangle = \alpha^2|00\rangle + \alpha\beta^*|01\rangle + \alpha^*\beta|10\rangle + \beta^2|11\rangle.$$

Como as amplitudes $\alpha\beta^*$ e $\alpha^*\beta$ não podem ser nulas (pois isto implicaria α e β nulos, não existindo portanto o estado $|\psi\rangle$), concluímos que $|\Phi\rangle$ não pode ser escrito como combinação linear (com a operação de produto tensorial) de dois qubits. Estados com esta propriedade são chamados *emaranhados* ou *enlaçados*. Em particular, estados na forma

$$\frac{1}{\sqrt{2}}(|x\rangle \pm |\neg x\rangle),$$

onde $\neg$ denota o complemento binário, são conhecidos como *estados de Bell* ou *pares de Einstein - Podolsky - Rosen (EPR)*.

De forma mais geral, o Teorema 2.1.1 nos garante que não é possível copiar um estado quântico perfeitamente, ou seja, todas as configurações de um estado não podem ser totalmente transferidas para uma cópia. Isto implica que não é possível medi-lo sem deteriorar suas informações. Porém, determinados arranjos dos estados quânticos podem ser associados a processos de medida que extraíam apenas uma determinada informação, sem alterar a configuração do arranjo. A escolha deste arranjo e do processo de medida é um procedimento específico e meticuloso, fundamental para transmissão de informações.

Veremos mais adiante que o processo de medida só fará sentido para o contexto de correção de erros se o mesmo for realizado sobre subespaços ortogonais do espaço vetorial complexo. Entretanto, faremos uma breve discussão da necessidade de efetuar medidas em autoespaços ortogonais no próximo parágrafo.

Já vimos que as transformações em Mecânica Quântica são unitárias. Adiante veremos que elas são diagonalizáveis sobre o espaço vetorial complexo. Assim, podemos dizer que este espaço pode ser decomposto em autoespaços ortogonais. Desta forma, as transformações, ou erros no nosso caso, levam estados quânticos de um autoespaço inicial a outro ortogonal ao primeiro. Efetuando a medida sobre estes espaços ortogonais, identificamos o estado final e o distinguimos do erro que corrompeu o espaço inicial. Como os erros são dados por operadores unitários, então os mesmos podem ser invertidos, o que restituirá o estado inicial. Este processo será discutido com detalhes no Exemplo 2.4.1.

Por enquanto, estudaremos os elementos e a estrutura do conjunto de erros, no qual nos basearemos para gerar os códigos corretores quânticos.

2.2 O Grupo de Pauli

O *grupo de Pauli* é um grupo de operadores muito utilizado na Mecânica Quântica. Consiste de três operadores, além da identidade. Estes operadores serão alvo de estudo nesta subseção.

O primeiro operador age sobre a base habitual binária trocando $|0\rangle$ por $|1\rangle$ e vice-versa. Esta transformação é análoga ao erro introduzido no caso clássico: se a entrada do canal for o bit 1, a saída do canal, para a mesma informação, é o bit 0 e vice-versa. Como os estados quânticos (que são agora ferramentas para a codificação) têm como referencial os estados $|0\rangle$ e $|1\rangle$, faz sentido considerarmos esta transformação um erro de transmissão a ser estudado. Este erro é chamado *bit flip* e denotado por X . A notação equivalente na forma matricial é

$$X = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}.$$

O segundo operador age de forma trivial sobre a base binária se a entrada for $|0\rangle$, e inverte a fase se a entrada for $|1\rangle$. Ou seja, essa operação leva $|0\rangle$ em $|0\rangle$ e $|1\rangle$ em $-|1\rangle$ (e vice-versa). No caso quântico, este operador rotaciona de 180 graus o vetor *spin* do estado e, desta forma faz sentido considerá-lo como um erro. Este erro é chamado *phase flip* e

denotado por Z . A notação equivalente na forma matricial é

$$Z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}.$$

O terceiro operador é uma composição dos dois primeiros. Como ambos foram considerados erros e como um não pode, aplicado ao outro, restituir a identidade (o que anularia o erro), então a composição deles também deve ser considerada um erro. Este erro é chamado *bit and phase flip* e denotado por Y . A matriz correspondente é

$$Y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix},$$

onde $i = \sqrt{-1}$.

Vale comentar que o elemento i não tem sentido físico e por ser uma constante, na maioria das vezes, deixaremos de considerá-lo.

Uma vez que os erros foram identificados na forma de operadores, o *grupo de Pauli* fica estabelecido como

$$G_1 = \{\pm I, \pm X, \pm Y, \pm Z\}. \quad (2.3)$$

Note que, como $Y = -XZ$, os operadores X e Z geram G_1 , que é fechado e que, por sua vez, gera a estrutura algébrica usual dos Quatérnios. No sentido de ocorrência de erros, o sinal negativo em cada um dos operadores de G_1 não tem significado. Entretanto, sob o ponto de vista algébrico eles são necessários para garantir as propriedades de grupo.

Considerando a base $v_0 = |0\rangle = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ e $v_1 = |1\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$, para $l \in \mathbb{Z}_2$, as expressões para os efeitos dos erros X e Z são, respectivamente

$$X : v_l \longrightarrow v_{1+l} \quad (2.4)$$

e

$$Z : v_l \longrightarrow (-1)^{(1,l)} v_l, \quad (2.5)$$

para $l \in \mathbb{Z}_2$.

Assim, temos:

- $l = 0$, $X(v_0) = v_1$ e $Z(v_0) = (-1)^{1,0} v_0 = v_0$.
- $l = 1$, $X(v_1) = v_0$ e $Z(v_1) = (-1)^{1,1} v_1 = -v_1$.

Na notação de Dirac, as eqs. (2.4) e (2.5) são escritas da seguinte forma:

$$X : |l\rangle \longrightarrow |l+1\rangle$$

e

$$Z : |l\rangle \longrightarrow (-1)^{(1,l)} |l\rangle.$$

Passamos agora a considerar a ação dos operadores do grupo de Pauli sobre n qubits, para n um inteiro positivo.

Seja G_n o grupo de Pauli agindo sobre n qubits. Definimos G_n como o grupo dos produtos tensoriais dos elementos do grupo de Pauli G_1 , de forma que

$$G_n = \{\pm \omega_1 \otimes \omega_2 \otimes \dots \otimes \omega_n\}, \quad (2.6)$$

onde cada $\omega_j \in \{I, X, Z, Y\}$, para $j \in \{1, \dots, n\}$.

Uma forma alternativa para escrevermos a atuação de erros do tipo X e Z sobre n qubits é considerar a extensão de v_l , citado em (2.4) e (2.5), na forma:

$$v_l = v_{l_1} \otimes v_{l_2} \otimes \dots \otimes v_{l_n},$$

que na notação de *bra-ket* resulta em:

$$v_l = |l\rangle = |l_1\rangle \otimes \dots \otimes |l_n\rangle,$$

para $l = (l_1, \dots, l_n)$ um vetor de $\mathbb{Z}_2^n = V$.

Considerando u_j o j -ésimo vetor da base canônica geradora do espaço vetorial V , definimos erros sobre o j -ésimo qubit de $|v\rangle$ como [22]

$$X(u_j) : |v\rangle \longrightarrow |v + u_j\rangle$$

e

$$Z(u_j) : |v\rangle \longrightarrow (-1)^{v \circ u_j} |v\rangle,$$

onde $v \circ u_j$ denota do produto interno usual no espaço vetorial binário.

De uma maneira mais simplificada, definimos $X(a), Z(b)$ para quaisquer vetores $a, b \in V$ por

$$X(a) : |v\rangle \longrightarrow |v + a\rangle, \quad (2.7)$$

que indica ocorrência de erro do tipo *bit flip* onde $a_j \neq 0$ e

$$Z(b) : |v\rangle \longrightarrow (-1)^{v \circ b} |v\rangle, \quad (2.8)$$

que indica ocorrência de erro do tipo *phase flip* onde $b_j \neq 0$.

Assim, $X(a)$ produz um erro do tipo *bit flip* no j -ésimo qubit toda vez que a j -ésima coordenada de a é 1, e $Z(b)$ produz um erro de fase no j -ésimo qubit toda vez que a j -ésima coordenada de b é 1. Quando a j -ésima coordenada de a_j e a j -ésima componente de b_j são iguais a 1, temos um erro do tipo Y .

Como os operadores X e Z são geradores de G_n , qualquer elemento deste grupo pode ser escrito como

$$e = \pm X(a)Z(b), \quad (2.9)$$

para algum a e b vetores em $\mathbb{Z}_2^n$.

Na forma da eq. (2.6), definimos, para $j = 1, \dots, n$:

$$a_j = \begin{cases} 1 & \text{se } \omega_j = X \text{ ou } Y \\ 0 & \text{se } \omega_j = Z \text{ ou } I \end{cases}$$

e

$$b_j = \begin{cases} 1 & \text{se } \omega_j = Z \text{ ou } Y \\ 0 & \text{se } \omega_j = X \text{ ou } I \end{cases}$$

Exemplo 2.2.1

Para $n = 4$, considere o subconjunto de erros de G_n dado por $XZXY$ e um elemento $|v\rangle$ na forma $|0110\rangle$. $XZXY$ é a descrição resumida da ação de um operador X sobre o primeiro e o terceiro qubit, de um operador Z sobre o segundo e de um operador Y sobre o quarto qubit. Assim, $a = (1011)$ e $b = (0101)$ e segue que

$$X(a)Z(b)(|v\rangle) = X(1011)Z(0101)(|0110\rangle).$$

Pelas eqs. (2.7) e (2.8)

$$Z(0101)|0110\rangle = (-1)^{(0101) \circ (0110)} |0110\rangle = -|0110\rangle,$$

e

$$X(1011)(-|0110\rangle) = -X(1011)|0110\rangle = -|0110 + 1011\rangle = -|1101\rangle.$$

Portanto, $XZXY$ operando sobre $|0110\rangle$ resulta em $-|1101\rangle$.

2.3 Grupos de Clifford

Segundo o Postulado citado na Seção 2.1, todas as transformações em Mecânica Quântica são unitárias, de forma a considerarmos que todos os possíveis erros são operadores unitários. O que conhecemos da literatura, entretanto, é que todos os possíveis erros podem ser considerados composições dos operadores de Pauli. O que apresentamos nesta seção é a justificativa teórica, através de conceitos e propriedades associadas a *grupos de Clifford*, para a redução do grupo de erros quânticos geral (grupo das transformações unitárias) ao grupo de Pauli.

O principal conceito envolvido na construção dos Grupos de Clifford é o de normalizador.

Definição 2.3.1 [16] *O normalizador L de um subgrupo H de G é dado por*

$$L = L_G(H) = \{g \in G : g^{-1}Hg = H\}.$$

Os *subgrupos de Clifford* são obtidos como *normalizador* L de G_n em $O(2^n)$, grupo das matrizes ortogonais $2^n \times 2^n$ sobre $\mathbb{Z}_2^n$.² Tais subgrupos possuem as seguintes propriedades [5]:

- L/G_n é isomorfo ao grupo ortogonal $O(2n)$,
- L age como grupo ortogonal $O(2n)$ sobre G_n ,

onde $O(2n)$ indica o grupo de matrizes ortogonais $2n \times 2n$ sobre $\mathbb{Z}_2^n$.

Estas propriedades garantem que todos os possíveis erros em $U(2^n)$, podem ser reduzidos às ações do grupo G_n . Os erros que não são elementos de G_n e não podem ser escritos como composição destes elementos encontram equivalentes no grupo de Pauli via operação de conjugação. Graças a esta característica, que decorre do fato de utilizarmos o normalizador, a redução do grupo de erros ao grupo de Pauli mantém a ortogonalidade proveniente do grupo de matrizes $O(2n)$.³

Segundo um postulado da Mecânica Quântica, dois estados quânticos só podem ser distinguidos se estes forem ortogonais. Assim, tendo a redução preservado a ortogonalidade

²Na verdade, deveríamos tomar $U(2^n)$, grupo das matrizes unitárias $2^n \times 2^n$ sobre $\mathbb{C}_2^n$. Porém, como foi discutido em [5], considerar os *subgrupos de Clifford* sobre $O(2^n)$ não acarreta perdas em teoria e simplifica nossa discussão.

³Este é o caso análogo ao que fazemos quando olhamos um sinal expandido em série de Fourier. O que em geral nos interessa é um período dessa expansão. Como o sinal nos períodos restantes são ortogonais ao do período de interesse, podemos desconsiderar as demais partes sem perdas.

proveniente de $O(2n)$, mantemos a possibilidade de distinguir os estados quânticos, o que no nosso contexto implica em detecção (e possível correção) de erros.

Com isso, temos uma resposta algébrica para o porquê de se considerar o grupo de Pauli sendo equivalente ao grupo de erros quando *a priori* o caso geral dos possíveis erros, seria o $U(2^n)$, para $G_n \subset U(2^n)$.

2.4 Estrutura Matemática do Grupo de Erros

Uma vez definido o grupo de erros a ser utilizado, apresentaremos a sua estrutura matemática. Para facilitar a notação, denotaremos G_n por E .

- O grupo E é gerado pelos operadores X e Z aplicados em vetores de $\mathbb{Z}_2^n$, conforme as eqs. (2.7) e (2.8). Assim, o número de elementos, ou **ordem de E** , é dado por $2 \cdot 2^n \cdot 2^n$, ou seja, $o(E) = 2^{2n+1}$, o que nos leva a afirmar que E é um 2-grupo.

A notação *p-grupo*, para p um número primo, indica que a ordem do grupo é potência do primo p . Esta é a abreviação para dizer, no nosso caso, que E é um *p-grupo de Sylow*. Esta propriedade nos leva a um resultado importante, [16]:

Teorema 2.4.1 (Sylow) *Considere G um grupo. Se p é primo e $p^\alpha | o(G)$, então G possui um subgrupo de ordem p^α .*

Como veremos adiante, este subgrupo de ordem p^α , que denotaremos por S , poderá gerar um código corretor de erros. Os *p-grupos de Sylow* e os subgrupos de ordem p^α possuem estrutura interessante para a codificação: p caracteriza o alfabeto que utilizaremos, no nosso caso o binário, e α representa a dimensão das palavras código associadas a S .

- E não é abeliano pois, por exemplo, $XZ = -ZX$.

Das equações (2.7) e (2.8), podemos concluir que

$$X(a)X(a') = X(a + a') \quad (2.10)$$

$$Z(b)Z(b') = Z(b + b'), \quad (2.11)$$

para a, b, a' e b' pertencentes a $\mathbb{Z}_2^n$.

Isto conduz a uma possível identificação entre as transformações unitárias do grupo de Pauli com o grupo aditivo de $\mathbb{Z}_2^n$. Entretanto, o grupo aditivo de $\mathbb{Z}_2^n$ é comutativo e o grupo de operadores de Pauli E é não comutativo. Portanto, é necessário que se construa um subgrupo do grupo E que seja um p -grupo como E e que contorne a não-comutatividade do grupo de erros, isto é, que seja abeliano. Se conseguirmos tal estrutura matemática, nosso trabalho será favorecido pois poderemos fazer uso da teoria clássica (que opera exatamente sobre o grupo aditivo de $\mathbb{Z}_2^n$) e seus procedimentos.

Equivalentemente, poderemos construir códigos sobre $\mathbb{Z}_2^n$ munido da soma usual tal que corrija os erros gerados por E .

Num trabalho datado de 1954, Hall e Higman, [15], propuseram a solução para este problema. Apresentaremos a seguir os passos que conduzem à solução do problema em questão.

Definição 2.4.1 [16]:

1. O centralizador C do grupo G em L é dado por

$$C = \{h \in L : h^{-1}gh = g, \quad \forall g \in G\}.$$

2. O comutador de dois elementos de G , digamos g e h é dado por

$$[g, h] = ghg^{-1}h^{-1}.$$

3. O comutador do grupo G , que denotaremos de $\Xi(G)$, é o menor subgrupo de G que contem todos os comutadores de G .

De outra forma, podemos dizer que o *normalizador* é o conjunto que age em G por conjugação e fixa o *centralizador*.

No caso de interesse, $G = E$, e o *centralizador* C é $\{\pm I\}$. Note que o *comutador* $\Xi(E)$ é $\{\pm I\}$. Assim, podemos dizer que o *comutador* e o *centralizador* do grupo gerador de erros (que é um p -grupo) coincidem.

O fato do *comutador* $\Xi(E)$ ser $\{\pm I\}$ está intimamente ligado a propriedade de E ser um grupo p -solúvel.

De fato, todo grupo pode ser caracterizado em termos de seus *subgrupos derivados* $D_i, i > 1$, que são definidos recursivamente pela regra, [8]:

$$D_0(G) = G, \quad D_1(G) = [G, G], \quad D_i(G) = [D_{i-1}(G), D_{i-1}(G)],$$

onde $D_1(G)$ é o *comutador* do grupo G .

A solubilidade de G é garantida por

Definição 2.4.2 [16]: Considere G um grupo. G é solúvel se, e somente se,

$$D_n(G) = I,$$

para algum n .

Como vimos que $\Xi(E) = \{\pm I\}$, temos por definição que $D_1(G) = \{\pm I\}$. Isso implica que $D_n(G) = I, \quad \forall n > 1$. Conseqüentemente, E é um 2- grupo solúvel.

Os próximos resultados mostram porque nosso estudo tomou a direção apresentada, lembrando que procurávamos um subgrupo de E que pudesse reproduzi-lo via um *isomorfismo* e que fosse abeliano.

Teorema 2.4.2 [16] Um grupo solúvel sempre tem um subgrupo normal e abeliano.

Este teorema garante a existência do subgrupo de G com a propriedade de ser abeliano. O Teorema 2.4.3 exhibe exatamente a forma desse subgrupo que, como propriedade, reconstrói o grupo G por *isomorfismo*.

Teorema 2.4.3 [15] Seja G um p -grupo que admite um p -grupo S como um grupo de operadores tal que um certo elemento $k \in S$ opera não trivialmente sobre G e opera trivialmente sobre todo subgrupo próprio de G que admite S . Então, para algum q , G é um q -grupo abeliano ou é um grupo cujo centro e comutador coincidem e cujo grupo quociente $G/\Xi(G)$ é abeliano; S transforma $G/\Xi(G)$ irredutivelmente.

O fato de um grupo ser não-comutativo não impede que um código seja gerado sobre ele, resultado que se conhece do caso clássico. No contexto quântico, porém, a não-comutatividade entre um operador de medida adequado (considerado a partir do subgrupo S de E) e um estado quântico transmitido indicará a ocorrência de erros no processo. Por

isso, sob o ponto de vista de detecção de erros, é necessário que o grupo de erros E (ou um isomorfo a ele) seja abeliano, o que fará o algoritmo de reconhecimento de erros que será definido eficaz.

Quando falamos do grupo de erros E , a existência do subgrupo S é garantida pelo *Teorema de Sylow*; S é o conjunto de erros que, por exemplo, o canal introduz na transmissão. É a partir do subgrupo S que iremos gerar o código. No caso quântico, ao contrário do caso clássico, os códigos são particulares para cada tipo de situação-problema. É preciso estabelecer com precisão de qual conjunto de erros queremos proteger (e corrigir) as informações que estarão sendo transmitidas.

Com isso, concluímos a nossa procura. O subgrupo $E/\Xi(E)$ reconstitui E por um *isomorfismo*; é *abeliano* pelo Teorema 2.4.3 e tem ordem 2^n , ou seja, é também um *2-grupo*. Portanto, o *grupo de erros* que iremos utilizar para alicerçar o processo de codificação será o $E/\Xi(E)$, que denotaremos por $\bar{E}$.

$\bar{E}$ possibilita a descrição de E como espaço vetorial de dimensão $2n$ (como será apresentado na Subseção 2.4.1). Associada a um espaço vetorial existe uma norma e, a partir dela, os conceitos de ortogonalidade e de geometria, fundamentais para o contexto de detecção/correção de erros segundo a caracterização da Mecânica Quântica.

O fato de S agir irreduzivelmente sobre $\bar{E}$ garante a identificação completa do código com o *grupo de erros*.

Segundo Félix Klein, “*uma geometria é o estudo das propriedades de um conjunto A que permanecem invariantes quando se submetem os elementos de A às transformações geradas por algum grupo*”. Baseados nesta idéia, definiremos o conceito de *grupos extra-especiais*, que gerarão as transformações sobre o espaço vetorial complexo.

Definição 2.4.3 [15] *Um p -grupo especial, (p primo), é um p -grupo abeliano cujo centro e comutador coincidem e cujo grupo obtido a partir do quociente pelo comutador gera um subgrupo abeliano.*

Definição 2.4.4 [15] *Um p -grupo extra-especial é um grupo não-abeliano especial cujo centro é cíclico.*

Como E satisfaz as condições das Definições 2.4.3 e 2.4.4, afirmamos que E é um *2-grupo extra-especial*.

Por esta propriedade de E , podemos associar a E via $\bar{E}$ e operações deste novo grupo, que mostraremos na próxima seção, a geometria ortogonal que nos permitirá encontrar códigos que corrigirão erros quânticos nos autoespaços ortogonais determinados pela decomposição do espaço vetorial complexo.

2.4.1 Considerações sobre $\bar{E}$

Sendo *abeliano* e de ordem 2^{2n} , temos que $\bar{E}$ é isomorfo ao grupo aditivo do espaço vetorial de dimensão $2n$ sobre $\mathbb{Z}_2^n$. Como $\bar{E}$ é o subgrupo das classes laterais de E módulo $\{\pm I\}$, o fato do grupo de Clifford agir transitivamente sobre E comprova que as classes laterais em $\bar{E}$ esgotam todas as possíveis transformações unitárias que possam gerar erros quânticos.

Pela eq. (2.6), todo elemento de E pode ser escrito como $e = \pm \omega_1 \otimes \omega_2 \otimes \cdots \otimes \omega_n$, onde $\omega_j \in \{I, X, Y, Z\}$ para $j \in \{1, \dots, n\}$, que pode ser descrito a partir de $X(a)Z(b)$, conforme a Seção 2.2. Associamos a cada elemento e sua imagem $\bar{e}$ em $\bar{E}$ e a representamos por um vetor de comprimento $2n$ da forma $(a | b)$, onde para $i \in \{1, \dots, n\}$,

$$a_i = \begin{cases} 1 & \text{se } \omega_j = X \text{ ou } Y \\ 0 & \text{se } \omega_j = Z \text{ ou } I \end{cases}$$

e

$$b_i = \begin{cases} 1 & \text{se } \omega_j = Z \text{ ou } Y \\ 0 & \text{se } \omega_j = X \text{ ou } I \end{cases}$$

Desta forma, a em $(a | b)$ representa a ocorrência de erros do tipo X ou Y sobre os n qubits considerados. Por sua vez, b representa a ocorrência de erros do tipo Z ou Y . Tanto a quanto b tem comprimento n e, por isso, $(a | b)$ tem comprimento $2n$. A barra é utilizada para facilitar a interpretação do vetor: as entradas à esquerda associamos erros X nos qubits e, à direita, erros do tipo Z . O caso em que as posições em a e b coincidem implica que um erro do tipo Y ocorreu.

Por exemplo, o erro $X \otimes I \otimes X \otimes I \otimes Z \otimes Y \otimes Z \otimes Y$ é denotado como vetor $(a | b)$ por $(10100101 | 0001111)$ [6].

Precisamos definir a operação usual ou a forma fundamental sobre $\bar{E}$.

Começamos definindo a forma quadrática q como

$$q(\bar{e}) = \sum_{j=0}^n a_j b_j,$$

onde $e = \pm X(a)Z(b)$ e $\bar{e}$ sua imagem em $\bar{E}$.

Temos então que:

$$e^2 = (\pm I)^{q(\bar{e})}$$

e

$$q(\bar{e}) = \begin{cases} 0, & \text{se } X(a)Z(b) \text{ e } X(a')Z(b') \text{ comutam} \\ 1, & \text{caso contrário} \end{cases},$$

para a, b, a' e b' vetores de $\mathbb{Z}_2^n$.

Como no grupo unitário $U(2^n)$ não podemos definir $q(\bar{e}) = e^2$ pois $(ie)^2 \neq e^2$, então uma forma linear não singular alternativa associada à forma quadrática q é definida como sendo

$$((X(a)Z(b), X(a')Z(b'))) = ab' + a'b.$$

Denotaremos esta operação por $*$ de tal forma que para $(a | b)$ e $(a' | b') \in \bar{E}$:

$$(a | b) * (a' | b') = ab' + a'b. \quad (2.12)$$

A forma $*$ é *simplética* pois para todo $(a | b) \in \bar{E}$ temos $(a | b) * (a | b) = 0$.

Dizemos que $\bar{e}$ e $\bar{f} \in \bar{E}$ são *perpendiculares* se, e somente se, $\bar{e} * \bar{f} = 0$.

Um subespaço $\bar{S}$ de $\bar{E}$ é dito *totalmente isotrópico* se $q(\bar{s}) = 0$, para todo $\bar{s} \in \bar{S}$. Isto significa que para $\bar{s} = (a | b)$, $\bar{s}' = (a' | b')$ em $\bar{S}$ o produto interno $\bar{s} * \bar{s}' = 0$, e segue que $\bar{S}$ é *totalmente isotrópico* se, e somente se, S é *abeliano*.

Assim, a conexão entre a construção algébrica de $\bar{E}$ e a geometria ortogonal binária está completa, graças às operações que foram definidas sobre $\bar{E}$.

Exemplo 2.4.1 [22]:

Para $n=3$, tomamos vetores em $\mathbb{Z}_2^3$. Então, se queremos corrigir erros do tipo bit flip em um qubit, podemos utilizar o seguinte subconjunto S de E :

$$S = \left\{ \pm I, \pm Z(110), \pm Z(011), \pm Z(101) \right\}.$$

O vetor (110) é uma notação alternativa para indicar um gerador que possibilita verificar se um erro do tipo X ocorreu no último qubit. (011) e (101) fazem a mesma checagem sobre o primeiro e o segundo qubits, respectivamente. Se utilizássemos (001) , (100) e (010) , teríamos o mesmo efeito descrito, mantida esta ordem. Note que queremos corrigir erros

X e no entanto efetuamos medidas do tipo Z. Isto ocorre devido à operação fundamental $$, definida em (2.12), que “compara” a parte relativa a erros bit flips, representada pelo vetor a em $(a | b)$, que no caso queremos corrigir, sobre a parte relativa aos operadores de medidas, representada pelo vetor b em $(a | b)$, que devem ser escolhidos para que o processo de reconhecimento de erros seja eficaz.*

Discutimos a seguir como a escolha ocorreu.

No nosso caso, o erro que queremos corrigir é um bit flip em algum dos três qubits. As possibilidades na forma vetorial são: $(100 | 000)$, $(010 | 000)$ e $(001 | 000)$. Afirmamos que o conjunto escolhido para gerar as medidas que indicarão a ocorrência de erros é dado, na forma vetorial por

$$\bar{S} = \{(000 | 000), (000 | 110), (000 | 101), (000 | 011)\}. \quad (2.13)$$

Isto porque o operador $(000 | 110)$ identifica um erro na forma $(010 | 00)$ pois

$$(000 | 110) * (010 | 000)$$

resulta em $ab' + a'b = 1$. Isto implica que $q(\bar{e}) = 1$, ou seja, os operadores de medidas e os de erros anticomutam. O mesmo ocorre para os pares $(001 | 000)$, $(000 | 011)$ e $(100 | 000)$, $(000 | 110)$.

*Vemos que para todo $(a | b)$ e $(a' | b') \in \bar{S}$, definido segundo (2.12) temos $(a | b) * (a' | b') = 0$. Então, $\bar{S}$ é totalmente isotrópico, o que implica que S é abeliano. Portanto, no contexto de correção de erros induz à identificação destes, uma vez que o espaço é abeliano e os operadores de medida e erros anticomutam.*

Concluimos também, com cálculos diretos, que $(000 | 110)$ e $(000 | 011)$ geram $\bar{S}$ como espaço binário. Queremos agora determinar $\bar{S}^\perp$, o conjunto geral de operadores que detectam um erro de bit flip. A condição que temos para $(a | b) \in \bar{S}$ estar contido em $\bar{S}^\perp$ é

$$(a | b) * (000 | 110) = 0$$

$$(a | b) * (000 | 011) = 0,$$

de onde resulta que

$$a(110) = 0 \Rightarrow a = (000) \quad \text{ou} \quad a = (111)$$

e que

$$a(011) = 0 \Rightarrow a = (000) \quad \text{ou} \quad a = (111).$$

Assim, concluímos que

$$S^\perp = \{\pm Z(xyz), \pm X(111)Z(xyz), (xyz) \in \mathbb{Z}_2^3\},$$

e, conseqüentemente,

$$\bar{S}^\perp = \{(000 | xyz), (111 | xyz), (xyz) \in \mathbb{Z}_2^3\}.$$

Portanto, qualquer operador na forma de elementos de $\bar{S}^\perp$ é eficaz na detecção de um erro de bit flip.

O grupo L (normalizador) age transitivamente sobre subespaços *totalmente isotrópicos* de uma dada dimensão. Assim, todo espaço *totalmente isotrópico* k -dimensional está contido no mesmo número de espaços *totalmente isotrópicos* maximais n -dimensionais. Se $\bar{T}$ é um subespaço *totalmente isotrópico* maximal, então o grupo $\bar{T}$ tem 2^n autovalores distintos e os autoespaços correspondentes determinam uma *base ortonormal* de $\mathbb{R}^{2^n}$ ou $\mathbb{C}^{2^n}$ [22].

Isto acontece porque utilizamos um subgrupo *abeliano* de E para realizar a *decomposição ortogonal* do espaço complexo. Como cada um dos operadores de Pauli X, Y, Z é diagonalizável sobre $\mathbb{C}^2$, então cada elemento de E é diagonalizável sobre $\mathbb{C}^{2^n}$. Como transformações num subgrupo *abeliano* S de E são *simultaneamente diagonalizáveis*, então $\mathbb{C}^{2^n}$ tem uma base ortonormal constituída dos autovetores dos elementos de S , isto é, os autoespaços das transformações em S são mutuamente ortogonais. Considerando um desses autoespaços, temos o *código corretor de erros quânticos* que mapeia $n - k$ qubits sobre n qubits. Os 2^{n-k} vetores da base ortonormal nesse autoespaço constituem as palavras-código, [22]. Portanto estados afetados por operadores unitários do grupo de Pauli são levados de um autoespaço inicial a outro ortogonal a este, de forma que a “mudança” pode ser detectada pois os estados inicial e o modificado podem ser distinguidos por serem ortogonais [22].

Para um estudo sistemático da construção de códigos é necessário analisar algumas propriedades que se aliam à Álgebra Linear e a Teoria de Grupos. Citaremos algumas delas a seguir, [22].

Se Q é um autoespaço para $S \subseteq E$ *abeliano*, definimos a função λ sobre S por:

$$s(\mathbf{q}) = \lambda(s)\mathbf{q},$$

para todo $s \in S$, $q \in Q$ e $\lambda \in \mathbb{C}$; as identificações são: $\lambda \neq 0$; $\lambda(I) = +1$ e $\lambda(-I) = -1$.

λ denota os autovalores associados aos elementos do grupo de ação $\bar{S}$ sobre cada um dos autoespaço Q . Dizemos que λ é a representação dos *caracteres* de S . Assim, λ é um *homomorfismo* de S para o grupo multiplicativo dos complexos não nulos. *Um subgrupo abeliano do grupo de erros é isomorfo ao grupo de seus caracteres.*

A partir do conceito de *caracteres*, obtemos as propriedades que os identificam com os autoespaços do espaço vetorial complexo, [22]. Essas propriedades são discriminadas a seguir.

A. Autoespaços distintos de S correspondem a representações distintas de $\bar{S}$ (caracteres). Então o número de autoespaços não excede $|\bar{S}|$, o número de elementos em $\bar{S}$.

B. Elementos de E levam autoespaços em autoespaços; elementos de $S^\perp$, e só eles, fixam autoespaços. Isto implica que $e \in E$ fixa $Q \Leftrightarrow e$ comuta com todo $s \in S$, ou seja, $e \in S^\perp$.

C. E é transitivo sob a ação nos autoespaços de S , isto é, para quaisquer dois espaços Q e Q' de S , há um elemento $e \in E$ tal que $Q' = eQ$.

D. Os autoespaços de S têm uma correspondência bijetora com as classes laterais de E módulo $S^\perp$. Existem $|\bar{S}|$ autoespaços.

Estas propriedades justificam comentários anteriores a respeito da particularidade de cada problema em codificação quântica. Cada subespaço S abeliano escolhido como base do código, pode ser identificado pelo seu grupo de caracteres dado pelos seus autoespaços. Esta identificação é importante pela facilidade algébrica e geométrica, uma vez que é muito mais simples trabalharmos com números inteiros e suas operações do que com a ação de subgrupos sobre grupos.

2.5 A Idéia dos Códigos

Continuação do Exemplo 2.5.1:

Consideremos os vetores fixados por S . Como $Z(b)|v\rangle = (\pm 1)^{b \circ v}|v\rangle$, então $|v\rangle$ é fixado por $S \Leftrightarrow v \circ b = 0$ para $b \in \{(110), (011), (101)\}$. As únicas possibilidades são $v = (000)$ ou $v = (111)$. Então, um autoespaço complexo de S pode ser definido na forma:

$$Q = \left\{ |000\rangle, |111\rangle \right\}_v \subseteq \mathbb{C}^2 \otimes \mathbb{C}^2 \otimes \mathbb{C}^2.$$

Se utilizarmos o autoespaço Q para codificar, temos a seguinte identificação:

$$0 \longrightarrow |000\rangle \quad \text{e} \quad 1 \longrightarrow |111\rangle.$$

Com isso, a informação “0” ou “1” é codificada através do correspondente mapeamento na base de vetores associada ao código. O ruído do canal pode alterar a informação codificada, levando-a a outro autoespaço. Utilizamos então os elementos do subgrupo específico S para fazer a medida e determinar para qual autoespaço a informação foi levada e como será decodificada.

Continuação do Exemplo 2.5.1 :

Suponha que $|v\rangle \in Q$. Como ocorreu um erro $e \in E$, então na saída deste canal temos $e|v\rangle$. Em que autoespaço a informação será decodificada?

Para identificar o outro autoespaço para onde a informação foi levada é preciso saber quais são as possibilidades. Os elementos de

$$S^\perp = \{\pm Z(b); \pm X(111)Z(b), \quad b \in \mathbb{Z}_2^3\}$$

são os que mapeiam Q sobre ele mesmo, ou seja, fixa S . Assim, os elementos de uma classe lateral $eS^\perp$ são levados para eQ .

Precisamos descrever E em função das suas classes laterais. Para isso, sabemos que $o(E) = 2^7$ e $o(S^\perp) = 2^5$. Portanto, o número de classes laterais é $2^2 = 4$, representadas por (000) , (100) , (010) e (001) .

Segue que:

$$E = S^\perp \cup X(100)S^\perp \cup X(010)S^\perp \cup X(001)S^\perp,$$

onde $\{X(100), X(010), X(001)\}$ é o conjunto gerador para o grupo que identifica o erro bit flip sobre um único qubit.

Os subespaços correspondentes são imagens de Q pela ação dos erros bit flip. Cada um dos Q_ω é um autoespaço de S . Tais autoespaços são dados por

$$Q_{100} = X(100)Q = \{(100), (011)\}$$

$$Q_{010} = X(010)Q = \{(010), (101)\}$$

$$Q_{001} = X(001)Q = \{(001), (110)\}$$

Como todo elemento de E pertence a uma das classes laterais de $S^\perp$, então $e|v\rangle \in Q_\omega$ para algum $\omega \in \{(100), (010), (001), (000)\}$. Deste modo, determinamos em qual Q_ω o estado $e|v\rangle$ está contido. Como resultado, conhecemos o valor de ω . Assim, aplicamos $X(\omega)$ a $e|v\rangle$ cujo resultado pertence a Q , restituindo a informação que foi alterada.

Os três elementos de $\bar{S}$, que não a identidade têm ordem 2 e cada um deles é a soma dos outros dois. Então, λ deve associar a cada um dos dois geradores os valores $+1$ ou -1 . Existem quatro possibilidades para os caracteres e estes são mostrados na tabela a seguir.

caracter	(000 110)	(000 101)	(000 011)
λ_{000}	+1	+1	+1
λ_{100}	-1	-1	+1
λ_{010}	-1	+1	-1
λ_{001}	+1	-1	-1

A determinação dos caracteres do grupo de ação será muito importante na escolha do autoespaço que irá gerar os *códigos estabilizadores* que serão apresentados.

Capítulo 3

Códigos Corretores de Erros Quânticos

Neste capítulo, apresentamos aspectos gerais da correção de erros.

Na Seção 3.1, estudaremos um caso particular da correção e detecção de erros do tipo *bit flip* sobre os autoespaços determinados pelos erros considerados. Introduzimos a questão da eficiência do processo de recobrimento de erros via o conceito de fidelidade. Na Seção 3.2, faremos o mesmo, considerando erros do tipo *phase flip*.

Na Seção 3.3, mostraremos as características de construção e de correção do código de repetição de nove qubits proposto por Shor. Baseados neste exemplo, na Seção 3.4 apresentaremos as condições e as propriedades de correção de erros gerais, a definição de códigos não degenerados, o conceito de peso quântico e distância. Também, discutiremos as restrições para a construção do conjunto de erros corrigíveis de um código dado. Ainda na Seção 3.4, apresentamos o código de Shor sob o ponto de vista de classes laterais do grupo de erros e seus geradores, abordando conceitos que serão aprimorados no Capítulo 4.

Na Seção 3.5 mencionaremos o formalismo dos códigos clássicos que será importante para o estudo da classe de códigos quânticos conhecida como CSS. Apresentamos como exemplo desta classe o código $CSS(7,1,3)$.

Na Seção 3.6 apresentamos os limitantes para os códigos quânticos: o limitante de Hamming, o de Gilbert-Varshamov e o de Knill-Laflamme, que nos orientam na procura de novos códigos.

Começamos com alguns comentários sobre conceitos da correção clássica com o objetivo de relacionarmos estes com nosso caso de interesse: a correção de erros quânticos.

É conhecido do caso clássico que o ruído é o responsável pela alteração na informação

nos sistemas de processamento. Como não é possível construir sistemas sem interferência, o que se procura fazer é criar esquemas que protejam a informação das ações do ruído e interferências do meio.

Para compreender como os esquemas de correção agem no caso clássico, suponha que queremos enviar um bit por um canal ruidoso. O efeito do ruído é trocar o valor do bit que está sendo enviado. Considere a probabilidade deste erro acontecer sendo p , com $p > 0$. Consequentemente, a probabilidade do bit ser transmitido sem erro é $1 - p$. Um canal que possui esta característica é conhecido como *canal binário simétrico*, ou simplesmente *BSC*.

Como poderíamos proteger o bit do erro que pode ser introduzido? Uma forma usual é fazer três cópias do bit da seguinte forma:

$$0 \longrightarrow 000 \quad \text{e} \quad 1 \longrightarrow 111,$$

e, em seguida, enviá-lo pelo canal.

Os bits 000 e 111 são chamados *bit lógico 0*, denotado por 0_L e *bit lógico 1*, denotado por 1_L , respectivamente. O receptor dos três bits enviados tem que decidir qual o valor do bit original. Suponha que a saída do canal tenha sido 001. Ainda que a probabilidade p de ocorrer erro seja pequena, o receptor, conhecendo as possíveis entradas, aceita que o terceiro bit foi “flipado” e que o bit enviado foi 0 ao invés de 1. Este critério de decisão é o de *lógica majoritária*.

O processo de codificação descrito é definido como sendo um *código de repetição*. Técnica similar tem sido utilizada como uma parte da conversação cotidiana: se temos dificuldade em entender alguém que nos fala, pedimos que repita o que foi dito. Esta interação fonte - receptor pode ser repetida até que a mensagem nos seja coerente. Este processo, do ponto de vista de minimizar o tempo de transmissão, não é considerado bom pela grande quantidade de dígitos que precisamos introduzir para codificar um único bit de informação. A procura de bons códigos, ou seja, códigos que tornem menor possível a quantidade de bits que transmitam com segurança certa quantidade de bits de informação é um ramo de estudo da Teoria da Informação.

3.1 Analogias na Teoria Quântica de Correção

Embora a teoria clássica de correção de erros seja bem fundamentada e muito desenvolvida, ela não pode ser aplicada diretamente no caso quântico. Os conceitos da teoria quântica

que não nos permitem analogias já foram mencionados: a impossibilidade de se copiar um estado quântico (Teorema da Não - Clonagem), as possibilidades de erros que podem ser introduzidos no processo correspondem a um contínuo; também, o fato da medida colapsar a informação contida num qubit, no caso do processo utilizado para fazê-la não tenha sido adequado.

Mas, conforme discutiremos a seguir, estes aspectos não comprometem a pesquisa em métodos de correção quântica, que é apenas uma parcela do vasto campo da computação quântica com relação a tolerância à falha introduzida por Von Neumann.

Um código corretor de erros quânticos pode ser entendido como uma função de k -qubits (um espaço vetorial complexo de dimensão 2^k) sobre n -qubits (um espaço vetorial complexo de dimensão 2^n), onde $n > k$. Estes k qubits são chamados *qubits de informação* e são eles que queremos proteger dos erros. Os $n - k$ qubits adicionais formam a redundância necessária para a proteção da informação.

Suponha que enviamos qubits através de um canal que transmite sem introduzir erros com probabilidade $1 - p$ e “flipa” os qubits com probabilidade p . Isto é, com probabilidade p o estado $|\psi\rangle$ é trocado pelo estado $X|\psi\rangle$, onde X denota a transformação que irá resultar no erro *bit flip*. Um canal com tal característica é denominado *canal bit flip*.

Suponha que codificamos o estado $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ em $|\psi\rangle = \alpha|000\rangle + \beta|111\rangle$. Como este estado será transmitido pelo canal bit flip, admitimos que, ou não aconteça erro algum, ou que um erro bit flip corrompeu algum dos qubits. Há dois estágios no procedimento de correção que utilizamos para reverter o erro e reconstituir o estado correto.

A primeira parte consiste em realizar uma medida que identifica se algum erro ocorreu e, em caso positivo, qual qubit foi corrompido. O resultado desta medida é chamado *síndrome de erro*. Para o canal bit flip, são possíveis quatro síndromes de erros e cada uma destas corresponde a um operador de projeção na seguinte forma:

$$P_0 \equiv |000\rangle\langle 000| + |111\rangle\langle 111|,$$

$$P_1 \equiv |100\rangle\langle 100| + |011\rangle\langle 011|,$$

$$P_2 \equiv |010\rangle\langle 010| + |101\rangle\langle 101|,$$

$$P_3 \equiv |001\rangle\langle 001| + |110\rangle\langle 110|.$$

A estes operadores de projeção associa-se as seguintes afirmações com respeito à localização de possíveis erros:

- $P_0 \longrightarrow$ nenhum erro foi introduzido pelo canal;
- $P_1 \longrightarrow$ o canal introduziu um erro sobre o primeiro qubit;
- $P_2 \longrightarrow$ o canal introduziu um erro sobre o segundo qubit;
- $P_3 \longrightarrow$ o canal introduziu um erro sobre o terceiro qubit.

Suponha que o estado obtido após a transmissão seja $|\psi'\rangle = \alpha |100\rangle + \beta |011\rangle$. Portanto, na base $B = \{|000\rangle, |001\rangle, |010\rangle, |011\rangle, |100\rangle, |101\rangle, |110\rangle, |111\rangle\}$ temos que

$$|\psi'\rangle = \begin{pmatrix} 0 \\ 0 \\ 0 \\ \beta \\ \alpha \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

e

$$\langle\psi'| = (|\psi'\rangle)^\dagger = (0 \ 0 \ 0 \ \beta \ \alpha \ 0 \ 0 \ 0).$$

A medida da síndrome é dada pelo operador $\langle\psi'| P_i |\psi'\rangle$, para P_i , para $i = 0, 1, 2, 3$, são matrizes de ordem 8 (na mesma base), que são construídas de acordo com a definição dos projetores associados. Por exemplo, a matriz que representa o projetor P_1 na base B é dada por

$$P_1 = \begin{bmatrix} 00000000 \\ 00000000 \\ 00000000 \\ 00010000 \\ 00001000 \\ 00000000 \\ 00000000 \\ 00000000 \end{bmatrix}.$$

Assim,

$$\langle \psi' | P_1 | \psi' \rangle = \begin{pmatrix} 0 & 0 & 0 & \beta & \alpha & 0 & 0 & 0 \end{pmatrix} \begin{bmatrix} 00000000 \\ 00000000 \\ 00000000 \\ 00010000 \\ 00001000 \\ 00000000 \\ 00000000 \\ 00000000 \end{bmatrix} \begin{pmatrix} 0 \\ 0 \\ 0 \\ \beta \\ \alpha \\ 0 \\ 0 \\ 0 \end{pmatrix} = |\beta|^2 + |\alpha|^2 = 1,$$

que é a síndrome associada ao erro ocorrido no primeiro qubit. Efetuando o cálculo para P_0 , P_2 e P_{+4} , obtemos síndrome 0, que era esperado pelo fato ocorrer erro sobre um único qubit.

Ressaltamos que a medida da síndrome não altera o estado que está sendo observado pois ela exibe apenas a informação de que ocorreu erro sem interferir nas constantes α e β que determinam o estado inicial. Isto é, a síndrome não contém informações sobre o estado a ser protegido dos erros.

A segunda parte é utilizar o valor da síndrome obtido para escolher o procedimento que reconstituirá o estado inicial. No nosso caso, a síndrome de erro 1, medida sobre o projetor P_1 , indica que um erro *bit flip* corrompeu o primeiro qubit. Basta-nos “flipar” novamente tal qubit para restituirmos o estado inicial com exatidão.

Esta análise de erros não é completamente adequada pois os erros em Mecânica Quântica não são igualmente gerados.

Um exemplo disto é que um erro *bit flip* não afeta o estado $|\gamma\rangle = \frac{(|0\rangle + |1\rangle)}{\sqrt{2}}$, um estado normalizado, mas afeta o estado $|0\rangle$. De fato, $X|\gamma\rangle = \frac{1}{\sqrt{2}}(X|0\rangle + X|1\rangle) = |\gamma\rangle$, enquanto $X|0\rangle = |1\rangle$.

Isto dificulta o processo de recobrimento dos erros e, conseqüentemente sua correção. Com o objetivo de contornar este obstáculo, introduzimos o conceito de *operadores observáveis* e da medida associada a cada um deles.

Definição 3.1.1 Um observável M é um operador hermitiano ($M = M^\dagger$) sobre o espaço que contém o sistema sobre o qual operamos. Pode ser escrito na forma

$$M = \sum_m m P_m,$$

onde P_m é o projetor sobre o autoespaço de M com caracter m . Os possíveis resultados das medidas são os caracteres m dos observáveis.

Assim, ao invés de utilizarmos diretamente os P_i 's para a obtenção da síndrome, podemos calcular os caracteres associados aos observáveis. Considerando o caso em que a transmissão de $|\psi\rangle = \alpha|000\rangle + \beta|111\rangle$ resultou em $|\psi'\rangle = \alpha|100\rangle + \beta|011\rangle$, podemos descobrir qual erro foi introduzido se medirmos os observáveis Z_1Z_2 e Z_2Z_3 . Z_1Z_2 compara o primeiro e o segundo qubit: se eles forem iguais, o resultado da medida é 1, caso contrário, a medida resulta em -1. Z_2Z_3 faz o mesmo no segundo e no terceiro qubit. Existe quatro possíveis resultados para a síndrome, como no primeiro estudo em que utilizamos diretamente os projetores. As quatro possibilidades são dadas pelas combinações dos valores possíveis de Z_1Z_2 e Z_2Z_3 :

- Se o autovalor de Z_1Z_2 é +1 e o de Z_2Z_3 é +1, concluímos que os três qubits são iguais e portanto não houve erro;
- Se o autovalor de Z_1Z_2 é +1 e o de Z_2Z_3 é -1, concluímos que o erro ocorreu no terceiro qubit;
- Se o autovalor de Z_1Z_2 é -1 e o de Z_2Z_3 é +1, concluímos que o erro ocorreu no primeiro qubit;
- Se o autovalor de Z_1Z_2 é -1 e o de Z_2Z_3 é -1, concluímos que o erro ocorreu no segundo qubit.

É necessário analisar quão eficiente é este processo de correção sobre os erros introduzidos na transmissão. Para isto, introduzimos o conceito de *fidelidade*.

Suponha que um estado $|\psi\rangle$ tenha sido enviado através de um canal que introduz erros do tipo *bit flip*. Sem utilizar um código de correção, o estado do qubit após a transmissão é prevista por

$$\rho = (1 - p) |\psi\rangle\langle\psi| + pX |\psi\rangle\langle\psi| X. \quad (3.1)$$

A *fidelidade* entre o resultado da transmissão e o estado original é dada por

$$F(|\psi\rangle, \rho) = \sqrt{\langle\psi|\rho|\psi\rangle}, \quad (3.2)$$

que, para um operador na forma de (3.1), resulta em

$$F = \sqrt{(1-p) + p\langle\psi|X|\psi\rangle\langle\psi|X|\psi\rangle}. \quad (3.3)$$

O objetivo de um processo de correção de erros é aumentar a fidelidade com a qual a informação quântica é transmitida e trazê-la o mais próxima possível de 1, valor que indica que o resultado da transmissão é o estado enviado.

Com o intuito de tornar o conceito mais claro, observamos que o segundo termo da raiz quadrada em (3.3) é não negativo e igual a 0 quando $|\psi\rangle = |0\rangle$. Portanto, a fidelidade mínima é $F = \sqrt{1-p}$. Suponha que o código que definia os operadores P_0, P_1, P_2, P_3 seja utilizado para proteger o estado $|\psi\rangle = \alpha|000\rangle + \beta|111\rangle$. O operador que prevê a forma do estado quântico depois de transmitido e corrigido pelo código em questão é

$$\rho = [(1-p)^3 + 3p(1-p)^2] |\psi\rangle\langle\psi| + \dots,$$

em que os termos omitidos representam erros do tipo *bit flip* sobre 2 ou 3 qubits. Todos são positivos e então a expressão que obtivemos em (3.2) é um limitante inferior para a fidelidade real.

Mas, $F = \sqrt{\langle\psi|\rho|\psi\rangle} \geq \sqrt{(1-p)^3 + 3p(1-p)^2}$. Isto significa que a fidelidade é pelo menos $\sqrt{13p^2 + 2p^3}$, o que leva a concluirmos que a transmissão é confiável apenas se $p \leq \frac{1}{2}$.

Esta análise, apesar de válida, não representa inovação alguma com respeito à teoria clássica. Por isso, na próxima seção descreveremos um canal que introduz

erros do tipo *phase flip*, que é um caso particular de erro quântico.

3.2 Modelo para Erros de Fase

Neste novo modelo de canal quântico, consideramos que o qubit é transmitido sem a introdução de erros com probabilidade $1-p$, ou que o mesmo tenha inversão de fase com probabilidade p . Mais precisamente, dizemos que o canal irá aplicar o operador Z , correspondente ao erro *phase flip* no estado inicial $|\psi\rangle$ com probabilidade p , $p > 0$. Se

$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ então o estado após a transmissão por este canal será $|\psi'\rangle = \alpha|0\rangle - \beta|1\rangle$. Por esta característica, este modelo é denominado *canal phase flip*.

Com o propósito de simplificar a notação, definiremos uma nova base $\{|+\rangle, |-\rangle\}$, onde

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

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

e também os qubits lógicos:

$$|0_L\rangle = |+++\rangle,$$

$$|1_L\rangle = |--\rangle.$$

Esta mudança de base é obtida aplicando

$$H = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}.$$

sobre a base $\{|0\rangle, |1\rangle\}$. Este operador é conhecido como *operador de Hadamard*¹.

A codificação para o canal *phase flip* é construída em dois passos:

- O primeiro passo é codificar cada bit de informação em outros três, de forma a gerar os operadores lógicos $|0_L\rangle$ e $|1_L\rangle$, como no canal *bit flip*.
- O segundo passo é aplicar o operador de Hadamard em cada qubit.

A identificação de erros ocorre pela aplicação das mesmas medidas projetivas utilizadas para detectar erros do tipo *bit flip*, P_0 , P_1 , P_2 e P_3 , agora conjugadas pelo operador de Hadamard. Assim, não utilizaremos diretamente P_i , mas sim $P'_i = H^{\otimes 3} P_i H^{\otimes 3}$, para $i = 0, 1, 2, 3$, onde $H^{\otimes 3} = H \otimes H \otimes H$.

Equivalentemente, a medida da síndrome é dada pela ação do operador $H^{\otimes 3} Z_1 Z_2 H^{\otimes 3} = X_1 X_2$ e $H^{\otimes 3} Z_2 Z_3 H^{\otimes 3} = X_2 X_3$. A medida dos observáveis $X_1 X_2$ e $X_2 X_3$ corresponde à comparação dos sinais do primeiro e do segundo qubits e do segundo e do terceiro qubits, respectivamente. Neste caso, também temos quatro possibilidades para a síndrome. A localização dos erros se faz de modo análoga ao apresentado para o canal *bit flip*. Por exemplo, suponha que efetuamos a medida de $X_1 X_2$ e o resultado foi -1 e a medida de $X_2 X_3$ resultou em

¹Esta mudança de base equivale a rotacionar de 45 graus no sentido anti-horário os eixos coordenados da esfera do Bloch.

+1. Logo, temos um erro de sinal no primeiro qubit. Restituímos o estado inicial aplicando $HX_1H = Z_1$ no estado corrompido.

O processo de correção e detecção para erros do canal *phase flip* tem as mesmas características dos procedimentos utilizados para o canal *bit flip*. Por isso, a fidelidade mínima para os dois canais é a mesma.

Temos, portanto, um procedimento que consegue detectar e corrigir erros *bit flip* e outro, equivalente, que detecta e corrige erros *phase flip*. Uma pergunta natural seria sobre a possibilidade de se gerar um único procedimento que tivesse as características de ambos. Discutiremos a resposta desta questão na próxima seção.

3.3 Código de Repetição de Shor com Nove Qubits

O código que apresentamos é uma concatenação dos dois procedimentos de detecção e correção exibidos nas seções anteriores. Portanto, corrige erros *bit flip* e *phase flip*. Esta concatenação garante esta propriedade de correção devido ao fato dos processos citados corrigirem separadamente erros do tipo X e Z , que são os geradores do grupo de erros, como vimos no Capítulo 2.

No procedimento apresentado por Shor, a codificação do qubit é feita utilizando o critério

$$|0\rangle \longrightarrow |+++ \rangle$$

$$|1\rangle \longrightarrow |-- - \rangle,$$

como foi feito quando queríamos corrigir erros do tipo *phase flip*. Em seguida, codificamos cada um dos qubits na forma

$$|+\rangle \longrightarrow \frac{(|000\rangle + |111\rangle)}{\sqrt{2}}$$

$$|-\rangle \longrightarrow \frac{(|000\rangle - |111\rangle)}{\sqrt{2}},$$

como fizemos quando queríamos corrigir erros *bit flip*. O resultado deste processo é um código de nove qubits com os qubits lógicos dados por:

$$|0\rangle \equiv |0_L\rangle = \frac{(|000\rangle + |111\rangle)(|000\rangle + |111\rangle)(|000\rangle + |111\rangle)}{2\sqrt{2}}$$

$$|1\rangle \equiv |1_L\rangle = \frac{(|000\rangle - |111\rangle)(|000\rangle - |111\rangle)(|000\rangle - |111\rangle)}{2\sqrt{2}}.$$

Afirmávamos que o código de nove qubits proposto por Shor protege contra erros *bit flip* e *phase flip* em qualquer um dos qubits. Justificamos a seguir esta afirmação.

Suponha que queiramos analisar primeiro a existência de um erro *bit flip* no primeiro bloco de três qubits. Realizando as medidas de Z_1Z_2 e Z_2Z_3 identificamos com exatidão em qual qubit ocorreu o erro e então podemos revertê-lo. O mesmo processo se aplica aos outros blocos: a medida de Z_4Z_5 e Z_5Z_6 identifica um erro *bit flip* no segundo bloco, da mesma forma que a medida de Z_7Z_8 e Z_8Z_9 o faz no terceiro bloco.

Queremos também analisar se algum erro do tipo *phase flip* ocorreu em algum qubit. Faremos isto estudando o sinal de cada um dos blocos de qubits. Por exemplo, se fizermos a medida do observável $X_1X_2X_3X_4X_5X_6$, seguido da medida de $X_4X_5X_6X_7X_8X_9$, identificamos um erro *phase flip* em qualquer um dos blocos e com isso podemos revertê-lo utilizando o processo adequado.

Uma característica interessante relacionada ao conceito de emaranhamento mencionado no Capítulo 1 é que não é possível identificar erros *phase flip* em cada qubit dos blocos, enquanto o fazemos com exatidão no bloco como um todo. Isso ocorre porque o qubit lógico não é codificado localmente. Esta propriedade de codificação aumenta a proteção contra o ruído, uma vez que este possui característica local, ou seja, age independentemente sobre diferentes qubits do bloco.

E se ocorrem dois erros na transmissão?

Quanto a esta questão, temos duas possibilidades a serem estudadas.

Se ocorrer um erro *bit flip* no primeiro qubit e um erro *phase flip* no primeiro bloco, a medida de $X_1X_2X_3X_4X_5X_6$ e $X_4X_5X_6X_7X_8X_9$ indicará o erro *phase flip* sobre o primeiro bloco e a medida de Z_1Z_2 e Z_2Z_3 indica em qual dos qubits do primeiro bloco ocorreu o erro *bit flip* (no caso, o primeiro). Assim, o operador $X_1X_2X_3Z_1$ restitui o estado inicial. Os processos de correção de erros *bit flip* e erros *phase flip* são independentes. Isto não é uma propriedade particular do código de Shor, segundo mostra o seguinte resultado:

Teorema 3.3.1 [11] *Se um código quântico corrige erros do tipo A e do tipo B então ele corrige qualquer combinação linear de A e B.*

Se ocorrer um erro do tipo *bit flip* no primeiro qubit e um outro de mesmo tipo sobre o segundo qubit simultaneamente, a detecção por lógica majoritária indicará erro sobre o terceiro qubit. Na tentativa de corrigi-lo pelo processo usual, iremos “flipar” o terceiro qubit,

o que destruirá a informação contida no estado como um todo. O mesmo vale para erros do tipo *phase flip* se considerarmos os blocos de qubits.

Desta forma, podemos afirmar que o código de repetição de Shor corrige um erro geral (do tipo *bit flip*, *phase flip*, ou qualquer combinação linear destes) sobre um qubit qualquer.

3.4 Critérios Gerais para Correção de Erros

Como foi mencionado no Capítulo 2, cada problema no contexto de correção de erros quânticos resulta em uma formulação particular do código a ser utilizado. No Exemplo 2.5.1, discutimos a correção de um erro *bit flip* sobre elementos de $\mathbb{Z}_2^3$ e geramos então o código. Para tanto, definimos um subgrupo S e a partir do estudo de sua estrutura geramos um procedimento adequado que garante a correção dos erros indicados.

A questão que discutiremos é o caso recíproco do Exemplo 2.5.1 : dado um código corretor de erros quântico C , contido num subgrupo S , quantos erros ele pode corrigir?

O seguinte teorema nos auxilia no início deste estudo:

Teorema 3.4.1 [21] *Seja C um código corretor de erros quânticos e P um operador de projeção sobre C . Seja ε um conjunto de operadores quânticos E_i . A condição necessária e suficiente para a existência de uma operação de correção sobre ε em C é*

$$PE_i^\dagger E_j P = \alpha_{ij} P, \quad (3.4)$$

onde α_{ij} é uma matriz de números complexos. Se tal operação existe, $\{E_i\}$ constitui um conjunto de erros corrigíveis.

Em geral, definimos o conjunto $\varepsilon \subseteq E$ de erros corrigíveis satisfazendo

Lema 3.4.1 [22] *Se $e, f \in \varepsilon$ com $e \neq f$ então $f^{-1}e \notin S^\perp - S$.*

Esta condição decorre da decomposição ortogonal de $\mathbb{C}^{2^n}$.

Uma outra questão é: qual procedimento iremos utilizar para corrigir os erros e decodificar a informação?

Suponha que começamos o processo de transmissão enviando uma palavra-código $q \in C$ e que um erro seja introduzido pelo canal, levando q a $z = e(q)$ para algum $e \in \varepsilon$. O objetivo é restituir o estado q sem destruir a informação que ele contém. Conforme mencionado na

Seção 2.6, os elementos de E permutam os autoespaços de S (propriedade **B**) e então podemos assumir que z está em algum autoespaço C' de S . Fisicamente, podemos fazer uma sequência de medidas sobre observáveis para identificar o subespaço C' . Esta sequência de medidas é interpretada matematicamente da seguinte forma: pela propriedade **A**, os diferentes autoespaços de S correspondem a 2^{n-k} caracteres distintos de $\bar{S} = S/\{\pm I\}$, de maneira que identificar o autoespaço C' é equivalente a identificar o caracter λ correspondente a C' . Suponha que $\{s_1, \dots, s_{n-k}\}$ seja o conjunto de geradores de S módulo $\{\pm I\}$. A sequência de medidas ² que identifica C' corresponde ao cálculo de

$$s_j(z) = \lambda(s_j)z,$$

para $1 \leq j \leq n-k$.

Entretanto, admitimos que um erro $e \in \mathcal{E}$ havia ocorrido, condição que nos leva a concluir (pela propriedade **C**) que $C' = f(C)$ para algum $f \in \mathcal{E}$. Então, decodificamos z por

$$f^{-1}(z) = f^{-1}e(q).$$

Note que o passo de decodificação não necessitou de informações de q e e , sendo este o processo geral de correção e decodificação.

Para garantir que este processo é válido, é necessário que se $e(C) = f(C)$ para $e, f \in \mathcal{E}$, então $f^{-1}e(q) = q$.

De fato, se $e(C) = f(C)$ temos $f^{-1}e(C) = C$, o que resulta da propriedade **B**, que $f^{-1}e \in S^\perp$. Mas, pela condição de construção de $\mathcal{E}$, $f^{-1}e$ está em S . Como q é um autovetor de S , concluímos que o processo de decodificação utilizado apresenta como estado corrigido o estado inicial.

Algumas hipóteses adicionais sobre a construção de $\mathcal{E}$ geram classificações para os códigos corretores de erros quânticos. Por exemplo.

Definição 3.4.1 *Se é possível escolher um $\mathcal{E}$ que satisfaz: se $e, f \in \mathcal{E}$ com $e \neq f$ então $f^{-1}e \notin S^\perp$ para um código C , dizemos que C é não degenerado. Caso contrário, C é degenerado.*

Observamos que erros em S não têm efeito sobre estados q , para $q \in C$, (ou seja, elementos de S estabilizam todos elementos de C) e então podemos considerar o quociente $S^\perp - S$ para o estudo.

²Estes valores são as síndromes que considerávamos ao utilizar projetores. Por estas medidas, identificamos λ e, conseqüentemente, o autoespaço C' contendo z .

Há muitas maneiras de especificar um conjunto ε satisfazendo a condição da Definição 3.4.1:

Continuação do Exemplo 2.5.1

Como $S = \{\pm I, \pm Z(110), \pm Z(101), \pm Z(011)\}$ temos:

$$S^\perp = S \cup Z(100)S \cup X(111)S \cup X(111)Z(100)S$$

e

$$E = S^\perp \cup X(100)S^\perp \cup X(010)S^\perp \cup X(001)S^\perp,$$

algumas possibilidades para ε são:

$$\varepsilon_1 = \{X(100), X(010), X(001)\},$$

$$\varepsilon_2 = \{X(011), X(101), X(110)\},$$

$$\varepsilon_3 = \{X(100), X(101), X(001)Z(100)\}.$$

Neste contexto com muitas possibilidades para ε é necessário obtermos um critério para a escolha do conjunto: quais propriedades do subgrupo S indicam se o conjunto de erros ε é abrangente o bastante para ser vantajosa sua utilização?

Considere um erro $e = X(a)Z(b) \in E$, com $a, b \in \mathbb{Z}_2^n$, correspondente a um vetor $(a|b) \in \bar{E}$. O número de qubits afetados por e é o número de qubits que sofreram a ação de X , Z , ou ambos. Esta quantidade é definida como *peso quântico de e* e a notação é $q\text{-wt}(e)$.

Por exemplo, se $e = X \otimes Y \otimes Z \otimes I$, então $(a|b) = (1100|0110)$, segundo a definição do Capítulo 2. Concluimos que $q - \text{wt}(e) = 3$. Por este exemplo, observamos a diferença entre o peso quântico e o peso clássico de Hamming, que é definido pelo número de entradas não nulas de $(a|b)$, que neste caso seria 4.

Fixando um inteiro positivo d e considerando

$$\varepsilon = \{e \in E : q - \text{wt}(e) \leq \lfloor (d-1)/2 \rfloor\},$$

temos que se $e, f \in \varepsilon$, então o peso quântico de $f^{-1}e$ é o mesmo da soma dos vetores correspondentes em $\bar{E}$ e é no máximo $2\lfloor (d-1)/2 \rfloor \leq (d-1) < d$. Então, sob o ponto

de vista de pesos quânticos dos operadores de erros, dizemos que o conjunto $\mathcal{E}$ é formado por erros corrigíveis se a condição de que o peso mínimo de $S^\perp - S$ é d , para d fixo, for satisfeita. Isto resulta em uma propriedade de correção análoga a que conhecemos do caso clássico [22].

Teorema 3.4.2 [22] *Se o peso mínimo de $\bar{S}^\perp - \bar{S}$ é d , então C corrige erros em $\lfloor (d-1)/2 \rfloor$ qubits.*

Os conceitos definidos nesta seção nos possibilita reescrever o Teorema 3.3.1 na seguinte forma:

Teorema 3.4.3 *Se um código corretor de erros quânticos C corrige todos os operadores de Pauli de peso t então este código corrige todos os erros sobre t qubits.*

Definimos um outro parâmetro dos códigos corretores de erros: a *distância*. Ela é tão importante para o estudo de

um código quanto o comprimento das palavras-códigos n ou o número de qubits que contêm informação k .

Definição 3.4.2 *Um código corretor de erros quânticos C tem distância d se o peso quântico mínimo entre os operadores não nulos de $S^\perp - S$ for d .*

Observando o Teorema 3.4.2, entendemos qual a importância do parâmetro *distância*: ela caracteriza a capacidade de correção do código correspondente.

Exemplo 3.4.1 *Se $n = 9$, temos que os operadores de erros $e = \pm X(a)Z(b)$, para $a, b \in \mathbb{Z}_2^9$, escritos na forma $(a|b)$, são vetores com 18 componentes. Suponha que $\bar{S}$ seja gerado por:*

$$\begin{array}{ll} (000000000|011000000) & (000000000|000000110) \\ (000000000|000011000) & (111111000|000000000) \\ (000000000|000000011) & (000000000|000110000) \\ (000000000|110000000) & (000111111|000000000), \end{array}$$

a representação na forma vetorial $(a | b)$ dos operadores definidos para a identificação de erros no código de repetição com nove qubits de Shor.

Assim, temos que $\bar{S}$ é um espaço 8-dimensional. Resulta que a dimensão de $\bar{S}^\perp$ é 10. Podemos escolher para geradores de $\bar{S}^\perp$ os vetores:

$$(000000111|000000000)$$

$$(111000000|000000000).$$

O peso quântico mínimo de $\bar{S}^\perp - \bar{S}$ é 3 e então, pelo Teorema 3.4.2, este código corrige 1 erro. Como $\bar{S}$ contém um vetor de peso 2, concluímos também que ele é degenerado pela Definição 3.4.1.

O código cujos geradores foram exibidos no Exemplo 3.4.1 é o código de repetição de Shor com nove qubits. Iremos discutir este exemplo com mais detalhes quando estudarmos o formalismo estabilizador dos códigos quânticos no Capítulo 4.

Apresentamos a seguir um resumo das características de correção dos códigos quânticos.

Mensagem	n qubits em $\mathbb{C}^{2^n}$
Grupo de Erros	E com $e = \pm X(a)Z(b)$
Código	$C \subseteq \mathbb{C}^{2^n}$ com $\dim(C) = 2^k$
Erros Detectáveis	$S^\perp \subseteq E$
Erros sem efeito	$S \subseteq S^\perp$
Erros Corrigíveis	$\varepsilon \subseteq E$ tal que $e, f \in \varepsilon \implies f^{-1}e \notin S^\perp/S$
Pesos	mínimo de $S^\perp/S$ é d ; máximo de ε é $\lfloor (d-1)/2 \rfloor$

3.5 Exemplo de um Código Corretor de Erros Quânticos

Princípios da teoria de códigos lineares clássicos podem ser adaptados para a construção de códigos corretores de erros quânticos. Um exemplo disto são os códigos CSS, abreviação de Calderbank - Shor - Steane, que exploram o conceito de códigos duais que apresentaremos nesta primeira subseção. Em seguida, trataremos, de fato, do processo de construção de um código CSS.

3.5.1 Formalismo dos códigos clássicos lineares

Em um código binário, a partir de 2^n vetores no espaço vetorial complexo de comprimento n identificamos um subconjunto contendo 2^k vetores, que são exatamente as palavras-código de um código. Elas formam um subespaço linear fechado k -dimensional do espaço vetorial binário F_2^n , chamado espaço C do código se for gerado pela base de k vetores

$v_1, v_2, \dots, v_k$, isto é, uma palavra-código arbitrária pode ser escrita como uma combinação linear destes vetores da base:

$$v(\alpha_1, \dots, \alpha_k) = \sum_i \alpha_i v_i, \quad (3.5)$$

para cada $\alpha_i \in \{0, 1\}$, e a operação de adição é módulo 2. Podemos dizer que o vetor de comprimento n codifica a mensagem de k -bits $\alpha = (\alpha_1, \dots, \alpha_k)$.

A base de k vetores $v_1, \dots, v_k$ pode ser representada por uma matriz $k \times n$

$$G = \begin{bmatrix} v_1 \\ \cdot \\ \cdot \\ \cdot \\ v_k \end{bmatrix},$$

chamada *matriz geradora* do código. A eq.(3.6), na notação matricial, é escrita como

$$v(\alpha) = \alpha G; \quad (3.6)$$

Uma outra maneira de caracterizar o subespaço de codificação k -dimensional de F_2^n é especificar as $n - k$ posições restantes. Existe uma matriz H $(n - k) \times n$ tal que

$$Hv = 0,$$

para todos os vetores $v \in G$. A matriz H é chamada *matriz verificação de paridade* do código C . As linhas de H são constituídas pelos $n - k$ vetores linearmente independentes e o espaço do código é o espaço dos vetores que são *ortogonais* a todos os $(n - k)$ vetores (a ortogonalidade é definida com respeito ao produto interno usual). Com isso, temos

$$HG^T = 0, \quad (3.7)$$

onde G^T denota a transposta de G . De outra forma, podemos dizer que as linhas de G são ortogonais às linhas de H [19].

Um conceito muito útil na teoria de codificação clássica é o de *código dual*. Note que se G é a matriz geradora $k \times n$ e H é a matriz verificação de paridade $(n - k) \times n$ do código C , a relação entre elas, segundo a eq. (3.8), é $HG^T = 0$. Tomando a transposta, segue que $GH^T = 0$. Então, podemos considerar H^T como uma matriz geradora e G como a matriz verificação de paridade de um código $(n - k)$ -dimensional, denotado por $C^\perp$, chamado

código dual de C . Em outras palavras, $C^\perp$ é o complemento ortogonal de C em F_2^n . Um código contém seu dual se todas as suas palavras-códigos tem peso par e são mutuamente ortogonais. Se $n = 2k$ é possível que $C = C^\perp$ e, neste caso, C é dito um código *auto-dual*.

3.5.2 Códigos CSS

Trataremos da construção do código $CSS(7, 1, 3)$, que foi proposto por Steane e é o código mais simples desta classe de códigos corretores de erros quânticos [22].

Considere o código clássico de Hamming com parâmetros $[7, 4, 3]$ que denotamos por H_3 . Este código é um subespaço 4-dimensional de um espaço binário 7-dimensional com peso clássico mínimo 3³. Para o código H_3 , a matriz verificação de paridade é dada por

$$H = \begin{bmatrix} 1 & 1 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 1 \end{bmatrix}.$$

A partir de H , note que todas as palavras-código têm peso par e que as linhas de H são mutuamente ortogonais (no produto interno usual sobre $\mathbb{Z}_2^7$). Da subseção anterior, concluímos que H_3 contém seu dual. Assim, escolhe-se para geradores de $\bar{S}$ os seis vetores obtidos das linhas de H e suas combinações:

$$\begin{aligned} s_1 &= (0000000|1110100) & s_2 &= (0000000|0111010) \\ s_3 &= (0000000|1101001) & s_4 &= (1110100|0000000) \\ s_5 &= (0111010|0000000) & s_6 &= (1101001|0000000). \end{aligned}$$

Temos então que $\bar{S}$ tem dimensão 6 em $\bar{E}$ que, por sua vez, tem dimensão 14. Isto implica que $\bar{S}^\perp$ tem dimensão 8. Note que os geradores de $\bar{S}^\perp$ são $s_1, \dots, s_6$ e também $s_7 = (0000000|0101100)$ e $s_8 = (0101100|0000000)$. Nesta descrição de $\bar{S}^\perp$, via geradores de H_3 , vemos que o peso quântico mínimo coincide com o peso clássico de Hamming de H_3 que é 3. Este código é não degenerado e corrige um erro em um qubit, segundo o Teorema 3.4.2. Continuaremos a discutir este exemplo no item 2 da Seção 4.2.

3.6 Regras para a Construção de Códigos

A utilização de conceitos teóricos já conhecidos para a obtenção de novos códigos quânticos corretores de erros tem sido um assunto muito explorado pelos pesquisadores. Em [12],

³Esta notação $[7, 4, 3]$ reflete o fato que os códigos de Hamming H_r com parâmetros $[2^r - 1, 2^r - 1 - r, 3]$ existem para $r \geq 3$.

foram apresentados alguns códigos quânticos que foram gerados a partir de outros já conhecidos por métodos heurísticos. O propósito da obtenção de novos códigos é bastante simples: que o produto final de uma transmissão associada a um código corretor seja o máximo de informação segura com a menor taxa possível de erros, fazendo o melhor uso (no sentido de aproveitamento) do canal que utilizamos para transmitir.

Entretanto, existem algumas “regras” que devem ser obedecidas por todos os códigos quânticos: todos os códigos citados em [12], assim como os que mencionamos no presente trabalho, têm em comum a característica de satisfazer alguns limitantes. Estes limitantes serão estudados na próxima seção.

O modelo de canal que é considerado para a transmissão de dados codificados no caso quântico é o *canal de depolarização*. Falaremos brevemente de suas características.

O canal de depolarização pode ser descrito por um único parâmetro: a probabilidade p de ocorrer erro sobre algum qubit. A ação deste canal sobre cada qubit é definida por:

- Nada acontece no qubit com probabilidade $1 - p$;
- X é aplicado no qubit com probabilidade $p/3$;
- Y é aplicado no qubit com probabilidade $p/3$;
- Z é aplicado no qubit com probabilidade $p/3$.

Se $|\rho\rangle$ é o estado a ser transmitido, a ação do canal de depolarização é da forma:

$$A(\rho) = (1 - p)\rho + p/3(X\rho X + Y\rho Y + Z\rho Z). \quad (3.8)$$

A fidelidade mínima para os estados enviados através de um canal de depolarização é dada por [21]:

$$F = \sqrt{1 - \frac{2p}{3}}. \quad (3.9)$$

3.6.1 Limitantes para a correção quântica

Sendo já conhecidas as principais características do canal pelo qual transmitiremos as informações, associamos a questão relativa à eficiência de um código corretor de erros quânticos que atua depois da ação do canal. O estudo baseia-se no número de qubits codificados e na distância mínima de um código, como na teoria clássica de correção. Apresentaremos os três limitantes mais utilizados.

1)- Limitante quântico de Hamming

Este é o limitante mais simples que só pode ser aplicado a códigos não degenerados.

Suponha que um código não degenerado seja usado para codificar k qubits de informação em n qubits de forma que possa corrigir erros sobre qualquer subconjunto de t ou menos qubits. Considere que j erros ocorreram, onde $j \leq t$. Existem $\binom{n}{j}$ conjuntos de localizações onde os erros podem ocorrer. A cada um dos conjuntos de localizações associamos três erros possíveis (os operadores de Pauli X , Y e Z) que podem corromper cada qubit, gerando um total de 3^j possíveis erros. O número total de erros que podem ocorrer sobre t ou menos qubits é, então,

$$\sum_{j=0}^t 3^j \binom{n}{j}. \quad (3.10)$$

No caso de codificar k qubits de forma não degenerada, cada erro é um subespaço 2^k -dimensional ortogonal aos outros. Todos estes subespaços devem estar contidos no espaço total 2^n -dimensional, gerado pelos n -qubits das palavras-códigos. Então, concluímos que

$$\sum_{j=0}^t 3^j \binom{n}{j} 2^k \leq 2^n, \quad (3.11)$$

que é o *limitante quântico de Hamming*.

Se queremos codificar um qubit em n de forma que um erro possa ser corrigido, temos

$$2(1 + 3n) \leq 2^n. \quad (3.12)$$

Esta condição não é satisfeita para $n \leq 4$.

Portanto, o menor código quântico não degenerado que pode corrigir um erro em qualquer qubit tem $n = 5$.

Sabemos que não existem só códigos quânticos não degenerados. Na busca por limitantes que sejam mais gerais, os conceitos utilizados baseiam-se no limitante de Hamming que acabamos de discutir. Este é um motivo que faz deste limitante alvo de estudos e aplicações frequentes. Até nossos dias, conhecemos a validade da cota de Hamming para códigos não degenerados e para todos os degenerados que se conhece. Entretanto, não existe prova geral que garante a sua validade sobre todos os códigos degenerados.

2)- Limitante de Gilbert-Varshamov

Este limitante nos diz quantos dígitos de informação podemos alocar num código com palavras-código de tamanho n fixo, de forma que possam ser manipuladas com segurança. Considere um código quântico com parâmetros (n, k) , tal que o comprimento das palavras-código n é tão grande quanto se queira tomar. Se este código protege t qubits de erros que possam ocorrer, então existe um k tal que

$$\frac{k}{n} \geq 1 - 2H\left(\frac{2t}{n}\right),$$

sendo $H(x) = -x \log_2 x - (1-x) \log_2 (1-x)$, a entropia binária de Shannon.

A prova da validade da desigualdade que define este limitante é feita via propriedades dos códigos da classe CSS e é bastante interessante apesar de sua complexidade [21]. Este limitante se aplica, comprovadamente, a todos os códigos quânticos.

3)- Limitante de Knill-Laflamme

Este limitante surgiu na busca de cotas mais gerais que, assim como o limitante de Gilbert-Varshamov, pudessem ser aplicadas a todos os códigos. É análoga ao limitante de Singleton do caso clássico.

Este limitante garante que qualquer código quântico que codifica k qubits de informação em n qubits a serem transmitidos e corrige erros sobre quaisquer t qubits deve satisfazer

$$n \geq 4t + k. \quad (3.13)$$

Concordando com o que concluimos quando estudamos o limitante de Hamming, a informação de que o menor código que pode corrigir erros sobre um qubit tem $n = 5$ é uma consequência também deste limitante. Este código tem parâmetros $(5, 1, 3)$ e foi proposto por Bennett *et al.* em [3]. Construiremos esse código quando tivermos em mãos a Teoria de Estabilizadores que apresentaremos no próximo Capítulo.

Capítulo 4

Códigos Estabilizadores

Neste capítulo, trataremos de uma classe importante dos códigos corretores de erros quânticos : os *códigos estabilizadores*¹. Os códigos do tipo *CSS*, mencionados no Capítulo 3, são casos particulares desta classe.

Na Seção 4.1, descreveremos o formalismo dos códigos estabilizadores e a importância desta teoria quanto à caracterização dos erros que podem ser corrigidos por um código considerado.

Na Seção 4.2, apresentaremos o código de repetição de nove qubits proposto por Shor e o *CSS*(7, 1, 3) descritos via teoria de estabilizadores. Também na Seção 4.2, consideramos a construção de códigos da classe *CSS* mais gerais e o código de cinco qubits.

Na Seção 4.3, apresentaremos as condições para a obtenção de códigos clássicos a partir de códigos quânticos estabilizadores e as operações que proporcionam tal construção.

Na Seção 4.4, trataremos de códigos gerados sobre $GF(4)$ e das identificações entre estes e os códigos quânticos estabilizadores. Tais identificações possibilitam a aproximação da teoria clássica e da teoria quântica, conforme será visto.

Antes de estudarmos de fato os códigos estabilizadores, justificamos a atenção particular reservada a eles. Para isto, destacamos dois pontos principais:

- Os códigos estabilizadores têm construção análoga à dos códigos clássicos;
- Gottesman, [12], apresentou o seguinte resultado: o aproveitamento da capacidade do canal pelos códigos estabilizadores é relevante, como podemos observar na Figura 4.1

¹Os primeiros códigos estabilizadores foram apresentados independentemente por Gottesman, [13], e por Calderbank, Rains, Shor e Sloane, [4, 5].

Em [12], encontramos exemplos de códigos quânticos que foram gerados a partir de códigos clássicos e comparações entre os códigos quânticos estabilizadores e os outros que não compõem esta classe. Por exemplo, para códigos estabilizadores, encontra-se que a capacidade do canal de depolarização com parâmetro p (ver Seção 3.6) é limitada superiormente por $1 - H(p)$. Este é exatamente o limitante para o canal clássico *BSC*.

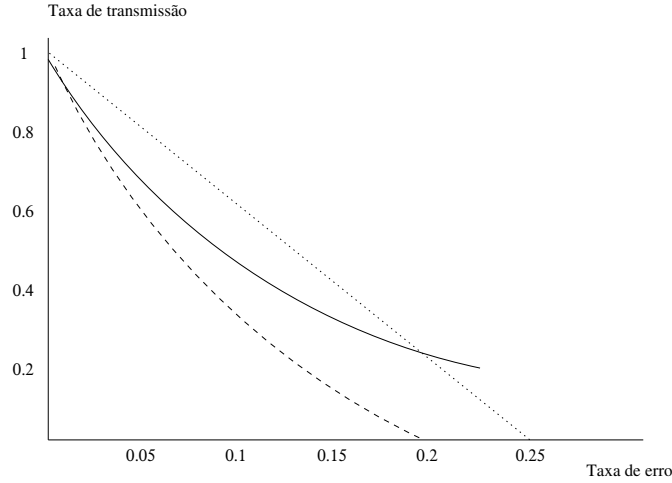


Figura 4.1: Eficiência de um código estabilizador.

Na Figura 4.1, a linha tracejada representa o limitante quântico de Hamming, a linha pontilhada, o limitante de Knill-Laflamme e a contínua representa

$$\frac{k}{n} \leq 1 - \frac{x}{2} \log_2 3 - \frac{1}{2} H(x),$$

que é a expressão obtida em [12] para a taxa de transmissão de códigos estabilizadores.

Os cálculos que levaram a esta expressão baseiam-se em conceitos clássicos.

A possibilidade de utilizar conceitos da teoria de codificação clássica para a formalização de características dos códigos quânticos estabilizadores (como a capacidade de canal, por exemplo) é um fato relevante e que justifica a atenção especial reservada a esta classe dos códigos quânticos.

4.1 O Formalismo Estabilizador

Considere o estado *EPR* de dois qubits discutido na Seção 2.1

$$|\Phi\rangle = \frac{|00\rangle + |11\rangle}{\sqrt{2}}. \quad (4.1)$$

Pode-se verificar que ele satisfaz:

$$X_1 X_2 |\psi\rangle = |\psi\rangle$$

e

$$Z_1 Z_2 |\psi\rangle = |\psi\rangle.$$

Assim, dizemos que o estado $|\psi\rangle$ é *estabilizado* pelo operadores $X_1 X_2$ e $Z_1 Z_2$.

A idéia básica do formalismo estabilizador é que muitos estados quânticos podem ser facilmente descritos pelos operadores que o estabilizam. Conforme veremos, exemplos desta descrição simplificada são os códigos de repetição de Shor e o CSS.

Toda a teoria deste formalismo é baseada na estrutura de grupos. O nosso interesse reside no grupo de erros sobre n qubits $G_n = E$, definido no Capítulo 2. Todos os elementos deste grupo têm as seguintes características:

- Cada $M \in E$ é uma matriz unitária, $M^{-1} = M^\dagger$, onde $\dagger$ representa o transposto conjugado complexo.
- Para cada elemento $M \in E$, $M^2 = \pm I$: $M^2 = I$ se o número de Y 's no produto tensorial é par e $M^2 = -I$ se o número de Y 's é ímpar.
- Se $M^2 = I$, então M é Hermitiana ($M = M^\dagger$); se $M^2 = -I$, então M é anti-Hermitiana ($M = -M^\dagger$).
- Quaisquer dois elementos M e $N \in E$ comutam ou anticomutam: $MN = \pm NM$.

A seguir apresentamos a definição de *estabilizador* de um grupo.

Definição 4.1.1 *Seja S um subgrupo abeliano de E e seja V_S o conjunto de estados com n qubits que são fixados por todos os elementos de S . Dizemos que V_S é o espaço vetorial estabilizado por S , de forma que S é dito estabilizador de V_S .*

Com isso, estamos aptos a definir o código estabilizador associado a S .

Definição 4.1.2 *Um código estabilizador H_S é definido como um espaço vetorial V_S que é estabilizado por um subgrupo abeliano S de E tal que $-I \notin S$.*

Como mencionado no Capítulo 2, o grupo de Pauli G_n caracteriza o código quântico da seguinte forma: denote um subgrupo abeliano de um grupo de Pauli sobre n -qubits. Todos os elementos de S agindo sobre $\mathbb{C}^{2^n}$ podem ser simultaneamente diagonalizados. Por isso, o código estabilizador H_S associado a S é um autoespaço correspondente ao autovalor +1, obtido pela ação de todos os elementos de S sobre os vetores de H_S . Isto é,

$$|\psi\rangle \in H_S \quad \text{se, e somente se,} \quad M|\psi\rangle = |\psi\rangle \quad \text{para todo} \quad M \in S. \quad (4.2)$$

O subgrupo S , assim como qualquer outro, pode ser caracterizado em função de seus geradores. Há elementos $\{M_i\} \in E$ que são *independentes* e tais que cada elemento de S pode ser escrito como um produto de M_i 's. Logo, se S tem $n - k$ geradores, o espaço de codificação H_S tem dimensão 2^k .

Para verificar isto, considere que cada $M \in S$ deve satisfazer $M^2 = I$, pois se $M^2 = -I$, então M não tem autovalor +1. Para cada um dos $M \neq \pm I$ restantes em E , os autovalores +1 e -1 têm igual probabilidade de ocorrer. Isto decorre do fato que para cada $M \neq \pm I$ existe pelo menos um $N \in E$ que anticomuta com M , de forma que

$$MN = -NM. \quad (4.3)$$

Assim, $M|\psi\rangle = |\psi\rangle$ se, e somente se, $M(N|\psi\rangle) = -N|\psi\rangle$. A ação de N estabiliza autoestados associados aos autovalores +1 e autoestados associados aos autovalores -1 em igual proporção. Desta forma, são $\frac{1}{2}(2^n) = 2^{n-1}$ estados mutuamente ortogonais que satisfazem

$$M_i|\psi\rangle = |\psi\rangle, \quad (4.4)$$

onde M_i é um dos geradores de S .

Considere agora M_2 um outro elemento de E que comute com M_1 , tal que $M_2 \neq \pm I, \pm M_1$. Podemos encontrar um $N \in E$ que comuta com M_1 mas anticomuta com M_2 , de forma que N preserva o autoestado +1 de M_1 mas altera os autoestados associados aos autovalores +1 e -1 de M_2 . Resulta que o estado $|\psi\rangle \in H_S$ satisfaz

$$M_1|\psi\rangle = M_2|\psi\rangle = |\psi\rangle, \quad (4.5)$$

e que H_S tem dimensão 2^{n-2} [23].

Continuando este processo, temos que se M_j é independente de $\{M_1, M_2, \dots, M_{j-1}\}$, então há um N que comuta com $M_1, \dots, M_{j-1}$ e anticomuta com M_j . Assim, restrito ao espaço com $M_1 = M_2 = \dots = M_{j-1}$, M_j tem tanto autovetores associados ao autovalor $+1$ quanto autovetores associados aos autovalores -1 . Pela Definição 4.1.2, apenas devem ser considerados autovetores relacionados aos autovalores positivos. Portanto, a dimensão do espaço passa a ser $2^n(1/2)^{n-k} = 2^k$. ■

A linguagem de estabilizadores é importante porque facilita a caracterização dos erros que um código é capaz de detectar e corrigir. Para tal, podemos associar os $n - k$ geradores do estabilizador $M_1, \dots, M_{n-k}$ aos operadores de verificação de paridade de um código. Se a informação codificada sofreu um erro que não pode ser detectado, então encontraremos que $M_i = 1$ para cada um dos geradores. Mas, se $M_i = -1$, para algum i , então o estado da informação está contido em um autoespaço ortogonal ao subespaço do código e um erro será detectado.

Sabemos que o operador erro pode ser expandido em termos de elementos E_a do grupo de Pauli. Cada E_a comuta ou anticomuta com um gerador do estabilizador M . Se E_a e M comutam então

$$ME_a|\psi\rangle = E_aM|\psi\rangle = E_a|\psi\rangle, \quad (4.6)$$

para $|\psi\rangle \in H_S$ e o erro preserva o valor $M = 1$. Mas, se E_a e M anticomutam, então

$$ME_a|\psi\rangle = -E_aM|\psi\rangle = -E_a|\psi\rangle, \quad (4.7)$$

de forma que o erro “flipa” o valor de M , podendo ser detectado através da medida de M , [23].

Em função dos geradores do estabilizador $\{M_i\}$ e dos erros E_a , considere a relação

$$M_i E_a = (-1)^{s_{ia}} E_a M_i. \quad (4.8)$$

Os s_{ia} 's, $i = 1, \dots, n - k$ constituem a *síndrome* para o erro E_a , pois $(-1)^{s_{ia}}$ é o resultado da medida M_i , se o erro E_a ocorre.

No caso de um *código não degenerado*, cada erro tem síndrome distinta, de forma que a medida dos $n - k$ geradores indica com precisão o erro que ocorreu. Entretanto, se o

código é *degenerado*, existem erros distintos com a mesma síndrome e as medidas não geram conclusões confiáveis a respeito do erro introduzido no processo.

Para cada $E_a, E_b \in \mathcal{E}$, considere as seguintes afirmações [23]:

1. $E_a^\dagger E_b \in S$,
2. existe um $M \in S$ que anticomuta com $E_a^\dagger E_b$

Se uma delas for satisfeita, garante-se que um erro pode ser detectado e corrigido. Esta é uma condição suficiente.

Assim, um *código estabilizador* que corrige $\{\mathcal{E}\}$ é um espaço H_S fixado por um subgrupo abeliano S do grupo de Pauli, onde 1) ou 2) é satisfeita para cada $E_a^\dagger E_b$, com $E_a, E_b \in \mathcal{E}$. Um código é não degenerado se a condição 1) não é satisfeita para qualquer $E_a^\dagger E_b$.

Podemos escolher o subespaço de codificação como sendo qualquer um dos 2^{n-k} autoespaços de $n - k$ elementos que comuta com E . Todavia, todos os códigos que são gerados por estas escolhas são *equivalentes*. Por códigos equivalentes entendemos códigos que diferem somente na localização dos qubits e na escolha da base utilizada para a descrição dos mesmos. Assim, o estabilizador de um código é transformado no estabilizador de um outro código através da permutação dos qubits. Esta permutação é definida como um produto tensorial de transformações sobre os qubits que compõem o estabilizador.

De fato, se dividirmos os geradores do estabilizador em dois conjuntos $\{M_1, \dots, M_j\}$ e $\{M_{j+1}, \dots, M_{n-k}\}$, então existe um $N \in E$ que comuta com os elementos do primeiro conjunto e anticomuta com os elementos do segundo. Aplicando N a $|\psi\rangle \in H_S$, preservam-se os autovalores do primeiro conjunto, enquanto os sinais dos autovalores do segundo são trocados.

O processo de recobrimento de erros pode falhar se há um $E_a^\dagger E_b$ que comuta com o estabilizador mas não pertence a este conjunto. Um operador com esta propriedade preserva o subespaço de codificação H_S mas age não trivialmente sobre ele, de forma que pode corromper a informação codificada. Se isto ocorre, o erro E_a pode ser interpretado como um erro E_b . Na tentativa de correção, a aplicação de $E_b^\dagger E_a$ causará a perda do dado enviado, pois modificará as características do estado transmitido.

Um código estabilizador com distância d tem a propriedade que cada $E_a \in E$, com peso menor que d , está no estabilizador ou anticomuta com algum elemento deste grupo. Sob este

ponto de vista, o código é *não degenerado* se o estabilizador não contém elementos de peso menor que d . Um código com distância $d = 2t + 1$ pode corrigir até t erros, assim como um código com distância $s + 1$ pode detectar até s erros, [23].

Em síntese, um código corretor de erros quânticos pode ser caracterizado pelos seus geradores e os $n - k$ geradores podem ser representados por uma matriz $(n - k) \times 2n$

$$H = (H_X | H_Z), \quad (4.9)$$

onde cada linha é um operador de Pauli expresso na notação $(a|b)$. A síndrome de um erro $E_a = (a_a|b_a)$ é determinada por suas propriedades de comutação com os geradores $M_i = (a'_i|b'_i)$. Isto é,

$$s_{ia} = (a_a|b_a) \cdot (a'_i|b'_i) = a_a \cdot b'_i + a'_i \cdot b_a, \quad (4.10)$$

que é exatamente a definição da operação fundamental sobre $\bar{E}$, eq. (2.14).

4.2 Alguns Exemplos de Códigos Estabilizadores

1) Código de nove qubits de Shor [23], [22], [11].

Como mencionado na Seção 3.3, o código de repetição de nove qubits de Shor pode detectar e corrigir erros do tipo *bit flip* e *phase flip*.

Os operadores lógicos para este código são:

$$|0\rangle \equiv |0_L\rangle = \frac{(|000\rangle + |111\rangle)(|000\rangle + |111\rangle)(|000\rangle + |111\rangle)}{2\sqrt{2}}$$

$$|1\rangle \equiv |1_L\rangle = \frac{(|000\rangle - |111\rangle)(|000\rangle - |111\rangle)(|000\rangle - |111\rangle)}{2\sqrt{2}}.$$

Para detectar um erro do tipo *bit flip* sobre o primeiro bloco de três qubits, compara-se o primeiro e o segundo qubits e em seguida, o segundo e o terceiro. Isto é equivalente a medir os autovalores de Z_1Z_2 e de Z_2Z_3 .

Se a mesma análise for feita sobre o segundo bloco, a medida que identifica um erro *bit flip* sobre um dos qubits é o autovalor dos operadores Z_4Z_5 e Z_5Z_6 . Para o terceiro bloco, a medida corresponde aos autovalores dos operadores Z_7Z_8 e Z_8Z_9 .

Para detectar um erro *phase flip* sobre um dos três blocos a medida a ser efetuada é o autovalor de $X_1X_2X_3X_4X_5X_6$ e $X_4X_5X_6X_7X_8X_9$.

Assim, os operadores que detectam erros no caso do código de Shor são:

$$M_1 = Z_1 Z_2$$

$$M_2 = Z_2 Z_3$$

$$M_3 = Z_4 Z_5$$

$$M_4 = Z_5 Z_6$$

$$M_5 = Z_7 Z_8$$

$$M_6 = Z_8 Z_9$$

$$M_7 = X_1 X_2 X_3 X_4 X_5 X_6$$

$$M_8 = X_4 X_5 X_6 X_7 X_8 X_9.$$

Cada um dos M_i , para $i \in \{1, \dots, 8\}$ tem a seguinte propriedade:

$$M_i |0_L\rangle = |0_L\rangle$$

$$M_i |1_L\rangle = |1_L\rangle.$$

Portanto, os operadores M_i , para $i \in \{1, \dots, 8\}$ são os geradores do estabilizador associado ao código de Shor.

Na notação de (4.9), os geradores podem ser representados por

$$\left[\begin{array}{cccccccc|cccccccc} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 \\ \hline 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{array} \right],$$

matriz que caracteriza o código no formalismo estabilizador.

Note que X_1 e X_2 anticomutam com o operador M_1 e por isso ambos são erros detectados pela medida de M_1 .

De fato, X_1 corresponde ao vetor $(100000000|000000000) = (a|b)$ e X_2 corresponde a $(010000000|000000000) = (a'|b')$. M_1 equivale a $(000000000|110000000) = (\alpha|\beta)$.

Portanto, o autovalor de $M_1 X_1$ é calculado por

$$(a|b) * (\alpha|\beta) = 1,$$

o que implica que a síndrome tem valor -1, o que indica que um erro foi detectado.

Os operadores da forma X_j , para $j \in \{1, \dots, 9\}$ não podem ser detectados pela medida do gerador M_1 pois ambos comutam.

Por exemplo, X_9 equivale ao vetor $(000000001|000000000) = (\bar{a}|\bar{b})$. Portanto,

$$(\bar{a}|\bar{b}) * (\alpha|\beta) = 0,$$

o que resulta em síndrome 1, que indica que não houve erro detectado.

2) Código de sete qubits [23], [22].

Neste exemplo, estudaremos a construção do código $CSS(7, 1, 3)$, também conhecido como código de Steane, sob o ponto de vista dos geradores do estabilizador, com o objetivo de compará-la à forma apresentada na Seção 3.5.2.

Considere um código de Hamming clássico com $n = 7$ e $k = 4$. Sua matriz de verificação de paridade é dada por

$$H = \begin{bmatrix} 1 & 1 & 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{bmatrix}.$$

A matriz geradora do estabilizador do código quântico é dada por:

$$\bar{H} = \begin{bmatrix} H & 0 \\ 0 & H \end{bmatrix}, \quad (4.11)$$

ou seja,

$$\bar{H} = \left[\begin{array}{ccccccc|cccccccc} 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{array} \right].$$

O processo para relacionar os geradores do estabilizador com os operadores de erros X e Z é transformar $\bar{H}$ em uma matriz da forma $[H_X|H_Z]$, ou seja, identificar cada entrada 1 da matriz H da esquerda em (4.11) com o operador X e cada entrada 1 da matriz H da direita em (4.11) com o operador Z . As entradas 0 são identificadas com o operador I .

A matriz resultante é

$$\bar{H} = \left[\begin{array}{cccccc|cccccccc} X & X & X & X & I & I & I & I & I & I & I & I & I \\ X & X & I & I & X & X & I & I & I & I & I & I & I \\ X & I & X & I & X & I & X & I & I & I & I & I & I \\ I & I & I & I & I & I & I & Z & Z & Z & Z & I & I & I \\ I & I & I & I & I & I & I & Z & Z & I & I & Z & Z & I \\ I & I & I & I & I & I & I & Z & I & Z & I & Z & I & Z \end{array} \right].$$

Este processo introduz três operadores que implementam medidas na verificação de paridade e por isso não alteram as características do código inicial.

Os geradores $X_1X_2X_3X_4$, $X_1X_2X_5X_6$ e $X_1X_3X_5X_7$, correspondentes às linhas 1, 2 e 3 da matriz $\bar{H}$, respectivamente, e os geradores $Z_1Z_2Z_3Z_4$, $Z_1Z_2Z_5Z_6$ e $Z_1Z_3Z_5Z_7$, correspondentes às linhas 4, 5 e 6, respectivamente, foram obtidos a partir do mesmo código clássico, o código de Hamming com parâmetros $[7, 4, 3]$, que contém o seu dual de parâmetros $[7, 3, 3]$.

Em [21] e [13] menciona-se o seguinte resultado:

Suponha que C_1 e C_2 sejam códigos clássicos lineares com parâmetros $[n, k_1]$ e $[n, k_2]$, respectivamente, tal que $C_2 \subset C_1$ e $C_1 \subset C_2^\perp$ com capacidade de correção t , onde $t = \lfloor (d - 1)/2 \rfloor$. O código quântico resultante (pela construção apresentada) tem parâmetros $(n, k_1 - k_2)$ e corrige t erros.

No caso do código de Steane, C_1 tem $n = 7$ e $k = 4$ e C_2 tem $n = 7$ e $k = 3$. Portanto, isto justifica os parâmetros do código de Steane: $n = 7$, $k = 1$ e $d = 3$.

Uma outra maneira de estudar o CSS(7, 1, 3) é considerar os seguintes operadores:

$$\begin{aligned} M_1 &= Z_1Z_2Z_3Z_4 \\ M_2 &= Z_1Z_2Z_5Z_6 \\ M_3 &= Z_1Z_3Z_5Z_7, \end{aligned} \tag{4.12}$$

e

$$\begin{aligned} M_4 &= X_1X_2X_3X_4 \\ M_5 &= X_1X_2X_5X_6 \\ M_6 &= X_1X_3X_5X_7. \end{aligned} \tag{4.13}$$

Todos os operadores M_i , para $i = 1, 2, 3, 4, 5, 6$ verificam a paridade das palavras-código. Os operadores M_i , para $i = 1, 2, 3$, detectam erros do tipo *bit flip*. Já os operadores M_i , para $i = 4, 5, 6$ detectam erros do tipo *phase flip*.

Se a medida de M_i , para $i = 1, 2, 3$ resulta em 1, o espaço gerado por eles é tal que as palavras-código satisfazem a matriz verificação de paridade do código de Hamming na base X . Equivalentemente, se a medida de M_i , para $i = 4, 5, 6$ resulta em 1, as palavras-código no espaço gerado satisfazem a matriz verificação de paridade na base Z , que corresponde a efetuar a transformada de Hadamard sobre a base X .

Assim, constrói-se o código de sete qubits sob a condição de que a matriz verificação de paridade de Hamming deve ser satisfeita nas bases X e Z . Os geradores comutam porque o código de Hamming contém o seu dual.

Esta descrição via geradores do estabilizador mostra de maneira objetiva como foi feita a escolha para os geradores s_i , para $i = 1, 2, 3, 4, 5, 6$ apresentada na Seção 3.5.2.

3) Códigos CSS mais gerais [23].

No exemplo anterior, consideramos um estabilizador cujos geradores podem ser escolhidos como um produto de X 's ($a|0$) ou um produto de Z 's ($0|b$) e que foram obtidos a partir do mesmo código clássico. Tais geradores foram descritos na forma

$$\bar{H} = \left[\begin{array}{c|c} H_X & 0 \\ \hline 0 & H_Z \end{array} \right]. \quad (4.14)$$

O que ocorre em casos mais gerais quando os geradores não são produtos de Z ou X , ou se eles forem obtidos de códigos clássicos diferentes?

Veremos então uma outra forma de descrever a classe de códigos CSS. Seja C_1 o código clássico linear com matriz verificação de paridade $H_1(n - k_1) \times n$ e seja C_2 um *subcódigo* de C_1 , com matriz verificação de paridade $H_2(n - k_2) \times n$, onde $k_2 < k_1$. As primeiras $n - k_1$ linhas de H_2 coincidem com as de H_1 , mas existe $k_1 - k_2$ linhas adicionais linearmente independentes em H_2 . Assim, cada palavra-código em C_2 está contida em C_1 , mas C_2 possui palavras-código a mais.

O subcódigo C_2 define uma relação de equivalência em C_1 ; dizemos que $u, v \in C_1$ são *equivalentes* ($u \equiv v$) se, e somente se, existe um w em C_2 tal que $u = v + w$. As classes de equivalência são as *classes laterais* de C_2 em C_1 .

Um código CSS é um código quântico com $k = k_1 - k_2$ que associa um código a cada

classe de equivalência. Cada elemento de uma base para o subespaço de codificação pode ser escrito na forma

$$|\bar{w}\rangle = \frac{1}{\sqrt{2^{k_2}}} \sum_{v \in C_2} |v + w\rangle, \quad (4.15)$$

uma superposição com pesos iguais de todas as palavras-código na classe lateral representada por w . Existe $2^{k_1 - k_2}$ palavras-código linearmente independentes. Os estados $|\bar{w}\rangle$ são normalizados e mutuamente ortogonais: $\langle \bar{w} | \bar{w}' \rangle = 0$, se w e w' pertencem a classes laterais diferentes, [23].

4) O Código de cinco qubits [23].

O código de cinco qubits é um exemplo de um código estabilizador não-CSS (pois os geradores do estabilizador não são produtos de X ou Z), que é perfeito, não degenerado e tem parâmetros $(5, 1, 3)$. Conforme mencionado na Seção 3.6.1, $n = 5$ é o menor comprimento das palavras-código de um código capaz de corrigir um erro em um qubit (limitante de Knill-Laflamme) e por isso dizemos que este é um código perfeito.

Os quatro geradores do estabilizador podem ser descritos por

$$\begin{aligned} M_1 &= X_1 Z_2 Z_3 X_4, \\ M_2 &= X_2 Z_3 Z_4 X_5, \\ M_3 &= X_1 X_3 Z_4 Z_5, \\ M_4 &= Z_1 X_2 X_4 Z_5, \end{aligned} \quad (4.16)$$

onde M_2, M_3, M_4 são obtidos a partir de M_1 por permutações cíclicas dos qubits. $M_5 = M_1 M_2 M_3 M_4$ não é independente dos outros quatro. Mas, se uma permutação cíclica de um gerador também é gerador, então trata-se de um código *cíclico*. Como nenhum dos geradores têm componentes Y , o quadrado de qualquer M_i é I (ver Seção 4.1). Para cada par de geradores, há duas posições coincidentes para um X e um Z , de forma que os geradores comutam. Cada operador de Pauli de peso 1 ou 2 anticomuta com pelo menos um gerador do estabilizador. Portanto, a distância do código é 3.

Considere os operadores de erros de peso 1. Estes operadores são da forma X_i , $i \in \{1, \dots, 6\}$ (um operador X na posição i e I nas restantes), ou da forma Z_i . Todos os operadores de erros com esta característica anticomutam com pelo menos um dos geradores. Por exemplo, operador de erro X_1 anticomuta com o gerador M_4 e Z_1 anticomuta com M_1 e M_3 .

No caso de operadores de erros com peso 2, considere, sem perda de generalidade, aqueles que atuam sobre os dois primeiros qubits. Suponha que exista algum que comute com todos os geradores M_i , $i = 1, 2, 3, 4$. Para comutar com a componente X_2 de M_2 e com a componente X_1 de M_3 , o operador deve ser da forma X_1X_2 . Todavia, X_1X_2 anticomuta com X_1Z_2 de M_1 e com Z_1X_2 de M_4 .

Portanto, o peso mínimo para o operador que comuta com todos os geradores é 3.

Na notação simplética, o estabilizador é representado pela matriz

$$\tilde{H} = \left[\begin{array}{ccccc|ccccc} 1 & 0 & 0 & 1 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 \end{array} \right]. \quad (4.17)$$

$\tilde{H}$ tem uma interpretação interessante. Cada uma das suas colunas pode ser vista como a *síndrome* de um erro sobre um qubit. Por exemplo, o operador *bit flip* X_j , comuta com M_i se M_i tem I ou X na posição j , e anticomuta se M_i tem um operador Z na posição j . A tabela

	X_1	X_2	X_3	X_4	X_5
M_1	0	1	1	0	0
M_2	0	0	1	1	0
M_3	0	0	0	1	1
M_4	1	0	0	0	1

lista a saída da medida $M_{1,2,3,4}$ para o evento de um *bit flip*. Por exemplo, se $M_1 = M_2 = M_3 = 1, M_4 = -1$ temos a identificação do erro, segundo a eq. (4.8).

Também, a submatriz da esquerda de $\tilde{H}$ pode ser vista como a tabela de síndromes para os erros *phase flip*.

	Z_1	Z_2	Z_3	Z_4	Z_5
M_1	1	0	0	1	0
M_2	0	1	0	0	1
M_3	1	0	1	0	0
M_4	0	1	0	1	0

Como Y anticomuta com X e Z , obtemos a síndrome para o erro Y_i somando a i -ésima coluna da tabela dos X_i 's e a i -ésima coluna da tabela dos Z_i 's.

	Y_1	Y_2	Y_3	Y_4	Y_5
M_1	1	1	1	1	0
M_2	0	1	1	1	1
M_3	1	0	1	1	1
M_4	1	1	0	1	1

Encontramos, por inspeção, que as 15 colunas das tabelas de síndrome de X , Y , e Z são todas distintas e com isso verificamos que o código é não degenerado e corrige um erro.

Por causa da ciclicidade do código, caracterizamos todos os 15 elementos não triviais do seu estabilizador. Além de $M_1 = XZZXI$ e dos quatro operadores obtidos por permutações cíclicas do qubit, o estabilizador também contém

$$M_3 M_4 = -YXXYI, \quad (4.18)$$

mais suas permutações cíclicas, e

$$M_2 M_5 = -ZYYZI, \quad (4.19)$$

e suas permutações cíclicas. Todos os elementos do estabilizador são operadores de Pauli de peso 4.

5)- Códigos detectores de erros

O estado de Bell

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

é estabilizado pelos geradores

$$Z_1 Z_2 \quad e \quad X_1 X_2.$$

O código associado a estes geradores possui $n = 2$, $k = 0$ e tem distância mínima 2 porque não contém operadores de peso 1 que comutem com X e Z . Como mencionado no Capítulo 2, um erro do tipo *bit flip*, *phase flip*, ou ambos, levam o estado $|\Phi\rangle$ a outro autoespaço que é ortogonal ao primeiro.

Podemos estender estes códigos para $n = 4$ e $k = 2$, de tal forma que os geradores do estabilizador sejam

$$Z_1 Z_2 Z_3 Z_4 \quad e \quad X_1 X_2 X_3 X_4.$$

O código estendido tem distância 2 como o código original. O subespaço do código é gerado pelos estados de paridade par que são invariantes sob a ação simultânea em quatro qubits. Uma base para o $(4, 2, 2)$ é:

$$|0000\rangle + |1111\rangle, |0011\rangle + |1100\rangle, |0101\rangle + |1010\rangle, |0110\rangle + |1001\rangle.$$

Evidentemente, X ou Z agindo sobre qualquer qubit leva um estado a outro ortogonal ao subespaço de codificação, de tal forma que qualquer erro sobre um qubit pode ser detectado.

A generalização desta classe é o código $(2m, 2m - 2, 2)$ com geradores:

$$Z_1 Z_2 \cdots Z_{2m}$$

e

$$X_1 X_2 \cdots X_{2m},$$

com comprimento sempre par para que os operadores comutem.

O subespaço de codificação é gerado pelos estados dados por

$$\frac{1}{\sqrt{2}}(|x\rangle + |\neg x\rangle), \quad (4.20)$$

onde x é um vetor de peso par $n = 2m$ e $\neg x$ denota o complementar binário de x .

4.3 Códigos Clássicos Construídos a partir de Códigos Estabilizadores

A construção de códigos quânticos a partir de códigos clássicos foi citada e exemplificada no Capítulo 3. Todavia, existem condições para a construção de códigos estabilizadores a partir de códigos clássicos?

Embora possa parecer elementar, apresentamos como obter de uma matriz $(n - k) \times 2n$ na forma $(H_X | H_Z)$, associada a um código estabilizador, a matriz geradora de um código clássico com taxa de correção de erros similar à do código quântico original.

Nesta direção, temos o seguinte resultado:

Teorema 4.3.1 [6] *Se existe um código quântico estabilizador (n, k) que corrige t erros, então existe um código binário clássico $[n - 1, k]$ que corrige pelo menos t erros.*

A construção envolve uma transformação da matriz $(H_X | H_Z)$ para a sua forma padrão, através da aplicação de operações elementares nas linhas. Essas operações elementares mantêm as características de correção do código. A primeira é a *adição de linhas*, onde a j -ésima linha é adicionada à i -ésima linha, para $i \neq j$. Isto corresponde a trocar o i -ésimo gerador pelo produto do i -ésimo pelo j -ésimo gerador, fixando sua fase em $+1$. A segunda é a *transposição de colunas*, onde a i -ésima coluna é trocada com a j -ésima coluna nas submatrizes

H_X e H_Z simultaneamente. Isto corresponde à transposição da i -ésima posição do qubit com a j -ésima posição.

Aplicando as transformações do dois tipos na matriz $(H_X \mid H_Z)$ obtemos:

$$\left[\begin{array}{c|c|c|c} I_{(s \times s)} & A_{(s \times (n-s))} & E_{1(s \times s)} & E_{2(s \times (n-s))} \\ \hline 0_{(r \times s)} & 0_{(r \times (n-s))} & E_{3(r \times s)} & E_{4(r \times (n-s))} \end{array} \right] \quad (4.21)$$

onde s é o posto da submatriz H_X , e $r = n - k - s$ ². Agora, somando as linhas dos últimos r geradores e transpondo colunas nas últimas $n - s$ posições dos qubits, a matriz pode ser reescrita na forma:

$$\left[\begin{array}{c|c|c|c|c|c|c} I_{(s \times s)} & A_{1(s \times t)} & A_{2(s \times r_1)} & B_{1(s \times s)} & B_{2(s \times t)} & B_{3(s \times r_1)} \\ \hline 0_{(r_1 \times s)} & 0_{(r_1 \times t)} & 0_{(r_1 \times r_1)} & C_{1(r_1 \times s)} & C_{2(r_1 \times t)} & I_{(r_1 \times r_1)} \\ \hline 0_{(r_2 \times s)} & 0_{(r_2 \times t)} & 0_{(r_2 \times r_1)} & D_{(r_2 \times s)} & 0_{(r_2 \times t)} & 0_{(r_2 \times r_1)} \end{array} \right] \quad (4.22)$$

onde r_1 é o posto de E_4 , $r_2 = r - r_1$ e $t = n - s - r_1$ ³. Note que se $r_2 > 0$ e $D \neq 0$ então um dos últimos r_2 geradores não deve comutar com um dos primeiros s geradores. Assim, podemos considerar $r_2 = 0$, $r_1 = r$ e $t = k$. Então, (4.21) pode ser substituída por

$$\left[\begin{array}{c|c|c|c|c|c|c} I_{(s \times s)} & A_{1(s \times k)} & A_{2(s \times r)} & B_{1(s \times s)} & B_{2(s \times k)} & B_{3(s \times r)} \\ \hline 0_{(r \times s)} & 0_{(r \times k)} & 0_{(r \times r)} & C_{1(r \times s)} & C_{2(r \times k)} & I_{(r \times r)} \end{array} \right] \quad (4.23)$$

onde $s + k + r = n$. Qualquer conjunto de geradores na forma (4.23) é dito estar na forma *padrão* [6].

Para uma matriz geradora na forma padrão, considere o código binário clássico gerado pela matriz $k \times (n - r)$:

$$\left[A_{\text{I}}^T_{(k \times s)} \mid I_{(k \times k)} \right]. \quad (4.24)$$

O próximo passo é verificar se o código clássico tem característica de correção de erros similares à do código quântico original.

O código clássico, cuja matriz geradora é dada por (4.24), consiste de 2^k palavras-código no espaço F_2^{n-r} . Com isso, procuramos construir um *isomorfismo* entre este código e a sua versão quântica especificada em (4.23), consistindo de 2^k palavras-códigos contidas num conjunto especial S de 2^{n-r} estados distintos de n -qubits. Este conjunto S tem a propriedade de ser fechado com respeito aos erros X nas primeiras $n - r$ posições do qubit. O efeito dos erros *bit flip* sobre as palavras-código do código clássico no espaço F_2^{n-r} é equivalente

²É como aplicar a eliminação de Gauss na submatriz H_X .

³É como aplicar a eliminação de Gauss agora na submatriz E_4 de (4.26).

ao efeito dos erros X nas primeiras $n - r$ posições do qubit nas palavras-código do código quântico sobre o espaço S .

Formalmente, o *isomorfismo* que queremos construir é uma função $\tau : F_2^{n-r} \longrightarrow S$, tal que [6]:

1. τ é bijetora;
2. Para cada $(n - r) - \text{upla } (y_1 \cdots y_{n-r}) \in F_2^{n-r}$, palavra-código do código clássico, $\tau(y_1 \cdots y_{n-r})$ é uma palavra-código do código quântico;
3. Para cada palavra-código $(y_1 \cdots y_{n-r}) \in F_2^{n-r}$ do código clássico e para cada vetor erro $(E_1 \cdots E_{n-r}) \in F_2^{n-r}$, temos

$$\tau(y_1 \cdots y_{n-r} \oplus E_1 \cdots E_{n-r}) = X^{E_1} \otimes \cdots \otimes X^{E_{n-r}} \otimes \overbrace{I \otimes \cdots \otimes I}^r \tau(y_1 \cdots y_{n-r}).$$

Se tal isomorfismo existir, um erro no i -ésimo bit do código clássico, para $i \in \{1, \dots, n - r\}$, corresponde a um erro X no i -ésimo qubit do código quântico. Assim, se o código quântico pode corrigir quaisquer t erros então ele pode corrigir quaisquer t erros X sobre as primeiras $n - r$ posições do qubit e o código clássico associado corrige quaisquer t erros. O procedimento é o seguinte: dada uma palavra-código $y_1 \cdots y_{n-r}$ sujeita a um erro $E_1 \cdots E_{n-r}$ de peso limitado por t , primeiro aplicamos a função τ sobre ela. Por 2) e 3), o resultado é $\tau(y_1 \cdots y_{n-r})$ sujeito a no máximo t erros Z sobre as primeiras $n - r$ posições dos qubits, que podem ser corrigidos. Por 1), τ^{-1} pode ser aplicada para corrigir a palavra-código quântica, implicando na correção da palavra-código original. Assim, se estabelecermos uma função τ que satisfaça as três propriedades, o código clássico especificado por (4.24) deve corrigir pelo menos a mesma quantidade de erros que o código quântico especificado por (4.23) pode corrigir.

Para o caso do canal de depolarização com parâmetro p , se o código quântico possui fidelidade $1 - \varepsilon$, terá também fidelidade $1 - \varepsilon$ para um canal que aplica X independentemente nas primeiras $n - r$ posições do qubit independentemente com probabilidade p . Assim, o código clássico correspondente corrige o erro com probabilidade $1 - \varepsilon$ sobre o canal BSC com parâmetro p .

A bijeção a qual nos referimos nesta seção foi construída por Cleve, [6]. Neste trabalho, o autor demonstra também que a bijeção τ satisfaz as propriedades descritas em 1), 2) e 3).

Ilustraremos os resultados com um exemplo.

Exemplo 4.3.1 [6] Considere o código quântico com parâmetros $(8, 3)$ que corrige um erro e cujo estabilizador é dado por

$$\left[\begin{array}{cccccccc|cccccccc} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{array} \right] \quad (4.25)$$

Colocando na forma padrão

$$\left[\begin{array}{cccc|cccc|cccc|cccc|c} 1 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 & 1 & 1 & 0 & 1 & 1 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 \\ \hline 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{array} \right], \quad (4.26)$$

temos um código equivalente a $(8, 3)$ na forma (4.23). Pelo exposto, o código linear binário resultante é gerado pela matriz na forma (4.24):

$$\left[\begin{array}{cccc|ccc} 1 & 1 & 1 & 0 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 1 & 0 \\ 1 & 0 & 1 & 1 & 0 & 0 & 1 \end{array} \right], \quad (4.27)$$

que é a matriz verificação de paridade do código de Hamming $[7, 3, 3]$.

Mencionamos que códigos clássicos podem ser utilizados para gerar códigos quânticos como é o caso dos códigos CSS. Concluimos que a recíproca também vale, ou seja, existe uma perfeita identificação quanto a possibilidade de construção entre códigos clássicos e códigos quânticos.

Cabe-nos a pergunta do que mais pode ser identificado entre as duas teorias de codificação. É de conhecimento geral que várias propriedades dos códigos quânticos coincidem com as dos códigos clássicos. A fundamentação teórica que garante tal afirmação está relacionada com o conceito de geração de códigos sobre $GF(4)$, o corpo de Galois com 4 elementos.

4.4 Códigos sobre $GF(4)$

Considere o caso do código de cinco qubits, Seção 4.2. Tal código quântico com parâmetros $(5, 1, 3)$ foi caracterizado a partir da escolha dos geradores do estabilizador M_i , para

$i = 1, 2, 3, 4$. Disso, concluímos que a distância mínima é $d = 3$. É com fundamentos envolvidos na construção de códigos sobre $GF(4)$ que justificaremos como foram obtidos tais geradores.

As formas tradicionais utilizadas para gerar códigos clássicos são:

1. utilização da estrutura de espaço vetorial k -dimensional cujos vetores têm comprimento n e cujas componentes pertencem a $\mathbb{Z}_2^n$,
2. utilização de corpos finitos com q elementos que em geral denominamos $GF(q) = F_q$ ou *corpo de Galois com q elementos*, onde $q = p^m$, para p primo.

Para códigos não binários, modelamos o erro por um deslocamento cíclico dos q símbolos do código. Existe, desta forma, $q - 1$ erros não triviais. O peso de Hamming de um vetor em $GF(q)$ é o número de elementos que são diferentes de zero. A distância de Hamming entre dois vetores é o peso da diferença entre eles. Um código clássico $[n, k, d]_q$ consiste de q^k palavras-código em F^n , onde a distância mínima entre um par delas é d . O limitante superior do número de palavras-código encontrado pelo método de *empacotamento de esferas* e que deve ser satisfeito pelo código $[n, k, d]$ com $d = 3$, é

$$1 + (q - 1)n \leq q^{n-k}. \quad (4.28)$$

O código de Hamming perfeito binário com parâmetros $n = 2^m - 1$, $k = n - m$ satura este limitante para $q = 2$ mas, admite generalização em qualquer F_q que pode ser construída a partir de

$$n = \frac{q^m - 1}{q - 1}, \quad (4.29)$$

para $k = n - m$.

O código $(5, 1, 3)$ pode ser gerado a partir de um código de Hamming clássico na forma $[5, 3, 3]_4$. Observamos assim uma identificação entre os códigos clássicos sobre $GF(4)$ e os códigos quânticos estabilizadores binários. Isto é possível porque o estabilizador pode ser associado a um conjunto de vetores sobre F_4 , fechado com relação à adição. Apresentaremos a seguir como é possível tal associação.

O $GF(4)$ possui quatro elementos que denotaremos por $0, 1, \omega, \bar{\omega}$, onde

$$\begin{aligned} 1 + 1 &= \omega + \omega = \bar{\omega} + \bar{\omega} = 0, \\ 1 + \omega &= \bar{\omega}, \end{aligned} \quad (4.30)$$

e $\omega^2 = \bar{\omega}$, $\omega\bar{\omega} = 1$.

Se identificarmos X com ω , Z com 1 , Y com $\omega + 1 = \bar{\omega}$ e I com $\omega + \bar{\omega} + 1 = 0$, notamos que a estrutura aditiva de $GF(4)$ está associada ao grupo de Pauli gerado por X e Z . Mas, como mencionado, o grupo de Pauli é não-abeliano. É preciso então estabelecer uma estrutura, como fizemos ao construir $\bar{E}$, que torne o grupo de Pauli abeliano quando restrito a tal estrutura.

Com este propósito, definimos a função linear *traço*, $tr : GF(4) \rightarrow \mathbb{Z}_2$ tal que [22, 5]:

$$tr(0) = tr(1) = 0 \quad tr(\omega) = tr(\bar{\omega}) = 1, \quad (4.31)$$

ou, na forma mais geral, como

$$tr(x) = x + \bar{x}$$

para $x \in GF(4)$.

A partir da função traço, é possível obter um espaço vetorial n -dimensional sobre $GF(4)$. Como ω e 1 formam uma base para o espaço vetorial bidimensional com relação a $GF(2) \equiv \mathbb{Z}_2$, podemos escrever qualquer elemento de $GF(4)^n$ como $\omega a + 1b$, onde $a, b \in \mathbb{Z}_2^n$.

Definindo a função $\xi : \bar{E} \rightarrow GF(4)^n$ como

$$\xi(a \mid b) = \omega a + 1b, \quad (4.32)$$

temos que, quando $n = 1$, ξ leva a classe $(1 \mid 0)$, contendo X , em ω e a classe $(0 \mid 1)$, contendo Z , em 1 , como necessário. A função ξ , além de ser uma bijeção, é um isomorfismo entre os espaços vetoriais binários, quando consideramos $GF(4)^n$ como um espaço vetorial binário sob a restrição de escalares.

Como $\omega a_j + b_j \neq 0$ se, e somente se, $a_j \neq 0$ ou $b_j \neq 0$, o peso quântico de um vetor $(a \mid b) \in \bar{E}$ coincide com o peso clássico de Hamming de sua imagem $\xi(a \mid b) = \omega a + 1b \in GF(4)^n$.

Tendo em vista que o produto interno usual $\circ$ sobre $GF(4)^n$ é dado por

$$v \circ \omega = v \circ \bar{\omega} = \sum_{j=1}^n v_j \bar{\omega}_j, \quad (4.33)$$

é possível mostrar que

$$tr[(\omega a + 1b) \cdot \overline{(\omega a' + 1b')}] = a \cdot b' + a' \cdot b = (a \mid b) * (a' \mid b'). \quad (4.34)$$

Isso implica que a forma $\circ$ é identicamente nula sobre um subespaço de $GF(4)^n$ se, e somente se, a forma $*$ é identicamente nula sobre o subespaço correspondente de $\bar{E}$. Isto é, o espaço $GF(4)^n$, sob a restrição de escalares, é isomorfo a $\bar{E}$ pela função ξ que identifica o peso quântico com o peso clássico. Também, um subespaço de $GF(4)^n$ é a imagem, sob ξ , de um subespaço totalmente isotrópico de $\bar{E}$ (com respeito a $*$) se, e somente se, é totalmente isotrópico (com respeito a $\circ$).

Segue um resultado importante:

Teorema 4.4.1 [22] *Se W é um subespaço l -dimensional totalmente isotrópico de $GF(4)^n$ (com respeito ao produto interno Hermitiano) tal que $W^\perp - W$ tem peso mínimo de Hamming d , então a restrição de escalares a $\mathbb{Z}_2$ associa um código estabilizador $(n, n - 2l, d)$ a W .*

Exemplificaremos os resultados discutidos nesta seção aplicando-os na construção do código quântico de parâmetros $(5, 1, 3)$.

Exemplo 4.4.1 *Mostraremos que o código $(5, 1, 3)$ pode ser gerado a partir do código $[5, 3, 3]_4$.*

A matriz verificação de paridade do código clássico estendido sobre F_4 é dada por

$$H = \begin{bmatrix} 1 & \omega & \omega & 1 & 0 \\ 0 & 1 & \omega & \omega & 1 \end{bmatrix}$$

cujas colunas geram os cinco espaços 1-dimensionais distintos de $GF(4)^2$. Agora,

$$(1\omega\omega 10) \leftrightarrow (ZXXZI) \leftrightarrow (01100 \mid 10010) = (a_2 \mid b_2),$$

e

$$(01\omega\omega 1) \leftrightarrow (IZXXZ) \leftrightarrow (00110 \mid 01001) = (a_3 \mid b_3).$$

De maneira similar, multiplicando a primeira linha de H por ω e somando a ela a segunda temos $(a_1 \mid b_1)$; multiplicando a segunda por ω e somando a ela a primeira, encontramos $(a_4 \mid b_4)$.

Com isto, concluímos que podemos obter os geradores do estabilizador do código $(5, 1, 3)$ a partir do código clássico com parâmetros $[5, 3, 3]$ estendido sobre $GF(4)$.

Conclusão

Por se tratar de um trabalho com abordagem teórica, pudemos analisar conceitos e resultados envolvidos na formalização do processo de transmissão de informações codificadas em estados quânticos.

A apresentação de propriedades físicas que definem as características da Mecânica Quântica, tais como o emaranhamento e o Teorema da Não-Clonagem, permitiu uma idéia geral das condições que devem ser satisfeitas pelo procedimento de codificação sob o ponto de vista quântico. Por se tratar, em geral, de conceitos desconhecidos em nosso ambiente clássico, optamos por apresentá-los no Capítulo 1, onde citamos também avanços relevantes em Teoria de Informação Quântica, com o intuito de que o leitor conheça as estratégias nesta área de pesquisa.

No Capítulo 2, baseados em resultados da Teoria de Grupos, desenvolvemos os passos que levaram à definição do grupo de operadores de erros, a partir do qual gera-se códigos corretores. A teoria dos Grupos de Clifford possibilitou uma descrição matricial desta definição, fazendo com que as justificativas pudessem ser citadas de maneira concisa, como de fato ocorre em Mecânica Quântica. Em seguida, a análise da estrutura matemática do grupo de erros e das consequências da mesma no estudo levou-nos à discussão da não-comutatividade do grupo de Pauli e da necessidade que esta característica fosse contornada no contexto de codificação. A procura por estratégias que permitissem associar condições necessárias para a criação de códigos neste ambiente não abeliano foi concluída com resultados conhecidos matematicamente, tais como Teorema de Sylow e suas propriedades, solubilidade de grupos e decomposição dos mesmos em classes laterais. Consequentemente, estudamos a estrutura do grupo quociente, definimos grupos especiais e extra-especiais e, em seguida, citamos conceitos da Geometria Ortogonal associada ao problema em estudo. Esta descrição explicitou a importância da decomposição ortogonal do espaço vetorial complexo, de forma que os erros gerados pelo grupo de Pauli pudessem ser distinguidos. Todavia, a ortogonalidade dos

autoespaços foi garantida pela propriedade abeliana que o grupo quociente (o grupo de erros quocientado pelo seu comutador) proporcionou. Por sua vez, os autovalores associados aos autovetores geradores dos autoespaços ortogonais descreveram a ação dos operadores de erros no subespaço escolhido para a criação do código. Propriedades da Álgebra Linear possibilitaram uma descrição simplificada destas ações.

Esta abordagem descritiva construiu relações entre conceitos de Teoria de Grupos e Álgebra Linear, propriedades físicas relacionadas às características do espaço quântico e necessárias observações no contexto de codificação e criação de códigos. Optamos por assim fazê-la afim de induzir o leitor a entender o processo teórico, reconhecendo as limitações impostas pela Física Quântica e compreendendo as conseqüências no processo de transmissão de informações.

No Capítulo 3, iniciamos a descrição das propriedades dos códigos corretores de erros quânticos. Salientamos as dificuldades associadas ao processo de medida dos estados quânticos, o que não nos permite criar um código de repetição análogo ao clássico que seja confiável. Os projetores (e observáveis) e as operações associadas possibilitaram uma interpretação fundamental no entendimento de como a detecção de erros quânticos ocorre. Por isso, optamos em mencionar estes conceitos e operações com uma quantidade razoável de detalhes. Em seguida, a descrição dos canais *bit flip* e *phase flip* e do procedimento de codificação utilizando nove qubits proposto por Shor exemplificou a detecção e correção de erros através de medidas de observáveis, o que torna o formalismo simples e elegante. A condição de distinção entre estados quânticos em espaços mutuamente ortogonais foi descrita em função de resultados relacionados diretamente a escolha de subespaços de codificação, de forma a identificar conceitos físicos e restrições fundamentais para a eficiência do processo de transmissão.

Tendo definido um código corretor de erros quânticos, apresentamos o parâmetro distância d a ele associado e as relações existentes entre d , n e k . O que se concluiu é que embora as definições sejam diferentes da clássica, as condições mantidas entre os parâmetros de um código quântico são similares e muitas vezes idênticas às do caso clássico. Visando descrever maiores correspondências, estudamos o código de Steane, que pertence a classe CSS, e que é obtido a partir do código clássico de Hamming com parâmetros $[7, 4, 3]$.

No Capítulo 4, descrevemos a classe dos códigos estabilizadores, que contém a classe dos códigos CSS, que comentamos no Capítulo 3. O formalismo dos códigos estabilizadores

foi apresentado segundo definições de teoria de grupos. A possibilidade de uma descrição simplificada (em termos de geradores do estabilizador do subgrupo de erros que pode ser corrigido por cada código) é uma característica relevante desta classe. Contudo, identificações quanto à construção com códigos clássicos também são fatores que justificam a atenção especial designada. Por fim, mencionamos a correspondência entre os códigos corretores de erros quânticos e os códigos clássicos gerados sobre $GF(4)$. Com o objetivo de exemplificar o quão valiosa é esta identificação, optamos em apresentar a construção do código quântico com parâmetros $(5, 1, 3)$ a partir dos elementos de $GF(4)$ e manipulações algébricas bastante simples.

No geral, este trabalho foi realizado com o intuito de inserir o leitor em um novo ambiente de transmissão de informações, abordando conceitos necessários para o entendimento dos processos teóricos envolvidos. Acreditamos que, pela descrição adotada, esta dissertação possa ser utilizada para um maior entendimento das aplicações da Mecânica Quântica e estimular novas pesquisas, principalmente no campo de codificação quântica, que ainda não conta com muitos partidários em nosso país.

Trabalhos Futuros

Foi estudado nesta dissertação o grupo de Pauli e suas aplicações em codificação de informações em estados quânticos. Sendo este grupo um caso particular dos Quatérnios, seria interessante expandir o estudo na tentativa de unir propriedades quânticas utilizadas para a codificação às características estruturais dos Quatérnios, o que pode favorecer o entendimento geométrico do espaço para a codificação quântica.

O mesmo poderia ser realizado com os Octônios, ainda que esta teoria não esteja totalmente desenvolvida e acessível. Neste caso, estaríamos considerando um grupo de erros com oito operadores, que não encontraria correspondência no grupo de Pauli ou no $GF(4)$. Por exemplo, o código de repetição com nove qubits de Shor não seria referencial, pois os geradores do grupo de erros não corresponderiam a X e Z .

Outro tema interessante associado ao presente trabalho é a imersão de nossos estudos (que foram todos realizados nas condições do espaço de Hilbert) no espaço hiperbólico. Estudos considerando a transmissão de dados pela codificação clássica sob as condições do plano hiperbólico mostraram-se valiosos em relação aos estudos sob o plano euclidiano. Por isso, espera-se que os resultados obtidos sobre espaço hiperbólico também tragam novas dimensões para o estudo do problema de codificação quântica. Existe ainda a identificação do plano hiperbólico com os Quatérnios, o que não deixa de ser um estímulo para a procura de correlações entre o espaço hiperbólico e o grupo de Pauli.

Além destas propostas, a procura de novos códigos, o que implicaria na maior compreensão das características dos estados quânticos e da teoria matemática envolvida no processo de codificação-detecção e decodificação de erros introduzidos pelo canal que transmitirá as informações na forma de *qubit*, também é um campo de pesquisa muito promissor.

Bibliografia

- [1] “Why God plays dice,” *New Scientist* 22 de agosto, 27 (1998).
- [2] Jornal Folha de São Paulo, *Folha Ciência e Folha Informática*, 27 de setembro de 2001, 28 de junho e 23 de agosto de 2000 .
- [3] C. H. Bennett, D.P. DiVincenzo, J. A. Smolin, e W. K. Wootters, “Mixed state entanglement and quantum error correction,” *Phys. Rev. A* **54**, 3824 (1996) ou www.arxiv.org/abs/quant-ph/9604024.
- [4] A. R. Calderbank, E. M. Rains, P. W. Shor, e N. J. A. Sloane, “Quantum error correction and orthogonal geometry,” *Phys. Rev. Lett.* **78**, 405 (1997) ou www.arxiv.org/abs/quant-ph/9605005.
- [5] A. R. Calderbank, E. M. Rains, P. W. Shor, e N. J. A. Sloane, “Quantum error correction via codes over GF(4),” *IEEE Trans. Inform. Theory* **44**, 1369 (1996) ou www.arxiv.org/abs/quant-ph/9608006.
- [6] R. Cleve, “Quantum stabilizer codes and classical linear codes,” www.arxiv.org/abs/quant-ph/9612048.
- [7] D. Deutsch and A. Ekert, “Quantum computation,” www.qubit.org .
- [8] K. Doerk and T. Hawkes, *Finite Soluble Groups*, Walter de Gruyter, Nova York (1992).
- [9] A. Ekert, “Quantum cryptography based on Bell’s theorem,” *Phys. Rev. Lett.* **67**, 661 (1991).
- [10] R. P. Feynman, “Simulating physics with computers,” *Int. J. Theor. Phys.* **21**, 467 (1982).

- [11] D. Gottesman, “An introduction to quantum error correction,” www.arxiv.org/abs/quant-ph/0004072.
- [12] D. Gottesman, *Stabilizer Code and Quantum Error Correction*, www.arxiv.org/abs/quant-ph/9705052.
- [13] D. Gottesman, “Class of quantum error-correcting codes saturating the quantum Hamming bound,” *Phys. Rev A* **54**, 1862 (1996).
- [14] L. K. Grover, “A fast quantum mechanical algorithm for database search,” *Proceedings of the 28th Annual Symposium on the Theory of Computing*, 212 (1996).
- [15] P. Hall, e G. Higman, “On the length of p-soluble groups and reduction theorems for Burnside’s problem,” *Proc. London Math. Soc.* **3**, 1 (1956).
- [16] I. N. Heirstein, *Topics in Algebra*, 2^a ed., Wiley, Nova York (1975).
- [17] E. Knill, e R. Laflamme, “A theory of quantum error correction codes,” *Phys. Rev. A* **55**, 900 (1997) ou www.arxiv.org/abs/quant-ph/9604034.
- [18] C. Lavor, L. R. U. Manssur, e R. Portugal, “Shor’s algorithm for factoring large integers”, Material baseado em notas de aula dos cursos de Pós-Graduação ministrados no Laboratório Nacional de Computação Científica, Petrópolis (2003).
- [19] F. J. MacWilliams, e N. J. Sloane, *The Theory of Error-Correcting Codes*, North-Holland, Amsterdam (1977).
- [20] M. A. Nielsen, “Regras para um mundo quântico complexo,” *Scientific American Brasil*, dezembro, 80 (2002).
- [21] M. A. Nielsen, e I. R. Chuang, *Quantum Computation and Quantum Information*, Cambridge University Press, Nova York (2000).
- [22] H. Pollatsek, “Quantum error correction: classic group theory meets a quantum challenge,” *The Mathematical Association of America (MAA)* **108**, 932 (2001).
- [23] J. Preskill, Notas de Aula, www.theory.caltech.edu/people/preskill/ph229/.

- [24] P. W. Shor, “Algorithms for quantum computation: Discrete logarithms and factoring,” *Proceedings of the 35th Annual Symposium on Foundations of Computer Science*, 124 (1994).
- [25] T. Siegfried, *O Bit e o Pêndulo. A Nova Física da Informação*, Campus, Rio de Janeiro (2000). (Este livro apresenta conceitos físicos em linguagem e exemplos bastante simples. A leitura é recomendada para uma visão geral das aplicações da Mecânica Quântica.)
- [26] S. Singh, *O Livro dos Códigos*, Record, Rio de Janeiro (2001). (Este livro apresenta a história dos códigos desde o Egito Antigo até nossos dias, citando aspectos da teoria de codificação quântica e os avanços que ela proporciona).
- [27] A. Steane, “Multiple particle interference and quantum error correction,” *Proc. Roy. Soc. London A* **452**, 2551 (1996) ou www.arxiv.org/abs/quant-ph/9601029.
- [28] H. Weyl, *The Theory of Groups and Quantum Mechanics*, Dover Publisher, Nova York (1950).