

UNIVERSIDADE FEDERAL DE JUIZ DE FORA
INSTITUTO DE CIÊNCIAS EXATAS
PROGRAMA DE PÓS-GRADUAÇÃO EM MATEMÁTICA

Rafaela Cristina Oliveira da Cunha

Um estudo sobre códigos quânticos topológicos

Juiz de Fora

2025

Rafaela Cristina Oliveira da Cunha

Um estudo sobre códigos quânticos topológicos

Dissertação apresentada ao Programa de Pós-Graduação em Matemática da Universidade Federal de Juiz de Fora como requisito parcial à obtenção do título de Mestre em Matemática. Área de concentração: Matemática Pura.

Orientadora: Dra. Catarina Mendes de Jesus Sanchez

Coorientadora: Dra. Beatriz Casulari da Motta Ribeiro

Juiz de Fora

2025

Ficha catalográfica elaborada através do Modelo Latex do CDC da UFJF
com os dados fornecidos pelo(a) autor(a)

da Cunha, Rafaela Cristina Oliveira.

Um estudo sobre códigos quânticos topológicos / Rafaela Cristina
Oliveira da Cunha. – 2025.

89 f. : il.

Orientadora: Catarina Mendes de Jesus Sanchez

Coorientadora: Beatriz Casulari da Motta Ribeiro

Dissertação (Mestrado) – Universidade Federal de Juiz de Fora, Instituto
de Ciências Exatas. Programa de Pós-Graduação em Matemática, 2025.

1. Palavra-chave. 2. Palavra-chave. 3. Palavra-chave. I. Sanchez,
Catarina Mendes de Jesus, orient. II. Ribeiro, Beatriz Casulari da Motta
Ribeiro, coorient. III. Um estudo sobre códigos quânticos topológicos.

Rafaela Cristina Oliveira da Cunha

Um estudo sobre códigos quânticos topológico

Dissertação apresentada ao Programa de Pós-graduação em Matemática da Universidade Federal de Juiz de Fora como requisito parcial à obtenção do título de Mestre em Matemática. Área de concentração: Matemática Pura

Aprovada em 11 de março de 2025.

BANCA EXAMINADORA

Prof^a. Dr^a. Beatriz Casulari da Motta Ribeiro - Coorientadora

Universidade Federal de Juiz de Fora

Prof. Dr. Antonio Falcó Montesinos

Universidad CEU Cardenal Herrera

Prof^a. Dr^a. Clarice Dias de Albuquerque

Universidade Federal do Cariri

Juiz de Fora, 02/04/2025.



Documento assinado eletronicamente por **Beatriz Casulari da Motta Ribeiro, Professor(a)**, em 03/04/2025, às 11:22, conforme horário oficial de Brasília, com fundamento no § 3º do art. 4º do [Decreto nº 10.543, de 13 de novembro de 2020](#).



Documento assinado eletronicamente por **Clarice Dias de Albuquerque, Usuário Externo**, em 12/05/2025, às 17:22, conforme horário oficial de Brasília, com fundamento no § 3º do art. 4º do [Decreto nº 10.543, de 13 de novembro de 2020](#).



Documento assinado eletronicamente por **Antonio Falco Montesinos, Usuário Externo**, em 20/05/2025, às 08:29, conforme horário oficial de Brasília, com fundamento no § 3º do art. 4º do [Decreto nº 10.543, de 13 de novembro de 2020](#).



A autenticidade deste documento pode ser conferida no Portal do SEI-Uff (www2.uff.br/SEI) através do ícone Conferência de Documentos, informando o código verificador **2329106** e o código CRC **B6404C25**.

Dedico este trabalho à minha mãe, à minha irmã, ao meu grande amor e aos meus amigos.

AGRADECIMENTOS

À minha família Isabel, minha mãe, e Júlia, minha irmã, Júlia, meu mais sincero e profundo agradecimento, por todo apoio incondicional, pelos incentivos constantes, pelo suporte nos momentos difíceis e pela confiança depositada em mim. Em especial à minha mãe, agradeço por cada gesto de companheirismo e compreensão, por cada sacrifício silencioso e por acreditar em mim.

Aos meus familiares, em especial, aos meus tios Helieber, Luciene, Paulo e Léia, à minha prima Carol e a Luciana e Wagner, por todo apoio dado.

Ao meu grande amor, Biel, meu mais profundo e sincero agradecimento, por durante toda essa jornada de pós-graduação, estar ao meu lado, me apoiando nos momentos difíceis e celebrando comigo nos bons momentos. Também, agradeço pelas conversas, por todo amor e carinho, pelas diversões, experiências únicas inesquecíveis que me possibilita diariamente e por todo apoio e companhia nos momentos bons e ruins.

Às professoras Catarina Mendes de Jesus Sanchez e orientadora Beatriz Casulari da Motta Ribeiro, por aceitarem me orientar. Em especial à professora Beatriz, por ter me guiado, em meio ao caos diário, com a maior dedicação possível, para me propiciar o maior entendimento quanto à teoria de Códigos Corretores de Erros clássica e quântica.

Aos amigos e colegas de UFJF, pelo apoio, incentivo e momentos de leveza em meio aos desafios. Em especial, agradeço aos meus amigos Thayonara, Vinícius Sangi Gustavo Dutra, Milena, Thiago Evangelista e Raphael. Além disso, pelos ensinamentos aprendidos e frutos colhidos durante essa trajetória.

Aos amigos Thayonara e Vinícius, pela amizade, pelo apoio incondicional, por me ajudarem no meu crescimento como pessoa, por serem meus companheiros desde 2019, pelas mais icônicas conversas sobre a vida e pelas melhores chamadas em grupo.

Ao Gustavo Dutra, que mesmo que temos tido pouco contato durante a pós-graduação, foi meu parceiro de graduação, entendendo minhas indecisões durante essa caminhada.

A Milena, por sua amizade, por ser minha parceira de praia/evento em Porto de Galinhas e por me ajudar a aguentar os problemas de pós-graduação.

A Sheu, pela amizade e por me aceitar como hóspede na casa dela, quando precisava ficar por São Pedro, sendo minha companheira de pós passeios.

Ao Thiago Evangelista, pelas conversas, pelos momentos engraçadíssimos e por me ajudar desde o finalzinho da graduação.

Ao Raphael, pela amizade e companheirismo e por, assim como o Biel, ser meu companheiro de Yu-Gi-Oh, propiciando momentos divertidíssimos.

Agradeço também, à Letianne, à Larissa, à Gabi, à Fran, ao Walter, à Sthefy, à

Camila, à Vanessa (minha parceira de AL no verão), à Luma e ao Vinícius Rinco (meus parceiros ubaenses e de RPG), ao Augusto, ao Paulo, ao Christian, a Letícia Naves, ao Heitor, ao Matheus e ao Daniel de Souza, por me acompanharem e me apoiarem nessa jornada.

Aos membros da banca examinadora, Dra. Beatriz Casulari da Motta Ribeiro, Dra. Catarina Mendes de Jesus Sanchez, Dra. Clarice Dias Albuquerque e Dr. Antonio Falcó Montesinos, por dedicarem seu tempo, conhecimento e atenção à avaliação deste trabalho.

À Universidade Federal de Juiz de Fora e ao Departamento de Matemática.

À CAPES, pela bolsa de pós-graduação.

RESUMO

Nesse trabalho, estamos interessados em explorar os fundamentos no estudo dos códigos corretores de erros, tanto no contexto clássico quanto no quântico. Para isso, buscamos contextualizar os desafios da informação quântica em relação à informação clássica e a necessidade de técnicas específicas para lidar com a fragilidade dos estados quânticos. Apresentamos os códigos quânticos coloridos, uma classe promissora de códigos quânticos que combina elementos topológicos para oferecer maior robustez e eficiência na correção de erros.

Palavras-chave: códigos corretores de erro; códigos quânticos corretores de erros; códigos estabilizadores; códigos quânticos coloridos.

ABSTRACT

In this work, we are interested in exploring the fundamentals in the study of error correcting codes, both in the classical and quantum contexts. To do this, we seek to contextualize the challenges of quantum information in relation to classical information and the need for specific techniques to deal with the fragility of quantum states. We will present the color quantum codes, a promising class of quantum codes that combine topological elements to offer greater robustness and efficiency in error correction.

Keywords: error correcting codes; quantum error correcting codes; stabilizer codes; color quantum codes.

LISTA DE ILUSTRAÇÕES

Figura 1 – Sistema de Comunicação Típico.	15
Figura 2 – Esfera de Bloch.	26
Figura 3 – Representação dos qubits e operadores X_f e Z_v	68
Figura 4 – Operadores face B_f^σ agindo nos qubits da face verde f	70

LISTA DE TABELAS

Tabela 1	– Ações das portas lógicas clássicas AND, OR, NAND e NOR. Fonte: [1].	31
Tabela 2	– Ações das portas X , Z e H sobre os qubits $ 0\rangle$ e $ 1\rangle$	34
Tabela 3	– Ação da porta CNOT.	35
Tabela 4	– Síndromes de erros para o caso <i>bit flip</i>	45
Tabela 5	– Resultado das medidas de Z_1Z_2 e Z_2Z_3 sobre $ \psi\rangle$	46
Tabela 6	– Resultado das medidas de Z_1Z_3 e Z_2Z_3 sobre $ \psi\rangle$	47
Tabela 7	– Resultado das medidas de Z_1Z_2 e Z_1Z_3 sobre $ \psi\rangle$	47
Tabela 8	– Resultados das medidas de X_1X_2 e X_2X_3	49

SUMÁRIO

1	INTRODUÇÃO	11
2	CÓDIGOS CORRETORES DE ERROS CLÁSSICOS	13
2.1	CÓDIGOS CORRETORES DE ERROS	14
2.1.1	Códigos lineares	16
2.2	MATRIZ GERADORA DE UM CÓDIGO	16
2.3	MATRIZ TESTE DE PARIDADE	18
2.4	DECODIFICAÇÃO POR SÍNDROME	20
2.5	CÓDIGOS DUAIS	21
3	INFORMAÇÃO E COMPUTAÇÃO QUÂNTICA	23
3.1	NOTAÇÃO DE DIRAC	23
3.2	POSTULADOS DA MECÂNICA QUÂNTICA E MEDIÇÕES	24
3.3	A UNIDADE DE INFORMAÇÃO QUÂNTICA	25
3.4	EMARANHAMENTO QUÂNTICO	27
3.5	PORTAS QUÂNTICAS	30
4	CÓDIGOS QUÂNTICOS CORRETORES DE ERROS	37
4.1	DIFERENÇAS ENTRE INFORMAÇÃO QUÂNTICA E CLÁSSICA	37
4.2	CÓDIGO <i>BIT FLIP</i>	39
4.3	CÓDIGO <i>PHASE FLIP</i>	47
4.4	CÓDIGO DE SHOR	50
4.5	CÓDIGOS DEGENERADOS	52
4.6	LIMITANTES DE HAMMING E SINGLETON	53
4.7	CÓDIGOS CSS	58
4.8	CÓDIGOS ESTABILIZADORES	61
4.8.1	O grupo de Pauli	61
4.8.2	Código estabilizador	62
5	CÓDIGOS QUÂNTICOS COLORIDOS TOPOLÓGICOS	67
5.1	CONCEITOS PRELIMINARES	67
5.2	CÓDIGOS COLORIDOS	69
6	CONCLUSÃO	72
	REFERÊNCIAS	73
	APÊNDICE A – Álgebra Linear e a Notação de Dirac	75

1 INTRODUÇÃO

Os trabalhos de Richard W. Hamming e C. E. Shannon no Laboratório Bell de Tecnologia na década de 1940 marcam o início da Teoria de Codificação, que é um campo de estudo em amplo desenvolvimento, tanto do ponto de vista teórico quanto tecnológico. Cotidianamente, nos mais diversos meios, circulam informações que carregam consigo possíveis erros gerados na transmissão ou armazenamento dos dados. De maneira geral, o objetivo do uso de códigos corretores de erros é identificar e corrigir o máximo possível tais erros a fim de que a informação esteja o mais próxima da original.

No caso dos códigos corretores de erros clássicos, os métodos para melhorar a confiabilidade da transmissão estão intrinsecamente ligados às propriedades dos corpos finitos. A ideia básica é transmitir informação extra junto com a mensagem original, isto é, estender a sequência de símbolos da mensagem para uma sequência mais longa de forma sistemática, adicionando a redundância. A redundância, por sua vez, permite que o receptor seja capaz de detectar uma certa quantidade de erros (que depende dos parâmetros do código) e, por vezes, corrigi-los. Dessa forma, um código corretor de erros clássico é essencialmente, uma maneira organizada de acrescentar algum dado a cada informação que se queira transmitir ou armazenar, de modo que permita, ao recuperar a informação, detectar e corrigir erros presentes na transmissão da informação.

Mais recentemente, na década de 1990, surgiu uma classe especial de códigos corretores de erros baseada em propriedades da mecânica quântica, chamados códigos corretores de erros quânticos (QECC). O primeiro código quântico foi introduzido por Shor em 1995, sendo uma adaptação quântica do código de repetição clássico. Esse código, de nove qubits, ficou conhecido como código de Shor em homenagem ao seu autor. Em 1996, Steane propôs outro código quântico $[[7, 3, 4]]$. Esses desenvolvimentos abriram caminho para a criação dos códigos Calderbank-Shor-Steane, conhecidos simplesmente por códigos CSS, que foram formulados por Calderbank e Shor, bem como por Steane. No mesmo ano, Gottesman apresentou uma nova classe mais abrangente de códigos, conhecida como códigos estabilizadores, que inclui os códigos CSS como um caso particular. Os códigos estabilizadores têm sua construção baseada nos códigos lineares clássicos e compreendem os códigos de Shor e CSS, além de nos fornecer um método de construção de códigos quânticos. Nos últimos anos, avanços contínuos na produção e aplicação de códigos quânticos cada vez mais eficientes têm impulsionado o progresso da computação quântica.

No presente trabalho, vamos explorar os fundamentos dos códigos corretores de erros clássicos bem como dos códigos quânticos, culminando em uma breve apresentação dos códigos quânticos topológicos, em especial os códigos coloridos como um exemplo bastante promissor de códigos quânticos.

Começaremos, no Capítulo 2, abordando os códigos corretores de erros clássicos,

destacando suas principais estruturas e aplicações na detecção e correção de erros em sistemas de comunicação e armazenamento de dados. Para isso, faremos uma breve apresentação de alguns conceitos básicos da Teoria de Codificação, definiremos o que são os códigos corretores de erros, a métrica de Hamming e os parâmetros de um código. Em seguida, apresentaremos uma classe especial de códigos, denominada códigos lineares, que consiste na classe de códigos mais utilizada na prática. Também, apresentaremos as matrizes geradora e teste de paridade de um código linear, assim como uma forma de decodificação denominada decodificação por síndrome. Finalizaremos esse capítulo apresentando brevemente os códigos duais que além de sua importância no contexto clássico, também serão de suma importância para a construção dos códigos CSS.

No Capítulo 3, introduzimos a Teoria da Informação e Codificação Quântica, contextualizando os desafios da computação quântica e a necessidade de técnicas específicas para lidar com a fragilidade dos estados quânticos. Começaremos apresentando a notação de Dirac, uma notação bem incomum para matemáticos, mas eficiente para trabalhar com estados quânticos. Essa notação será utilizada no restante do trabalho. Em seguida, apresentaremos os postulados da Mecânica Quântica, que são um conjunto de princípios fundamentais que definem o comportamento dos sistemas quânticos. Para entrarmos nesse “mundo quântico” faz-se necessário a definição da unidade de informação quântica, o *qubit* e o conceito de superposição quântica. Como último assunto do capítulo, apresentaremos as portas quânticas, que estão para a computação quântica, assim como as portas lógicas clássicas estão para os computadores convencionais, e que realizam operações em qubits, manipulando seus estados através de transformações unitárias.

No Capítulo 4, introduziremos a teoria dos QECC, que estendem os conceitos clássicos para o domínio quântico, com o principal objetivo de proteger a informação quântica do ruído. Dedicaremos a segunda seção desse capítulo para destacarmos as diferenças entre a informação quântica e clássica. Em seguida, apresentaremos três exemplos de códigos quânticos. Dentre esses três códigos, o mais simples é o *bit flip*. Ele codifica três qubits e assemelha-se ao código clássico de repetição. Já o código *phase flip* não possui nenhum código equivalente nos códigos clássicos. O código de Shor é um código concatenado dos códigos *bit flip* e *phase flip* de nove qubits, e protege estados quânticos contra a ação de erros completamente arbitrários, desde que apenas um qubit seja afetado, conforme [20]. Por fim, apresentaremos os códigos CSS e os códigos estabilizadores.

Por fim, no Capítulo 5, apresentaremos os códigos quânticos coloridos, uma classe de códigos quânticos topológicos muito promissora, que combinam elementos topológicos para oferecer maior robustez e eficácia na correção de erros, sendo de particular interesse para a implementação prática da computação quântica tolerante a falhas.

2 CÓDIGOS CORRETORES DE ERROS CLÁSSICOS

Este capítulo está baseado nas referências [7] e [14] e é dedicado a apresentação dos conceitos básicos da Teoria de Códigos Corretores de Erros. Para mais informações, veja [1], [7] ou [14].

Consideremos os seguintes elementos de um código:

- $\mathcal{A}$ um **alfabeto**;
- **Letras** são os elementos de $\mathcal{A}$;
- $\mathcal{A}^n$ denota que $\mathcal{A}$ é um alfabeto cuja maior palavra tem **comprimento** n ;
- **Palavras** são os elementos de $\mathcal{A}^n$;
- Os elementos de $\mathcal{A}^n$ poderão ser representados por $a_1 \cdots a_n$ em vez de $(a_1, \dots, a_n)$.

Exemplo 2.0.1 (Código robô). Suponhamos que temos um robô que se move sobre um tabuleiro quadriculado, de modo que, ao darmos um dos comandos (Leste, Oeste, Norte ou Sul), o robô se desloca do centro de uma casa para o centro da casa adjunta indicada pelo comando.

Os quatro comandos acima descritos podem ser codificados como elementos de $\{0, 1\} \times \{0, 1\}$, como segue-se:

$$\mathcal{C}_1 = \begin{cases} \text{Leste} \mapsto 00 \\ \text{Oeste} \mapsto 01 \\ \text{Norte} \mapsto 10 \\ \text{Sul} \mapsto 11 \end{cases}.$$

A representação $\mathcal{C}_1$ dos comandos dados pelos elementos de $\{0, 1\} \times \{0, 1\}$ é chamada de **Código Fonte**.

Suponhamos, que esses pares ordenados devam ser transmitidos via radio e que o sinal no caminho sofra interferências. Imaginemos que a mensagem 00 possa, na chegada, ser recebida como 01, o que faria com que o robô, em vez de ir para Leste, fosse para Oeste. Desse modo, o erro ocorrido durante a transmissão não seria percebido, apesar de ter comprometido a mensagem original. Esse problema ocorre devido à “proximidade” das palavras. Para diminuir a possibilidade de que casos como esse ocorram, o que se faz é recodificar as palavras, de modo a introduzir **redundâncias**¹ que permitam detectar e corrigir erros.

¹ As redundâncias são utilizadas basicamente para a proteção de uma mensagem contra os efeitos do ruído.

Por exemplo, podemos modificar o nosso código da seguinte forma:

$$\mathcal{C}_2 = \begin{cases} \text{Leste} \mapsto 000 \\ \text{Oeste} \mapsto 010 \\ \text{Norte} \mapsto 101 \\ \text{Sul} \mapsto 111 \end{cases}.$$

Nesta nova codificação, as duas primeiras posições reproduzem o código da fonte, enquanto na terceira posição é introduzida uma redundância. O novo código introduzido na recodificação é chamado de **código de canal**. Por exemplo, ao receber a mensagem 011, é possível perceber que a mensagem está errada, pois ela não aparece no código. Porém, não há maneira de saber se a mensagem correta era 010 ou 111. Logo, esse código detecta erros, mas não os corrige. Isso acontece porque as palavras do código ainda estão muito próximas. Isso pode ser corrigido com mais redundâncias.

Por exemplo:

$$\mathcal{C}_3 = \begin{cases} \text{Leste} \mapsto 00000 \\ \text{Oeste} \mapsto 01011 \\ \text{Norte} \mapsto 10110 \\ \text{Sul} \mapsto 11101 \end{cases}.$$

Agora, caso haja erro em uma letra, será possível não somente detectar, mas também corrigir erros. De fato, basta escolhermos no código a palavra que tem apenas uma letra diferente da recebida. Por exemplo, suponhamos que se tenha introduzido um erro ao transmitirmos a palavra 10110, de modo que a mensagem recebida seja 11110. Comparando essa mensagem com as demais palavras do código, notamos que essa mensagem não pertence ao código e, portanto, detectamos erro. A palavra do código mais próxima da referida mensagem (a que tem menor número de letras diferentes) é 10110, que é precisamente a palavra transmitida.

O procedimento de transmissão de informação pode ser esquematizado como na Figura 1.

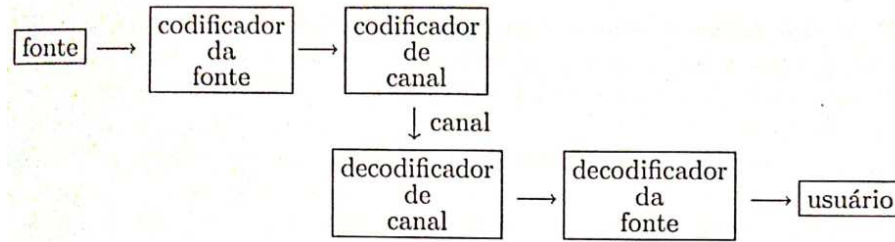
2.1 CÓDIGOS CORRETORES DE ERROS

Sejam $\mathbb{N} = \{1, 2, 3, \dots\}$ o conjunto dos números naturais e $\mathbb{K}$ um corpo finito com q elementos.

Agora, vamos definir formalmente os elementos necessários para a construção de um código corretor de erros.

Definição 2.1.1. Um conjunto finito $\mathcal{A} = \{a_1, a_2, \dots, a_q\}$, com $q \in \mathbb{N}$, é dito um **alfabeto**. Os elementos $a_i \in \mathcal{A}$, onde $i \in \{1, 2, \dots, q\}$, são chamados de **letras**. O **número de**

Figura 1 – Sistema de Comunicação Típico.



Fonte: [14]

elementos de $\mathcal{A}$, será denotado por $|\mathcal{A}|$ e simbolizado por q . A sequência $a_{k_1}a_{k_2}\cdots a_{k_j}$ ou $(a_{k_1}, a_{k_2}, \dots, a_{k_j})$, com $a_{k_i} \in \mathcal{A}$ para cada $i \in \{1, 2, \dots, j\}$, é denominada por **palavra** do código $\mathcal{A}$, ou simplesmente **palavra-código**. O **comprimento de uma palavra** é o número de letras que a forma. Além disso, $\mathcal{A}^n$ denota que $\mathcal{A}$ é um alfabeto cuja **maior palavra** tem comprimento n .

Agora, conhecidos os elementos de um código corretor de erros, definamos o que é tal código.

Definição 2.1.2. Um **código corretor de erros** $\mathcal{C}$ é um subconjunto próprio qualquer de $\mathcal{A}^n$, para algum $n \in \mathbb{N}$.

Definição 2.1.3. Dados dois elementos $u, v \in \mathcal{A}^n$ e $\mathcal{C} \subset \mathcal{A}^n$ um código, a **distância de Hamming** entre u e v é dada por

$$d(u, v) = \#\{i; u_i \neq v_i, 1 \leq i \leq n\},$$

onde $\#$ denota a quantidade de elementos do conjunto $\{i; u_i \neq v_i, 1 \leq i \leq n\}$. Definimos a **distância mínima** de $\mathcal{C}$ como o número

$$d = \min\{d(u, v); u, v \in \mathcal{C} \text{ e } u \neq v\}.$$

Definição 2.1.4. Dado $x \in \mathbb{K}^n$, definimos o **peso** de x como o número inteiro

$$\omega(x) = \#\{i; x_i \neq 0\},$$

isto é, a distância de x ao vetor todo nulo 0 de $\mathbb{K}^n$. O **peso mínimo** de um código $\mathcal{C}$ é dado por

$$\omega(\mathcal{C}) = \min\{\omega(x); x \in \mathcal{C} \setminus \{0\}\}.$$

Teorema 2.1.5. *Seja $\mathcal{C}$ um código com distância mínima d . Então $\mathcal{C}$ pode corrigir até $\kappa = \left\lfloor \frac{d-1}{2} \right\rfloor$ erros e detectar até $d-1$ erros, onde $\lfloor a \rfloor$ denota o maior número inteiro menor ou igual a a .*

Demonstração. Ver Teorema 3.12, capítulo 3 de [7]. $\square$

Exemplo 2.1.6. No código robô (Exemplo 2.0.1), como $d = 3$, então pelo Teorema 2.1.5, podemos corrigir até

$$\kappa = \left\lfloor \frac{d-1}{2} \right\rfloor = \left\lfloor \frac{3-1}{2} \right\rfloor = \left\lfloor \frac{2}{2} \right\rfloor = \lfloor 1 \rfloor = 1 \text{ erro}$$

e detectar até $d - 1 = 3 - 1 = 2$ erros.

2.1.1 Códigos lineares

A classe de códigos mais utilizada na prática é a dos códigos lineares, à qual definimos a seguir.

Definição 2.1.7. Um código $\mathcal{C} \subset \mathbb{F}_q^n$ é dito um **código linear** se for um subespaço vetorial de $\mathbb{F}_q^n$.

Proposição 2.1.8. *Seja $\mathcal{C} \subset \mathbb{F}_q^n$ um código linear com distância mínima d . Então, as seguintes afirmações são válidas:*

1. Para todos $x, y \in \mathbb{F}_q^n$ vale $d(x, y) = \omega(x - y)$;
2. $d = \omega(\mathcal{C})$.

Demonstração. Ver Proposição 4.5, capítulo 4 de [7]. $\square$

2.2 MATRIZ GERADORA DE UM CÓDIGO

Definição 2.2.1. Seja $\mathcal{C} \subset \mathbb{F}_q^n$ um código linear. Chamaremos de **parâmetros do código linear** $\mathcal{C}$ a terna $[n, k, d]_q$, onde:

- n é o comprimento do código;
- k é a **dimensão** de $\mathcal{C}$ como subespaço vetorial sobre $\mathbb{F}_q$;
- d é a distância mínima de $\mathcal{C}$, que é igual ao peso $\omega(\mathcal{C})$ do código $\mathcal{C}$.

Perceba que o número de elementos de $\mathcal{C}$ é igual a $M = q^k$.

Definição 2.2.2. Seja $\mathcal{B} = \{v_1, v_2, \dots, v_k\}$ uma base ordenada de $\mathcal{C}$ e considere a matriz G , cujas linhas são os vetores $v_i = (v_{i1}, v_{i2}, \dots, v_{in})$, $i \in \{1, 2, \dots, k\}$ da base, isto é,

$$G = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_k \end{bmatrix} = \begin{bmatrix} v_{11} & v_{12} & \cdots & v_{1n} \\ v_{21} & v_{22} & \cdots & v_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ v_{k1} & v_{k2} & \cdots & v_{kn} \end{bmatrix}_{k \times n}.$$

A matriz G é chamada de **matriz geradora** de $\mathcal{C}$ associada a base $\mathcal{B}$.

Exemplo 2.2.3. Considere $\mathcal{C}$ o código linear sobre $\mathbb{F}_2$ gerado pela matriz

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

Como G é 4×7 , pela forma que definimos matriz geradora (Definição 2.2.2) e pela Definição 2.2.1, podemos perceber que comprimento do código é $n = 7$, a dimensão é $k = 4$.

Vimos que as linhas da matriz G são formadas pelos vetores de uma base (fixada) do código $\mathcal{C}$. Assim, as palavras de $\mathcal{C}$ são dadas através das combinações lineares das linhas de G . As palavras de $\mathcal{C}$ e seus respectivos pesos são:

- $c_0 = (0, 0, 0, 0, 0, 0, 0)$
 $\Rightarrow \omega(c_0) = 0;$
- $c_1 = (1, 0, 0, 0, 1, 1, 1)$
 $\Rightarrow \omega(c_1) = 4;$
- $c_2 = (0, 1, 0, 0, 1, 1, 0)$
 $\Rightarrow \omega(c_2) = 3;$
- $c_3 = (0, 0, 1, 0, 1, 0, 1)$
 $\Rightarrow \omega(c_3) = 3;$
- $c_4 = (0, 0, 0, 1, 0, 1, 1)$
 $\Rightarrow \omega(c_4) = 3;$
- $c_5 = c_1 + c_2 = (0, 1, 1, 0, 0, 1, 1)$
 $\Rightarrow \omega(c_5) = 4;$
- $c_6 = c_1 + c_3 = (0, 0, 1, 1, 1, 1, 0);$
 $\Rightarrow \omega(c_6) = 4;$
- $c_7 = c_1 + c_4 = (1, 0, 0, 1, 1, 0, 0)$
 $\Rightarrow \omega(c_7) = 3;$
- $c_8 = c_2 + c_3 = (0, 0, 1, 1, 1, 1, 0)$
 $\Rightarrow \omega(c_8) = 4;$
- $c_9 = c_2 + c_4 = (0, 1, 0, 1, 1, 0, 1)$
 $\Rightarrow \omega(c_9) = 4;$
- $c_{10} = c_1 + c_2 + c_3 = (0, 1, 1, 1, 0, 0, 0)$
 $\Rightarrow \omega(c_{10}) = 3.$
- $c_{11} = c_1 + c_2 + c_4 = (0, 1, 1, 1, 0, 0, 0)$
 $\Rightarrow \omega(c_{11}) = 3;$
- $c_{12} = c_2 + c_3 + c_4 = (0, 0, 1, 0, 1, 0, 1)$
 $\Rightarrow \omega(c_{12}) = 3;$
- $c_{13} = c_1 + c_2 + c_3 + c_4 = (0, 1, 1, 0, 0, 1, 1)$
 $\Rightarrow \omega(c_{13}) = 4.$

Logo, a distância mínima de $\mathcal{C}$ é $d = 3$ (pois é o peso mínimo de uma palavra não nula do código). Portanto, $[7, 4, 3]_2$ são os parâmetros de $\mathcal{C}$.

Neste caso, se queremos, por exemplo, codificar $(1, 0, 0, 1)$, temos:

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

Ou seja, ao codificar a informação $(1, 0, 0, 1)$ deverá ser recebida a mensagem $(1, 0, 0, 1, 1, 0, 0)$, a menos de ruído.

Definição 2.2.4. Dizemos que uma matriz geradora G de um código linear $\mathcal{C}$ está na **forma padrão** quando ela se apresenta na forma

$$G = [Id_k | A],$$

em que Id_k é a matriz identidade $k \times k$ e A é uma matriz qualquer de ordem $k \times (n - k)$.

2.3 MATRIZ TESTE DE PARIDADE

Aqui, veremos uma outra maneira de definir um código linear, agora como a solução de um sistema linear de equações.

Definição 2.3.1. Uma **matriz teste de paridade** para um código linear $\mathcal{C} \subset \mathbb{F}_q^n$ é uma matriz H , $m \times n$, com entradas em $\mathbb{F}_q$ tal que

$$\mathcal{C} = \{u \in \mathbb{F}_q^n; uH^t = 0\},$$

onde H^t é a transposta da matriz H .

Proposição 2.3.2. *Seja $\mathcal{C}$ um código linear com matriz teste de paridade H . Se todo conjunto com $d-1$ colunas de H é linearmente independente e existem d colunas linearmente dependentes, então a distância mínima de $\mathcal{C}$ é d .*

Demonstração. Ver Teorema 4.17, capítulo 4 de [7]. □

Exemplo 2.3.3. Seja $\mathbb{F}_2$ o alfabeto. Consideremos o código $\mathcal{C}$ do Exemplo 2.2.3, gerado pela matriz

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

O conjunto

$$\mathcal{B} = \{1000111, 0100110, 0010101, 0001011\}$$

é uma base de $\mathcal{C}$. Ou seja, dado $v \in \mathcal{C}$, v é escrito de forma única como

$$v = (x_1, x_2, x_3, x_4, x_1 + x_2 + x_3, x_1 + x_2 + x_4, x_1 + x_3 + x_4),$$

onde $x_i \in \mathbb{F}_2$, $i \in \{1, 2, 3, 4\}$. Vamos mostrar que a matriz

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

é a matriz teste de paridade para o código $\mathcal{C}$.

De fato, seja $u \in \mathbb{F}_2^7$. Temos:

$$uH^t = 0 \Leftrightarrow \begin{bmatrix} u_1 & u_2 & u_3 & u_4 & u_5 & u_6 & u_7 \end{bmatrix} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 0 \\ 1 & 0 & 1 \\ 0 & 1 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \end{bmatrix}$$

$$\Leftrightarrow \begin{bmatrix} u_1 + u_2 + u_3 + u_5 & u_1 + u_2 + u_4 + u_6 & u_1 + u_3 + u_4 + u_7 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \end{bmatrix}$$

$$\Leftrightarrow \begin{cases} u_1 + u_2 + u_3 + u_5 = 0 \\ u_1 + u_2 + u_4 + u_6 = 0 \\ u_1 + u_3 + u_4 + u_7 = 0 \end{cases} \Leftrightarrow \begin{cases} u_1 = u_1 \\ u_2 = u_2 \\ u_3 = u_3 \\ u_4 = u_4 \\ u_5 = u_1 + u_2 + u_3 \\ u_6 = u_1 + u_2 + u_4 \\ u_7 = u_1 + u_3 + u_4 \end{cases}$$

Ou seja, $u = (u_1, u_2, u_3, u_4, u_1 + u_2 + u_3, u_1 + u_2 + u_4, u_1 + u_3 + u_4)$. Portanto, $\mathcal{C} = \{u \in \mathbb{F}_2^7; uH^t = 0\}$, isto é, H é a matriz teste de paridade para $\mathcal{C}$.

Olhando para a matriz H , vemos que suas colunas são linearmente independentes duas a duas e que há um conjunto de exatamente três colunas linearmente dependentes (por exemplo, a primeira coluna é a soma da segunda coluna e da sétima coluna). Logo, pela Proposição 2.3.2 a distância mínima é 3.

Proposição 2.3.4. *Seja G uma matriz geradora de um código linear $\mathcal{C}$ de dimensão k . Uma matriz H , $m \times n$, é uma matriz teste de paridade para $\mathcal{C}$ se, e somente se, $GH^t = 0$ e o posto de H é $n - k$.*

Demonstração. Ver Proposição 4.19, capítulo 4 de [7]. □

Proposição 2.3.5. *Seja $\mathcal{C} \subset \mathbb{F}_q^n$ um código de dimensão k com matriz geradora $G = [Id_k | A]$ na forma padrão. Então, $H = [-A^t | Id_{n-k}]$ é uma matriz teste de paridade para $\mathcal{C}$.*

Demonstração. Ver Proposição 4.21, capítulo 4 de [7]. □

Exemplo 2.3.6. Considere o código linear $\mathcal{C}$ sobre $\mathbb{F}_5 = \mathbb{Z}_5$ gerado pela matriz

$$G = \begin{bmatrix} 1 & 0 & 0 & 1 & 1 & 2 \\ 0 & 1 & 0 & 1 & 2 & 1 \\ 0 & 0 & 1 & 2 & 1 & 1 \end{bmatrix}.$$

Perceba que G está na forma padrão, $n = 6$ e $k = 3$. Além disso, como estamos trabalhando em $\mathbb{Z}_5$

$$A = \begin{bmatrix} 1 & 1 & 2 \\ 1 & 2 & 1 \\ 2 & 1 & 1 \end{bmatrix} \Rightarrow -A^t = \begin{bmatrix} -1 & -1 & -2 \\ -1 & -2 & -1 \\ -2 & -1 & -1 \end{bmatrix} = \begin{bmatrix} 4 & 4 & 3 \\ 4 & 3 & 4 \\ 3 & 4 & 4 \end{bmatrix}$$

Logo, pela Proposição 2.3.5,

$$H = [-A^t | Id_{n-k}] = \begin{bmatrix} 4 & 4 & 3 & 1 & 0 & 0 \\ 4 & 3 & 4 & 0 & 1 & 0 \\ 3 & 4 & 4 & 0 & 0 & 1 \end{bmatrix}$$

é uma matriz teste de paridade para $\mathcal{C}$.

2.4 DECODIFICAÇÃO POR SÍNDROME

Dada uma matriz geradora G de um código linear $\mathcal{C}$, e dado $v \in \mathcal{C}$, vimos que podemos codificar este vetor da seguinte forma:

$$v \mapsto vG.$$

Além disso, se a matriz geradora está na forma padrão $G = [Id_k | A]$, podemos codificar anexando as $n - k$ coordenadas de vA a v .

Agora, para fazer o caminho inverso, isto é, para decodificar uma mensagem, precisamos de um trabalho que é um pouco mais árduo. Para decodificar escolhendo o vetor mais próximo no código, devemos encontrar uma palavra do código de comprimento n que está mais próxima a n -upla recebida. Para um código sem estrutura óbvia, isso só pode ser feito calculando a distância entre cada palavra do código e a n -upla recebida, o que pode ser bem trabalhoso.

Vamos ver então um algoritmo de decodificação que explora a linearidade dos códigos lineares.

Definição 2.4.1. Seja $\mathcal{C}$ um código linear com matriz teste de paridade H . Dado um vetor $v \in \mathbb{F}_q^n$, definimos a **síndrome** de v como sendo o vetor $s(v) = vH^t$.

Quando uma mensagem c é transmitida e um vetor v é recebido, a diferença entre os dois vetores é chamada **erro** e denotada por e . Dessa forma:

$$v = c + e.$$

Se H é uma matriz teste de paridade para o código linear $\mathcal{C}$, então, como $\mathcal{C} = \{v \in \mathbb{F}_q^n; vH^t = 0\}$, temos:

$$\begin{aligned} s(v) &= vH^t = (c + e)H^t = cH^t + eH^t \\ &= 0 + eH^t, \quad \text{pois } c \in \mathcal{C} \\ &= eH^t \\ &= s(e), \end{aligned}$$

ou seja, a síndrome de v é a mesma do vetor erro. Ademais,

$$s(v) = 0 \Leftrightarrow vH^t = 0 \Leftrightarrow v \in \mathcal{C}.$$

Se $\omega(e) \leq 1$, então a síndrome $s(v) = s(e) = eH^t$ é só um múltiplo escalar de uma coluna de H .

Isso nos dá um algoritmo de decodificação simples para códigos lineares que corrijam 1 erro. Primeiro, calculamos a síndrome de v . Assim, temos dois casos:

- $s(v) = 0$.

Neste caso, $v \in \mathcal{C}$, e então não houve erro na transmissão.

- $s(v) \neq 0$.

Neste caso, procure a coluna de H que é um múltiplo escalar de $s(v)$. Se essa coluna não existir, então houve mais de um erro e o código não será capaz de corrigir. Agora, se $s(v) = \lambda \cdot \text{coluna } j$, para algum $\lambda \in \mathbb{N}$, então, para corrigir o erro, basta somar a v o vetor com $-\lambda$ na j -ésima posição e 0 nas demais posições para corrigir o erro.

2.5 CÓDIGOS DUAIS

Definição 2.5.1. Seja $\mathcal{C} \subset \mathbb{F}_q^n$ um código linear. Definimos

$$\mathcal{C}^\perp = \{v \in \mathbb{F}_q^n; \langle u, v \rangle = 0, \text{ para todo } u \in \mathcal{C}\},$$

em que $\langle \cdot, \cdot \rangle$ significa o produto interno.

Noutros termos, $\mathcal{C}^\perp$ é o subespaço ortogonal a $\mathcal{C}$ com respeito ao produto interno.

O subespaço $\mathcal{C}^\perp$ é o conjunto de soluções de um sistema linear de equações de posto k e n indeterminadas. Então, o código dual $\mathcal{C}^\perp$ é um código linear de dimensão $n - k$ e comprimento n sobre $\mathbb{F}_q$.

Proposição 2.5.2. *Se H é uma matriz $(n - k) \times n$ que é uma matriz teste de paridade para um código linear $\mathcal{C}$ de dimensão k , então H^t é uma matriz geradora para $\mathcal{C}^\perp$. Ainda, se G é uma matriz geradora para $\mathcal{C}$, então G^t é uma matriz teste de paridade para $\mathcal{C}^\perp$.*

Demonstração. Ver Proposição 4.19, capítulo 4 de [7]. $\square$

Definição 2.5.3. Dizemos que um código linear $\mathcal{C}$ é **auto-dual** quando $\mathcal{C} = \mathcal{C}^\perp$.

Seja $\mathcal{C}$ um código um código auto-dual. Sejam G e H , respectivamente, as matrizes geradora e teste de paridade de $\mathcal{C}$. Então, pela Proposição 2.3.4, vale $GH^t = 0$. Sabemos, pela Proposição 2.5.2, que H é uma matriz geradora de $\mathcal{C}^\perp$ e G é uma matriz teste de paridade para $\mathcal{C}^\perp$, donde, pela Proposição 2.3.4, $HG^t = 0$. Agora, sendo $\mathcal{C}$ auto-dual, $\mathcal{C} = \mathcal{C}^\perp$ e assim G é matriz teste de paridade para $\mathcal{C}$. A proposição nos dá uma condição para que um código linear seja auto-dual.

Proposição 2.5.4. *Um código linear $[n, k, d]_q$ é auto-dual se, e somente se, sua matriz geradora G é tal que $GG^t = 0$ e $\text{posto}(G) = n - k$.*

Exemplo 2.5.5. Consideremos $\mathcal{C}$ o $[4, 2, d]_7$ código linear gerado pela matriz

$$G = \begin{bmatrix} 1 & 0 & 2 & 3 \\ 0 & 1 & 3 & 5 \end{bmatrix}.$$

Então, $\mathcal{C}$ é um código auto-dual.

De fato, como $\mathcal{C} \subset \mathbb{F}_7^4$, devemos utilizar as operações módulo 7, assim temos:

$$GG^t = \begin{bmatrix} 1 & 0 & 2 & 3 \\ 0 & 1 & 3 & 5 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 2 & 3 \\ 3 & 5 \end{bmatrix} = \begin{bmatrix} 14 & 21 \\ 21 & 35 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = 0.$$

Ademais, como posto de G é 2, segue da Proposição 2.3.4, que G é também uma matriz teste de paridade para G . Logo, $\mathcal{C}^\perp$ é gerado por G e, então, $\mathcal{C} = \mathcal{C}^\perp$.

Sendo $\mathcal{C}$ um código auto-dual de dimensão k , G e H suas matrizes geradora e teste de paridade, respectivamente, então $G = H^t$ ou $H = G^t$.

Observe que um código auto-dual tem parâmetros $[n, n/2, d]_q$.

3 INFORMAÇÃO E COMPUTAÇÃO QUÂNTICA

A Mecânica Quântica é a formalização matemática da Física Quântica, desenvolvida para descrever o comportamento de sistemas em escalas microscópicas, como átomos, partículas subatômicas e fótons. Seus princípios — como superposição, emaranhamento e a natureza probabilística das medições — desafiam as noções clássicas de realidade e provocaram uma profunda reformulação na compreensão dos fenômenos físicos. Ao longo do século XX, os fundamentos da Mecânica Quântica revolucionaram o entendimento da natureza sob a ótica da física clássica.

A Computação Quântica, por sua vez, baseia-se nessas propriedades fundamentais para propor um novo paradigma de processamento de informações, distinto e potencialmente mais poderoso que o modelo clássico. A partir desses fundamentos, surgiu a área de Informação e Computação Quântica, que explora como as propriedades quânticas podem ser utilizadas para representar, processar e transmitir informações de forma mais eficiente do que os métodos tradicionais permitem. Esse campo propõe uma nova forma de computação, com o potencial de resolver certos problemas complexos de maneira exponencialmente mais rápida, além de oferecer avanços em segurança da informação e comunicação.

Trata-se de um tema em pleno desenvolvimento, com grande interesse da comunidade científica e tecnológica. As pesquisas em Informação e Computação Quântica vêm crescendo de forma acelerada, impulsionadas por desafios teóricos profundos e por avanços experimentais significativos. Esta dissertação se insere nesse contexto dinâmico, buscando contribuir para a compreensão e o avanço desse fascinante e promissor ramo da ciência moderna. Para mais informações sobre o assunto, recomendamos a referência [20].

3.1 NOTAÇÃO DE DIRAC

Cabe destacarmos que, no Anexo, são apresentadas algumas teorias relevantes que podem servir de apoio à compreensão dos conceitos que exploraremos, especialmente no contexto da Álgebra Linear.

Apresentaremos nessa seção uma notação bem incomum para matemáticos, mas eficiente para os físicos na denotação de estados quânticos - *A notação de Dirac*. Para isso, tomamos como fonte base [24].

Antes de iniciar o estudo sobre os códigos quânticos, é necessário compreender o significado da notação “ $|\rangle$ ”. Qual é sua leitura? Como deve ser pronunciada? Nesta seção, será apresentada a interpretação e o uso dessa notação, fundamental para a descrição de estados quânticos no formalismo da Mecânica Quântica.

Na física quântica, o estado físico de um sistema é dado por todas as informações

possíveis de obter-se de tal sistema, podendo ser representado por uma função complexa de onda ou por um vetor de estado contido num espaço vetorial complexo. Porém, prova-se que a representação de espaços vetoriais em três dimensões de alguns elementos, como os spins, não é possível.

Paul Dirac, pioneiro da Física e Mecânica Quântica, foi responsável por introduzir a representação dos estados quânticos por meio de espaços vetoriais complexos, conforme apresentado em sua obra clássica [9]. Nessa formalização, tais vetores passaram a ser denominados por **kets**, os quais representam, de forma vetorial, um estado físico da Mecânica Quântica, contendo todas as informações associadas a esse estado. Os kets são simbolicamente representados por “ $| \rangle$ ” e possuem um elemento dual correspondente, chamado **bra**, representado por “ $\langle |$ ”. O produto interno entre um *bra* e um *ket* é denotado por “ $\langle | \rangle$ ”, sendo denominado **braket**. Essa notação, introduzida por Dirac, ficou amplamente conhecida como **notação bra-ket** ou **notação de Dirac**.

3.2 POSTULADOS DA MECÂNICA QUÂNTICA E MEDIÇÕES

Dessa seção em diante, até o final desse capítulo, tomaremos como base [1] e [5].

O objetivo desta seção será apresentarmos brevemente os postulados da mecânica quântica.

Definição 3.2.1. Um **espaço de Hilbert** é um espaço vetorial complexo com produto interno.

Em essência, a mecânica quântica é uma estrutura matemática para o desenvolvimento de uma teoria física. A conexão entre o mundo físico e o formalismo da mecânica quântica é dada pelos *Postulados da Mecânica Quântica*. Apresentaremos esses postulados, que estão presentes em [1], como forma de conhecimento, em [20] estão disponíveis mais informações

- **Postulado 1:** *A qualquer sistema físico isolado existe um **espaço de Hilbert** associado, conhecido como **espaço de estados** do sistema. O estado de sistema é totalmente descrito pelo seu **vetor de estado**, um vetor unitário no espaço de estados.*
- **Postulado 2:** *A evolução de um sistema quântico fechado é descrita por uma transformação unitária. Ou seja, o estado $|\psi\rangle$ de um sistema em um tempo t_1 está relacionado ao estado $|\psi'\rangle$ do sistema no tempo t_2 por um operador unitário U que depende somente de t_1 e t_2 . Simbolicamente, temos:*

$$|\psi'\rangle = U|\psi\rangle.$$

- **Postulado 2':** A evolução temporal do estado de um sistema quântico fechado é descrita pela equação de Schrödinger:

$$i\hbar \frac{d|\psi\rangle}{dt} = H|\psi\rangle$$

Nessa equação, $\hbar$ é uma constante física chamada **constante de Planck**, cujo valor é determinado experimentalmente e H é um operador Hermitiano conhecido como **Hamiltoniano** do sistema. Na prática, é comum absorver o fator $\hbar$ dentro de H , efetivamente convencendo-se $\hbar = 1$.

- **Postulado 3:** As medidas quânticas $|0\rangle$ e $|1\rangle$ são descritas por determinados operadores de medida $\{M_m\}$. Esses operadores atuam sobre o espaço de estados do sistema. O índice m refere-se aos possíveis resultados da medida.
- **Postulado 4:** O espaço de estados de um sistema físico composto é o produto tensorial dos espaços de estados dos sistemas individuais.

3.3 A UNIDADE DE INFORMAÇÃO QUÂNTICA

A unidade fundamental da informação clássica é o *bit*, que pode assumir dois estados; a saber, 0 ou 1. Chamamos de **qubit** (ou **bit quântico**) a unidade de informação quântica.

Definição 3.3.1. Definimos o **estado** de um qubit como um vetor em um espaço vetorial complexo bidimensional $\mathbb{C}^2$, que é um espaço de Hilbert.

Diferentemente do bit clássico, um qubit pode assumir, além dos dois estados $|0\rangle$ e $|1\rangle$, todos os outros estados correspondentes às combinações lineares de $|0\rangle$ e $|1\rangle$, as quais chamamos **superposições**, do seguinte modo:

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle, \quad (3.1)$$

em que $\alpha, \beta \in \mathbb{C}$. Neste caso, dizemos que o qubit $|\psi\rangle$ está em uma **superposição de dois estados**.

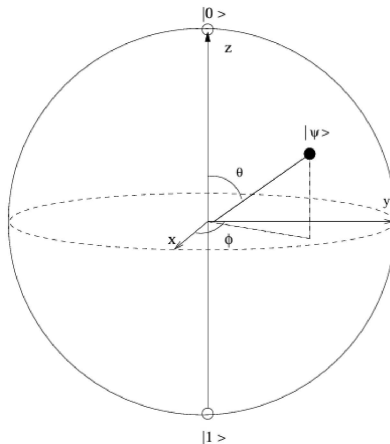
Os estados $|0\rangle$ e $|1\rangle$ formam a **Base Computacional**, que é uma base ortonormal para o **espaço de estados**. Assim, em particular, $|\psi\rangle$ é unitário, isto é, $\langle\psi|\psi\rangle = 1$, o que equivale à $|\alpha|^2 + |\beta|^2 = 1$. Valendo $|\alpha|^2 + |\beta|^2 = 1$, quando medirmos um qubit a probabilidade de obtermos $|0\rangle$ como resultado será $|\alpha|^2$ e a probabilidade de obtermos $|1\rangle$ será de $|\beta|^2$.

Reescrevendo a expressão 3.1 em função de dois parâmetros θ e ϕ , obtemos

$$|\psi\rangle = \cos\left(\frac{\theta}{2}\right)|0\rangle + e^{i\phi}\sin\left(\frac{\theta}{2}\right)|1\rangle,$$

tal representação permitiu uma representação geométrica para o qubit. A partir dos parâmetros θ e ϕ podemos definir uma esfera de raio unitário, denominada **Esfera de Bloch**, como na seguinte figura:

Figura 2 – Esfera de Bloch.



Fonte: [1]

Nos polos da esfera encontram-se os estados $|0\rangle$ e $|1\rangle$, e qualquer combinação linear que represente um estado em superposição, estará representada univocamente em um ponto na casca desta esfera.

Um sistema de 2 qubits, os **estados na base computacional** são $|00\rangle$, $|01\rangle$, $|10\rangle$ e $|11\rangle$, os quais pertencem ao espaço de Hilbert $\mathbb{C}^2 \otimes \mathbb{C}^2 = \mathbb{C}^4$. A superposição desses quatro estados é dada por

$$|\psi\rangle = \alpha_{00}|00\rangle + \alpha_{01}|01\rangle + \alpha_{10}|10\rangle + \alpha_{11}|11\rangle.$$

A probabilidade da medição resultar em $|ij\rangle$, $i, j = \{0, 1\}$, será de $|\alpha_{ij}|^2$.

Mais geralmente, um estado φ com n qubits é uma superposição dos 2^n estados

$$|00 \cdots 0\rangle, |00 \cdots 1\rangle, \dots, |11 \cdots 1\rangle,$$

onde a sequência dentro de cada *ket* corresponde à representação binária dos números $0, 1, \dots, 2^n - 1$. Esse estado pode ser expresso como

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

sujeito à condição

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

3.4 EMARANHAMENTO QUÂNTICO

Nesta seção, definiremos formalmente um emaranhamento quântico.

Um estado de dois qubits pode ou não ser resultado do produto tensorial de dois estados de um qubit. Consideremos os estados de um qubit $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ e $|\varphi\rangle = \gamma|0\rangle + \delta|1\rangle$, com $\alpha, \beta, \gamma, \delta \in \mathbb{C}$. O estado definido pelo produto tensorial de $|\psi\rangle$ e $|\varphi\rangle$ é:

$$\begin{aligned}
 |\psi\rangle \otimes |\varphi\rangle &= (\alpha|0\rangle + \beta|1\rangle) \otimes (\gamma|0\rangle + \delta|1\rangle) \\
 &= (\alpha|0\rangle \otimes \gamma|0\rangle) + (\alpha|0\rangle \otimes \delta|1\rangle) + (\beta|1\rangle \otimes \gamma|0\rangle) + (\beta|1\rangle \otimes \delta|1\rangle) \\
 &= \alpha\gamma(|0\rangle \otimes |0\rangle) + \alpha\delta(|0\rangle \otimes |1\rangle) + \beta\gamma(|1\rangle \otimes |0\rangle) + \beta\delta(|1\rangle \otimes |1\rangle) \\
 &= \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle
 \end{aligned} \tag{3.2}$$

Sendo que, na última linha apenas reescrevemos a expressão da linha anterior utilizando a notação $|v\rangle \otimes |w\rangle = |vw\rangle$.

Um estado de dois qubits genérico

$$a|00\rangle + b|01\rangle + c|10\rangle + d|11\rangle$$

é da forma (3.2) se, e somente se, vale

$$\begin{cases} a = \alpha\gamma, \\ b = \alpha\delta, \\ c = \beta\gamma, \\ d = \beta\delta, \end{cases}$$

o que implica que

$$\frac{b}{d} = \frac{\alpha}{\beta} = \frac{a}{c} \Rightarrow ad = bc.$$

Dessa forma, um estado de dois qubits, em geral, não pode ser expresso como um produto tensorial de estados individuais de um qubit. Com isso, surge a seguinte definição:

Definição 3.4.1. Dizemos que um estado de dois qubits é **emaranhado** quando ele não é produto tensorial de estados de um qubit.

Exemplo 3.4.2. Conforme disponível no anexo, da definição de produto tensorial, o estado $|10\rangle$ pode ser escrito como o produto tensorial dos estados $|1\rangle$ e $|0\rangle$,

$$|1\rangle \otimes |0\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix} = |10\rangle$$

Dessa forma, o estado $|10\rangle$ não é um estado emaranhado.

Por outro lado, o estado

$$|v\rangle = \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix}$$

é um estado emaranhado. De fato, vamos mostrar que $|v\rangle$ não pode ser escrito como produto tensorial de estados de um qubit. Supondo o contrário, consideremos os estados de um qubit $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ e $|\varphi\rangle = \gamma|0\rangle + \delta|1\rangle$, com $\alpha, \beta, \gamma, \delta \in \mathbb{C}$, de modo que $|v\rangle = |\psi\rangle \otimes |\varphi\rangle$. Deveríamos ter

$$\begin{aligned} \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix} &= \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle \\ \Rightarrow \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix} &= \alpha\gamma \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) + \alpha\delta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) + \beta\gamma \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) + \beta\delta \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) \\ &\Rightarrow \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix} = \alpha\gamma \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix} + \alpha\delta \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix} + \beta\gamma \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix} + \beta\delta \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix} \\ &\Rightarrow \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix} = \begin{bmatrix} \alpha\gamma \\ \alpha\delta \\ \beta\gamma \\ \beta\delta \end{bmatrix} \Rightarrow \begin{cases} \alpha\gamma = 0 \\ \alpha\delta = 1 \\ \beta\gamma = 1 \\ \beta\delta = 0 \end{cases} \Rightarrow \begin{cases} \alpha = 0 \text{ ou } \gamma = 0 \\ \alpha\delta = 1 \\ \beta\gamma = 1 \\ \beta = 0 \text{ ou } \delta = 0 \end{cases}. \end{aligned}$$

Mas, como $\alpha\delta = \beta\gamma = 1$, deveríamos ter $\alpha, \beta, \gamma, \delta \neq 0$, contradizendo $\alpha\gamma = 0$ e $\beta\delta = 0$. Portanto, o estado $|v\rangle$ é emaranhado.

Os estados emaranhados de dois qubits mais conhecidos são:

$$\beta_{|00\rangle} = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle), \quad (3.3)$$

$$\beta_{|10\rangle} = \frac{1}{\sqrt{2}}(|00\rangle - |11\rangle), \quad (3.4)$$

$$\beta_{|01\rangle} = \frac{1}{\sqrt{2}}(|01\rangle + |10\rangle), \quad (3.5)$$

$$\beta_{|11\rangle} = \frac{1}{\sqrt{2}}(|01\rangle - |10\rangle). \quad (3.6)$$

Esses estados são conhecidos como *estados de Bell* ou *pares EPR* (Einstein - Podolsky - Rosen). Eles desempenham um papel fundamental em diversas áreas da informação quântica, atuando como unidade básica de emaranhamento em muitas aplicações de criptografia quântica. Além disso, são amplamente utilizados em outros contextos, como na teleportação quântica, bem explorada em [3].

Exemplo 3.4.3. Vamos mostrar que os estados de Bell realmente são emaranhados. Para isso, sejam $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ e $|\varphi\rangle = \gamma|0\rangle + \delta|1\rangle$, $\alpha, \beta, \gamma, \delta \in \mathbb{C}$.

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

Suponhamos que existam estados de um único qubit $|\psi\rangle$ e $|\varphi\rangle$ tais que $\beta_{|00\rangle} = |\psi\rangle \otimes |\varphi\rangle$.

Assim:

$$\beta_{|00\rangle} = \frac{|00\rangle + |11\rangle}{\sqrt{2}} \stackrel{(3.2)}{=} \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle,$$

donde $\alpha\delta = 0$, o que implicaria que $\alpha = 0$ ou $\delta = 0$, contradizendo o fato de que $\alpha\gamma = \frac{1}{\sqrt{2}} = \beta\delta$.

2. $\beta_{|10\rangle} = \frac{1}{\sqrt{2}}(|00\rangle - |11\rangle)$:

Suponhamos que existam estados de um único qubit $|\psi\rangle$ e $|\varphi\rangle$ tais que $\beta_{|10\rangle} = |\psi\rangle \otimes |\varphi\rangle$.

Assim:

$$\beta_{|10\rangle} = \frac{|00\rangle - |11\rangle}{\sqrt{2}} \stackrel{(3.2)}{=} \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle,$$

donde $\alpha\delta = 0$, o que implicaria que $\alpha = 0$ ou $\delta = 0$, contradizendo o fato de que $\alpha\gamma = \frac{1}{\sqrt{2}}$ e $\beta\delta = -\frac{1}{\sqrt{2}}$.

3. $\beta_{|01\rangle} = \frac{1}{\sqrt{2}}(|01\rangle + |10\rangle)$:

Suponhamos que existam estados de um único qubit $|\psi\rangle$ e $|\varphi\rangle$ tais que $\beta_{|01\rangle} = |\psi\rangle \otimes |\varphi\rangle$.

Assim:

$$\beta_{|01\rangle} = \frac{|01\rangle + |10\rangle}{\sqrt{2}} \stackrel{(3.2)}{=} \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle,$$

donde $\alpha\gamma = 0$, o que implicaria que $\alpha = 0$ ou $\gamma = 0$, contradizendo o fato de que $\alpha\delta = \frac{1}{\sqrt{2}} = \beta\gamma$.

4. $\beta_{|11\rangle} = \frac{1}{\sqrt{2}}(|01\rangle - |10\rangle)$:

Suponhamos que existam estados de um único qubit $|\psi\rangle$ e $|\varphi\rangle$ tais que $\beta_{|11\rangle} = |\psi\rangle \otimes |\varphi\rangle$.

Assim:

$$\beta_{|11\rangle} = \frac{|01\rangle - |10\rangle}{\sqrt{2}} \stackrel{(3.2)}{=} \alpha\gamma|00\rangle + \alpha\delta|01\rangle + \beta\gamma|10\rangle + \beta\delta|11\rangle,$$

donde $\alpha\gamma = 0$, o que implicaria que $\alpha = 0$ ou $\gamma = 0$, contradizendo o fato de que $\alpha\delta = \frac{1}{\sqrt{2}}$ e $\beta\gamma = -\frac{1}{\sqrt{2}}$.

Exemplo 3.4.4. O estado

$$|\xi\rangle = \begin{bmatrix} \frac{1}{2} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \frac{1}{2} \end{bmatrix}$$

é um estado emaranhado, enquanto o estado

$$|\zeta\rangle = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

é não emaranhado porque pode ser decomposto em

$$|\zeta\rangle = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} \otimes \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 0 & 1 \end{bmatrix}.$$

Exemplo 3.4.5. Para n qubits, um estado emaranhado, conhecido como o **estado GHZ**, é definido por

$$\frac{1}{\sqrt{2}}(|00 \cdots 0\rangle + |11 \cdots 1\rangle).$$

3.5 PORTAS QUÂNTICAS

Um computador quântico é construído a partir de um circuito quântico que inclui fios e portas lógicas quânticas, responsáveis por transportar e manipular a informação quântica. Essa construção é semelhante à dos circuitos de computadores clássicos, que também possuem fios e portas lógicas para transportar e processar a informação.

A computação quântica começa com a preparação de qubits iniciais que contêm os dados de entrada do problema. As portas lógicas quânticas operam sequencialmente nos qubits conforme o algoritmo, modificando as amplitudes que representam a superposição. Após a aplicação dessas portas, o resultado da computação é obtido por meio de medições em qubits previamente selecionados.

Para um qubit clássico, a única porta não trivial é a **NOT**, que inverte o estado, transformando 0 em 1 e 1 em 0. Para múltiplos bits clássicos, existem outras portas lógicas conhecidas como **AND**, **OR**, **NAND** e **NOR**, cujas ações são dadas na tabela abaixo.

bit de entrada ab	bit de saída			
	AND	OR	NAND	NOR
00	0	0	1	1
01	0	1	1	0
10	0	1	1	0
11	1	1	0	0

Tabela 1 – Ações das portas lógicas clássicas AND, OR, NAND e NOR. Fonte: [1].

Definição 3.5.1. Dizemos que uma porta lógica é **universal** quando ela pode ser utilizada para implementar qualquer função lógica (através da composição), sem a necessidade de utilizarmos outras portas.

Observação 3.5.2. O termo universal é utilizado na Definição 3.5.1 justamente pela versatilidade dessas portas, já que através dessas portas é possível criar qualquer outro tipo de porta lógica.

As portas universais clássicas são formadas pela combinação das portas AND, OR e NOT. A partir dessas combinações, é possível obter as portas NAND e NOR.

A porta lógica NAND é construída a partir das portas NOT e AND. Sua operação consiste essencialmente na multiplicação dos bits de entrada, seguida da negação do resultado. Essa operação é dada da seguinte forma:

$$a \text{ AND } b = a \cdot b = c.$$

Aplicando sob o bit de saída c da porta NOT, teremos então como saída $\bar{c}$, completando assim ação da porta NAND que é dada por $\text{NAND} = \text{NOT}(\text{AND})$.

Por outro lado, a porta NOR é derivada da combinação das portas OR e NOT. Se considerarmos a e b como as duas entradas da porta OR e c como a saída, a operação lógica da porta OR é dada por:

$$a \text{ OR } b = a \oplus b = c,$$

onde $\oplus$ denota a soma módulo 2. Aqui, análogo ao caso da porta AND, basta aplicar a operação NOT à saída da OR, resultando na seguinte expressão lógica:

$$a \text{ NOR } b = \text{NOT}(a \text{ AND } b) = \text{NOT}\{a \oplus b = c\} = \bar{c}.$$

Agora, vejamos como se comportam as portas lógicas no caso quântico, começando com o seguinte exemplo:

Exemplo 3.5.3. No caso quântico, também temos uma porta equivalente à porta NOT. Essa porta é dada pela matriz de Pauli X , denominada por **bit flip**. Temos:

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

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

Logo, o *bit flip* leva $|0\rangle$ em $|1\rangle$ e $|1\rangle$ em $|0\rangle$. Além disso, a ação de X no estado de superposição $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, $\alpha, \beta \in \mathbb{C}$, é linear, de modo que,

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

ou matricialmente

$$X \begin{bmatrix} \alpha \\ \beta \end{bmatrix} = \begin{bmatrix} \beta \\ \alpha \end{bmatrix}.$$

Uma maneira de descrever as portas quânticas sobre um qubit, se dá através de matrizes 2×2 . Para que isso aconteça, é importante lembrarmos que para um estado quântico qualquer $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, a condição $|\alpha|^2 + |\beta|^2 = 1$ também deve ser satisfeita pelo estado $|\psi'\rangle = \alpha'|0\rangle + \beta'|1\rangle$ após a aplicação da porta lógica, isto é, $|\alpha'|^2 + |\beta'|^2 = 1$, de modo que a matriz que representa a porta lógica deve ser unitária. Essa é a única condição imposta sobre as portas quânticas. Dessa forma, *qualquer matriz unitária especifica uma porta lógica quântica*.

Exemplo 3.5.4. Denotemos por $A^\dagger$ a matriz complexa conjugada de A . Vamos mostrar que as matrizes de Pauli são unitárias. De fato:

$$I^\dagger I = II = II^\dagger = I$$

$$\begin{aligned} X^\dagger X &= \begin{bmatrix} 0^* & 1^* \\ 1^* & 0^* \end{bmatrix}^T \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \\ &= \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \\ &= I. \end{aligned} \qquad \begin{aligned} XX^\dagger &= \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 0^* & 1^* \\ 1^* & 0^* \end{bmatrix}^T \\ &= \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \\ &= I. \end{aligned}$$

$$\begin{aligned} Y^\dagger Y &= \begin{bmatrix} 0^* & (-i)^* \\ i^* & 0^* \end{bmatrix}^T \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \\ &= \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \\ &= I \end{aligned} \qquad \begin{aligned} YY^\dagger &= \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \begin{bmatrix} 0^* & (-i)^* \\ i^* & 0^* \end{bmatrix}^T \\ &= \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \\ &= I. \end{aligned}$$

$$\begin{aligned} Z^\dagger Z &= \begin{bmatrix} 1^* & 0^* \\ 0^* & (-1)^* \end{bmatrix}^T \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \\ &= \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \\ &= I \end{aligned} \qquad \begin{aligned} ZZ^\dagger &= \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 1^* & 0^* \\ 0^* & (-1)^* \end{bmatrix}^T \\ &= \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \\ &= I. \end{aligned}$$

Logo, as matrizes de Pauli são unitárias e assim, são portas lógicas quânticas.

Exemplo 3.5.5. A porta Z , conhecida como *operador phase flip*, é tal que

$$Z|0\rangle = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix} = |0\rangle,$$

$$Z|1\rangle = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \end{bmatrix} = -|1\rangle,$$

ou seja, o operador *phase flip* age trocando o sinal do estado $|1\rangle$ e preservando o estado $|0\rangle$.

Exemplo 3.5.6. Uma outra porta muito utilizada é conhecida como *porta de Hadamard*, onde sua representação matricial é dada por

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

Perceba que

$$H^\dagger H = \left(\frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \right) \cdot \left(\frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}^2 = \frac{1}{2} \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix} = I. \quad (3.7)$$

Mais ainda, H é hermitiano¹

O exemplo a seguir mostra como a porta de Hadamard age sobre os qubits $|0\rangle$ e $|1\rangle$.

Exemplo 3.5.7. Temos:

$$\begin{aligned} H|0\rangle &= \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \frac{1}{\sqrt{2}} \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \\ H|1\rangle &= \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix} = \frac{1}{\sqrt{2}} \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} - \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle). \end{aligned}$$

Denotemos

$$H|0\rangle = |+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \quad \text{e} \quad H|1\rangle = |-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) \quad (3.8)$$

Os estados $|+\rangle$ e $|-\rangle$ são superposições emaranhados.

Vamos mostrar que $\{|+\rangle, |-\rangle\}$ é uma base de $\mathbb{C}$. Para isso, seja $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ um estado qualquer.

¹ A definição de operador hermitiano encontra-se disponível no anexo na Definição .0.25.

Primeiramente, vamos mostrar que $\{|+\rangle, |-\rangle\}$ gera $\mathbb{C}$, mostrando que $|\psi\rangle$ pode ser escrito como combinação linear de $|+\rangle$ e $|-\rangle$. Tomando $a = \frac{1}{\sqrt{2}}(\alpha + \beta)$ e $b = \frac{1}{\sqrt{2}}(\alpha - \beta)$, $a, b \in \mathbb{C}$, temos

$$|\psi\rangle = \underbrace{\frac{1}{\sqrt{2}}(\alpha + \beta)}_a |+\rangle + \underbrace{\frac{1}{\sqrt{2}}(\alpha - \beta)}_b |-\rangle.$$

De fato:

$$\begin{aligned} a|+\rangle + b|-\rangle &= \frac{1}{\sqrt{2}}(\alpha + \beta)|+\rangle + \frac{1}{\sqrt{2}}(\alpha - \beta)|-\rangle \\ &= \frac{1}{\sqrt{2}}(\alpha + \beta) \left(\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \right) + \frac{1}{\sqrt{2}}(\alpha - \beta) \left(\frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) \right) \\ &= \frac{1}{2}[(\alpha + \beta) + (\alpha - \beta)]|0\rangle + \frac{1}{2}[(\alpha + \beta) - (\alpha - \beta)]|1\rangle \\ &= \frac{1}{2}2\alpha|0\rangle + \frac{1}{2}2\beta|1\rangle = \alpha|0\rangle + \beta|1\rangle \\ &= |\psi\rangle. \end{aligned}$$

Ademais, se existem $a, b \in \mathbb{C}$ tais que $|\psi\rangle = a|+\rangle + b|-\rangle = 0$, então $a \left(\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \right) + b \left(\frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) \right) = 0$, donde

$$\begin{aligned} \left(\frac{a}{\sqrt{2}} + \frac{b}{\sqrt{2}} \right) |0\rangle + \left(\frac{a}{\sqrt{2}} - \frac{b}{\sqrt{2}} \right) |1\rangle = 0 &\Rightarrow \frac{a}{\sqrt{2}} + \frac{b}{\sqrt{2}} = 0 \quad \text{e} \quad \frac{a}{\sqrt{2}} - \frac{b}{\sqrt{2}} = 0 \\ &\Rightarrow a + b = 0 \quad \text{e} \quad a - b = 0 \Rightarrow a = b = 0. \end{aligned}$$

Logo, o conjunto $\{|+\rangle, |-\rangle\}$ é LI.

Portanto, $\{|+\rangle, |-\rangle\}$ é uma base de $\mathbb{C}$ a qual denominamos por **base conjugada**. Os vetores dessa base podem ser escritos matricialmente como

$$|+\rangle = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix} \quad \text{e} \quad |-\rangle = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix}.$$

A tabela a seguir resume como as portas lógicas que temos até dado momento agem sobre os bits $|0\rangle$ e $|1\rangle$.

bit de entrada	X	Z	H
$ 0\rangle$	$ 1\rangle$	$ 0\rangle$	$ +\rangle$
$ 1\rangle$	$ 0\rangle$	$- 1\rangle$	$ -\rangle$

Tabela 2 – Ações das portas X , Z e H sobre os qubits $|0\rangle$ e $|1\rangle$.

Uma matriz unitária A , por definição, é inversível, já que

$$A^\dagger A = AA^\dagger = I.$$

Mais ainda, sua inversa $A^\dagger$ também é uma matriz unitária. Assim, o fato de qualquer porta quântica ser representada por uma matriz unitária implica que sua inversa também será uma porta quântica. Logo, as portas quânticas são reversíveis, isto é, é sempre possível inverter uma porta quântica utilizando outra porta quântica.

Exemplo 3.5.8. Uma outra porta lógica muito importante, agora para mais de um qubit, é conhecida como **NOT-controlada**, ou simplesmente **CNOT**. Essa porta tem dois qubits de entrada, denominados **qubit de controle** e **qubit alvo**, respectivamente. Essa porta lógica age da seguinte forma:

- Se o qubit controle for $|1\rangle$, então a porta CNOT age no qubit alvo, trocando-o de $|0\rangle$ para $|1\rangle$, e de $|1\rangle$ para $|0\rangle$.
- Se o qubit controle for $|0\rangle$, então nada acontece com o qubit alvo.

Assim, temos:

bit de entrada	CNOT
00	00
01	01
10	11
11	10

Tabela 3 – Ação da porta CNOT.

Outra maneira de descrevermos a porta CNOT se dá em vê-la como uma generalização da porta clássica XOR, pondo $|a, b\rangle \mapsto |a, b \oplus a\rangle$, onde $\oplus$ denota a adição módulo 2. De fato:

- $|0, 0\rangle \mapsto |0, 0 \oplus 0\rangle = |0, 0\rangle$;
- $|0, 1\rangle \mapsto |0, 1 \oplus 0\rangle = |0, 1\rangle$;
- $|1, 0\rangle \mapsto |1, 0 \oplus 1\rangle = |1, 1\rangle$;
- $|1, 1\rangle \mapsto |1, 1 \oplus 1\rangle = |1, 0\rangle$.

Ainda, outra maneira de representarmos a porta CNOT se dá matricialmente por

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

que é uma matriz unitária, já que é uma matriz cujas entradas são número reais, sua transposição é igual a matriz original (ou seja, $(CNOT)^\dagger = CNOT$) e

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

A porta CNOT é uma das mais interessantes, pois qualquer operação lógica quântica envolvendo múltiplos qubits pode ser construída a partir dela e de portas de um único qubit. Ou seja, a porta CNOT é universal.

Vejamos no exemplo a seguir, um outro fato interessante a se destacar sobre a porta CNOT, que consiste no fato de que essa porta pode produzir um estado emaranhado.

Exemplo 3.5.9. Através da porta CNOT, podemos fazer o estado não emaranhado $|00\rangle$ tornar-se um estado emaranhado. De fato:

1. Consideremos o estado $|00\rangle$;
2. Apliquemos o operador $H \otimes I$ em $|00\rangle$. Temos:

$$\begin{aligned} (H \otimes I)|00\rangle &= H|0\rangle \otimes I|0\rangle \\ &= \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \otimes |0\rangle \\ &= \frac{1}{\sqrt{2}}(|0\rangle \otimes |0\rangle + |1\rangle \otimes |0\rangle) \\ &= \frac{1}{\sqrt{2}}(|00\rangle + |10\rangle). \end{aligned}$$

3. Agora, aplicando nesse estado a porta CNOT, obtemos:

$$\begin{aligned} CNOT((H \otimes I)|00\rangle) &= CNOT\left(\frac{1}{\sqrt{2}}(|00\rangle + |10\rangle)\right) \\ &= \frac{1}{\sqrt{2}}(CNOT(|00\rangle) + CNOT(|10\rangle)) \\ &= \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle) \\ &\stackrel{(3.3)}{=} \beta_{|00\rangle}, \end{aligned}$$

que conforme o Exemplo 3.4.3, é um estado emaranhado.

Repetindo o processo anterior para cada $i, j \in \{0, 1\}$, isto é, considerando o estado não emaranhado $|ij\rangle$, em seguida aplicando o operador $H \otimes I$ em $|ij\rangle$ e, por fim, aplicando a porta CNOT nesse estado, vamos obter o estado de Bell $\beta_{|ij\rangle}$, que é um estado emaranhado.

4 CÓDIGOS QUÂNTICOS CORRETORES DE ERROS

Neste capítulo introduziremos a *Teoria dos Códigos Quânticos Corretores de Erros (QECC)*, que tem como objetivo a proteção da informação quântica contra a ação de ruído, para isso usaremos como as referências [1] e [5]. A informação quântica é processada por meio de operações unitárias seguidas de medições, realizadas em sistemas físicos específicos. Interações com o ambiente que rodeiam os sistemas podem provocar imperfeições nas operações e até mesmo erros no processamento da informação. Para superar esses problemas e garantir um bom funcionamento de um computador quântico, o uso de QECC tem sido uma das principais estratégias.

Os códigos quânticos corretores de erros funcionam codificando estados quânticos para que eles possam permanecer estáveis, protegidos, em relação à ação de ruídos, e depois para recuperarmos os estados basta uma decodificação.

4.1 DIFERENÇAS ENTRE INFORMAÇÃO QUÂNTICA E CLÁSSICA

Assim como no caso de códigos clássicos, nos códigos quânticos também acrescentamos redundâncias, para que seja possível detectar e corrigir erros, mesmo que haja corrompimento de alguma informação pelo ruído.

Os QECC têm sua construção baseada nos Códigos Corretores de Erros Clássicos, apesar das diferenças fundamentais na informação quântica.

A primeira diferença consiste na *impossibilidade de copiar um qubit*. Afim de mostrarmos isso, vejamos o seguinte teorema.

Teorema 4.1.1 (Teorema da não-clonagem). *Nenhum dispositivo quântico pode ser construído que produza $|\psi\rangle|\psi\rangle$, dado $|\psi\rangle$ arbitrário.*

Demonstração. Ver página 48, capítulo 3 de [1]. □

Definição 4.1.2. Definimos um **código de repetição** como um código simples para proteção de bit contra erros que podem ser introduzidos no canal.

Apresentaremos agora um exemplo de informação clássica, utilizando como princípio de codificação a cópia de bits.

Exemplo 4.1.3. Aqui vamos apresentar um código de repetição de 3 bits, e

$$\begin{aligned} 0 &\mapsto 000 \\ 1 &\mapsto 111. \end{aligned}$$

Admitindo que o receptor saiba que o processo de codificação utilizado foi a repetição, de posse dos 3 dígitos recebidos, ele deverá decidir qual mensagem (bit) foi enviada

originalmente. Suponhamos que o receptor tenha recebido a mensagem “101”, a qual foi codificada e enviada seguindo a regra acima mencionada. Sem dificuldades, o receptor entenderá que o erro ocorreu no segundo bit e decidirá que a mensagem enviada foi “111”. Por outro lado, se dois bits forem trocados, isto é, se ocorrerem dois erros no processo de transmissão da mensagem, o receptor provavelmente decidirá pelo bit de informação errado. Isso deixa transparecer a “fragilidade” desse tipo de codificação.

Uma outra diferença entre as informações clássica e quântica é que *em um único qubit podem ocorrer erros distintos de forma contínua*. Isso significa que um erro num estado $\alpha|0\rangle + \beta|1\rangle$ pode mudar α e β por uma quantidade pequena ε , onde esses pequenos erros podem acumular-se durante o tempo, corrompendo a informação.

No Exemplo 3.5.5, vimos que o operador Z , que age trocando o sinal do estado $|1\rangle$. Além disso, sabemos que um qubit pode estar em uma superposição de dois estados, isto é,

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle, \quad \alpha, \beta \in \mathbb{C}.$$

Exemplo 4.1.4. O erro *phase flip*, por exemplo, faria com que o estado

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$$

mudasse para

$$|\psi\rangle = \alpha|0\rangle - \beta|1\rangle.$$

Assim, outra diferença é que *erros do tipo phase flip não ocorrem no caso clássico*. Erros como esse são bem graves, pois podem transformar o estado $|+\rangle$ no estado $|-\rangle$, conforme o exemplo a seguir.

Exemplo 4.1.5. Vamos entender melhor o comportamento do erro *phase flip* nos elementos da base conjugada. Consideremos os estados da base conjugada $\{|+\rangle, |-\rangle\}$. Aplicando o operador Z em cada elemento dessa base, temos:

$$Z|+\rangle = Z\left(\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)\right) = \frac{1}{\sqrt{2}}(Z|0\rangle + Z|1\rangle) = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) = |-\rangle.$$

$$Z|-\rangle = Z\left(\frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)\right) = \frac{1}{\sqrt{2}}(Z|0\rangle - Z|1\rangle) = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) = |+\rangle.$$

Logo, Z leva o estado $|+\rangle$ no estado $|-\rangle$ e o estado $|-\rangle$ no estado $|+\rangle$.

Outra diferença fundamental entre as informações quânticas e clássica é que *as medidas destroem a informação quântica*, como vimos no Postulado 3. Na correção clássica de erros, é possível observar e corrigir erros na saída do canal e decidir qual o procedimento a ser adotado, ou seja, esse problema não ocorre. Em contrapartida, a observação no caso quântico geralmente destrói o estado quântico sob observação, assim impossibilitando sua recuperação.

Mesmo com essas diferenças mencionadas, é possível desenvolver a teoria de correção quântica de erros com adaptações para que os códigos quânticos de correção de erros possam funcionar.

Na próxima seção, veremos alguns exemplos de códigos quânticos de erros.

Definição 4.1.6. Um **código quântico corretor de erros (QECC)** é uma função de um espaço de Hilbert de dimensão 2^k em um espaço de Hilbert de dimensão 2^n , com $k < n$. Os vetores no espaço 2^n dimensional são denominadas **palavras-código**.

Definição 4.1.7. Seja $\mathcal{C}$ um QECC. Definimos a **distância mínima** d de $\mathcal{C}$ como a menor distância de Hamming entre duas palavras-código distintas, isto é, o número

$$d = \min\{d(u, v); u, v \in \mathcal{C} \text{ e } u \neq v\}.$$

Um QECC $\mathcal{C}$ com comprimento n , dimensão k e distância mínima d é chamado um código $[[n, k, d]]$. Um código com distância mínima d pode corrigir até t erros ocorridos nos qubits de uma palavra-código, onde

$$t = \left\lfloor \frac{d-1}{2} \right\rfloor.$$

4.2 CÓDIGO BIT FLIP

Apresentaremos agora um primeiro exemplo de código quântico, chamado *código bit flip*. Este código será o mais simples que veremos e codifica três qubits, por isso, muitas vezes também é denominado em literaturas como *código bit flip de três qubits*.

Observação 4.2.1. Como vimos no Exemplo 4.1.3, a codificação por repetição falha quando mais de um bit dos três enviados seja invertido, o que ocorre com a probabilidade de $3p^2 - 2p^3$. Sem adicionarmos as redundâncias, a probabilidade de erro é p , de modo que o código de repetição garante a confiabilidade da transmissão se $p < \frac{1}{2}$.

De forma análoga ao que fizemos no Exemplo 4.1.3, agora consideremos o envio do qubit $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, $|\alpha|^2 + |\beta|^2 = 1$, através do canal *bit flip*. Neste canal o qubit é transmitido sem erros com uma probabilidade $1 - p$ e o sofre a inversão de bit com probabilidade p . Em outras palavras, caso o canal atue sobre o qubit $|\psi\rangle$, o estado transmitido será $X|\psi\rangle = \beta|0\rangle + \alpha|1\rangle$, onde X representa a matriz de Pauli, conhecida como **operador bit flip**. Esse canal é chamado **canal bit flip** ou de **inversão de bit**.

Exemplo 4.2.2. Mostraremos que o código *bit flip* protege qubits contra efeitos de ruídos deste canal, descrevendo o processo realizado pelo código *bit flip* para proteção, identificação e correção dos erros do canal *bit flip*.

Consideremos $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ o estado de um qubit. Suponhamos que $|\psi\rangle$ seja codificado com três qubits, da seguinte forma:

$$\alpha|000\rangle + \beta|111\rangle.$$

Podemos reescrever essa operação convenientemente como

$$|0\rangle \rightarrow |0_L\rangle \equiv |000\rangle \quad (4.1)$$

$$|1\rangle \rightarrow |1_L\rangle \equiv |111\rangle \quad (4.2)$$

onde $|0_L\rangle$ e $|1_L\rangle$ são chamados de **estados lógicos**.

Suponhamos que o estado inicial $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ tenha sido perfeitamente codificado como $\alpha|000\rangle + \beta|111\rangle$.

Ao codificarmos $|\psi\rangle$ para $\alpha|000\rangle + \beta|111\rangle$, cada qubit é codificado pelo canal *bit flip* de maneira independente, de modo que, caso ocorra erro em um deles durante a codificação, isso não influenciará a codificação dos demais.

Analogamente ao caso clássico, no código *bit flip* conseguimos identificar e corrigir erros, caso tenham ocorrido em no máximo um qubit. Assim, suponhamos que, ao codificarmos $|\psi\rangle$ para $\alpha|000\rangle + \beta|111\rangle$ um erro tenha ocorrido em no máximo um 1 qubit. Desejamos corrigir esse erro, mas sem perder a superposição $\alpha|000\rangle + \beta|111\rangle$. Para isso, é necessário entendermos sobre o funcionamento do procedimento de correção de erros para recuperarmos o estado original. Este procedimento é dividido nos seguintes passos:

1. No primeiro passo, queremos detectar se ocorreu erro em algum qubit do nosso estado quântico. Para isso, devemos realizar uma medida sobre o estado quântico, a qual irá detectar erro (caso haja algum) nos dizendo onde foi que ele ocorreu (caso tenha ocorrido). O resultado dessa medida é chamado de **síndrome do erro**.
2. No segundo passo do processo de correção de erros, utilizamos o valor da síndrome obtido no passo anterior para determinar qual será o procedimento de recuperação do estado original.

Consideremos os seguintes operadores de projeção:

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

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

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

$$P_3 = |001\rangle\langle 001| + |110\rangle\langle 110|,$$

Para o canal *bit flip*, esses quatro operadores são correspondentes as quatro síndromes de erro, da seguinte forma:

- O operador P_0 corresponde a erro em nenhum qubit.

- Para cada $i \in \{1, 2, 3\}$, o operador P_i corresponde a erro no i -ésimo qubit.

Observação 4.2.3. Calculando os operadores projeções obtemos:

$$P_0 = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}; \quad P_1 = \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}; \quad (4.3)$$

$$P_2 = \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}; \quad P_3 = \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}. \quad (4.4)$$

Definição 4.2.4. Dado um estado $|\psi\rangle$, sendo $|\psi'\rangle$ o estado obtido após a transmissão, a *medida da síndrome* é dada pelo operador

$$|\psi'\rangle P_i |\psi'\rangle, \quad \text{para } i = 0, 1, 2, 3. \quad (4.5)$$

Ao realizarmos uma medição, os resultados possíveis são 0 ou 1. Caso o resultado seja 0, podemos afirmar que o erro apontado pelo operador não ocorreu naquele qubit. Por outro lado, caso o resultado seja 1, então ocorreu o erro apontado pelo operador.

Exemplo 4.2.5. Nos itens a seguir, analisaremos ocorrências do erro *bit flip*, sendo $|\psi\rangle = \alpha|000\rangle + \beta|111\rangle$ o estado que gostaríamos de receber.

1. Suponhamos que ocorreu um erro *bit flip* no primeiro qubit. Assim, ao invés de recebermos o estado $|\psi\rangle$, o resultado obtido após a transmissão foi $|\psi'\rangle =$

$\alpha|100\rangle + \beta|011\rangle$. Temos:

$$\begin{aligned}
 |\psi'\rangle &= \alpha|100\rangle + \beta|011\rangle = \alpha(|1\rangle \otimes |0\rangle \otimes |0\rangle) + \beta(|0\rangle \otimes |1\rangle \otimes |1\rangle) \\
 &= \alpha \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) + \beta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) \\
 &= \alpha \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix} \right) + \beta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix} \right) = \alpha \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} + \beta \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ \beta \\ \alpha \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}.
 \end{aligned}$$

Assim,

$$\langle\psi'| = [0 \ 0 \ 0 \ \bar{\beta} \ \bar{\alpha} \ 0 \ 0 \ 0].$$

Logo, calculando a medida da síndrome:

$$\begin{aligned}
 \langle\psi'|P_1|\psi'\rangle &\stackrel{(4.3)}{=} [0 \ 0 \ 0 \ \bar{\beta} \ \bar{\alpha} \ 0 \ 0 \ 0] \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ 0 \\ 0 \\ \beta \\ \alpha \\ 0 \\ 0 \\ 0 \end{bmatrix} \\
 &= |\beta|^2 + |\alpha|^2 = 1.
 \end{aligned}$$

Essa síndrome está associada ao erro no primeiro qubit. Perceba que, caso seja feito o cálculo da síndrome, nesse caso, para os demais projetores, os resultados serão 0.

2. Suponhamos que ocorreu um erro *bit flip* no segundo qubit, isto é, em vez do estado $|\psi\rangle$ o estado resultante da transmissão foi $|\psi'\rangle = \alpha|010\rangle + \beta|101\rangle$. Calculando $|\psi'\rangle$ e $\langle\psi'|$, obtemos:

$$|\psi'\rangle = \begin{bmatrix} 0 \\ 0 \\ \alpha \\ 0 \\ 0 \\ 0 \\ \beta \\ 0 \\ 0 \end{bmatrix} \quad \text{e} \quad \langle\psi'| = [0 \ 0 \ \bar{\alpha} \ 0 \ 0 \ \bar{\beta} \ 0 \ 0].$$

Portanto, a medida da síndrome é dada por

$$\begin{aligned}
 \langle \psi' | P_2 | \psi' \rangle &\stackrel{(4.4)}{=} \begin{bmatrix} 0 & \bar{\alpha} & 0 & 0 & 0 & 0 & \bar{\beta} & 0 \end{bmatrix} \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ \alpha \\ 0 \\ 0 \\ 0 \\ 0 \\ \beta \\ 0 \end{bmatrix} \\
 &= |\alpha|^2 + |\beta|^2 = 1.
 \end{aligned}$$

Essa síndrome está associada ao erro no segundo qubit. Caso seja feito o cálculo da síndrome, para os demais projetores, os resultados serão 0.

3. Por fim, suponhamos que ocorreu um erro *bit flip* no terceiro qubit. Então, em vez do estado $|\psi\rangle$, recebemos $|\psi'\rangle = \alpha|001\rangle + \beta|110\rangle$. Calculando $|\psi'\rangle$ e $\langle\psi'|$, obtemos:

$$|\psi'\rangle = \begin{bmatrix} 0 \\ \alpha \\ 0 \\ 0 \\ 0 \\ 0 \\ \beta \\ 0 \end{bmatrix} \quad \text{e} \quad \langle\psi'| = \begin{bmatrix} 0 & \bar{\alpha} & 0 & 0 & 0 & 0 & \bar{\beta} & 0 \end{bmatrix}.$$

Calculando a medida da síndrome:

$$\begin{aligned}
 \langle \psi' | P_3 | \psi' \rangle &\stackrel{(4.4)}{=} \begin{bmatrix} 0 & \bar{\alpha} & 0 & 0 & 0 & 0 & \bar{\beta} & 0 \end{bmatrix} \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ \alpha \\ 0 \\ 0 \\ 0 \\ 0 \\ \beta \\ 0 \end{bmatrix} \\
 &= |\alpha|^2 + |\beta|^2 = 1.
 \end{aligned}$$

Essa síndrome está associada ao erro no terceiro qubit. Caso seja feito o cálculo da síndrome, para os demais projetores, os resultados serão 0.

Definição 4.2.6. [20] Uma medida projetiva é descrita por um observável M , que é um operador no espaço de estados do sistema sendo observado. O observável tem uma

decomposição espectral

$$M = \sum_m m P_m,$$

onde P_m é projetor sobre o autoespaço de M com autovalor m . Os possíveis resultados da medida correspondem aos autovalores m do observável. Medindo-se o estado $|\psi\rangle$, a probabilidade de se obter o resultado m é dada por

$$p(m) = \langle \psi | P_m | \psi \rangle.$$

Obtido o resultado m , o estado do sistema logo após a medida será

$$\frac{P_m |\psi\rangle}{\sqrt{p(m)}}.$$

Utilizando a definição acima é possível mostrarmos que a medição da síndrome não altera o estado como mostraremos no exemplo a seguir. A síndrome de erro nos informa somente se ocorreu erro e caso ocorreu, onde ele ocorreu, não nos dizendo nada sobre o estado quântico, isto é, sobre os valores das amplitudes α e β .

Exemplo 4.2.7. Verifiquemos que a medição da síndrome não altera o estado, utilizando o Exemplo 4.2.5.

1. Se $m = 1$ ocorreu, então, pela Definição 4.2.6, o estado após a medição será dado por

$$\frac{P_1 |\psi'\rangle}{p(1)} = \frac{P_1 (\alpha |100\rangle + \beta |011\rangle)}{\sqrt{\langle \psi' | P_1 | \psi' \rangle}} = \frac{\alpha |100\rangle + \beta |011\rangle}{\sqrt{1}} = \alpha |100\rangle + \beta |011\rangle = |\psi'\rangle.$$

2. Se $m = 2$ ocorreu, então, pela Definição 4.2.6, o estado após a medição será dado por

$$\frac{P_2 |\psi'\rangle}{p(2)} = \frac{P_2 (\alpha |010\rangle + \beta |101\rangle)}{\sqrt{\langle \psi' | P_2 | \psi' \rangle}} = \frac{\alpha |010\rangle + \beta |101\rangle}{\sqrt{1}} = \alpha |010\rangle + \beta |101\rangle = |\psi'\rangle.$$

3. Se $m = 3$ ocorreu, então, pela Definição 4.2.6, o estado após a medição será dado por

$$\frac{P_3 |\psi'\rangle}{p(3)} = \frac{P_3 (\alpha |001\rangle + \beta |110\rangle)}{\sqrt{\langle \psi' | P_3 | \psi' \rangle}} = \frac{\alpha |001\rangle + \beta |110\rangle}{\sqrt{1}} = \alpha |001\rangle + \beta |110\rangle = |\psi'\rangle.$$

Após detectarmos se ocorreu erro em algum qubit, a recuperação é feita de maneira bem simples. Definimos os operadores X_1 , X_2 e X_3 como

$$X_1 = X \otimes I \otimes I, \quad X_2 = I \otimes X \otimes I, \quad X_3 = I \otimes I \otimes X. \quad (4.6)$$

Quando gostaríamos de receber um estado $|\psi\rangle$ e ocorrer um erro *bit flip* no i -ésimo qubit, $i \in \{1, 2, 3\}$, obtendo o estado $|\psi'\rangle$ após a transmissão, basta aplicar o operador X_i sobre o estado transmitido para recuperarmos o estado original $|\psi\rangle$. Em outras palavras, cada um dos operadores definidos em (4.6) tem a função de inverter o qubit com erro na transmissão.

Exemplo 4.2.8. Se a síndrome de erro for 1 para o operador P_2 , ou seja, se um erro ocorreu no segundo qubit, como no Exemplo 1, então devemos inverter o qubit novamente aplicando o operador X_2 , e assim obtemos o estado inicial. De fato, nesse caso, o estado inicial é $|\psi\rangle = \alpha|000\rangle + \beta|111\rangle$ e o estado transmitido foi $|\psi'\rangle = \alpha|010\rangle + \beta|101\rangle$. Temos:

$$\begin{aligned}
X_2|\psi'\rangle &= X_2(\alpha|010\rangle + \beta|101\rangle) \\
&= I \otimes X \otimes I[\alpha(|0\rangle \otimes |1\rangle \otimes |0\rangle) + \beta(|1\rangle \otimes |0\rangle \otimes |1\rangle)] \\
&= \alpha[I \otimes X \otimes I(|0\rangle \otimes |1\rangle \otimes |0\rangle)] + \beta[I \otimes X \otimes I(|1\rangle \otimes |0\rangle \otimes |1\rangle)] \\
&\stackrel{\text{Def. 0.27}}{=} \alpha(I|0\rangle \otimes X|1\rangle \otimes I|0\rangle) + \beta(I|1\rangle \otimes X|0\rangle \otimes I|1\rangle) \\
&= \alpha(|0\rangle \otimes |0\rangle \otimes |0\rangle) + \beta(|1\rangle \otimes |1\rangle \otimes |1\rangle) \\
&= \alpha|000\rangle + \beta|111\rangle \\
&= |\psi\rangle.
\end{aligned}$$

As possíveis síndromes de erro e seus correspondentes procedimentos de correção são apresentados na tabela a seguir.

síndrome m	procedimento
0	não faça nada
1	inverte o primeiro qubit, aplicando X_1 em $ \psi'\rangle$
2	inverte o segundo qubit, aplicando X_2 em $ \psi'\rangle$
3	inverte o terceiro qubit, aplicando X_3 em $ \psi'\rangle$

Tabela 4 – Síndromes de erros para o caso *bit flip*.

O procedimento de correção de erros descrito nessa tabela funciona perfeitamente, desde que a inversão ocorra em no máximo um dos três qubits. A probabilidade de que isso aconteça é de $1 - 3p^2 + 2p^3$. Logo, a probabilidade de erro é $3p^2 - 2p^3$, exatamente como no caso clássico de repetição. Assim como no caso clássico, desde que $p < \frac{1}{2}$, a codificação e a decodificação aumentam a confiança na proteção do estado quântico, de forma que o código *bit flip* será mais confiável do que enviarmos um único qubit.

Alguns erros podem corromper um estado quântico por completo. Isso pode dificultar o processo de recobrimento dos erros e, por conseguinte, sua correção. Vejamos um exemplo que nos mostra isso.

Exemplo 4.2.9. Consideremos o operador X . Temos:

$$\begin{aligned}
X|+\rangle &= X\left(\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)\right) = \frac{1}{\sqrt{2}}X|0\rangle + \frac{1}{\sqrt{2}}X|1\rangle \\
&= \frac{1}{\sqrt{2}}|1\rangle + \frac{1}{\sqrt{2}}|0\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) = |+\rangle.
\end{aligned}$$

Temos também:

$$\begin{aligned} X|-\rangle &= X \left(\frac{1}{\sqrt{2}}(|0\rangle - |1\rangle) \right) = \frac{1}{\sqrt{2}}X|0\rangle - \frac{1}{\sqrt{2}}X|1\rangle \\ &= \frac{1}{\sqrt{2}}|1\rangle - \frac{1}{\sqrt{2}}|0\rangle = \frac{1}{\sqrt{2}}(|1\rangle - |0\rangle) \neq |-\rangle. \end{aligned}$$

Assim, como o operador X inverte os estados $|0\rangle$ e $|1\rangle$, ele não afeta o estado $|+\rangle$, mas afeta o estado $|-\rangle$.

Problemas como os do Exemplo 4.2.9, podem dificultar o processo de recobrimento dos erros e, por conseguinte, suas correções. Todavia, existe uma interpretação alternativa para a medida de síndrome que faz uso do conceito de observável.

No modo alternativo de correção de erro *bit flip*, em vez de medirmos os projetores P_0 , P_1 , P_2 e P_3 , executamos duas medidas, primeiro através do observáveis Z_1Z_2 e depois através dos observáveis Z_2Z_3 , onde

$$Z_1Z_2 = Z \otimes Z \otimes I \quad \text{e} \quad Z_2Z_3 = I \otimes Z \otimes Z. \quad (4.7)$$

Esses observáveis cumprem o papel de comparação, com Z_1Z_2 comparando os dois primeiros qubits enquanto Z_2Z_3 compara o segundo e o terceiro qubits.

Cada um desses observáveis tem autovalores ± 1 e, assim, cada medida fornece um único bit de informação para um total de dois bits de informação.

Ao realizarmos a medição utilizando Z_1Z_2 , se $Z_1Z_2|\psi\rangle = |\psi\rangle$, então os dois primeiros qubits são iguais, e se $Z_1Z_2|\psi\rangle = -|\psi\rangle$ então, os dois primeiros qubits são diferentes. De forma análoga comparamos o segundo e o terceiro qubit. Na tabela a seguir apresentamos a listagem das possibilidades dessas medidas.

$Z_1Z_2 \psi\rangle$	$Z_2Z_3 \psi\rangle$	qubit onde ocorreu o erro <i>bit flip</i>
$+\psi\rangle$	$+\psi\rangle$	nenhum
$-\psi\rangle$	$+\psi\rangle$	primeiro qubit
$-\psi\rangle$	$-\psi\rangle$	segundo qubit
$+\psi\rangle$	$-\psi\rangle$	terceiro qubit

Tabela 5 – Resultado das medidas de Z_1Z_2 e Z_2Z_3 sobre $|\psi\rangle$.

Outros possíveis conjuntos de observáveis que podem ser utilizados para a medida das síndromes são:

$$\begin{aligned} Z_1Z_3 &= Z \otimes I \otimes Z \quad \text{e} \quad Z_2Z_3 = I \otimes Z \otimes Z \\ Z_1Z_2 &= Z \otimes Z \otimes I \quad \text{e} \quad Z_1Z_3 = Z \otimes I \otimes Z. \end{aligned}$$

As tabelas desses conjuntos de observáveis são dadas por:

$Z_1 Z_3 \psi\rangle$	$Z_2 Z_3 \psi\rangle$	qubit onde ocorreu o erro <i>bit flip</i>
$+\lvert\psi\rangle$	$+\lvert\psi\rangle$	nenhum
$-\lvert\psi\rangle$	$+\lvert\psi\rangle$	primeiro qubit
$+\lvert\psi\rangle$	$-\lvert\psi\rangle$	segundo qubit
$-\lvert\psi\rangle$	$-\lvert\psi\rangle$	terceiro qubit

Tabela 6 – Resultado das medidas de $Z_1 Z_3$ e $Z_2 Z_3$ sobre $|\psi\rangle$.

$Z_1 Z_2 \psi\rangle$	$Z_1 Z_3 \psi\rangle$	qubit onde ocorreu o erro <i>bit flip</i>
$+\lvert\psi\rangle$	$+\lvert\psi\rangle$	nenhum
$-\lvert\psi\rangle$	$-\lvert\psi\rangle$	primeiro qubit
$-\lvert\psi\rangle$	$+\lvert\psi\rangle$	segundo qubit
$+\lvert\psi\rangle$	$-\lvert\psi\rangle$	terceiro qubit

Tabela 7 – Resultado das medidas de $Z_1 Z_2$ e $Z_1 Z_3$ sobre $|\psi\rangle$.

Uma vez que foi detectado o erro, o passo seguinte será corrigi-lo. Se nenhum erro ocorreu, nada se tem a fazer. Porém, se as síndromes indicarem que ocorreu erro em algum qubit, devemos aplicar o operador X sobre o qubit onde esse erro ocorreu, fazendo assim a correção. Além disso, assim como no caso clássico de repetição de três bits, esse código tem distância mínima $d = 3$ e é capaz de corrigir até

$$t = \left\lfloor \frac{d-1}{2} \right\rfloor = \lfloor 1 \rfloor = 1 \text{ erro}$$

4.3 CÓDIGO *PHASE FLIP*

Vimos que o código *bit flip* não apresenta uma mudança significativa em relação aos códigos corretores de erros clássicos. Veremos agora um canal de ruído mais interessante, conhecido como ***phase flip*** ou ***inversão de fase*** para um qubit. Este código é mais um tipo de código quântico que envolve três qubits e, ao contrário do código *bit flip*, não tem nenhum código equivalente nos códigos clássicos.

Assim como fizemos para o canal *bit flip*, gostaríamos de enviar o qubit $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ mas agora através do canal *phase flip*. Esse canal também transmite um qubit sem erros com probabilidade $1 - p$, e inverte a fase dos estados $|0\rangle$ e $|1\rangle$ com probabilidade $p > 0$ levando o estado $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ para $Z|\psi\rangle = \alpha|0\rangle - \beta|1\rangle$.

A seguir, apresentaremos o funcionamento do código *phase flip*, responsável pela proteção, detecção e correção de erros no canal correspondente, utilizando uma forma simples para transformar o canal *phase flip* em um canal de *bit flip*. Para isso, utilizaremos a base conjugada $\{|+\rangle, |-\rangle\}$. Vale lembrar que, nessa base, o operador Z atua como um operador *bit flip*, pois troca $|+\rangle$ por $|-\rangle$ e vice-versa, conforme ilustrado no Exemplo 4.1.5.

Como já sabemos proceder com erros *bit flip*, então passaremos o qubit $|\psi\rangle$ para a base conjugada. Assim os estados lógicos (4.1) e (4.2) são escritos na base conjugado como

$$\begin{aligned} |0_L\rangle &\equiv |+++ \rangle \\ |1_L\rangle &\equiv |-- \rangle, \end{aligned}$$

onde

$$|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle) \quad \text{e} \quad |-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle).$$

As operações de codificação, detecção de erro e recuperação são realizadas da mesma forma que no caso do canal *bit flip*, porém em relação a base conjugada. Como pode ser observado em (3.7), a porta de Hadamard H pode ser invertida por ela própria. Em particular, temos:

$$H|+\rangle = H(H|0\rangle) = |0\rangle \quad \text{e} \quad H|-\rangle = H(H|1\rangle) = |1\rangle.$$

Assim, a porta de Hadamard realiza a mudança entre as bases $\{|0\rangle, |+\rangle\}$ e $\{|+\rangle, |-\rangle\}$. Portanto utilizaremos essa porta para mudança entre os canais.

O processo de codificação é dado da seguinte forma:

1. No primeiro passo codificamos cada qubit de informação em três qubits, similarmente ao que foi feito para o caso *bit flip*;
2. Depois de realizarmos o primeiro passo, aplicamos a porta de Hadamard sobre cada um dos qubits, $H^{\otimes 3} = H \otimes H \otimes H$, transformando $\alpha|000\rangle + \beta|111\rangle$ em $\alpha|+++ \rangle + \beta|-- \rangle$.

A operação de codificação pode ser escrita na base computacional como:

$$|0\rangle \rightarrow |000\rangle \rightarrow H^{\otimes 3}|000\rangle = |0_L\rangle \tag{4.8}$$

$$|1\rangle \rightarrow |111\rangle \rightarrow H^{\otimes 3}|111\rangle = |1_L\rangle \tag{4.9}$$

onde:

$$\begin{aligned} |0_L\rangle &\equiv \frac{1}{2\sqrt{2}}(|000\rangle + |001\rangle + |010\rangle + |011\rangle + |100\rangle + |101\rangle + |110\rangle + |111\rangle) \\ |1_L\rangle &\equiv \frac{1}{2\sqrt{2}}(|000\rangle - |001\rangle - |010\rangle + |011\rangle - |100\rangle + |101\rangle + |110\rangle - |111\rangle) \end{aligned}$$

Apesar do código de repetição de três bits não aparecer explicitamente nessa operação, ele está presente na distribuição dos sinais antes de cada ket.

A identificação e correção de um possível erro é realizada analogamente ao que foi feito antes, aplicando-se as mesmas medidas projetivas do canal *bit flip*, porém devemos

mudar a base dos operadores de projeção P_0, P_1, P_2 e P_3 para a base conjugada, o que pode ser feito, da seguinte maneira:

$$P_i \rightarrow P'_i \equiv H^{\otimes 3} P_i H^{\otimes 3}, \quad i = 0, 1, 2, 3.$$

Também, equivalentemente ao que fizemos para o caso *bit flip*, podemos utilizar os observáveis $Z_1 Z_2$ e $Z_2 Z_3$ (4.7), novamente mudado-os de base, para realizar as medidas de síndromes, de modo que

$$H^{\otimes 3} Z_1 Z_2 H^{\otimes 3} = X_1 X_2 = X \otimes X \otimes I \quad \text{e} \quad H^{\otimes 3} Z_2 Z_3 H^{\otimes 3} = X_2 X_3 = I \otimes X \otimes X.$$

A medida do observável $X_1 X_2$ compara os sinais do primeiro e segundo qubits, verificando se são “ambos positivos” ou “ambos negativos”, enquanto a medida de $X_2 X_3$ faz o mesmo que o observável $X_1 X_2$, mas comparando os sinais do segundo e terceiro qubits.

Exemplo 4.3.1. Considere os estados $|+\rangle|+\rangle \otimes (\cdot)$ ou $|-\rangle|-\rangle \otimes (\cdot)$, onde $(\cdot) = |+\rangle$ ou $(\cdot) = |-\rangle$, então o resultado da medida de $X_1 X_2$ é +1. Por outro lado, se considerarmos os estados como $|+\rangle|-\rangle \otimes (\cdot)$ ou $|-\rangle|+\rangle \otimes (\cdot)$, a medida desse observável será -1.

A tabela a seguir mostra todas as possibilidades de resultados das medidas e os possíveis erros detectados.

autovalores de $(X_1 X_2, X_2 X_3)$	bloco onde ocorreu o erro
(+1,+1)	nenhum
(-1,+1)	primeiro bloco
(-1,-1)	segundo bloco
(+1,-1)	terceiro bloco

Tabela 8 – Resultados das medidas de $X_1 X_2$ e $X_2 X_3$.

Também poderíamos utilizar os seguintes conjuntos de observáveis:

$$\begin{aligned} X_1 X_3 &= H^{\otimes 3} Z_1 Z_3 H^{\otimes 3} \quad \text{e} \quad X_2 X_3 = H^{\otimes 3} Z_2 Z_3 H^{\otimes 3} \\ X_1 X_2 &= H^{\otimes 3} Z_1 Z_2 H^{\otimes 3} \quad \text{e} \quad X_1 X_3 = H^{\otimes 3} Z_1 Z_3 H^{\otimes 3}. \end{aligned}$$

Para corrigir erros do tipo *phase flip*, podemos aplicar os mesmos operadores usados no canal *bit flip*, mas conjugados pela matriz de Hadamard. Por exemplo, se uma inversão de fase for detectada em um qubit na posição i , basta aplicar o operador $H X H$ sobre o i -ésimo qubit para realizar a correção.

Em resumo, o código *phase flip* tem as mesmas características do código *bit flip*. Em particular, os critérios pra o aumento da proteção, em comparação com caso em que o código não é aplicado, são os mesmos. Ademais, esse código também tem distância mínima $d = 3$ e também corrige até $t = 1$ erro.

4.4 CÓDIGO DE SHOR

Até dado momento, vimos que é possível construir um código que detecta e corrige erros *bit flip* e outro que detecta e corrige erros *phase flip*. Nesta seção veremos um código que é capaz de corrigir tanto erros *bit flip* como erros *phase flip*. Este será o **código de Shor**, que foi criado por Peter Shor em 1995 [23]. Esse código é um código com nove qubits, que é obtido por uma combinação de três qubits *bit flip* e *phase flip*¹.

A primeira coisa a se fazer é codificar $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ utilizando o código *phase flip*, o que, como vimos em (4.8) e (4.9), pode ser feito da seguinte forma:

$$\begin{aligned} |0\rangle &\rightarrow |+++\rangle \\ |1\rangle &\rightarrow |---\rangle. \end{aligned}$$

Em seguida, codificamos cada um desses qubits usando o código *bit flip*:

$$\begin{aligned} |+\rangle &\rightarrow \frac{1}{\sqrt{2}}(|000\rangle + |111\rangle) \\ |-\rangle &\rightarrow \frac{1}{\sqrt{2}}(|000\rangle - |111\rangle). \end{aligned}$$

Logo, o resultado dessa codificação será um código de nove qubits, com estados lógicos dados por:

$$\begin{aligned} |0\rangle &\rightarrow |0_L\rangle \equiv \frac{1}{2\sqrt{2}}(|000\rangle + |111\rangle)(|000\rangle + |111\rangle)(|000\rangle + |111\rangle) \\ |1\rangle &\rightarrow |1_L\rangle \equiv \frac{1}{2\sqrt{2}}(|000\rangle - |111\rangle)(|000\rangle - |111\rangle)(|000\rangle - |111\rangle). \end{aligned}$$

Observação 4.4.1. O código de Shor não contradiz o Teorema da Não-Clonagem, pois uma palavra-código arbitrária é uma superposição dos dois estados lógicos.

Devido à maneira como foi construído, o código de Shor tem a capacidade de proteger contra um único erro de inversão de bit, um único erro de inversão de fase, um erro de outro tipo, ou ambos simultaneamente, afetando o mesmo qubit ou qubits diferentes. Para verificar se ocorreu algum desses erros fazemos a seguinte verificação:

¹ Essa técnica de combinar dois ou mais códigos corretores de erros para formar um código mais robusto é conhecida como **concatenação**. Esta é uma maneira bastante útil de obter-se códigos novos a partir de códigos já conhecidos, podendo ainda aproveitarmos as melhores características de cada código individual, aumentando assim a capacidade de detecção e correção de erros.

1. Para verificarmos se ocorreu um erro *bit flip* em algum qubit, como, em geral, não sabemos exatamente em que bloco ocorreu o erro, calculamos a medida dos observáveis Z_1Z_2 , Z_2Z_3 , Z_4Z_5 , Z_5Z_6 , Z_7Z_8 e Z_8Z_9 ², realizando a comparação entre os qubits, assim como fizemos no código *bit flip*, mas agora analisando por blocos de três qubits. Caso tenha ocorrido um erro *bit flip*, recuperamos o estado original invertendo o qubit que sofreu o erro.
2. Agora, para verificarmos se ocorreu um erro *phase flip* em algum qubit, aplicamos os observáveis $X_1X_2X_3X_4X_5X_6$ e $X_4X_5X_6X_7X_8X_9$ comparando, respectivamente, os sinais do primeiro e segundo blocos de três qubits e do segundo e terceiro blocos de três qubits, da mesma forma que no caso *bit flip* se comparava o sinal de dois qubits. Uma vez calculada as medidas desses observáveis, é possível determinar em qual bloco de qubits ocorreu o erro de inversão de fase e depois corrigi-lo aplicando o operador Z sobre qualquer um dos qubits desse bloco.
3. Ainda, com esse código ainda é possível corrigir um erro *bit flip* e *phase flip* ocorrendo no mesmo qubit, aplicando o operador ZX a esse qubit corrompido. Para isso, basta aplicarmos os dois procedimentos descritos acima.

Observação 4.4.2. Vale destacarmos:

- Os procedimentos descritos acima são independentes, mas o código de Shor não corrige um erro de inversão de bit e um erro de inversão de fase se eles ocorrerem em qubits distintos.
- Se ocorrer um erro *phase flip*, não saberemos dizer em qual dos qubits ocorreu esse erro. Saberemos apenas em qual dos blocos o erro ocorreu.
- Para recuperarmos o estado original, sempre será realizado os dois primeiros procedimentos descritos acima, independente de ter ocorrido *bit flip*, *phase flip*, ou os dois.

Exemplo 4.4.3. Vejamos exemplos de possíveis erros *bit flip* para um código de Shor:

- (a) Suponhamos que um erro *bit flip* ocorreu no primeiro bloco de três qubits. Então, podemos detectá-lo através da medida dos observáveis Z_1Z_2 e Z_2Z_3 , e corrigi-lo aplicando o operador X no qubit onde ocorreu o erro.

- (b) Agora, supondo que tenha ocorrido um erro *bit flip* no segundo ou terceiro blocos do qubit, podemos detectar o qubit corrompido pela medida dos observáveis Z_4Z_5 e Z_5Z_6 ou Z_7Z_8 e Z_8Z_9 , respectivamente.

Exemplo 4.4.4. Vejamos exemplos de possíveis erros *phase flip* para um código de Shor:

- (a) Se ocorrer um único erro *phase flip* em qualquer um dos três primeiros qubits, o primeiro bloco inverterá o sinal, fazendo com o que o primeiro e o segundo blocos possuam sinais distintos e o segundo e terceiro blocos, sinais iguais.
- (b) Se ocorrer um erro *phase flip* em algum dos qubits do segundo bloco, esse bloco mudará de sinal e, conseqüentemente, o primeiro e o segundo bloco terão sinais distintos, assim como o segundo e terceiro blocos.
- (c) Se ocorrer um erro *phase flip* em um dos qubits do terceiro bloco, esse bloco terá sinal invertido, fazendo com que o primeiro e o segundo blocos tenham sinais iguais, e o segundo e terceiro, sinais diferentes.

Perceba que, seja qual for o qubit corrompido do bloco de três qubits, o efeito do erro é o mesmo e conseguimos recuperar o estado invertendo o sinal do bloco onde foi apontado o erro. Por exemplo, $(|000\rangle - |111\rangle)$ e $(|000\rangle - |111\rangle)$ têm o mesmo sinal em ambos os bloco, enquanto $(|000\rangle - |111\rangle)$ e $(|000\rangle + |111\rangle)$ têm sinais distintos.

Vale destacarmos que o código de Shor é superior aos códigos clássicos de repetição, pois realiza uma busca mais eficiente por erros. O código de Shor, além de proteger contra inversões de bit e de fase em um qubit, também protege estados quânticos contra erros completamente arbitrários, desde que apenas um qubit seja afetado [20]. Demonstra-se que, embora os erros possam ocorrer de forma contínua em um qubit, todos podem ser corrigidos ao corrigir apenas um subconjunto discreto de erros. Os demais erros serão automaticamente corrigidos. A discretização dos erros é crucial para entender a eficácia das correções quânticas.

4.5 CÓDIGOS DEGENERADOS

Vimos que em muitos aspectos, os códigos quânticos corretores de erros assemelham-se bem aos códigos corretores de erros clássicos. Por exemplo, no caso quântico um erro é identificado medindo a síndrome de um erro e, em seguida, corrigimos, assim como ocorre no caso clássico. Porém, vimos também que há distinções entre as propriedades desses tipos de código. Uma dessas diferenças reside em um código quântico, conhecido como *código degenerado*.

Definição 4.5.1. *Códigos quânticos degenerados* são códigos quânticos em que diferentes erros agem de forma equivalente sobre os estados codificados, de modo que a correção do erro não precisa distinguir entre eles.

A principal característica dos códigos degenerados é que não é possível determinar em qual qubit ocorreu um erro, uma vez que o efeito do erro é o mesmo em diferentes qubits. Essa ideia é ilustrada para o efeito dos erros *phase flip* sobre as palavras-código do código de Shor, conforme dissemos na Observação 4.4.2. Consideremos o efeito dos erros Z_1 e Z_2 nas palavras-código para o código Shor. Como vimos no Exemplo 4.4.4, o efeito desses erros é o mesmo em ambas as palavras-código. Por outro lado em códigos corretores de erros clássicos, erros em bits diferentes necessariamente levam a diferentes palavras-código corrompidas.

Assim, códigos quânticos degenerados têm vantagens e desvantagens. Uma desvantagem é que algumas das técnicas de demonstração utilizadas classicamente em limites para correção de erros falham, porque não podem ser aplicadas a códigos degenerados. Um exemplo disso, será o limitante de Hamming, que veremos na próxima seção.

4.6 LIMITANTES DE HAMMING E SINGLETON

Os limitantes de Hamming e de Singleton são um dos critérios para determinarmos a qualidade de um código. Começamos com o limitante (ou cota) de Singleton. Duas palavras de um $[[n, k, d]]$ código linear clássico devem diferir em qualquer conjunto de coordenadas $n - d + 1$, já que elas devem estar a uma distância de pelo menos d uma da outra.

Proposição 4.6.1 (Limitante de Singleton). *Os parâmetros $[n, k, d]_q$ de um código linear satisfazem a desigualdade*

$$d \leq n - k + 1.$$

Demonstração. Seja H é uma matriz teste de paridade do código $\mathcal{C}$. Então, pela Proposição 2.3.4, a matriz H tem posto $n - k$. Pela Proposição 2.3.2, se a distância mínima de $\mathcal{C}$ é d , então quaisquer $d - 1$ colunas de H devem ser linearmente independentes e devem existir d colunas de H linearmente dependentes. Ademais, como $\text{posto}(H) = n - k$, existe pelo menos um grupo de $n - k$ colunas de H linearmente independentes e qualquer quantidade maior de colunas é linearmente dependente. Logo, podemos ter $d - 1 = n - k$, se quaisquer $n - k$ colunas de H forem linearmente independentes, ou $d - 1 < n - k$ se existir um grupo de $n - k$ colunas de H linearmente dependentes. Por conseguinte, temos que $d - 1 \leq n - k$, isto é, $d \leq n - k + 1$, como queríamos. $\square$

Definição 4.6.2. Um código cujos parâmetros satisfazem uma igualdade no limitante de Singleton é dito um *código separável por distância máxima* ou simplesmente um *código MDS* (Maximum Distance Separable).

Dados um elemento $a \in \mathcal{A}^n$ e um número real $t \geq 0$, definimos o *disco* e a *esfera* de centro a e raio t como sendo, respectivamente, os conjuntos

$$D(a, t) = \{u \in \mathcal{A}^n; d(u, a) \leq t\},$$

$$S(a, t) = \{u \in \mathcal{A}^n; d(u, a) = t\}.$$

Estes conjuntos são finitos e o próximo lema nos fornecerá as suas cardinalidades.

Antes de enunciar o próximo lema, ao longo do texto usaremos a notação usual para a combinação de elementos, dois a dois, ou seja,

$$\binom{n}{i} = \frac{n!}{i!(n-i)!}$$

e denotaremos por $|B|$ o número de elementos de um conjunto finito B .

Lema 4.6.3. Para todo $a \in \mathcal{A}^n$ e todo número natural $r > 0$, temos que

$$|D(a, r)| = \sum_{i=0}^r \binom{n}{i} (q-1)^i.$$

Demonstração. Primeiramente, vamos mostrar que $|S(a, i)| = \binom{n}{i} (q-1)^i$. De fato, sendo

$$S(a, i) = \{u \in \mathcal{A}^n; d(u, a) = i\},$$

então se $u \in S(a, i)$, segue que u possui i entradas que são distintas das respectivas entradas de a . Assim, para determinarmos o número de elementos de $S(a, i)$, pelo princípio multiplicativo, basta conhecermos o número de possibilidades de se escolher i entradas de modo a garantirmos que todas sejam distintas das respectivas entradas em A . Isso nos dá $|S(a, i)| = \binom{n}{i} (q-1)^i$.

Agora, se $i \neq j$, então $S(a, i) \cap S(a, j) = \emptyset$. Com efeito, supondo que existisse $u \in S(a, i) \cap S(a, j)$, deveríamos ter $d(u, a) = i$ e $d(u, a) = j$, ou seja, $i = j$, o que é uma contradição. Além disso,

$$\bigcup_{i=0}^r S(a, i) = D(a, r).$$

Logo,

$$|D(a, r)| = \sum_{i=0}^r |S(a, i)| = \sum_{i=0}^r \binom{n}{i} (q-1)^i,$$

como queríamos. □

Note que, a cardinalidade de $D(a, r)$ depende apenas n (tamanho da maior palavra em $\mathcal{A}$), q (número de elementos de $\mathcal{A}$) e r (raio do disco ou esfera).

Dado um código $\mathcal{C} \subset \mathcal{A}^n$, com distância mínima d , defini-se

$$\kappa = \left\lfloor \frac{d-1}{2} \right\rfloor, \quad (4.10)$$

onde $\lfloor n \rfloor$ representa a **parte inteira** de $n \in \mathbb{R}$.

Lema 4.6.4. *Seja $\mathcal{C} \subset \mathcal{A}^n$ um código com distância mínima d . Se c e c' são palavras distintas de $\mathcal{C}$, então*

$$D(c, \kappa) \cap D(c', \kappa) = \emptyset. \quad (4.11)$$

Demonstração. De fato, suponhamos, por absurdo, que existisse $x \in D(c, \kappa) \cap D(c', \kappa)$. Desta forma, teríamos $d(x, c) \leq \kappa$ e $d(x, c') \leq \kappa$, donde seguiria

$$\begin{aligned} d(c, c') &\stackrel{\text{desigualdade triangular}}{\leq} d(c, x) + d(x, c') \\ &\stackrel{\text{simetria}}{=} \underbrace{d(x, c)}_{\leq \kappa} + \underbrace{d(x, c')}_{\leq \kappa} \\ &= 2\kappa \\ &\leq d-1. \end{aligned}$$

Absurdo, pois $d(c, c') \geq d$, já que d é a distância mínima de $\mathcal{C}$. □

O Teorema 2.1.5 juntamente com as definições acima nos permite traçar uma estratégia para detecção e correção de erros. De fato, sejam $\mathcal{C} \subset \mathcal{A}^n$ um código com distância mínima d e $\kappa = \lfloor \frac{d-1}{2} \rfloor$.

Quando o receptor recebe uma palavra r , uma, e somente uma, das seguintes afirmações é verificada:

1. Existe $c \in \mathcal{C}$ tal que $r \in D(c, \kappa)$, isto é, $d(r, c) \leq \kappa$. Neste caso, substitui-se r por c .
2. Não existe $c \in \mathcal{C}$ tal que $r \in D(c, \kappa)$, isto é, $d(r, c) > \kappa$, para todo $c \in \mathcal{C}$. Neste caso, não é possível decodificar r com uma boa margem de segurança.

Observe que em 1. não se pode ter certeza absoluta de que c tenha sido a palavra transmitida, pois poderíamos ter cometido mais do que κ erros, afastando assim r da palavra transmitida e aproximando-a a outra palavra do código. Esta questão deve ser encarada em termos probabilísticos.

Definição 4.6.5. Sejam $\mathcal{C} \subset \mathcal{A}^n$ um código com distância mínima d e $t = \left\lfloor \frac{d-1}{2} \right\rfloor$. Dizemos que o código $\mathcal{C}$ é **perfeito** quando

$$\bigcup_{c \in \mathcal{C}} D(c, t) = \mathcal{A}^n.$$

Um código perfeito que corrige até t erros, pode corrigir todos os erros de peso menor ou igual a t e nenhum de peso maior que t . Se imaginarmos uma esfera “pequena” (no sentido de raio pequeno) em torno de cada uma das palavras-código de um código, sendo cada esfera de mesmo raio (número inteiro positivo), então dizermos que um código é perfeito é equivalente a dizermos que as esferas de raio t em volta das palavras-código são disjuntas, e que a união de todas essas esferas contém todos os vetores de comprimento n .

Agora, vejamos como se comporta o limitante de Singleton no caso quântico. O limitante quântico de Singleton estabelece um limite para a capacidade dos códigos quânticos de correção de erro. Para a demonstração veja [20].

Proposição 4.6.6. *Qualquer código quântico que codifique k qubits em n qubits e que seja capaz de corrigir erros em quaisquer t qubits deve satisfazer*

$$n \geq 4t + k.$$

Segue-se que o menor código que codifique um único qubit e seja capaz de corrigir um erro arbitrário em um único qubit deve satisfazer

$$n \geq 4 + 1 = 5.$$

Teorema 4.6.7 (Limitante quântico de Singleton). *Seja $[[n, k, d]]$ um código quântico que codifica k qubits lógicos em n qubits físicos e corrige até $t = \left\lfloor \frac{d-1}{2} \right\rfloor$ erros, então os parâmetros do código devem satisfazer a seguinte desigualdade:*

$$n - k \geq 2(d - 1). \quad (4.12)$$

O limitante de Hamming é um critério para determinarmos se um dado código com certas características existe ou não, conforme o seguinte teorema:

Teorema 4.6.8 (Limitante de Hamming). *Se $\mathcal{C}$ um código com parâmetros $[n, M, d]$ que corrige até $t = \left\lfloor \frac{d-1}{2} \right\rfloor$, então*

$$M \left[1 + \binom{n}{1}(q-1) + \binom{n}{2}(q-1)^2 + \cdots + \binom{n}{t}(q-1)^t \right] \leq q^n,$$

em que

$$1 + \binom{n}{1}(q-1) + \binom{n}{2}(q-1)^2 + \cdots + \binom{n}{t}(q-1)^t = \sum_{i=0}^t \binom{n}{i}(q-1)^i.$$

Demonstração. Dado um código $\mathcal{C}$ em $\mathcal{A}^n$ de comprimento n contendo $M = q^k$ palavras-código e que corrige $t = \left\lfloor \frac{d-1}{2} \right\rfloor$ erros, pelo Lema 4.6.4, se $u, u' \in \mathcal{C}$, com $u \neq u'$, então

$$D(u, t) \cap D(u', t) = \emptyset.$$

Ademais, como

$$\bigcup_{u \in \mathcal{C}} D(u, t) \subset \mathcal{A}^n,$$

sendo $\bigcup_{u \in \mathcal{C}} D(u, t) \subset \mathcal{A}^n$ se, e somente se, $\mathcal{C}$ é perfeito, segue que

$$\sum_{u \in \mathcal{C}} |D(u, t)| \leq q^n.$$

Mas, pelo Lema 4.6.3, $|D(u, t)| = \sum_{i=0}^t \binom{n}{i} (q-1)^i$, logo, $\sum_{u \in \mathcal{C}} |D(u, t)| \leq q^n$ implica que

$$\begin{aligned} \sum_{u \in \mathcal{C}} \sum_{i=0}^t \binom{n}{i} (q-1)^i &\leq q^n \Rightarrow \underbrace{q^k}_{=M} \sum_{i=0}^t \binom{n}{i} (q-1)^i \leq q^n \\ \Rightarrow M \left[1 + \binom{n}{1} (q-1) + \binom{n}{2} (q-1)^2 + \cdots + \binom{n}{t} (q-1)^t \right] &\leq q^n, \end{aligned}$$

como queríamos. $\square$

Surge então a seguinte pergunta: como se comporta o limitante de Hamming no caso quântico?

Um ponto negativo é que o limitante quântico de Hamming é aplicável somente a códigos não-degenerados, porém ele aponta para algumas ideias de como devem ser os limites gerais.

Suponhamos que um código não-degenerado seja utilizado para codificar k qubits em n qubits, de modo que ele possa corrigir erros em qualquer subconjunto com até t qubits. Suponhamos também que foram ocorridos j erros, $j \leq t$. Existem $\binom{n}{j}$ possíveis conjuntos de posições onde o erro pode ocorrer. Para cada um desses conjuntos existem três tipos de erros possíveis, X , Y e Z , onde cada um desses erros pode ocorrer em cada qubit. Logo, são ao todo, 3^j erros possíveis, e assim, o número total de erros que podem ocorrer em t ou menos qubits é:

$$\sum_{j=0}^t \binom{n}{j} 3^j.$$

Em particular, $j = 0$ corresponde ao caso em que nenhum erro ocorre em qualquer um dos qubits, que é o caso em que o “erro” I ocorre. Para codificar k qubits de forma não-degenerada, cada um desses erros deve corresponder a um subespaço ortogonal com 2^k dimensões. Como todos esses subespaços devem estar contidos em um espaço de n qubits, com 2^n dimensões, então

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

onde essa desigualdade é chamada **limitante quântico de Hamming**.

Em particular, para $k = 1$ e $t = 1$, o limitante de Hamming é dado por

$$2(1 + 3n) \leq 2^n$$

que só é satisfeita para $n \geq 5$. Portanto, para um código não-degenerado codificar um qubit contra qualquer tipo de erro, ele não pode ter menos de 5 qubits. De fato, existe um código quântico de 5 qubits que atinge o limitante quântico de Hamming.

Como nem todos os códigos quânticos são não-degenerados, o limitante quântico de Hamming é mais útil como uma regra prática para decidirmos sobre a existência de códigos quânticos. Até dado momento, não se conhece códigos que violem o limitante de Hamming, nem mesmo os códigos degenerados.

4.7 CÓDIGOS CSS

Nessa seção veremos um exemplo de códigos quânticos de correção de erros, os **códigos Calderbank-Shor-Steane**, conhecidos simplesmente por **códigos CSS**. Esses códigos são uma ampla classe de códigos quânticos que são construídos a partir dos códigos lineares clássicos, aos quais trabalhamos na Subseção 2.1.1

Definição 4.7.1. Sejam $\mathcal{C}_1$ e $\mathcal{C}_2$ códigos lineares clássicos $[n, k_1]$ e $[n, k_2]$, respectivamente, tais que $\mathcal{C}_2 \subset \mathcal{C}_1$ e $\mathcal{C}_1$ e $\mathcal{C}_2^\perp$ corrigem ambos t erros. Seja $x \in \mathcal{C}_1$ uma palavra-código qualquer. Definimo o estado quântico $|x + \mathcal{C}_2\rangle$ por

$$|x + \mathcal{C}_2\rangle \equiv \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |x + y\rangle,$$

onde $|\mathcal{C}_2|$ denota a cardinalidade do código $\mathcal{C}_2$ e $+$ é a adição realizada módulo 2, coordenada a coordenada. O espaço vetorial gerado pelos estados quânticos $|x + \mathcal{C}_2\rangle$, para todo $x \in \mathcal{C}_1$, é denominado **código CSS** de $\mathcal{C}_1$ sobre $\mathcal{C}_2$ e é denotado por $\text{CSS}(\mathcal{C}_1, \mathcal{C}_2)$.

Seja $x' \in \mathcal{C}_1$ tal que $x - x' \in \mathcal{C}_2$. Então, $|x + \mathcal{C}_2\rangle = |x' + \mathcal{C}_2\rangle$.

De fato, se $x - x' \in \mathcal{C}_2$, então $x' = x - z$, para algum $z \in \mathcal{C}_2$ e

$$|x' + \mathcal{C}_2\rangle = \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |x' + y\rangle = \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |x - z + y\rangle, \text{ com } z \in \mathcal{C}_2.$$

Agora, substituindo $y' = y - z$, como $y, z \in \mathcal{C}_2$, então $y - z \in \mathcal{C}_2$ e $y' = y - z$ também percorre todos os elementos de $\mathcal{C}_2$ à medida que y percorre $\mathcal{C}_2$. Assim,

$$\begin{aligned} |x' + \mathcal{C}_2\rangle &= \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |x - z + y\rangle, \text{ com } z \in \mathcal{C}_2 \\ &= \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y' \in \mathcal{C}_2} |x + y'\rangle = |x \in \mathcal{C}_2\rangle. \end{aligned}$$

Logo, o estado $|x + \mathcal{C}_2\rangle$ depende somente da classe lateral $\mathcal{C}_1/\mathcal{C}_2$ a qual x pertence. Além disso, se x e x' pertencem a diferentes classes laterais de $\mathcal{C}_2$, então para nenhum $y, y' \in \mathcal{C}_2$ vale $x + y = x' + y'$, e como consequência, $|x + \mathcal{C}_2\rangle$ e $|x' + \mathcal{C}_2\rangle$ são estados ortogonais.

O número de classes laterais de $\mathcal{C}_2$ em $\mathcal{C}_1$ é $|\mathcal{C}_1|/|\mathcal{C}_2|$. logo a dimensão de CSS $(\mathcal{C}_1, \mathcal{C}_2)$ é

$$[\mathcal{C}_1 : \mathcal{C}_2] = \frac{|\mathcal{C}_1|}{|\mathcal{C}_2|} = \frac{2^{k_1}}{2^{k_2}} = 2^{k_1 - k_2}. \quad (4.13)$$

Portanto, o código CSS $(\mathcal{C}_1, \mathcal{C}_2)$ têm parâmetros $[n, k_1 - k_2]$. Ademais, como $\mathcal{C}_1$ e $\mathcal{C}_2^\perp$ corrigem t erros, por definição, então o código CSS $(\mathcal{C}_1, \mathcal{C}_2)$ corrige erros em até t qubits, onde esses erros podem ser *bit flip* e *phase flip*. Para mais detalhes sobre a identificação e recuperação de erros, recomendamos [20].

Agora, vamos buscamos um exemplo de código CSS. Mas, para isso, precisamos definir o código de Hamming, o que pode ser visto na seguinte definição:

Definição 4.7.2. Sejam $n = 2^r - 1$, com $r \geq 2$ um inteiro, e H_r a matriz de ordem $r \times n$ cujas colunas são todos vetores não nulos de $\mathbb{F}_2^r$ (corpo finito). Definimos o **código de Hamming** como o código linear $\mathcal{H}_r$ com parâmetros $[2^r - 1, 2^r - r - 1, 3]$ que tem H_r como matriz teste de paridade, isto é,

$$\mathcal{H}_r = \{x \in F_2^n; xH_r^t = 0\}.$$

Exemplo 4.7.3 (Código de Steane). Vamos construir um código CSS usando o código de Hamming da Definição 4.7.2, considerando $r = 3$. Assim, temos um código de Hamming $[7, 3, 4]$, ao qual denotaremos por $\mathcal{C}$. Consideremos $\mathcal{C}_1 \equiv \mathcal{C}$ e $\mathcal{C}_2 \equiv \mathcal{C}^\perp$ e seja

$$H_{\mathcal{C}_1} = \begin{bmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{bmatrix}_{3 \times 7}$$

a matriz teste de paridade de $\mathcal{C}_1$. Como $\mathcal{C}_2 = \mathcal{C}_1^\perp$ e $H_{\mathcal{C}_1}$ é matriz teste de paridade de $\mathcal{C}_1$, então pela Proposição 2.5.2, $G_{\mathcal{C}_2} = H_{\mathcal{C}_1}^t$ é matriz geradora de $\mathcal{C}_2$. Além disso, pela Proposição 2.3.4, se $G_{\mathcal{C}_1}$ é a matriz geradora de $\mathcal{C}_1$, então $G_{\mathcal{C}_1}H_{\mathcal{C}_1}^t = 0$. Fazendo os cálculos, determinamos

$$G_{\mathcal{C}_1} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 \end{bmatrix}_3.$$

³ Perceba que, como estamos em $\mathbb{F}_2 = \{0, 1\}$, então

$$G_{\mathcal{C}_1}H_{\mathcal{C}_1}^t = \begin{bmatrix} 2 & 2 & 2 \\ 2 & 2 & 2 \\ 2 & 2 & 2 \\ 4 & 2 & 2 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

Também, pela Proposição, 2.5.2 como $\mathcal{C}_2 = \mathcal{C}_1^\perp$ e $G_{\mathcal{C}_1}$ é matriz geradora de $\mathcal{C}_1$, então $H_{\mathcal{C}_2} = G_{\mathcal{C}_1}^t$ é matriz teste de paridade de $\mathcal{C}_2$.

Utilizando as matrizes geradoras ou teste de paridade dos códigos $\mathcal{C}_1$ e $\mathcal{C}_2$, determinamos

$$\begin{aligned} \mathcal{C}_1 = \{ & (0, 0, 0, 0, 0, 0, 0), (1, 0, 1, 0, 1, 0, 1), (0, 1, 1, 0, 0, 1, 1), (1, 1, 0, 0, 1, 1, 0), \\ & (0, 0, 0, 1, 1, 1, 1), (1, 0, 1, 1, 0, 1, 0), (0, 1, 1, 1, 1, 0, 0), (1, 1, 0, 1, 0, 0, 1), \\ & (1, 1, 1, 1, 1, 1, 1), (1, 1, 1, 0, 0, 0, 0), (0, 0, 1, 1, 0, 0, 1), (0, 1, 0, 1, 0, 1, 0), \\ & (1, 0, 0, 1, 1, 0, 0), (1, 0, 0, 0, 0, 1, 1), (0, 1, 0, 0, 1, 0, 1), (0, 0, 1, 0, 1, 1, 0) \} \end{aligned}$$

e

$$\begin{aligned} \mathcal{C}_2 = \{ & (0, 0, 0, 0, 0, 0, 0), (1, 0, 1, 0, 1, 0, 1), (0, 1, 1, 0, 0, 1, 1), (1, 1, 0, 0, 1, 1, 0) \\ & (0, 0, 0, 1, 1, 1, 1), (1, 0, 1, 1, 0, 1, 0), (0, 1, 1, 1, 1, 0, 0), (1, 1, 0, 1, 0, 0, 1) \}. \end{aligned}$$

Portanto, $\mathcal{C}_2 \subset \mathcal{C}_1$. Ademais, $\mathcal{C} = (\mathcal{C}^\perp)^\perp = (\mathcal{C}_1^\perp)^\perp = \mathcal{C}_2^\perp$, portanto, ambos os códigos $\mathcal{C}_1$ e $\mathcal{C}_2^\perp$ têm distância mínima $d = 3$, sendo capazes de corrigir erros em até 1 bit.

Como $\mathcal{C}_1$ é um código $[7, 4]$ e $\mathcal{C}_2$ é um código $[7, 3]$, então CSS $(\mathcal{C}_1, \mathcal{C}_2)$ é um código quântico $[7, 1]$, capaz de corrigir erros em 1 único qubit, isto é, um código $[7, 1, 3]$. O código $\mathcal{C}_1$ é definido como sendo o espaço vetorial gerado pelos estados $|x + \mathcal{C}_2\rangle$, para todo $x \in \mathcal{C}_1$. Aqui, temos apenas dois estados, os quais são denotados pelos estados lógicos $|0_L\rangle$ e $|1_L\rangle$, onde

$$\begin{aligned} |0_L\rangle &= |0 + \mathcal{C}_2\rangle = \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |0 + y\rangle \\ &= \frac{1}{\sqrt{8}} [|0000000\rangle + |1010101\rangle + |0110011\rangle + |1100110\rangle + \\ &\quad |0001111\rangle + |1011010\rangle + |0111100\rangle + |1101001\rangle] \end{aligned}$$

e

$$\begin{aligned} |1_L\rangle &= |1 + \mathcal{C}_2\rangle = \frac{1}{\sqrt{|\mathcal{C}_2|}} \sum_{y \in \mathcal{C}_2} |1 + y\rangle \\ &= \frac{1}{\sqrt{8}} [|1111111\rangle + |0101010\rangle + |1001100\rangle + |0011001\rangle + \\ &\quad |1110000\rangle + |0100101\rangle + |1000011\rangle + |0010110\rangle] \end{aligned}$$

onde $0 = (0, 0, 0, 0, 0, 0, 0)$ e $1 = (1, 1, 1, 1, 1, 1, 1)$. Vale $0 \in \mathcal{C}_2$ e se considerarmos qualquer outro $x \in \mathcal{C}_2$, então teremos o mesmo estado $|0_L\rangle$. Analogamente, $1 \in \mathcal{C}_1$ e se considerarmos qualquer outro $x \in \mathcal{C}_1 - \mathcal{C}_2$, teremos o mesmo estado $|1_L\rangle$. O espaço vetorial gerado por $|0_L\rangle$ e $|1_L\rangle$ é chamado de **código de Steane**.

4.8 CÓDIGOS ESTABILIZADORES

Até dado momento, a teoria que apresentamos sobre correção quântica de erros nos fornece critérios para correção de erros. Agora, vamos apresentar um método de construção de tais códigos, vendo uma classe de códigos quânticos que compreende os códigos de Shor e CSS. Esses códigos, desenvolvidos por Gottesman, são chamados *códigos estabilizadores* tendo sua construção baseada nos códigos lineares clássicos, pois fazem uso de operadores que são analogias das matrizes teste de paridade. Para realizarmos esse estudo, esmerilaremos as propriedades de grupo fornecida pelo conjunto dos operadores de Pauli, precisando primeiro definir e entender melhor como funciona esse grupo.

4.8.1 O grupo de Pauli

Definição 4.8.1. Para 1 qubit, definimos o **grupo de Pauli** $\mathcal{G}$ como o conjunto

$$\mathcal{G} = \{\pm I, \pm iI, \pm X, \pm iX, \pm Y, \pm iY, \pm Z, \pm iZ\}$$

munido com a operação de multiplicação matricial. Para n qubits, o **grupo de Pauli geral** $\mathcal{G}_n$ é definido como todas as matrizes formadas pelo produtos tensoriais de ordem n das matrizes de Pauli, com fatores multiplicativos ± 1 e $\pm i$.

No grupo de Pauli geral, podemos desconsiderar os valores imaginário. De fato, apesar de $\sigma_y \in \mathcal{G}_n$ ter componentes imaginárias, se σ_y aparecer um número par de vezes no produto tensorial, os coeficientes serão reais, e se σ_y aparecer um número ímpar de vezes os coeficientes nas palavras-códigos pode ser imaginários. Porém prova-se em [13], que sempre que um código complexo existir, então também existe um código real com os mesmos parâmetros.

Para contornarmos esse problema, consideremos

$$Y \equiv ZX \equiv \sigma_y \equiv \begin{bmatrix} 0 & i \\ -i & 0 \end{bmatrix},$$

Perceba que $Y^2 = -I$. Dessa forma os operadores $\pm I, \pm X, \pm Y$ e $\pm Z$ formam o conjunto $\pm\{I, X, Y, Z\}$ que munidos da operação de multiplicação de matrizes, são um grupo de ordem 8. Logo, o grupo de Pauli de n qubits é o grupo

$$\mathcal{G}_n = \pm\{I, X, Y, Z\}^{\otimes n}$$

de ordem 2^{2n+1} .

Observação 4.8.2. Os elementos de $\mathcal{G}_n$ satisfazem as seguintes propriedades:

1. Todo $M \in \mathcal{G}_n$ é unitário, isto é, $M^{-1} = M^\dagger$;

2. Para todo $M \in \mathcal{G}_n$ vale $M^2 = \pm I \equiv \pm I^{\otimes n}$, onde:

- Se o número de Y 's no produto tensorial for par, então $M^2 = I$ e M é hermitiano ($M = M^\dagger$);
- Se o número de Y 's no produto tensorial for ímpar, então $M^2 = -I$ e M é anti-hermitiano ($M = -M^\dagger$).

3. Quaisquer $M, N \in \mathcal{G}_n$ comutam ou anticomutam, ou seja, $MN = \pm NM$.

4.8.2 Código estabilizador

O grupo de Pauli apresentado na subseção anterior é utilizado para caraterizar os códigos quânticos corretores de erros no seguinte sentido: sejam $\mathcal{H}$ um espaço de Hilbert de dimensão 2^n e $\mathcal{G}_n$ o grupo de Pauli de n qubits. Consideremos S um subgrupo abeliano de $\mathcal{G}_n$.

Definição 4.8.3. Definimos o **código estabilizador** $\mathcal{C}_S \subset \mathcal{H}$ associado a S como o autoespaço simultâneo com autovalor 1 de todos os elementos de S , isto é,

$$\mathcal{C}_S = \{|\psi\rangle; M|\psi\rangle = |\psi\rangle, \text{ para todo } M \in S\}.$$

Como o grupo S preserva todas as palavras-código, dizemos que S é **estabilizador** do código e chamamos seus elementos de **operadores estabilizadores**. Estados quânticos, e consequentemente, códigos quânticos, são descritos de forma mais compacta através dos operadores que o estabilizam.

Não é todo subgrupo S de $\mathcal{G}_n$ que pode ser utilizado como estabilizador para um espaço vetorial não-trivial. Para isso, S deve satisfazer as seguintes condições:

- Os elementos de S devem comutar;
- $-I$ não pode ser um elemento de S .

Vamos descrever o grupo S de maneiras mais compacta, através de seus geradores.

Definição 4.8.4. Dizemos que um conjunto $\{M_1, M_2, \dots, M_l\}$ de elementos independentes de S **gera** S , quando todo elemento de S pode ser escrito com um produto dos elementos $M_1, M_2, \dots, M_l$.

A partir da Definição 4.8.4, podemos concluir que o grupo S pode ser caracterizado por seus geradores $\{M_i\}$. Assim, para ver se um vetor em particular é estabilizado por S , é preciso apenas verificar se o vetor é estabilizado por seus geradores.

Proposição 4.8.5. Se S têm $n - k$ geradores, então o espaço do código $\mathcal{C}_S$ tem dimensão 2^k , ou seja, $\mathcal{C}_S$ codifica k qubits.

Demonstração. Seja $\{M_1, M_2, \dots, M_{n-k}\}$ o conjunto de todos os geradores de S . Para cada $M \in S$, vale $M^2 = I$, já que, caso contrário M teria autovalor diferente de ± 1 . Para todo $M \neq \pm I$ existem pelo menos um $N \in \mathcal{G}_n$ tal que N anticomuta com M , de modo que os autovalores $+1$ e -1 ocorrem com a mesma probabilidade. Tomemos $M = M_1$. Então, dado $|\psi\rangle \in \mathcal{C}_S$, temos $M_1|\psi\rangle = |\psi\rangle$ se, e somente se,

$$M_1 N |\psi\rangle = -N M_1 |\psi\rangle = -N |\psi\rangle.$$

Ou seja, N estabiliza autoestados associados aos autovalores $+1$ e autoestados associados aos autovalores -1 de M_1 com probabilidades iguais. Logo, temos $\frac{1}{2}(2^n) = 2^{n-1}$ estados mutualmente ortogonais tais que $M_1|\psi\rangle = |\psi\rangle$.

Agora, seja $M_2 \in \mathcal{G}_n$ tal que $M_2 \neq \pm I$, M_1 e comuta com M_1 . Existe um $N \in \mathcal{G}_n$ que comuta com M_1 e anticomuta com M_2 . Então, N preserva o autoespaço associado ao autovalor $+1$ de M_1 e, dentro desse espaço, altera na mesma proporção os autoestados associados aos autovalores $+1$ e -1 de M_2 . Logo, o autoespaço satisfazendo

$$M_1|\psi\rangle = M_2|\psi\rangle = |\psi\rangle$$

tem dimensão 2^{n-2} .

Procedendo analogamente, seja $M_j \in \mathcal{G}_n$ independente de $\{M_1, M_2, \dots, M_{j-1}\}$ e que comuta com todo M_i , com $i \in \{1, 2, \dots, j-1\}$. Então, existe um N que comuta com $M_1, M_2, \dots, M_{j-1}$ e anticomuta com M_j . Logo, no espaço $M_1 = M_2 = \dots = M_{j-1}$, M_j tem tanto autovetores associados ao autovalor $+1$ quanto autovalores associados ao autovalor -1 . A cada adição de um gerador, a dimensão dos autoespaços simultâneos é dividida ao meio. Portanto, com $n - k$ geradores, a dimensão do espaço é

$$\left(\frac{1}{2}\right)^{n-k} 2^n = 2^{n-(n-k)} = 2^k.$$

□

Os $n - k$ geradores do estabilizador S funcionam como os **operadores verificação de paridade** de um código. Ou seja, são observáveis que medimos para diagnosticar erros.

Exemplo 4.8.6. Suponhamos que uma informação foi codificada. Se $M_i = 1$, para todo $i \in \{1, 2, \dots, n - k\}$, então não é detectado nenhum erro. Porém, se para algum i , $M_i = -1$, então o estado da informação está contido em um autoespaço ortogonal ao subespaço do código, e como consequência um erro será detectado.

Seja $\varepsilon = \{E_a\} \subset \mathcal{G}_n$ o conjunto de erros que desejamos corrigir. Um operador erro E_a pode ser escrito em termos dos elementos do grupo de Pauli. Para um gerador M do estabilizador S , E_a comuta ou anticomuta com esse M .

- Se E_a comuta com M , então

$$ME_a|\psi\rangle = E_aM|\psi\rangle = E_a|\psi\rangle,$$

para todo $|\psi\rangle \in \mathcal{C}_S$. Logo, o erro E_a preserva o valor $M = 1$.

- Se E_a anticomuta com M , então

$$ME_a|\psi\rangle = -E_aM|\psi\rangle = -E_a|\psi\rangle,$$

para todo $|\psi\rangle \in \mathcal{C}_S$. Logo, o erro troca o valor de M , e pode ser detectado através de uma medida de M .

Por vezes, será bem útil utilizarmos a síndrome de erro para um código estabilizador. Então, vamos definir tal síndrome. Seja S o estabilizador do código $\mathcal{C}_S$ com $n - k$ geradores. Para geradores M e M_i , $i \in \{1, 2, \dots, n - k\}$, e erros E_a , definimos $f_M : \mathcal{G}_n \rightarrow \mathbb{Z}_2$ por

$$f_M(E_a) = \begin{cases} 0, & \text{se } [M, E_a] = 0 \\ 1, & \text{se } \{M, E_a\} = 0 \end{cases} \quad 4$$

e

$$f(E_a) = (f_{M_1}(E_a), f_{M_2}(E_a), \dots, f_{M_{n-k}}(E_a)).$$

Então, $f(E_a)$ é uma sequência binária de $(n - k)$ -bits. Como vimos, para fazer a operação de correção de erro para um código estabilizador, devemos medir os autovalores dos geradores do estabilizador.

Definição 4.8.7. Dado um gerador M do estabilizador S e erros E_a , definimos a **síndrome de erro** como o autovalor de M , isto é,

$$(-1)^{f_M(E_a)}.$$

Se o código for não-degenerado, então os $f_{M_i}(E_a)$ serão distintos, para todo $E_a \in \varepsilon$ e medindo os $n - k$ geradores é possível diagnosticar o erro completamente, isto é, cada erro tem uma síndrome distinta. Assim, nesse caso, a síndrome indica exatamente o erro que ocorreu. Por outro lado, se o código for degenerado, então existem erros distintos com síndromes iguais, e por conseguinte, as síndromes apontam qual conjunto de erros degenerados ocorreu.

Em geral, existe uma condição a ser satisfeita pelo estabilizador que é uma condição suficiente para que um erro seja detectado e corrigido. Essa condição é dada por uma generalização do Teorema 5 de [1], que veremos a seguir.

⁴ Perceba que, $f_M(E_a) = 0$, quando M e E_a comutam e $f_M(E_a) = 1$, quando M e E_a anticomutam.

Dado ε espaço de erros agindo em um espaço de Hilbert $\mathcal{H}$, então um subespaço $\mathcal{C}$ de $\mathcal{H}$ é um código quântico corretor de erros se, e somente se,

$$|\psi\rangle E_a^\dagger E_b |\psi\rangle = \alpha_{ab},$$

para todos $E_a, E_b \in \varepsilon$ e $|\psi\rangle \in \mathcal{C}$, onde α_{ab} é independente de $|\psi\rangle$. Essa afirmação é válida se uma das seguintes condições for satisfeita:

1. $E_a^\dagger E_b \in S$;
2. Existe um $M \in S$ que anticomuta com $E = E_a^\dagger E_b$.

De fato

- $E_a^\dagger E_b \in S$;

Se vale 1, então dado qualquer $|\psi\rangle \in \mathcal{C}_S$, vale $E_a^\dagger E_b |\psi\rangle = |\psi\rangle$, donde

$$\langle \psi | E_a^\dagger E_b | \psi \rangle = \langle \psi | | \psi \rangle = 1.$$

- Existe um $M \in S$ que anticomuta com $E = E_a^\dagger E_b$.

Como $M \in S$, vale $M|\psi\rangle = |\psi\rangle$, para todo $|\psi\rangle \in \mathcal{C}_S$. Temos:

$$\langle \psi | E_a^\dagger E_b | \psi \rangle = \langle \psi | EM | \psi \rangle = -\langle \psi | ME | \psi \rangle = -\langle \psi | E | \psi \rangle = -\langle \psi | E_a^\dagger E_b | \psi \rangle,$$

e assim, $\langle \psi | E_a^\dagger E_b | \psi \rangle = 0$.

Logo, através do que provamos acima, podemos caracterizar um código estabilizador que corrige o conjunto de erros $\{\varepsilon\}$ como um espaço $\mathcal{C}_S$ fixado por um subgrupo abeliano S do grupo de Pauli $\mathcal{G}_n$, onde as condições (1) e (2) são satisfeitas, para todo $E_a^\dagger E_b$, com $E_a, E_b \in \varepsilon$. Também, um o código é não-degenerado quando a condição (1) não é satisfeita para qualquer $E_a^\dagger E_b$.

O subgrupo abeliano S com $n - k$ geradores define 2^{n-k} autoespaços simultâneos, e qualquer um desses autoespaços pode ser escolhido como espaço de codificação. Cada um desses autoespaços gera um código, e todos esses códigos são equivalentes, no sentido de que diferem apenas no rotulamento dos qubits e na escolha da base usada para sua descrição. Portanto, o estabilizador de um código pode ser transformado no estabilizador de outro código por permutação dos qubits, juntamente com um produto tensorial de transformações de um único qubit.

Como o estabilizador S é um subgrupo abeliano de $\mathcal{G}_n$, os operadores de $\mathcal{G}_n$, os operadores de $\mathcal{G}_n$ que não comutam com todos os $M_i \in S$ levam $\mathcal{C}_S$ em seu complemento ortogonal, conseqüentemente são erros detectáveis. Todavia, existem muitos elementos de $\mathcal{C}_S$ que comutam com todo $M \in S$ mas não pertencem a S . Operadores com esta

características preservam o subespaço de codificação $\mathcal{C}_S$, porém não agem trivialmente sobre ele, e como consequência, podem corromper a informação. Por exemplo, suponhamos que $E_a^\dagger E_b$ comuta com todo elemento de S , mas que $E_a^\dagger E_b \notin S$. Como $E_a|\psi\rangle$ e $E_b|\psi\rangle$ têm a síndrome, podemos interpretar erroneamente um erro E_a como um erro E_b . Assim, o efeito do erro junto com a tentativa em corrigi-lo nos fazendo aplicar $E_a^\dagger E_b$ ao dado, o que pode causar dano.

Definição 4.8.8. Dado um operador de Pauli seu *peso* como o número de fatores diferentes de I no tensor.

Um código estabilizador com distância d possui a propriedade de que cada $E \in \mathcal{G}_n$ com peso menor que d pertence ao estabilizador ou anticomuta com algum elemento dele. Assim, se o estabilizador não contiver um elemento de peso menor do que d , o código é considerado não-degenerado. Para que um código seja capaz de corrigir t erros, sua distância mínima deve ser pelo menos $d = 2t + 1$. Além disso, um código com distância $s + 1$ pode detectar até s erros ou corrigir s erros em locais conhecidos [22].

5 CÓDIGOS QUÂNTICOS COLORIDOS TOPOLÓGICOS

Nosso principal objetivo neste capítulo será fazer uma apresentação de uma teoria que une códigos quânticos com topologia, apresentando uma classe de códigos quânticos topológicos, os códigos coloridos que é um exemplo de código de superfície que surgiu a partir do código tórico de Kitaev [18] e que é considerada um operador local, já que todos os geradores estabilizadores agem apenas em um pequeno número de qubits próximos.

5.1 CONCEITOS PRELIMINARES

Definição 5.1.1. Definimos uma **tesselação** sobre uma superfície M fixada como sendo a cobertura de todo M por figuras, de modo que não existam espaços e nem sobreposições entre essas figuras.

Definição 5.1.2. Dada uma superfície M , consideremos uma tesselação qualquer sobre M . Os vértices são ditos **0-cells**, as arestas são ditas **1-cells** e as faces são ditas **2-cells**.

Definição 5.1.3. Consideremos uma tesselação qualquer sobre uma superfície M e A_i o conjunto de todas as i -cells, $i = 0, 1, 2$. Para cada $A \subset A_i$, definimos uma **i -cadeia** em M como a soma formal finita

$$c = \sum_i c_i p_i; \quad c_i = \begin{cases} 0; & \text{se } p_i \notin A \\ 1; & \text{se } p_i \in A \end{cases}.$$

Para cada $i = 0, 1, 2$, somando as i -cadeias obtemos novas i -cadeias, onde esta soma é feita módulo 2. Assim, usando essa operação modular, formamos um grupo abeliano aditivo C_i formado pelas i -cadeias, onde o elemento neutro corresponde ao conjunto vazio. A partir desses grupos abelianos definimos três homomorfismos de grupos ∂_i , denominados por **operadores bordo**, que levam uma i -cadeia na soma das fronteiras de todas as i -cells que fazem parte dessa i -cadeia, onde

$$\partial_0 : C_0 \rightarrow \{0\}, \quad \partial_1 : C_1 \rightarrow C_2, \quad \partial_2 : C_2 \rightarrow C_1.$$

Definição 5.1.4. Definimos o **núcleo** de ∂_1 e a **imagem** de ∂_2 como os subgrupos de C_1 dados, respectivamente, por

$$Z_1 = \text{nuc } \partial_1 = \{z \in C_1; \partial_1 z = 0\}$$

$$B_1 = \text{Im } \partial_2 = \{b \in C_1; \exists c \in C_2 \text{ tal que } b = \partial_2 c\}.$$

Os elementos de Z_1 são chamados de **ciclos**.

Definição 5.1.5 (Característica de Euler). Dada uma tesselação qualquer, consideremos V , E e F o **número de vértices**, **arestas** e **faces**, respectivamente. Definimos a **característica de Euler** como a quantidade

$$\chi \equiv V - E + F. \tag{5.1}$$

O gênero g pode ser definido, de maneira informal, como o número de “alças” ou “buracos” de uma superfície. Trata-se de um invariante topológico fundamental, utilizado para distinguir diferentes classes de superfícies. Uma descrição mais rigorosa da forma normal das superfícies pode ser encontrada em [29].

Uma outra relação que utilizaremos é

$$\chi = 2(1 - g), \quad (5.2)$$

onde g é o gênero de uma superfície orientável e conexa.

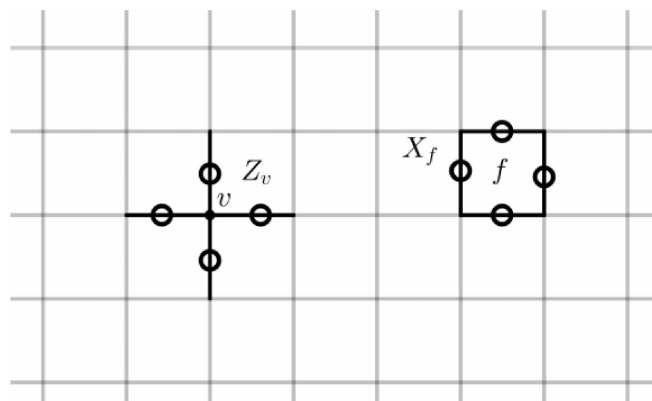
Definição 5.1.6. Sejam v um vértice e f uma face da tesselação fixada na superfície fechada. Definimos o **operador face** X_f e o **operador vértice** Z_v como os operadores de Pauli dados por

$$X_f \equiv \prod_{e \in \partial_2 f} X_e \quad \text{e} \quad Z_v \equiv \prod_{e; v \in \partial_1 e} Z_e,$$

onde X e Z são as matrizes de Pauli e consideramos $\partial_2 f$ e $\partial_1 e$ como conjuntos.

Observe que o operador face X_f consiste de um produto tensorial do operador de Pauli X agindo nos qubits que estão nas arestas pertencentes à face f , com a identidade agindo nos demais qubits. Já o operador vértice Z_v consiste de um produto tensorial do operador de Pauli Z agindo nos qubits que estão nas arestas adjacentes ao vértice v , com a identidade agindo nos demais qubits. Além disso, as arestas correspondem aos qubits. Na figura abaixo podemos ver a ação dos operadores face X_f e vértice Z_v agindo nos qubits que formam a face f e nos qubits adjacentes ao vértice v , respectivamente.

Figura 3 – Representação dos qubits e operadores X_f e Z_v .



Fonte: [5]

Os operadores face comutam entre si, assim como os operadores vértice também comutam entre si. Pelo Exemplo .034 do anexo, os operadores X e Z anticomutam.

Todavia, X_f e Z_v terão nenhuma ou terão duas arestas em comum, logo, os operadores face e vértice sempre comutarão, ou seja, $[X_f, Z_v] = 0^1$. Logo, pelo Teorema da diagonalização simultânea (Teorema .0.33), X_f e Z_v são simultaneamente diagonalizáveis, e portanto eles têm os mesmos autovalores iguais $+1$. Sendo assim, $-I \notin X_f$ e $-I \notin Z_v$.

5.2 CÓDIGOS COLORIDOS

Os **códigos coloridos** são uma classe de códigos quânticos topológicos desenvolvidos por Héctor Bombín e Miguel A. Martin-Delgado em 2007, para superfícies de gênero $g = 0$ e $g = 1$, o que pode ser visto melhor em [4]. Tais códigos são construídos sobre tesselações trivalentes e 3-colorível sobre uma superfície bidimensional.

Definição 5.2.1. Dizemos que uma tesselação é **trivalente** quando em cada vértice da tesselação encontram-incidem 3 arestas. Dizemos que uma tesselação é **3-colorível** quando é possível cobrir todas as faces da tesselação usando apenas 3 cores, de modo que as faces vizinhas que compartilham uma mesma aresta (adjacente) não tenham a mesma cor.

Para os códigos quânticos coloridos, as cores que utilizaremos serão vermelho (R), verde (G) e azul (B). Para cada aresta atribuímos uma cor, como por exemplo, as arestas verdes são as arestas adjacentes às faces azuis e vermelhas, isto é, arestas vermelhas não pertencem à faces verdes e, de modo análogo para arestas verdes e azuis.

Nos códigos coloridos os qubits são anexados aos vértices da tesselação. Os operadores do grupo estabilizador correspondem aos chamados operadores de face, sendo dois para cada face f : um do tipo X e outro do tipo Z . Cada um desses operadores atua sobre os qubits associados às arestas (ou vértices, dependendo da convenção adotada) pertencentes à face f . Para cada face f da tesselação, denotamos tais operadores por B_f^σ , com $\sigma = X, Z$. Na Figura 4 podemos visualizar uma tesselação hexagonal trivalente e 3-colorível no toro, onde os lados opostos são identificados.

Consideremos F_R , F_G e F_B os **conjuntos das faces vermelhas, verdes e azuis**, respectivamente. Temos

$$\prod_{f \in F_R} B_f^\sigma = \prod_{f \in F_G} B_f^\sigma = \prod_{f \in F_B} B_f^\sigma,$$

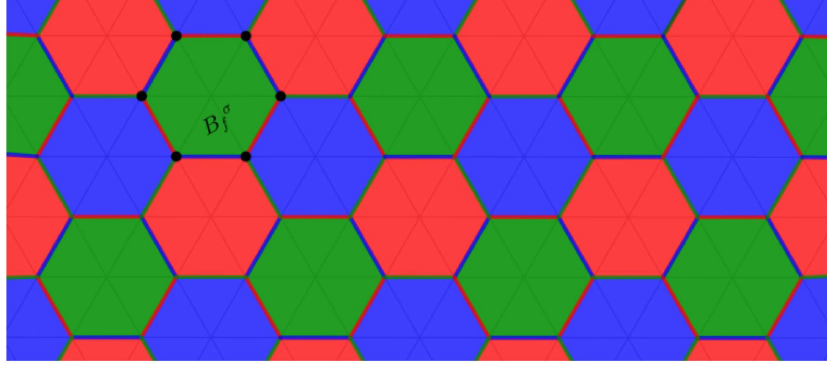
com $\sigma = X, Z$. Assim, quatro desses operadores não são independentes. Logo, o número de geradores estabilizadores independentes é dado por

$$r = 2(|F_R| + |F_G| + |F_B|) - 4, \quad (5.3)$$

pois temos dois operadores para cada face e um total de $|F_R| + |F_G| + |F_B|$ faces, e, como 4 operadores não são independentes, descartamos 4.

¹ Definimos o comutador no anexo, na Definição .0.31

Figura 4 – Operadores face B_f^g agindo nos qubits da face verde f .



Fonte: [5]

Definição 5.2.2. Chamamos de *tesselações reduzidas* as três tesselações construídas a partir de uma tesselação do código colorido, as quais são marcadas pela cor das faces reduzidas.

Obtemos a tesselação reduzida vermelha da seguintes forma:

1. Cada face vermelha corresponde a um vértice da tesselação reduzida, que em certo sentido, estamos reduzindo cada face vermelha a um único ponto.
2. As arestas da tesselação reduzida correspondem a dois vértices da tesselação e são as arestas que conectam faces vermelhas da tesselação, ou seja, arestas vermelhas.
3. As faces verdes e azuis da tesselação são as faces da tesselação reduzida.

Analogamente, obtemos a tesselação reduzida verde e azul.

Desse modo, a quantidade de faces vermelhas F_R da tesselação é igual a quantidade de vértices V da tesselação reduzida, ou seja, $|F_R| = V$. Ademais, $|F_G| + |F_B| = F$, onde F é a quantidade de faces da tesselação reduzida. Então, pela expressão 5.3, temos:

$$\begin{aligned} r &= 2(|F_R| + |F_G| + |F_B|) - 4 = 2(V + F) - 4 \\ &= 2(V + F - 2). \end{aligned} \tag{5.4}$$

O número de qubits é o dobro do número de arestas E da tesselação reduzida, isto é, $n = 2E$. Agora, como a quantidade de geradores independentes de um código estabilizador $[n, k]$ é igual a $n - k$, segue que $n - k = r$, ou seja, $k = n - r$. Logo, usando

as expressões da característica de Euler dadas em (5.1) e (5.2), e também (5.4), temos:

$$\begin{aligned}
 k &= n - r = \stackrel{n=2E \text{ e } (5.4)}{=} 2E - 2(V + F - 2) = 2(E - V - F + 2) \\
 &= 2[2 - (V - E + F)] \stackrel{(5.1)}{=} 2(2 - \chi) \stackrel{(5.2)}{=} 2(2 - (2(1 - g))) \\
 &= 4g.
 \end{aligned}$$

Logo, a quantidade k de qubits codificados depende apenas do gênero da superfície considerada, sendo independente da tesselação.

6 CONCLUSÃO

Neste trabalho, buscamos fornecer uma visão abrangente sobre os códigos corretores de erros, principalmente sobre os códigos quânticos corretores de erros, destacando sua importância fundamental na viabilização e no avanço das tecnologias quânticas. Exploramos como esses códigos desempenham um papel crucial na mitigação de erros inerentes aos sistemas quânticos, contribuindo para a estabilidade e confiabilidade dos processos computacionais e de comunicação quântica, onde todos esses fatores são essenciais para o desenvolvimento de aplicações inovadoras nessa área.

Em trabalhos futuros, com a teoria estudada ao longo desta dissertação, seria possível explorar uma construção dos códigos quânticos coloridos para superfícies compactas com $g \geq 2$. Isso foi feito por Waldir Silva Soares e Eduardo B. Silva em 2018, e pode ser visto em [5] e [25], onde [5] segue a ideia utilizada por [1], para gerar códigos quânticos topológicos em superfícies compactas com $g \geq 2$. Como essa é uma teoria muito recente, há ainda espaço para novas aplicações e/ou construções que refinem os códigos já conhecidos.

REFERÊNCIAS

- 1 ALBUQUERQUE, C. D. **Análise e construção de códigos quânticos topológicos sobre variedades bidimensionais**. Campinas: UNICAMP, 2009.
- 2 AMARAL, B.; BARAVIERA, A. T.; CUNHA, M. T. **Mecânica quântica para matemáticos em formação**. Impa-28th Colóquio Brasileiro de Matemática, 2011.
- 3 BENNETT, C. H. et al. **Teleporting an unknown quantum state via dual classical and Einstein–Podolsky–Rosen channels**. Physical Review Letters, v. 70, n. 13, p. 1895–1899, 1993.
- 4 BOMBIN, H.; MARTIN-DELGADO, M. A. **Topological Quantum Distillation**. Physical Review Letters, v. 97, n. 18, p.180501, 2006.
- 5 BRIZOLA, E. M. **Códigos quânticos coloridos em superfícies compactas com gênero $g \geq 2$** . Maringá: UEM, 2019.
- 6 CALDERBANK, A. R., SHOR, P. W. **Good quantum error-correcting codes exist**. Physical Review A, 54, pp. 1098-1105, 1996.
- 7 CUNHA, R. C. **Códigos corretores de erros e um exemplo de aplicação na biologia**. Trabalho de Conclusão de Curso de Bacharelado em Matemática - Universidade Federal de Juiz de Fora, Juiz de Fora. 2023.
- 8 DENNIS, E. et al. **Topological quantum memory**. Journal Of Mathematical Physics, p. 4452-4505. 2002
- 9 DIRAC, P. A. M. **The Principles of Quantum Mechanics**. 4th ed. Oxford: Clarendon Press, 1958.
- 10 GARCIA, A.; LEQUAIN, Y. **Elementos de Álgebra**. Projeto Euclides. Editora IMPA, Rio de Janeiro, 2014.
- 11 GONÇALVES, A. **Introdução à álgebra**. Instituto de Matemática Pura e Aplicada, 2017.
- 12 GOTTESMAN, D. **Class of quantum error-correcting codes saturating the quantum Hamming bound**. Physical Review A, 54, pp. 1862, 1996.
- 13 GOTTESMAN, D. **Stabilizer code and quantum error correction**, 1997.
- 14 HEFEZ, A.; VILLELA, M. L. T. **Códigos Corretores de Erros**. 2. ed. Rio de Janeiro: IMPA, 2008.
- 15 HERSTEIN, I. N. **Topics in algebra**. John Wiley e Sons, 1991.
- 16 JORIO, A.; FROSSARD, J. V. **Material de Estudos para Mecânica Quântica**. Programa de Pós-Graduação em Física, UFMG, 2019.
- 17 LEHNER, J. **A short course in automorphic functions**. Courier Corporation, 2014.

- 18 KITAEV, A. Yu. **Fault-tolerant quantum computation by anyons**. 1997, *Annals of Physics*, 303(1), 2–30. <https://doi.org/10.1006/aphy.2002.6227>
- 19 LEITÃO, M. R. **Tesselações no ensino de geometria Euclidiana**. 2015. Dissertação (Mestrado Profissional em Matemática em Rede Nacional) – Universidade Federal do Ceará, Juazeiro do Norte, 2015.
- 20 NIELSEN, M. A.; CHUANG, I. L. **Quantum Computation and Quantum Information**. Cambridge: Cambridge University Press, 2010.
- 21 PORTUGAL, R. **Códigos Quânticos**. Notas do minicurso do 1o Encontro de Teoria dos Códigos e Criptografia, UFABC, 2010. Disponível em: <https://www.lncc.br/portugal/CodigosQuanticos.pdf>. Acesso em: 4 fev. 2025.
- 22 PRESKILL, J. **Lectures notes for physics 219: quantum error correction**. California Institute of Technology, 1999.
- 23 SHOR, Peter W. **Scheme for Reducing Decoherence in Quantum Computer Memory**. *Physical Review A*, v. 52, n. 4, p.R2493, 1995.
- 24 SILVA, E. Elísio Física. Disponível em: <http://elisiofisica.blogspot.com/2011/03/notacao-de-dirac-introducao.html>. Acesso em: 5 fev. 2025
- 25 SOARES, W. S.; SILVA, E. B. **Hyperbolic quantum color codes**. 2018. Disponível em: <https://arxiv.org/abs/1804.06382>. Acesso em: 4 fev. 2025.
- 26 SOUZA, R. de C. **Portas lógicas quânticas universais para o grau de liberdade de caminho da luz**. Universidade Federal Fluminense, Volta Redonda, 2021. Disponível em: <https://www.portal.if.uff.br/posgrad/wp-content/uploads/sites/3/2021/10/Raiane-Carvalho-de-Souza-M.pdf>. Acesso em: 29 jan. 2025.
- 27 STEANE, A. M. **Error correcting codes in quantum theory**. *Phys. Rev. Letters*, 77, pp. 793, 1996.
- 28 STEANE, A. M. **Multiple particle interference and quantum error correction**. *Proc. R. Soc. Lond. A*, 452, pp. 2551-2577, 1996.
- 29 STILLWELL, J. **Geometry of Surfaces**. Springer-Verlag, 2000.

APÊNDICE A – Álgebra Linear e a Notação de Dirac

Os estados $|0\rangle$ e $|1\rangle$ também podem ser representados na forma matricial como

$$|0\rangle = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \quad \text{e} \quad |1\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix}.$$

Escrevendo dessa forma, a superposição $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$ pode ser escrita na forma

$$|\psi\rangle = \begin{bmatrix} \alpha \\ \beta \end{bmatrix}.$$

Essa representação matricial é muito conveniente para compreendermos mais facilmente as operações que estão sendo realizadas.

O objetivo aqui será apresentar alguns conceitos de álgebra linear (que serão utilizados ao longo do texto) já na Notação de Dirac. Apresentaremos os conceitos de espaço vetorial, base e dimensão e subespaços vetoriais. Ademais, também definiremos operadores lineares e as principais características relacionadas a eles, além das operações de produto interno e produto tensorial.

Definição .0.1. Um **espaço vetorial** sobre $\mathbb{K}$ (corpo) (ou um $\mathbb{K}$ -**espaço vetorial**) é um conjunto V munido das operações de adição e multiplicação por escalar

$$\begin{aligned} + : V \times V &\rightarrow V & \cdot : \mathbb{K} \times V &\rightarrow V \\ (|u\rangle, |v\rangle) &\mapsto |u\rangle + |v\rangle & (\alpha, |v\rangle) &\mapsto \alpha|v\rangle \end{aligned} \quad \text{e}$$

as quais devem satisfazer as seguintes propriedades:

$$(V1) \quad (|u\rangle + |v\rangle) + |w\rangle = |u\rangle + (|v\rangle + |w\rangle), \quad \forall |u\rangle, |v\rangle, |w\rangle \in V$$

$$(V2) \quad \exists 0 \in V; |v\rangle + 0 = 0 + |v\rangle = |v\rangle, \quad \forall |v\rangle \in V.$$

$$(V3) \quad \forall |v\rangle \in V, \exists -|v\rangle \in V; |v\rangle + (-|v\rangle) = 0.$$

$$(V4) \quad |u\rangle + |v\rangle = |v\rangle + |u\rangle, \quad \forall |u\rangle, |v\rangle \in V.$$

$$(V5) \quad \alpha(\beta|v\rangle) = (\alpha\beta)|v\rangle, \quad \forall \alpha, \beta \in \mathbb{K}, \forall |v\rangle \in V.$$

$$(V6) \quad 1|v\rangle = |v\rangle, \quad \forall |v\rangle \in V.$$

$$(V7) \quad \alpha(|u\rangle + |v\rangle) = \alpha|u\rangle + \alpha|v\rangle, \quad \forall \alpha \in \mathbb{K}, \forall |u\rangle, |v\rangle \in V.$$

$$(V8) \quad (\alpha + \beta)|v\rangle = \alpha|v\rangle + \beta|v\rangle, \quad \forall \alpha, \beta \in \mathbb{K}, \forall |v\rangle \in V.$$

Os elementos de $\mathbb{K}$ serão chamados de **escalares**, e os elementos do espaço vetorial V serão chamados de **vetores**. Vale ressaltarmos que os elementos $\lambda \cdot |v\rangle$ serão denotados simplesmente por $\lambda|v\rangle$, conforme a Definição .0.1.

Observação .0.2. Vale salientarmos que, no presente texto o corpo sempre será $\mathbb{C}$, o corpo dos números complexos, salvo menção explícita em contrário. Ademais, V denotará sempre um espaço vetorial sobre $\mathbb{C}$, a menos que mencionemos o contrário.

Definição .0.3. Uma expressão do tipo

$$\alpha_1|v_1\rangle + \alpha_2|v_2\rangle + \cdots + \alpha_n|v_n\rangle,$$

com $\alpha_1, \alpha_2, \dots, \alpha_n \in \mathbb{C}$ é chamada uma **combinação linear** dos vetores $|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle \in V$. Dado um conjunto S de vetores dizemos que ele **gera** V , e escrevemos $V = \text{ger}(S)$ se todo elemento de V pode ser escrito como combinação linear dos elementos de S .

Definição .0.4. Dizemos que um conjunto de vetores $S = \{|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle\} \subset V$ é **linearmente independente (LI)** se a equação

$$\alpha_1|v_1\rangle + \alpha_2|v_2\rangle + \cdots + \alpha_k|v_n\rangle = 0$$

admite somente a solução trivial $(\alpha_1, \alpha_2, \dots, \alpha_n) = (0, 0, \dots, 0)$. Caso contrário, dizemos que S é **linearmente dependente (LD)**.

Definição .0.5. Uma **base** para um espaço vetorial V é um conjunto LI $\mathcal{B} = \{|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle\}$ tal que todo vetor de V pode ser escrito como uma combinação linear dos elementos de $\mathcal{B}$. Definimos a **dimensão** de V como o número de vetores em uma base¹.

É possível provar que duas bases de V devem ter necessariamente o mesmo número de elementos, de modo que a dimensão está bem definida, isto é, não depende da base que escolhemos para V .

Da existência de uma base $\mathcal{B} = \{|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle\}$ do espaço V surge a notação para vetores mais utilizada: dado um vetor $|v\rangle$ sempre podemos escrevê-lo como

$$|v\rangle = a_1|v_1\rangle + a_2|v_2\rangle + \cdots + a_n|v_n\rangle$$

e de forma única. Com efeito, supondo que $|v\rangle = b_1|v_1\rangle + b_2|v_2\rangle + \cdots + b_n|v_n\rangle$, teremos

$$a_1|v_1\rangle + a_2|v_2\rangle + \cdots + a_n|v_n\rangle = b_1|v_1\rangle + b_2|v_2\rangle + \cdots + b_n|v_n\rangle$$

$$\Rightarrow (a_1 - b_1)|v_1\rangle + (a_2 - b_2)|v_2\rangle + \cdots + (a_n - b_n)|v_n\rangle = 0$$

e como $\mathcal{B}$ é base de V , da condição LI, devemos ter $a_i - b_i = 0$, ou seja, $a_i = b_i$, para cada $i \in \{1, 2, \dots, n\}$. Assim podemos representar o vetor por meio de seus coeficientes na base dada: $|v\rangle = (a_1, a_2, \dots, a_n)_{\mathcal{B}}$. Quando não houver confusão a respeito da base que está sendo utilizada denotaremos simplesmente $|v\rangle = [v]_{\mathcal{B}} = (a_1, a_2, \dots, a_n)$.

¹ A definição de dimensão aqui dada vale para espaços vetoriais que podem ser gerados por um número finito de vetores. Em outros casos, o espaço vetorial não pode ser gerado por nenhum conjunto finito e dizemos que esse espaço vetorial possui **dimensão é infinita**.

Definição .0.6. Sejam V um espaço vetorial sobre $\mathbb{K}$ e $W \subseteq V$. Dizemos que W é um **subespaço vetorial** (ou simplesmente **subespaço**) de V quando

- $W \neq \emptyset$.
- W é fechado para as operações de V , isto é,
 - (a) Se $|u\rangle, |v\rangle \in W$, então $|u\rangle + |v\rangle \in W$.
 - (b) Se $a \in \mathbb{K}$ e $|v\rangle \in W$, então $a|v\rangle \in W$.
- W é um espaço vetorial com as operações de V .

Proposição .0.7. *Seja $W \subseteq V$, sendo V um espaço vetorial. São equivalentes:*

1. W é subespaço de V .
2. $0 \in W$ e W é fechado para as operações de V .
3. $W \neq \emptyset$ e $a|u\rangle + |v\rangle \in W$ para todos $a \in \mathbb{K}$ e $|u\rangle, |v\rangle \in W$.

Definição .0.8. Uma aplicação $A : V \rightarrow W$ entre os espaços vetoriais V e W é dita um **operador linear** quando

$$A\left(\sum_i a_i |v_i\rangle\right) = \sum_i a_i A(|v_i\rangle),$$

para todos $s \in \mathbb{N}^*$, $a_1, a_2, \dots, a_s \in \mathbb{K}$ e $|v_1\rangle, |v_2\rangle, \dots, |v_s\rangle \in V$.

Dizemos que um operador linear A é **definido no espaço vetorial** V , se A é um operador linear de V em V .

Observação .0.9. Algumas observações importantes:

- Aqui, vamos denotar $A(|v\rangle)$ por $A|v\rangle$.
- Quando dissermos que $A : V \rightarrow W$ é um operador linear, estaremos deixando implícito que V e W são espaços vetoriais (sobre o mesmo corpo).
- Definimos o **operador identidade** $I_V : V \rightarrow V$ por $I_V|v\rangle \equiv |v\rangle$.
- A **composição** de dois operadores lineares $A : U \rightarrow V$ e $B : V \rightarrow W$ é dada por

$$(B \circ A)|v\rangle \equiv (BA)|v\rangle = B(A|v\rangle) = BA|v\rangle.$$

² Equivalentemente, $A(a_1|v_1\rangle + a_2|v_2\rangle + \dots + a_s|v_s\rangle) = a_1A(|v_1\rangle) + a_2A(|v_2\rangle) + \dots + a_sA(|v_s\rangle)$.

Uma maneira equivalente e muito útil para descrevermos os operadores lineares, se dá através de suas representações matriciais. Seja $A : V \rightarrow W$ um operador linear. Sejam $\{|v_1\rangle, |v_2\rangle, \dots, |v_m\rangle\}$ e $\{|w_1\rangle, |w_2\rangle, \dots, |w_n\rangle\}$ bases de V e W , respectivamente. Então, para cada $j \in \{1, 2, \dots, m\}$, existem $A_{1j}, A_{2j}, \dots, A_{nj} \in \mathbb{C}$ tais que

$$A|v_j\rangle = \sum_i A_{ij}|w_i\rangle. \quad (.1)$$

Assim, a **matriz** associada ao operador A , denotada por $A = [A_{ij}]$, é definida como a matriz $m \times n$ cujas entradas são os valores A_{ij} .

Definição .0.10. Um **funcional linear** $\langle \chi |$ é um operador linear $\chi : V \rightarrow \mathbb{C}$, onde V é um espaço vetorial. O **espaço dual** é o espaço de todos os funcionais lineares de V e é denotado por V^* .

Definição .0.11. Seja $|v\rangle$ um vetor no espaço vetorial $\mathbb{C}^n$. Definimos o **dual** de $|v\rangle$, denotado por $\langle v|$, como o vetor transposto de $|v\rangle$ cujos elementos são trocados pelos seus conjugados. Ou seja,

$$\langle v| = (|v\rangle)^\dagger,$$

onde $\dagger$ significa a conjugação transposta. Esse vetor de $(\mathbb{C}^n)^*$, na notação de Dirac, é conhecido como **bra**.

Definição .0.12. Dado um espaço vetorial V , um **produto interno** é uma aplicação

$$\begin{aligned} \langle \cdot | \cdot \rangle : V \times V &\rightarrow \mathbb{C} \\ (|u\rangle, |v\rangle) &\mapsto \langle u|v\rangle \end{aligned}$$

tal que, para todos $|u\rangle, |v\rangle, |w\rangle \in V$ e $\lambda, \mu \in \mathbb{C}$ valem:

1. $\langle \lambda u + \mu v | w \rangle = \lambda^* \langle u | w \rangle + \mu^* \langle v | w \rangle;$
2. $\langle u | v \rangle = (\langle v | u \rangle)^{*4};$
3. $\langle u | u \rangle \geq 0;$
4. Se $\langle u | u \rangle = 0$, então $|u\rangle = 0$.

Definição .0.13. Sejam V um espaço vetorial e $|v\rangle, |w\rangle \in V$. Definimos o **produto interno** $\langle \cdot | \cdot \rangle : V \times V \rightarrow \mathbb{C}$ por

$$\langle v | w \rangle = |v\rangle^\dagger |w\rangle. \quad (.2)$$

³ Perceba que aqui estamos aplicando A a cada vetor da base de V e escrevendo o resultado obtido como uma combinação linear dos vetores da base de W .

⁴ O asterisco denota a operação de tomar o complexo conjugado do número

Perceba que a expressão .2 equivale ao produto matricial do vetor $|v\rangle$ transposto-conjugado, denotado por $|v\rangle^\dagger$, por $\langle v|$.

O produto interno nos permite introduzir uma noção que generaliza a um espaço vetorial qualquer a ideia de perpendicularidade no espaço, com a qual já estamos familiarizados

Definição .0.14. Dizemos que dois vetores $|u\rangle$ e $|v\rangle$ são **ortogonais** quando $\langle u|v\rangle = 0$. Dizemos que um conjunto $E = \{|v_1\rangle, |v_2\rangle, \dots, |v_k\rangle\}$ é **ortogonal** se seus elementos são dois a dois ortogonais. Dizemos que um conjunto $E = \{|v_1\rangle, |v_2\rangle, \dots, |v_k\rangle\}$ é **ortonormal** se é ortogonal e $\langle v_i|v_i\rangle = 1$, para todo i .

Definição .0.15. A **norma** de um estado $|v\rangle$ é definida por

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

Definiremos agora os espaços de interesse para a computação gráfica e informação quântica.

Definição .0.16. Um espaço vetorial complexo de dimensão finita munido de produto interno é dito um **espaço de Hilbert**.

Observação .0.17. • Um **subespaço** W de um espaço de Hilbert V , também é um espaço de Hilbert.

- Definimos o **complemento ortogonal** de W em V , denotado por $W^\perp$ como o conjunto de vetores ortogonais a todos os vetores de W .
- O espaço vetorial V pode ser escrito como soma direta de W e $W^\perp$, isto é,

$$V = W \oplus W^\perp.$$

Operadores lineares também podem ser representados através de um produto interno usando a **representação de produto externo**, da seguinte maneira: sejam V e W espaços vetoriais com produto interno, $|v\rangle \in V$ e $|w\rangle \in W$. O **produto externo** de $|v\rangle$ e $|w\rangle$ é operador linear $|w\rangle\langle v| : V \rightarrow W$ dado por

$$(|w\rangle\langle v|)|v'\rangle \equiv |w\rangle\langle v|v'\rangle = \langle v|v'\rangle|w\rangle,$$

para cada $|v'\rangle \in V$.

Definição .0.18. Definimos um **autovetor** de um operador linear A em V como um vetor não-nulo $|v\rangle \in V$ tal que

$$A|v\rangle = \lambda|v\rangle,$$

onde $\lambda \in \mathbb{C}$ é chamado de **autovalor** de A , correspondente a $|v\rangle$.

O **autoespaço** correspondente a um autovalor λ é o conjunto de vetores com autovalores λ acrescido do vetor nulo⁵.

Observação .0.19. Algumas propriedades sobre autovetores e autovalores:

1. Qualquer múltiplo de um autovetor também é um autovetor;
2. Se dois autovetores estão associados ao mesmo autovalor, então qualquer combinação linear desses autovetores é um autovetor;
3. O número de autovetores linearmente independentes associados a um mesmo autovalor é a **multiplicidade** desses autovalor.

Definição .0.20. Seja A um operador no espaço vetorial V . Uma **representação diagonal** de A é uma representação

$$A = \sum_i \lambda_i |i\rangle \langle i|,$$

onde os vetores $|i\rangle$ formam um conjunto de autovetores ortogonais de A , com autovalores correspondentes λ_i . Dizemos que A é **diagonalizável** se A possui uma representação diagonal.

Seja V um espaço vetorial. Assumindo a existência de uma base qualquer para V , uma pergunta que surge é se existe uma base ortonormal de V . A resposta é afirmativa e é isso que vamos construir agora.

Fixemos uma base qualquer $\mathcal{B} = \{|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle\}$ de V . Vamos obter a partir de $\mathcal{B}$ uma base ortonormal $\{|u_1\rangle, |u_2\rangle, \dots, |u_n\rangle\}$ por meio de um procedimento conhecido como **ortogonalização de Gram-Schmidt**.

Para obtermos o vetor $|u_1\rangle$ tomamos

$$|u_1\rangle = \frac{|v_1\rangle}{\| |v_1\rangle \|}.$$

Para obtermos $|u_2\rangle$, gostaríamos que $|u_2\rangle$ tenha norma unitária e que seja ortogonal ao vetor já obtido $|u_1\rangle$. Para verificar essa segunda condição, procuramos um vetor na forma

$$|w_2\rangle = |v_2\rangle + \alpha_1 |u_1\rangle$$

⁵ É possível verificar que o autoespaço correspondente a um autovalor λ é um subespaço vetorial de V .

(que está no subespaço gerado por $|v_1\rangle, |v_2\rangle$) de forma que $\langle w_2|u_1\rangle = 0$. Assim,

$$\begin{aligned} 0 &= \langle w_2|u_1\rangle = \langle v_2 + \alpha_1 u_1|u_1\rangle \\ &= \langle v_2|u_1\rangle + \alpha_1 \underbrace{\langle u_1|u_1\rangle}_{=\| |u_1\rangle\|=1} \\ &= \left\langle v_2 \left| \frac{v_1}{\| |v_1\rangle\|} \right. \right\rangle + \alpha_1 \\ &= \frac{1}{\| |v_1\rangle\|} \langle v_2|v_1\rangle + \alpha_1. \end{aligned}$$

Então,

$$\alpha_1 = -\frac{1}{\| |v_1\rangle\|} \langle v_2|v_1\rangle$$

e o vetor $|u_2\rangle$ é então definido como sendo $|u_2\rangle = \frac{|w_2\rangle}{\| |w_2\rangle\|}$.

Para obtermos $|u_3\rangle$ procederemos de forma similar: primeiro procuramos

$$|w_3\rangle = |v_3\rangle + \alpha_1 |u_1\rangle + \alpha_2 |u_2\rangle$$

que deve ser simultaneamente ortogonal a $|u_1\rangle$ e $|u_2\rangle$, determinando assim α_1 e α_2 como sendo

$$\alpha_1 = -\langle v_3|u_1\rangle \quad \text{e} \quad \alpha_2 = -\langle v_3|u_2\rangle$$

e, por fim, determinamos $|u_3\rangle = \frac{|w_3\rangle}{\| |w_3\rangle\|}$.

Procedendo analogamente, encontramos o vetor auxiliar $|w_k\rangle$, para cada $k \in \{2, 3, \dots, n\}$, que é dado por

$$|w_k\rangle = |v_k\rangle - \sum_{i=1}^{k-1} \langle v_k|u_i\rangle |u_i\rangle,$$

donde, $|u_k\rangle = \frac{|w_k\rangle}{\| |w_k\rangle\|}$, para todo $k \in \{2, 3, \dots, n\}$. Dessa forma, podemos exibir todos os vetores $|u_1\rangle, |u_2\rangle, \dots, |u_n\rangle$, e, por construção, eles geram o mesmo espaço que $|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle$, já que cada $|u_k\rangle$ se escreve como combinação linear dos vetores $|v_1\rangle, |v_2\rangle, \dots, |v_n\rangle$. Ademais, os vetores $|u_1\rangle, |u_2\rangle, \dots, |u_n\rangle$ também são ortonormais, sendo assim $\{|u_1\rangle, |u_2\rangle, \dots, |u_n\rangle\}$ a base ortonormal procurada do espaço V .

Seja A um operador linear no espaço de Hilbert $\mathcal{H}$. Existe um único operador linear $A^\dagger$ em $\mathcal{H}$, chamado de **operador adjunto** ou **conjugado Hermitiano** de A em $\mathcal{H}$, satisfazendo a

$$(|v\rangle, A|w\rangle) = (A^\dagger|v\rangle, |w\rangle),$$

para todos $|v\rangle, |w\rangle \in \mathcal{H}$.

Por convenção, o adjunto de um vetor $|v\rangle$ é $|v\rangle^\dagger \equiv \langle v|$. Assim, as principais propriedades da operação **adaga** ou **transposta-conjugada** são:

1. $(AB)^\dagger = B^\dagger A^\dagger$;
2. $(A|v\rangle)^\dagger = \langle v|A^\dagger$;
3. $(|w\rangle\langle v|)^\dagger = |v\rangle\langle w|$;
4. $(A^\dagger)^\dagger = A$;
5. Se A é inversível, então $A^\dagger$ também é e $(A^\dagger)^{-1} = (A^{-1})^\dagger$.
6. $(\sum_i a_i A_i)^\dagger = \sum_i a_i^* A_i^\dagger$ ⁶.

Com respeito a representação matricial de um operador A , a conjugação Hermitiana transforma a matriz A em sua transposta-conjugada, isto é,

$$A^\dagger = (A^*)^T,$$

onde "*" denota conjugação complexa e "T" é a operação transposição.

Exemplo .0.21. Consideremos os seguintes exemplos:

1. $\begin{bmatrix} 1+2i & 2i \\ 2+i & -1 \end{bmatrix}^\dagger = \begin{bmatrix} (1+2i)^* & (2i)^* \\ (2+i)^* & (-1)^* \end{bmatrix}^T = \begin{bmatrix} 1-2i & -2i \\ 2-i & -1 \end{bmatrix}^T = \begin{bmatrix} 1-2i & 2-i \\ -2i & -1 \end{bmatrix}.$
2. $\begin{bmatrix} 0 & 7i \\ 11-i & -10i \end{bmatrix}^\dagger = \begin{bmatrix} (0)^* & (7i)^* \\ (11-i)^* & (-10i)^* \end{bmatrix}^T = \begin{bmatrix} 0 & -7i \\ 11+i & 10i \end{bmatrix}^T = \begin{bmatrix} 0 & 11+i \\ -7i & 10i \end{bmatrix}.$

Definição .0.22. Um operador linear A em V é dito um **operador normal** quando

$$A^\dagger A = AA^\dagger.$$

Teorema .0.23 (Teorema Espectral). *Um operador linear A em V é diagonalizável se, e somente se, A for normal.*

Definição .0.24. Um operador linear U em V é dito um **operador unitário** quando

$$U^\dagger U = UU^\dagger = I.$$

Operadores unitários são normais, logo, são diagonalizáveis com relação a uma base ortonormal. Autovetores de um operador unitário associados a autovalores distintos são ortogonais. Os autovalores têm módulo iguais a 1. A aplicação de um operador unitário sobre um vetor preserva a norma.

Definição .0.25. Um operador linear A é dito **hermitiano** ou **auto-adjunto** quando

$$A^\dagger = A.$$

⁶ Está propriedade mostra que a operação adaga é conjugada linear.

Operadores hermitianos são normais, portanto, são diagonalizáveis com relação a uma base ortonormal. Autovetores de um operador hermitiano associados a autovalores distintos dão ortogonais. Os autovalores de um operador hermitiano são reais. Uma matriz real simétrica é hermitiana.

Definição .0.26. Um operador A em V é dito **positivo** quando

$$\langle v|A|v\rangle \geq 0,$$

para todo $|v\rangle \in V$. Se a desigualdade acima for estrita para todo $|v\rangle \in V|0$, então o operador A é dito **positivo definido**.

Os operadores positivos são hermitianos, logo, são diagonalizáveis.

Destacamos uma importante classe de operadores Hermitianos, conhecida como **projetores**. Seja V um espaço vetorial com dimensão n e W um subespaço de V com dimensão k . Utilizando o processo de ortogonalização de Gram-Schmidt é possível encontrarmos uma base ortonormal $\{|1\rangle, |2\rangle, \dots, |n\rangle\}$ de V , de modo que $\{|1\rangle, |2\rangle, \dots, |k\rangle\}$ seja uma base ortonormal de W . Definimos o **projetor** sobre W por

$$P \equiv \sum_{i=1}^k |i\rangle\langle i|.$$

É possível verificar que esta definição independe da base escolhida para W .

O operador $|v\rangle\langle v|$ é hermitiano, para qualquer $|v\rangle$, de modo que P também é hermitiano.

O **complemento ortogonal** de P é o operador

$$Q \equiv I - P$$

que é um projetor sobre o espaço $|k+1\rangle, \dots, |n\rangle$. A ação de projeção de P resultará no processo de medição quântica.

Os principais operadores para um qubit são conhecidos como **matrizes de Pauli** e são representados respectivamente por

$$\begin{aligned} \sigma_0 &\equiv I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \\ \sigma_1 &\equiv \sigma_x \equiv X \equiv \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \\ \sigma_2 &\equiv \sigma_y \equiv Y \equiv \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}, \\ \sigma_3 &\equiv \sigma_z \equiv Z \equiv \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}. \end{aligned}$$

Esses operadores formam um grupo multiplicativo não abeliano $\mathcal{G}$. Também, são unitários e hermitianos, portanto seus autovalores são iguais a 1 ou -1. Ou seja:

$$\sigma_j^2 = I \text{ e } \sigma_j^\dagger = \sigma_j, \quad j = 0, 1, 2, 3.$$

Valem:

$$X|0\rangle = |1\rangle, \quad X|1\rangle = |0\rangle, \quad Z|0\rangle = |0\rangle \text{ e } Z|1\rangle = -|1\rangle.$$

As matrizes de Pauli formam uma base para o espaço vetorial das matrizes de dimensão 2×2 . Logo, um operador genérico que atua em 1 qubit pode ser escrito como uma combinação linear das matrizes de Pauli.

O produto tensorial é uma forma de juntar espaços vetoriais para obter um outro espaço vetorial maior, que é um procedimento muito propício para a descrição da mecânica quântica de sistemas com muitas partículas.

Sejam V e W espaços de Hilbert de dimensões m e n , respectivamente. O **produto tensorial** $V \otimes W$ é um espaço vetorial de dimensão mn , cujos elementos são combinações lineares de produtos tensoriais $|v\rangle \otimes |w\rangle$ do elementos $|v\rangle \in V$ e $|w\rangle \in W$. Em particular, o produto tensorial de base ortonormais de V e W é uma base de $V \otimes W$. Pode-se escrever $|vw\rangle$ ou $|v\rangle|w\rangle$ em vez de $|v\rangle \otimes |w\rangle$.

Sejam V e W espaços vetoriais. O produto tensorial deve satisfazer as seguintes propriedades:

1. $\alpha(|v\rangle \otimes |w\rangle) = (\alpha|v\rangle) \otimes |w\rangle = |v\rangle \otimes (\alpha|w\rangle)$, para todos $|v\rangle \in V$, $|w\rangle \in W$ e $\alpha \in \mathbb{C}$.
2. $(|v_1\rangle + |v_2\rangle) \otimes |w\rangle = |v_1\rangle|w\rangle + |v_2\rangle|w\rangle$, para todos $|v_1\rangle, |v_2\rangle \in V$ e $|w\rangle \in W$.
3. $|v\rangle \otimes (|w_1\rangle + |w_2\rangle) = |v\rangle|w_1\rangle + |v\rangle|w_2\rangle$, para todos $|v\rangle \in V$ e $|w_1\rangle, |w_2\rangle \in W$.

Definição .0.27. Se A e B são operadores em V e W , respectivamente, definimos o operador $A \otimes B$ em $V \otimes W$ por

$$(A \otimes B)(|v\rangle \otimes |w\rangle) \equiv A|v\rangle \otimes B|w\rangle,$$

para todo $|v\rangle \otimes |w\rangle \in V \otimes W$.

Qualquer operador $C : V \otimes W \rightarrow V' \otimes W'$ pode ser representado como uma combinação linear de produtos tensoriais de operadores $A : V \rightarrow V'$ e $B : W \rightarrow W'$.

Utilizando os produtos internos definidos nos espaços de Hilbert V e W , podemos definir o produto interno em $V \otimes W$ por

$$\left(\sum_i a_i |v_i\rangle \otimes |w_i\rangle, \sum_j b_j |v'_j\rangle \otimes |w'_j\rangle \right) = \sum_{ij} a_i^* b_j \langle v_i | v'_j \rangle \langle w_i | w'_j \rangle.$$

São naturalmente estendidas para $V \otimes W$ todas as outras noções como operador adjunto, unitário, hermitiano.

As notações $|\psi\rangle^{\otimes n}$ e $A^{\otimes n}$ significam os produtos tensoriais de $|\psi\rangle$ e de A , por eles mesmos n vezes.

Todas as discussões sobre produto tensorial podem se tornar mais evidentes utilizando uma representação matricial conhecida como **produto de Kronecker**, apresentada na seguinte definição.

Definição .0.28. Sejam $A_{m \times n}$ e $B_{p \times q}$ matrizes. O **produto tensorial** de A e B é dado pela matriz

$$A \otimes B = \begin{bmatrix} A_{11}B & A_{12}B & \cdots & A_{1n}B \\ A_{21} & A_{22}B & \cdots & A_{2n}B \\ \vdots & \vdots & \ddots & \vdots \\ A_{m1} & A_{m2}B & \cdots & A_{mn}B \end{bmatrix}_{mp \times nq},$$

onde A_{ij} é o elemento localizado na linha i e coluna j da matriz A e cada termo $A_{ij}B$ denota uma submatriz $p \times q$ cujos elementos são os produtos de A_{ij} por cada elemento de B , isto é,

$$A_{ij}B = \begin{bmatrix} A_{ij}B_{11} & A_{ij}B_{12} & \cdots & A_{ij}B_{1q} \\ A_{ij}B_{21} & A_{ij}B_{22} & \cdots & A_{ij}B_{2q} \\ \vdots & \vdots & \ddots & \vdots \\ A_{ij}B_{p1} & A_{ij}B_{p2} & \cdots & A_{ij}B_{pq} \end{bmatrix}_{p \times q},$$

para todos $i = 1, 2, \dots, m$ e $j = 1, 2, \dots, n$.

Exemplo .0.29. Sejam $A = \begin{bmatrix} 1 & 0 \\ 0 & 7 \end{bmatrix}$ e $B = \begin{bmatrix} 2 & 0 & 4 \\ 9 & 10 & 0 \\ 0 & 3 & 17 \end{bmatrix}$. Então:

$$A \times B = \begin{bmatrix} 1 & 0 \\ 0 & 7 \end{bmatrix} \otimes \begin{bmatrix} 2 & 0 & 4 \\ 9 & 10 & 0 \\ 0 & 3 & 17 \end{bmatrix} = \begin{bmatrix} 2 & 0 & 4 & 0 & 0 & 0 \\ 9 & 10 & 0 & 0 & 0 & 0 \\ 0 & 3 & 17 & 0 & 0 & 0 \\ 0 & 0 & 0 & 14 & 0 & 28 \\ 0 & 0 & 0 & 63 & 70 & 0 \\ 0 & 0 & 0 & 0 & 21 & 119 \end{bmatrix}.$$

Observação .0.30. O produto tensorial não é comutativo. Por exemplo,

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

mas, por outro lado,

$$|1\rangle \otimes |0\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0(1) \\ 0(0) \\ 1(1) \\ 0(0) \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}.$$

Definição .0.31. Sejam A e B dois operadores lineares em um mesmo espaço vetorial V . Definimos o **comutador** entre A e B como

$$[A, B] = AB - BA.$$

Dizemos que A e B **comutam** quando $[A, B] = 0$. Definimos o **anticomutador** entre A e B como

$$\{A, B\} = AB + BA.$$

Dizemos que A e B **anticomutam** quando $\{A, B\} = 0$.

Observação .0.32. Sejam A , B e I operadores, onde I é a identidade. Valem:

$$1. [A, A] = 0;$$

$$[A, A] = AA - AA = A^2 - A^2 = 0.$$

$$2. [A, B] = -[B, A];$$

$$[A, B] = AB - BA = -(BA - AB) = -[B, A].$$

$$3. [I, A] = 0 = [A, I].$$

$$[I, A] = IA - AI = A - A = 0.$$

$$[A, I] = AI - IA = A - A = 0.$$

$$4. \{A, A\} = 2A^2.$$

$$\{A, A\} = AA + AA = A^2 + A^2 = 2A^2.$$

$$5. \{A, B\} = \{B, A\}$$

$$\{A, B\} = AB + BA = BA + AB = \{B, A\}.$$

$$6. \{I, A\} = 2A = \{A, I\}$$

$$\{I, A\} = IA + AI = A + A = 2A.$$

$$\{A, I\} = AI + IA = A + A = 2A.$$

Uma das mais importantes propriedades advindas dos conceitos de comutador e anticomutador se dá através da propriedade de diagonalização simultânea de operadores com esses conceitos. Isso é visto no seguinte teorema:

Teorema .0.33 (Teorema da diagonalização simultânea). *Sejam A e B operadores hermitianos. Então, $[A, B] = 0$ se, e somente se, existir uma base ortonormal tal que A e B sejam ortonormais nesta base. Neste caso, dizemos que A e B são **simultaneamente diagonalizáveis**⁷.*

Exemplo .0.34 (Matrizes de Pauli comutam ou anticomutam). Pelos itens 1 e 3 da Observação .0.32, respectivamente, cada uma das matrizes de Pauli comuta com ela mesma e por I comuta com X , Y , e Z . Podemos verificar que as relações de comutação para os operadores de Pauli são:

$$[X, Y] = 2iZ, [Y, Z] = 2iX, [Z, X] = 2iY.$$

De fato:

$$\begin{aligned} [X, Y] &= XY - YX = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} - \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \\ &= \begin{bmatrix} i & 0 \\ 0 & -i \end{bmatrix} - \begin{bmatrix} -i & 0 \\ 0 & i \end{bmatrix} \\ &= \begin{bmatrix} 2i & 0 \\ 0 & -2i \end{bmatrix} = 2i \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} = 2iZ. \end{aligned} \quad (.3)$$

$$\begin{aligned} [Y, Z] &= YZ - ZY = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} - \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \\ &= \begin{bmatrix} 0 & i \\ i & 0 \end{bmatrix} - \begin{bmatrix} 0 & -i \\ -i & 0 \end{bmatrix} \\ &= \begin{bmatrix} 0 & 2i \\ 2i & 0 \end{bmatrix} = 2i \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} = 2iX. \end{aligned} \quad (.4)$$

$$\begin{aligned} [Z, X] &= ZX - XZ = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} - \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \\ &= \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix} - \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix} \\ &= \begin{bmatrix} 0 & 2 \\ -2 & 0 \end{bmatrix} = 2i \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} = 2iY. \end{aligned} \quad (.5)$$

⁷ Do Teorema .0.33, em particular, podemos entender que duas matrizes hermitianas que comutam entre si são simultaneamente diagonalizáveis.

Perceba que, pelas expressões (.3), (.4) e (.5), respectivamente, temos

$$\{X, Y\} = XY + YX = 0, \quad \{Y, Z\} = YZ + ZY = 0 \quad \text{e} \quad \{Z, X\} = ZX - XZ = 0,$$

isto é, as matrizes X e Y , Y e Z e ainda X e Z anticomutam.

Observação .0.35 (Um breve resumo sobre notação de Dirac:). Temos:

1. **Ket:** $|\alpha\rangle$, lê-se ket alfa;
2. **Bra:** $\langle\alpha|$, lê-se bra alfa;
3. **Produto escalar:** $\langle\alpha|\alpha\rangle$, lê-se braket alfa;
4. **Operador:** $A * |\alpha\rangle = A|\alpha\rangle$;
5. **Multiplicação por escalar (números complexos):** $a|\alpha\rangle = |\alpha\rangle a$.

Consideremos o vetor $\vec{v}$ o qual estamos acostumados. Vamos representar $\vec{v}$ por $|v\rangle$ e seu elemento dual por $\langle v|$ na notação de Dirac da seguinte maneira:

$$\vec{v} = (v_1, v_2) \Rightarrow |v\rangle = \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} \Rightarrow \langle v| = [\overline{v_1} \quad \overline{v_2}].$$

Agora, vejamos como ficam as principais operações vetoriais na notação de braket:

• **Adição:**

$$|u\rangle + |v\rangle = \begin{bmatrix} u_1 \\ u_2 \end{bmatrix} + \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \begin{bmatrix} u_1 + v_1 \\ u_2 + v_2 \end{bmatrix}.$$

• **Subtração:**

$$|u\rangle - |v\rangle = \begin{bmatrix} u_1 \\ u_2 \end{bmatrix} - \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \begin{bmatrix} u_1 - v_1 \\ u_2 - v_2 \end{bmatrix}.$$

• **Multiplicação por escalar:**

$$\alpha|v\rangle = \alpha \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \begin{bmatrix} \alpha v_1 \\ \alpha v_2 \end{bmatrix}.$$

Produto interno:

$$\langle u|v\rangle = [\overline{u_1} \quad \overline{u_2}] \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \overline{u_1}v_1 + \overline{u_2}v_2.$$

• **Produto Vetorial:**

$$|u\rangle\langle v| = \begin{bmatrix} u_1 \\ u_2 \end{bmatrix} [\overline{v_1} \quad \overline{v_2}] = \begin{bmatrix} u_1\overline{v_1} & u_1\overline{v_2} \\ u_2\overline{v_1} & u_2\overline{v_2} \end{bmatrix}.$$

- **Produto Tensorial:**

$$|u\rangle \otimes |v\rangle \equiv |uv\rangle = \begin{bmatrix} u_1 \\ u_2 \\ \vdots \\ u_m \end{bmatrix} \otimes \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{bmatrix} = \begin{bmatrix} u_1 v_1 \\ u_1 v_2 \\ \vdots \\ u_1 v_n \\ u_2 v_1 \\ u_2 v_2 \\ \vdots \\ u_2 v_n \\ \vdots \\ u_m v_1 \\ u_m v_2 \\ \vdots \\ u_m v_n \end{bmatrix}.$$