

UNIVERSIDADE ESTADUAL DE PONTA GROSSA
SETOR DE CIÊNCIAS EXATAS E NATURAIS
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIAS – FÍSICA

ROSILENE ALES

MATRÓIDES E CÓDIGOS QUÂNTICOS

PONTA GROSSA
2017

ROSILENE ALES

MATRÓIDES E CÓDIGOS QUÂNTICOS

Dissertação apresentada para obtenção do título de Mestre na Universidade Estadual de Ponta Grossa, no Programa de Pós Graduação em Ciências, área de concentração Física.

Orientador: Prof. Dr. Giuliano Gadioli La Guardia.

PONTA GROSSA

2017

Ficha Catalográfica
Elaborada pelo Setor de Tratamento da Informação BICEN/UEPG

Ales, Rosilene
A371 Matrôides e códigos quânticos/ Rosilene
Ales. Ponta Grossa, 2017.
87f.

Dissertação (Mestrado em Ciências -
Área de Concentração: Física),
Universidade Estadual de Ponta Grossa.
Orientador: Prof. Dr. Giuliano Gadioli
La Guardia.

1.Matrôides. 2.Códigos quânticos.
3.Dualidade. I.La Guardia, Giuliano
Gadioli. II. Universidade Estadual de
Ponta Grossa. Mestrado em Ciências. III.
T.

CDD: 511.6

TERMO DE APROVAÇÃO

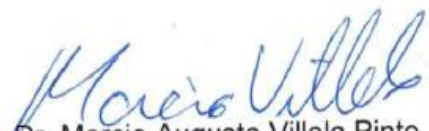
ROSILENE ALES

“MATRÓIDES E CÓDIGOS QUÂNTICOS”.

Dissertação aprovada como requisito parcial para obtenção do grau de Mestre no Programa de Pós-Graduação em Ciências - Física da Universidade Estadual de Ponta Grossa, pela seguinte banca examinadora.

Orientador:


Prof. Dr. Giuliano Gadioli La Guardia
Departamento de Matemática, UEPG/PR


Prof. Dr. Marcio Augusto Villela Pinto
Departamento de Engenharia Mecânica, UFPR/PR


Prof. Dr. Wanderley Aparecido Cerniauskas
Departamento de Matemática, UEPG/PR

Ponta Grossa, 29 de setembro de 2017.

Dedico aos meus pais Pedro e Marlene

AGRADECIMENTOS

- Primeiramente a Deus, pelo dom da vida, por estar ao meu lado e dado forças para nunca desistir.
- À minha família, pais, irmãos e em especial à minha irmã Rosana que me incentivou para entrar no mestrado e que nos momentos difíceis e de cansaço, soube me acolher e me levantar para não desanimar.
- Ao Prof. Dr. Giuliano Gadioli La Guardia, orientador desta dissertação, pela sabedoria, compreensão e empenho, pelas palavras amigas e exigências que me fizeram ser forte. Sem esquecer sua competência, participação na correção, revisão e discussões para a conclusão deste trabalho.
- Ao coordenador do Programa de Pós-graduação: Ciências em Física, José Danilo Szezech Júnior e ao ex-coordenador Sérgio da Costa Saab, pela oportunidade de crescimento, pelo conhecimento científico adquirido, pela amizade e confiança em mim depositada.
- A todos que estavam presentes: amigos, familiares, diretoras, professores e colegas de trabalho, pela amizade sincera, respeito e pelo incentivo na busca do conhecimento. Só tenho a dizer muito obrigado a todos.

“O que vale na vida não é o ponto de partida e sim a caminhada. Caminhando e semeando, no fim terás o que colher”.

(Cora Coralina)

RESUMO

Whitney identificou as propriedades fundamentais de dependência, que são comuns entre grafos e matrizes dando origem a Teoria de Matróides em 1935. Neste trabalho será apresentada a construção de novas famílias de Matróides e a resolução de teoremas de maneira ampla e de fácil compreensão, pois o tema é definido de forma matemática puramente abstrata. Assim, para obter um novo Matróide utiliza-se um já dado, sendo definido em termos de seus conjuntos independentes. Destaca-se que Matróides é encontrado nas seguintes abrangências: em espaços vetoriais, ciclos em grafos, funções afins, circuitos, bases, *rank*, fecho e dualidade. Deste modo, para construir um código quântico, precisa compreender a teoria de informação e codificação quântica como os Postulados da Mecânica Quântica, estados de um ou vários qubits, operadores unitários, portas lógicas, medidas de estados quânticos, códigos estabilizadores e a classe dos códigos de Calderbank-Shor-Steane (CSS). Os códigos lineares são os códigos da Álgebra Linear, subespaços vetoriais definidos sobre corpos finitos. Os códigos quânticos tem a finalidade de proteger possíveis erros de um canal para detectar e corrigir tais erros. Com o fundamento teórico adquirido, pode se verificar a possibilidade de construir a conexão da teoria de Matróides e a teoria de codificação quântica, por meio da matriz de verificação de paridade de um código CSS e a matriz que gera um dado Matróide vetorial.

Palavras-chave: Matróides. Códigos Quânticos. Dualidade.

ABSTRACT

Whitney identified the fundamental properties of dependence, which are common among graphs and matrices giving rise to Matroid Theory in 1935. In this work will be presented the construction of new families of Matroid and the resolution of theorems in a comprehensive and easy to understand, since the theme is defined in purely abstract mathematical form. Thus, to obtain a new Matroid one uses an already given one, being defined in terms of its independent sets. It should be noted that Matroid is found in the following ranges: in vector spaces, cycles in graphs, related functions, circuits, bases, *rank*, closure and duality. Thus, to construct a quantum code, one must understand quantum information and coding theory such as the Postulates of Quantum Mechanics, single- or multi-qubit states, unit operators, logic gates, quantum state measures, stabilizing codes, and the class of codes of Calderbank-Shor-Steane (CSS). The linear codes are the codes of Linear Algebra, vector subspaces defined on finite bodies. Quantum codes are intended to protect potential errors from a channel to detect and correct such errors. With the acquired theoretical foundation, the possibility of constructing the connection of the Matroid theory and the theory of quantum codification can be verified by means of the parity check matrix of a CSS code and the matrix that generates a given vector matroid.

Keywords: Matroids. Quantum Codes. Duality.

LISTA DE ILUSTRAÇÃO

Figura 1. O ciclo é um grafo que satisfaz o circuito eliminatório do axioma (C3).	24
Figura 2. Este Matróide ciclo gráfico é isomorfo para $M[A]$	25
Figura 3. Dois grafos não isomorfos com ciclo Matróides iguais.	31
Figura 4. (a) Um grafo de G , (b) Um subgrafo induzido de aresta $G[X]$ de G	37
Figura 5. K_5	42
Figura 6. Cinco Matróides e seus duais.	49
Figura 7. Um simples Matróide com um não simples dual.	50
Figura 8. Representação da proposição 3.1.3.....	51
Figura 9. Algumas portas clássicas de um ou mais bits (esquerda), e à direita o protótipo da porta quântica de muitos q-bits, o NÃO-controlado. A representação matricial do NÃO-controlado, U_{CN} , é escrita com relação às amplitudes de $ 00\rangle$, $ 01\rangle$, $ 10\rangle$ e $ 11\rangle$ nessa ordem.	65
Figura 10. Circuito para trocar os estados de dois qubits.....	80
Figura 11. Construção usada na demonstração de que as portas Hadamard, de fase de CNOT geram o normalizador $N(G_n)$	81

LISTA DE TABELA

Tabela 1. Geradores estabilizadores para o código de Steane de sete qubits.....75

Tabela 2. Propriedades de transformação dos elementos do grupo de Pauli sob conjugação por várias operações comuns. O NÃO CONTROLADO tem o q-bit 1 como controle e o q-bit 2 como alvo.....79

LISTA DE SIGLAS

CSS	Calderbank-Shor-Steane
CNOT	Porta quântica: não-controlado
EPR	Estado de Bell ou par de Einstein Podolsky-Rosen
HI	Hipótese de indução
MDS	<i>Maximum distance separable</i>

LISTA DE SÍMBOLOS

$\mathcal{I}(M)$	Coleção de conjuntos independentes de um Matróide M
$\mathcal{C}(M)$	Conjunto de circuitos de um Matróide M
$\mathcal{B}(M)$	Coleção das bases de um Matróide M
$\mathcal{B}$	Coleção de bases de um Matróide M
2^E	Conjunto das partes de E
$r(M)$	Função <i>Rank</i> de um Matróide M
$cl(M)$	Operador Fecho de um Matróide M
F_q	Corpo finito com q elementos
F_q^n	Conjunto das n -uplas de elementos definidos no corpo F_q
$\mathcal{C}^\perp$	Código Dual
$\dagger$	Operação linear sobre matrizes denominado adjunto ou conjugado Hermitiano

SUMÁRIO

1. INTRODUÇÃO	15
2. MATRÓIDES	19
2.1. DEFINIÇÃO DE MATRÓIDE	19
2.2. BASE	26
2.3 FUNÇÃO RANK	32
2.4 OPERADOR FECHO	38
3. DUALIDADE	46
3.1. AS DEFINIÇÕES E PROPRIEDADES BÁSICAS.	46
4. CONSTRUÇÃO DE CÓDIGOS QUÂNTICOS.....	56
4.1 CÓDIGOS LINEARES CLÁSSICOS.....	56
4.2 CÓDIGOS QUÂNTICOS.....	62
4.3 CÓDIGOS DE CALDERBANK–SHOR-STEANE (CSS)	67
4.4 CÓDIGOS ESTABILIZADORES	72
4.4.1. O FORMALISMO DE ESTABILIZADORES	73
4.5. PORTAS UNITÁRIAS E O FORMALISMO DE ESTABILIZADORES	78
4.6. MEDIDAS SEGUNDO O FORMALISMO DE ESTABILIZADORES	82
5. CONCLUSÃO	84
6. SUGESTÕES PARA TRABALHOS FUTUROS.....	85
7. REFERÊNCIAS	86

1. INTRODUÇÃO

O estudo da Teoria de Matróides (OXLEY, 1992) é bastante abstrato e generaliza diversas teorias disponíveis na literatura, tais como: a Teoria de Espaços Vetoriais, Módulos Vetoriais, Teoria de Grafos, o conceito de dependência algébrica, entre outros. Whitney (1935) identificou as propriedades fundamentais de dependência que são comuns a grafos e matrizes. Alguns autores interagiram junto a Whitney na teoria de Matróides tais como Van Der Waerden (VAN DER WAERDEN, 1937) que distinguiu as propriedades fundamentais que são comuns a álgebra e a dependência linear e o trabalho sobre Matróide de Birkoff (BIRKHOFF, 1935) e MacLane (MACLANE, 1936).

Uma das características da Teoria de Matróides é que um Matróide pode ser definido de várias maneiras diferentes, mas todas equivalentes. Mais especificamente, um Matróide pode ser definido mediante conjuntos independentes, conjuntos independentes maximais (bases), conjuntos dependentes minimais (circuitos), a função *rank* de um Matróide, função fecho, entre outras. Todas essas caracterizações serão tratadas neste trabalho.

As duas classes fundamentais de Matróides surgiram a partir do conceito de matrizes e de grafos.

Sendo $A_1, A_2, \dots, A_n$, colunas de uma matriz A , em um corpo finito, os subconjuntos de rótulos dessas colunas que são linearmente independentes (ou dependentes) em um dado espaço vetorial, caracterizam o que conhecemos como a classe dos Matróides Vetoriais, classe de Matróides que generaliza o conceito de espaço vetorial.

Whitney (1935) verificou que outros sistemas que não são representados por matrizes e que satisfazem estas propriedades, podem ser definidos de forma a generalizar certas propriedades satisfeitas por uma matriz.

Um Matróide é um par ordenado composto por um conjunto finito e por uma estrutura genérica do conceito de dependência geométrica; tal estrutura é uma generalização de noção de independência e dependência linear. No entanto, há diversas definições para Matróides que estão interligados com diversos ramos da Matemática tais como: na resolução de problemas algébricos, álgebra linear e no estudo dos espaços vetoriais. Mais precisamente, o conceito de Matróides tem sido encontrado em diversas áreas do conhecimento, tais como na teoria de grafos

(COSTALONGA, 2011), programação linear (SIMÃO, 2009) e teoria de redes elétricas (MOREIRA, 1993). Assim, é uma estrutura que generalizam os Matróides em diversas direções, conforme sua propriedade deve-se manter para que o objetivo seja alcançado.

O objetivo principal desta dissertação é introduzir os elementos básicos da Teoria de Matróides bem como os conceitos fundamentais da Teoria da Codificação Quântica. No que concerne à Teoria de Matróides abordaremos algumas demonstrações que não foram apresentadas em Oxley (1992), bem como abriremos algumas outras que foram apresentadas nos seguintes tópicos: conjuntos independentes e dependentes, conjuntos dependentes minimais (circuitos), conjuntos independentes maximais (bases), função *rank*, operador fecho, dualidade (Matróides duais). O livro texto utilizado na dissertação com relação à Teoria de Matróides foi o livro *Matroid Theory* (OXLEY, 1992) que, na nossa concepção, é o livro mais completo com relação a esse tema de estudo.

Com relação à Teoria da Codificação Quântica, foram abordados os seguintes tópicos: Postulados da mecânica Quântica, estados de um ou vários qubits (do inglês, *quantum bits*), operadores unitários, portas lógicas, medidas de estados quânticos, códigos estabilizadores, e a classe dos códigos de Calderbank-Shor-Steane (CSS). Para que fosse possível entender a fundo a classe dos códigos CSS, foi necessário entender um pouco sobre a Teoria dos Códigos de Bloco Lineares (MACWILLIAMS e SLOANE, 1977). Os códigos quânticos têm por finalidade proteger os estados quânticos de erros em um dado canal, corrigindo ou detectando esse erro.

Primeiramente, se faz necessário introduzir a teoria da Mecânica Quântica. Em 1982 (FEYNMAN, 1982), surgiu a conexão entre a Mecânica Quântica e a computação, iniciando uma nova pesquisa: a Computação Quântica e a Informação Quântica, que é a continuidade da evolução dos sistemas eletrônicos, computadores e das telecomunicações. De modo informal, a Teoria da Informação e Codificação Quântica é a teoria que investiga a criação de sistemas necessários para o envio da informação por canais ruidosos de comunicação, de forma a preservá-la de ruídos indesejáveis.

Os computadores quânticos utilizam componentes e conceitos relacionados à Mecânica Quântica. Na verdade, há grande dificuldade para se construir um computador quântico, porque a manipulação dos bits quânticos (qubits) é

extremamente complicada, devido ao fato de estados quânticos poderem colapsar para os estados zero ou um.

Na Computação Quântica, a unidade fundamental de informação é o qubit ou bit quântico, que será definido formalmente no Capítulo 4.

Sabe-se que há uma dificuldade intrínseca na computação clássica, ou seja, trabalhar com fatores primos de números de alta ordem de grandezas, ou seja, na resolução do algoritmo discreto. No caso de um computador quântico, esta tarefa não representaria quase nenhum esforço computacional, devido ao paralelismo quântico. Mais precisamente, o paralelismo quântico é uma característica do computador quântico efetuar diversas operações ao mesmo tempo, o que reduz drasticamente o tempo de execução do programa.

Com o advento do computador quântico, faz-se necessário estudar a Teoria de Codificação Quântica, que é a teoria natural para se entender como construir um código quântico capaz de proteger a informação quântica de erros do tipo *bit-flip*, que são erros do tipo mudança de bit, bem como erros do tipo *phase-shift*, que são erros do tipo mudança de fase de um dado estado quântico arbitrário.

Na primeira etapa, o estado quântico de informação é codificado por meio de operações unitárias por meio de um código corretor de erros quânticos. Este código quântico é um subespaço do espaço vetorial complexo gerado pelo produto tensorial de espaços vetoriais complexos munidos com produto interno, completo nesse produto (espaço de Hilbert). A segunda etapa, quando ocorre a atuação do ruído, é a identificação e a aplicação do operador correspondente para poder corrigir o erro, ou seja, fabricar o efeito inverso do causado pelo ruído. Portanto, para a identificação do erro, se deve ter atenção, pois o erro tem a capacidade de transportar o estado quântico codificado para um subespaço diferente do que foi inicialmente projetado. Para a identificação de erros no estado quântico é imprescindível que todos os espaços vetoriais de destino sejam ortogonais. As inversões de estados quânticos não deformam o espaço vetorial original de codificação, já que os diferentes processos de erros transportam os estados quânticos originais para outros espaços vetoriais que seja possível aplicar o processo de correção.

Na inversão de bit, realizando a comparação com o código de repetição clássica é possível identificar que as palavras-código, do código de repetição são os

rótulos para o código quântico em questão. O rótulo de um estado quântico é a representação binária $| \rangle$.

Para que seja possível entender profundamente a construção dos códigos quânticos CSS, é necessário o estudo aprofundado da construção dos códigos clássicos lineares, que serão utilizados para a construção CSS. Tais códigos (clássicos) lineares são os famosos códigos da Álgebra Linear, ou seja, subespaços vetoriais definidos sobre um corpo finito. Assim, também estudaremos nessa dissertação um pouco da Teoria dos Códigos Clássicos (veja Capítulo 4).

Deste modo, fornecer recursos para que, futuramente, seja possível construir conexões entre a Teoria de Matróides e a Teoria de Codificação Quântica, por meio da relação existente entre a matriz de paridade de um código CSS (4.3.15) e entre a matriz que gera um dado Matróide vetorial.

A dissertação está organizada da seguinte maneira. No Capítulo 2, falaremos de Matróide. No capítulo 3, proporcionaremos a Dualidade. No capítulo 4, apresentaremos a Construção de Códigos Quânticos, e finalizando com os seguintes capítulos 5. Conclusão do trabalho. 6. Sugestões de trabalhos futuros e, por fim no capítulo 7. Referências.

2. MATRÓIDES

Nesta seção, apresentaremos a definição e os conceitos referentes à Teoria de Matrôides, sendo que uma de suas características é que eles podem ser definidos de maneiras diferentes, porém todas equivalentes.

2.1. DEFINIÇÃO DE MATRÓIDE

Definição 2.1.1 Um Matrôide M é um par ordenado $(E, \mathcal{I})$ onde E é um conjunto finito e $\mathcal{I}$ uma coleção de subconjuntos de E que satisfaz as seguintes condições:

(I1) $\emptyset \in \mathcal{I}$.

(I2) Se $I \in \mathcal{I}$ e $I' \subseteq I$, então $I' \in \mathcal{I}$.

(I3) Se I_1 e $I_2 \in \mathcal{I}$ e $|I_1| < |I_2|$, então existe um elemento e de $I_2 - I_1$ tal que $I_1 \cup \{e\} \in \mathcal{I}$.

Se M é o Matrôide $(E, \mathcal{I})$, então dizemos que M é um Matrôide sobre E . O membro de $\mathcal{I}$ é chamado de conjunto independente de M e E é chamado conjunto base de M . Frequentemente escrevemos $\mathcal{I}(M)$ para denotarmos $\mathcal{I}$ e $E(M)$ para E , particularmente quando muitos Matrôides estão sendo considerados. O subconjunto de E que não está em $\mathcal{I}$ é chamado dependente.

O nome “Matrôide” foi criado por Whitney (1935), de uma classe fundamental, de modelos desta estrutura de dependência linear, de um conjunto de vetores que tem origem em matrizes como da seguinte maneira.

Proposição 2.1.1 Seja E o conjunto de colunas rotuladas de uma matriz $m \times n$ sobre um corpo F , e seja $\mathcal{I}$ o conjunto dos subconjuntos X de E para os quais multiconjuntos de colunas rotuladas por X é linearmente independente (L.I) no espaço vetorial de dimensão $V(m, F)$ sobre o corpo F . Então $(E, \mathcal{I})$ é um Matrôide.

Demonstração: Evidentemente, $\mathcal{I}$ satisfaz (I1), pois $\emptyset \in \mathcal{I}$ e (I2) também é satisfeito, pois, subconjuntos de conjuntos linearmente independentes também são linearmente independentes. Para provar (I3), sejam I_1 e I_2 conjuntos independentes

de $\mathcal{I}$. Supor que $|I_1| < |I_2|$, devemos provar que $\exists e \in I_2 - I_1$ com $I \cup \{e\} \in \mathcal{I}$. Por absurdo supor que $\forall e \in I_2 - I_1, I_1 \cup \{e\} \notin \mathcal{I} \Rightarrow I_1 \cup \{e\}$ gera vetores linearmente dependentes. Seja W subespaço de $V(m, F)$ gerado por $I_1 \cup I_2$. Logo a dimensão de W é no mínimo $|I_2|$, pois $I_2 \subset I_1 \cup I_2$. Portanto, W está contido no gerado de I_1 . Assim $|I_2| \leq \dim W \leq |I_1| < |I_2|$ é uma contradição. Portanto, concluímos que $I_2 - I_1$ contém um elemento e tal que $I_1 \cup \{e\} \in \mathcal{I}$ isto é, (I3) é válida. $\square$

Em que o multiconjunto é um conjunto cujos elementos podem ser repetidos (LA GUARDIA, PALAZZO e LAVOR, 2007).

Um Matróide obtido de uma matriz A como acima, é denotada por $M[A]$ *Matróide Vetorial* de A . Como um exemplo deste Matróide, considere o seguinte.

Exemplo 2.1.1 Seja A a matriz

$$A = \begin{matrix} & \begin{matrix} 1 & 2 & 3 & 4 & 5 \end{matrix} \\ \begin{pmatrix} 1 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 \end{pmatrix} \end{matrix}$$

sobre o corpo $\mathbb{R}$, cujas colunas são rotuladas por 1, 2, 3, 4 e 5. Obtemos.

$$E = \{1, 2, 3, 4, 5\} \text{ e}$$

$$\mathcal{I} = \{\emptyset, \{1\}, \{2\}, \{4\}, \{5\}, \{1, 2\}, \{1, 5\}, \{2, 4\}, \{2, 5\}, \{4, 5\}\}.$$

A coleção de subconjuntos dependentes é:

$$\{\{3\}, \{1, 3\}, \{1, 4\}, \{2, 3\}, \{3, 4\}, \{3, 5\}\} \cup \{X \subseteq E: |X| \geq 3\}.$$

Dentre os conjuntos dependentes existem os *minimais*, ou seja, conjuntos dependentes cujos subconjuntos próprios são independentes, e isso origina a seguinte definição:

Definição 2.1.2 Um conjunto dependente minimal de um Matróide M será denominado circuito de M , e o conjunto de todos os circuitos do Matróide M será denominado por $\mathcal{C}$ ou $\mathcal{C}(M)$.

Os circuitos gerados pelo Matróide do Exemplo 2.1.1 estão em $\mathcal{C}(M)$ dado por:

$$\mathcal{C}(M) = \{\{3\}, \{1, 4\}, \{1, 2, 5\}, \{2, 4, 5\}\}.$$

Um circuito de M com n elementos é denominado n - circuito. Evidentemente, tal como no exemplo anterior, toda vez que $\mathcal{C}(M)$ for especificado, M também o será: os membros de $\mathcal{I}(M)$ são os subconjuntos de M e E que não contêm nenhum membro de $\mathcal{C}(M)$. Assim, um Matróide é determinado exclusivamente por seu conjunto $\mathcal{C}$ de circuitos.

Vamos agora analisar algumas propriedades de $\mathcal{C}$, com vista a caracterizar esses dois subconjuntos de que pode ocorrer como o conjunto de circuitos de um Matróide sobre E . É fácil ver isso.

Lema 2.1.1 Seja $\mathcal{C}$ o conjunto dos circuitos de um Matróide M , então são válidas as seguintes propriedades:

(C1) $\emptyset \notin \mathcal{C}$.

(C2) Se C_1 e C_2 são membros de $\mathcal{C}$ e $C_1 \subseteq C_2$, então $C_1 = C_2$.

(C3) Se C_1 e C_2 são membros distintos de $\mathcal{C}$ e $e \in C_1 \cap C_2$. Então existe um membro C_3 de $\mathcal{C}$ tal que $C_3 \subseteq (C_1 \cup C_2) - e$.

Demonstração: (C1) $\emptyset \notin \mathcal{C}$. Por (I1) $\emptyset \in \mathcal{I}$, pois $\mathcal{I}$ não contém circuito.

Provar (C2) supor que $C_1, C_2 \in \mathcal{C}$ com $C_1 \subseteq C_2$ e $C_1 \neq C_2$. $\exists x \in C_2 - C_1$. Então $C_1 \subseteq C_2 - \{x\}$. Mas $C_2 - \{x\} \in \mathcal{I}$, por (I2) $C_1 \in \mathcal{I}$ é um absurdo.

Para verificar (C3). Supor que não existe um circuito contido em $(C_1 \cup C_2) - \{e\}$, ou seja, $(C_1 \cup C_2) - e$ é independente. Por (C2), $C_2 - C_1 \neq \emptyset$ então existe f em $C_2 - C_1$. Como C_2 é conjunto dependente minimal, $C_2 - f \in \mathcal{I}$. Seja I um subconjunto de $C_1 \cup C_2$ que é maximal com a propriedade que este contém a $C_2 - f$ é independente. Evidentemente $f \notin I$. Mas como C_1 é um circuito, algum elemento g de C_1 não está em $\mathcal{I}$. Como $f \in C_2 - C_1$, os elementos f e g são distintos. Então

$$|I| \leq |(C_1 \cup C_2) - \{f, g\}| = |(C_1 \cup C_2)| - 2 < |(C_1 \cup C_2) - \{e\}|.$$

Agora aplicando, (I3), tomando I_1 como I e I_2 como $(C_1 \cup C_2) - e$. O conjunto independente resultado contradiz com a maximalidade de I . $\square$

A condição (C3) é chamada de axioma da *eliminação do circuito*.

A seguir, apresentaremos um teorema que caracteriza um Matróide por meio das condições (C1) – (C3).

Teorema 2.1.1 Seja E um conjunto e $\mathcal{C}$ uma coleção de subconjuntos de E que satisfaz (C1) – (C3). Seja $\mathcal{I}$ uma coleção de subconjuntos de E que não contém membros de $\mathcal{C}$. Então $(E, \mathcal{I})$ é um Matróide que tem $\mathcal{C}$ como sua coleção de circuitos.

Demonstração: Primeiramente mostraremos que $\mathcal{I}$ satisfaz (I1) – (I3). Por (C1) $\emptyset$ não contém um membro de $\mathcal{C}$, por (I1) $\emptyset \in \mathcal{I}$. Se I não contém membro de $\mathcal{C}$ e I' está contido em I , então I' não contém membros de $\mathcal{C}$. Logo (I2) contém.

Para provar (I3), suponhamos que I_1 e I_2 são membros de $\mathcal{I}$ e $|I_1| < |I_2|$. Suponhamos que (I3) falha para o par (I_1, I_2) . Agora $\mathcal{I}$ contém um membro que é um subconjunto de $I_1 \cup I_2$ e contém mais elementos que I_1 . Escolha um subconjunto I_3 de $\mathcal{I}$ contido em $I_1 \cup I_2$ tal que a $|I_1 - I_3|$ é minimal. Como (I3) falha, $I_1 - I_3$ é não vazio, assim podemos escolher um elemento e deste conjunto. Para cada elemento f de $I_3 - I_1$, seja $T_f = (I_3 \cup e) - f$. Então $T_f \subseteq I_1 \cup I_2$ e $|I_1 - T_f| < |I_1 - I_3|$. Portanto $T_f \notin \mathcal{I}$, assim T_f contém um membro de $\mathcal{C}_f$ de $\mathcal{C}$. Evidentemente $f \notin \mathcal{C}_f$. Mas, $e \in \mathcal{C}_f$, pois caso contrário temos $\mathcal{C}_f \subseteq I_3$ contradiz o feito que $I_3 \in \mathcal{I}$.

Seja g um elemento de $I_3 - I_1$. Se $\mathcal{C}_g \cap (I_3 - I_1) = \emptyset$. Então $\mathcal{C}_g \subseteq (I_1 \cap I_3) \cup e - g \subseteq I_1$ o que dá uma contradição. Portanto, tem um elemento h em $\mathcal{C}_g \cap (I_3 - I_1)$. Então $e \in \mathcal{C}_g \cap \mathcal{C}_h$. Assim, (C3) implica que tem um membro C de $\mathcal{C}$ tal que $C \subseteq (\mathcal{C}_g \cup \mathcal{C}_h) - e$. Mas, ambos $\mathcal{C}_g$ e $\mathcal{C}_h$ são subconjuntos de $I_3 \cup e$ e, portanto $C \subseteq I_3$ é uma contradição. Concluimos que (I_3) não falha e consequentemente $(E, \mathcal{I})$ é um Matróide M . $\square$

Para provar que $\mathcal{C}$ é um conjunto $\mathcal{C}(M)$ de circuitos de M , notamos que os seguintes enunciados são equivalentes.

- (i) C é um circuito de M .
- (ii) $C \notin \mathcal{I}(M)$ e $C - x \in \mathcal{I}(M)$ para todo x em C .
- (iii) C contém um membro C' de $\mathcal{C}$ como um subconjunto, mas C' não é um subconjunto próprio de C .
- (iv) $C \in \mathcal{C}$.

Como consequência imediata do teorema anterior tem-se:

Corolário 2.1.1 Seja $\mathcal{C}$ uma coleção de subconjuntos de E . Então $\mathcal{C}$ é uma coleção de subconjuntos do circuito de um Matróide M se, e somente se, satisfaz as seguintes condições:

(C1) $\emptyset \notin \mathcal{C}$.

(C2) Se C_1 e C_2 são membros de $\mathcal{C}$ e $C_1 \subseteq C_2$. Então $C_1 = C_2$.

(C3) Se C_1 e C_2 são membros distintos de $\mathcal{C}$ e $e \in C_1 \cap C_2$, então, existe um membro C_3 de $\mathcal{C}$ tal que $C_3 \subseteq (C_1 \cup C_2) - e$.

Proposição 2.1.2 Seja I um conjunto independente em um Matróide M e e um elemento de M tal que $I \cup \{e\}$ é dependente. Então M tem um único circuito C contido em $I \cup \{e\}$, que contém e .

Demonstração: Evidentemente, $I \cup \{e\}$ contém um circuito, pois $I \cup \{e\}$ é um conjunto dependente. Portanto, todo circuito C contido em $I \cup \{e\}$ contém e , pois, caso contrário, $C \subset I$, um absurdo. Supor que existe um circuito C' , distinto de C , contido em $I \cup \{e\}$. Então, por (C3), $(C \cup C') - e$ contém um circuito C' , mas $(C \cup C') - e \subseteq I$, e assim $C' \subset I$, o que é um absurdo, pois I é independente. Portanto, C é o único circuito contido em $I \cup \{e\}$ e que contém e . $\square$

Dois importantes classes fundamentais de Matróides são: Matróides vetoriais, apresentadas anteriormente, e Matróides derivados de grafos. Em que o grafo G é definido pelo conjunto finito de vértices (pontos) e um conjunto finito de arestas (linhas), cada uma destas arestas está conectada dois vértices ao qual exibiremos na sequência.

Proposição 2.1.3. Seja E o conjunto das arestas de um grafo G e $\mathcal{C}$ a família dos conjuntos das arestas dos ciclos de G . Então $\mathcal{C}$ é a família dos circuitos de um Matróide sobre E .

Demonstração: (C1) $\emptyset \notin \mathcal{C}$. Por (I1), $\emptyset \in \mathcal{I}$.

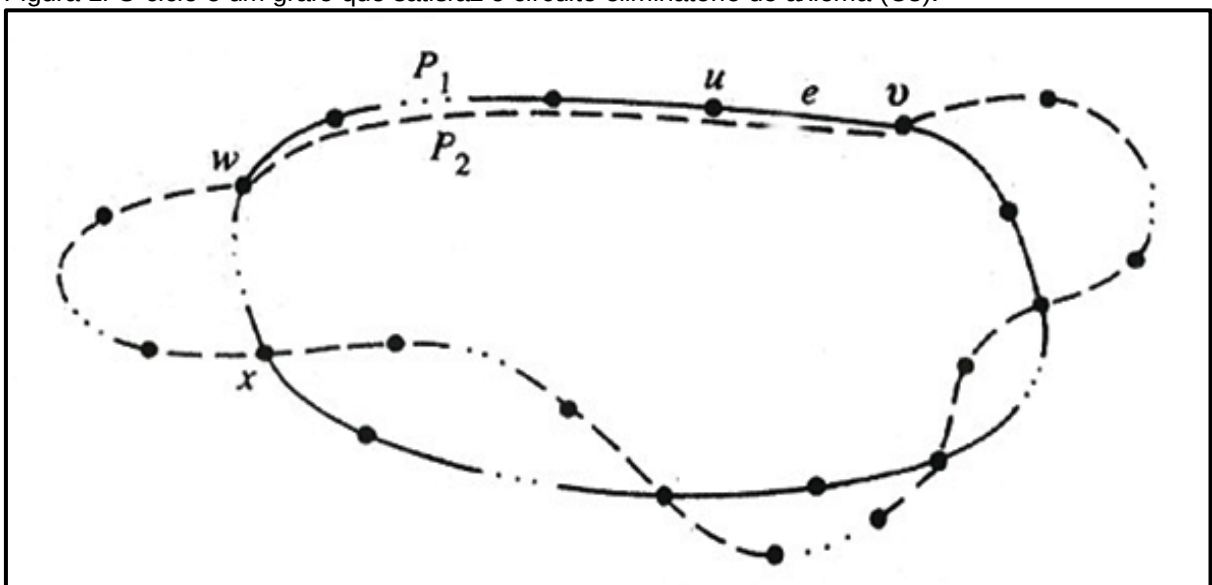
(C2) Supor C_1 e $C_2 \in \mathcal{C}$ com $C_1 \subseteq C_2$ e $C_1 \neq C_2$. Então $\exists e \in C_2 - C_1$, e assim $C_1 \subseteq C_2 - \{e\}$. Mas $C_2 - \{e\} \in \mathcal{I}$, donde, por (I2), $C_1 \in \mathcal{I}$, uma contradição.

(C3) Sejam C_1 e C_2 os conjuntos das arestas de dois ciclos distintos de G que possuem e como uma aresta em comum. Considere a Fig. (1). Sejam u e v os vértices da aresta e . Vamos agora construir um ciclo de G cujas arestas estão contidas em $C_1 \cup C_2 - e$, para $i = 1, 2$, seja P_1 um caminho de u a v em G cujas arestas estão em $C_1 - e$. Começando em u , percorrendo P_1 em direção de v , seja w o primeiro vértice cuja próxima aresta de P_1 não pertence a P_2 . Continuando a percorrer P_1 de w em direção v , encontramos um vértice x distinto de w que também está em P_2 . Como P_1 e P_2 terminam em v tal vértice existe. Unindo o subcaminho de P_1 que vai de w a x ao subcaminho de P_2 vai de x até w encontramos um ciclo cujas arestas estão contidas em $C_1 \cup C_2 - e$. Portanto, $\mathcal{C}$ satisfaz a propriedade (C3), $C_3 \subseteq C_1 \cup C_2 - e$. Logo $\mathcal{C}$ é a família dos circuitos de um Matróide sobre E . $\square$

Assim, um ciclo é uma sequência alternada de arestas e vértices, de um grafo G , formando um ciclo fechado.

Definição 2.1.3 Dois Matróides M_1 e M_2 são isomorfos, escreve-se $M_1 \cong M_2$, se existir uma bijeção (correspondente a injetiva e sobrejetiva) ψ de $E(M_1)$ para $E(M_2)$ tal que para todo $X \subseteq E(M_1)$, $\psi(X)$ é independente em M_2 se, e somente se, X é independente em M_1 .

Figura 1. O ciclo é um grafo que satisfaz o circuito eliminatório do axioma (C3).



FONTE: Oxley, 1992.

Nesta Proposição 2.1.3 observamos que um ciclo gráfico é fechado.

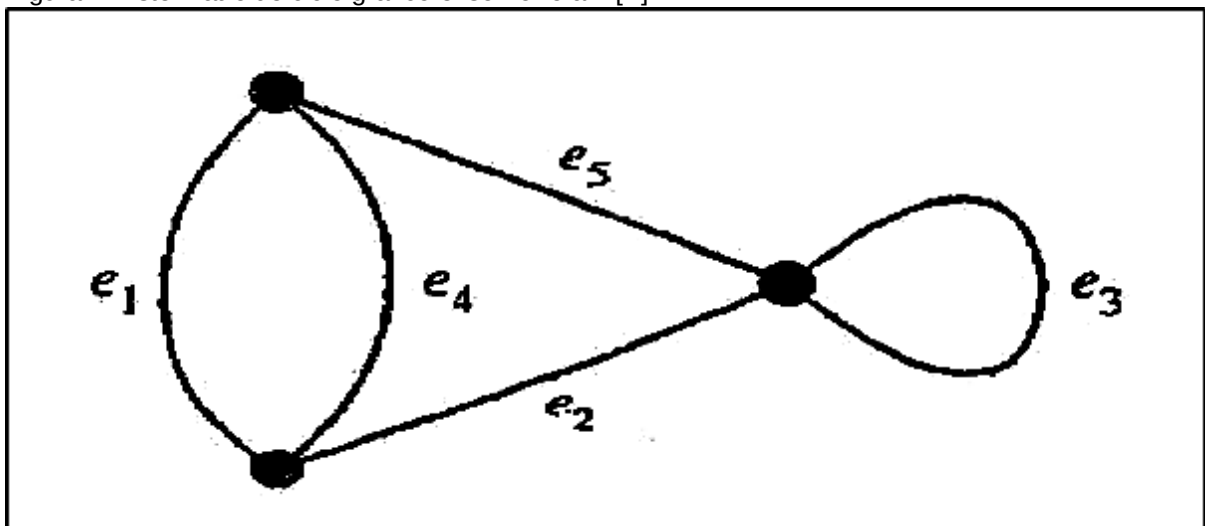
Matróides construídos através das arestas de um determinado grafo H , com o grafo G , denomina-se Matróide-ciclo e denotado por $M(H)$.

Um Matróide que é isomorfo a um Matróide-ciclo originado de um grafo G é denominado *gráfico*. Denotaremos este Matróide por $M(G)$. Podemos concluir, pela proposição acima, que é um subconjunto X de arestas de G é um conjunto independente em $M(G)$ se, e somente se, X não contém nenhum conjunto de arestas de ciclos, ou equivalentemente, o subgrafo (é o subconjunto de arestas e vértices do grafo original) induzido $G[X]$ pelas arestas de X é uma floresta (não contém circuitos).

Exemplo 2.1.2 Seja G o grafo na Fig (2), e seja $M = M(G)$. Então $E(M) = \{e_1, e_2, e_3, e_4, e_5\}$ e $\mathcal{C}(M) = \{\{e_3\}, \{e_1, e_4\}, \{e_1, e_2, e_5\}, \{e_2, e_4, e_5\}\}$. Comparando M com o Matróide $M[A]$ no Exemplo 2.1.1. Uma bijeção ψ de $\{1, 2, 3, 4, 5\}$ para $\{e_1, e_2, e_3, e_4, e_5\}$ é definida por $\psi(i) = e_i$. Um conjunto X é um circuito em $M[A]$ se, e somente se, $\psi(X)$ é um circuito em M . Equivalentemente, um conjunto Y é independente em $M[A]$ se $\psi(Y)$ é independente em M . Assim, os Matróides $M[A]$ e M têm a mesma estrutura, ou seja, são isomorfos.

Por exemplo, o Matróide $M[A]$ no Exemplo 2.1.2 é gráfico. Se M é isomorfo ao Matróide vetorial da matriz A sobre um corpo F então M é dito *representável*. Nesse caso, A é uma representação para M sobre F ou F – *representável* para M .

Figura 2. Este Matróide ciclo gráfico é isomorfo a $M[A]$.



FONTE: Oxley, 1992.

Na Fig (2), o laço e_3 é o par $\{e_1, e_4\}$, são arestas paralelas do circuito $M(G)$.

Chamamos a um elemento e um ciclo de um Matróide arbitrário M , se $\{e\}$ é um circuito de M , mas se f e g são elementos de M tal que $\{f, g\}$ é um circuito, então f e g se diz paralelo em M . Uma classe paralela de M é um subconjunto maximal X de $E(M)$ tal que quaisquer dois membros distintos de X são paralelos e não membros de X é um ciclo. Uma classe paralela é trivial se esta contém somente um elemento. Se M não contém ciclos e não é uma classe trivial é chamado de um Matróide simples ou geometria combinatória.

Sabe-se que um Matróide M pode ser especificado por sua coleção de conjuntos independentes ou pela sua coleção de circuitos. Veremos, na seção seguinte, que M também poderá ser determinado por sua coleção de conjuntos independentes maximais de um Matróide.

2.2. BASE

Nesta seção estudaremos as bases de um Matróide, mostrando que esses conjuntos contêm muito em comum com o conceito de base de um espaço vetorial.

Definição 2.2.1 Seja $(E, \mathcal{I})$ um Matróide M . Um conjunto independente *maximal*, ou seja, um conjunto independente tal que, qualquer elemento adicionado a este faz com que o novo conjunto se torne dependente, é denominado *base* de M . O conjunto de todas as bases de M será denotado por $\mathcal{B}$.

Como um exemplo, o conjunto $\mathcal{B}$ do Exemplo 2.1.1 é dado por $\mathcal{B} = \{\{1,2\}, \{1,5\}, \{2,4\}, \{2,5\}, \{4,5\}\}$. Observe que todas as bases possuem a mesma cardinalidade. Esse fato é demonstrado a seguir.

Teorema. 2.2.1 De fato, sejam B_1 e B_2 bases de um Matróide M . Então $|B_1| = |B_2|$.

Demonstração: Sejam B_1 e B_2 bases de um Matróide M . Então, por definição de base, B_1 e $B_2 \in \mathcal{I}$. Suponha por absurdo que $|B_1| \neq |B_2|$; então, sem perda de generalidade, podemos supor que $|B_1| < |B_2|$. Por (I3), existe um elemento e de $B_2 - B_1$ tal que $B_1 \cup \{e\} \in \mathcal{I}$. Logo, B_1 não é conjunto independente maximal, um absurdo.

Portanto, $|B_1| \geq |B_2|$. Analogamente, segue-se que $|B_2| \geq |B_1|$. Portanto, $|B_1| = |B_2|$, como requerido. $\square$

Todas as bases de uma matróide M contém o mesmo número de elementos. E tem a mesma cardinalidade, isto é são equicardinais.

Denotamos por $\mathcal{B}$ ou $\mathcal{B}(M)$ a coleção de bases de um Matróide M . Em que o próximo resultado descreve algumas propriedades das bases $\mathcal{B}$.

Lema 2.2.1 Se M é um Matróide e $\mathcal{B}$ é a coleção de bases de M , então $\mathcal{B}$ satisfaz as seguintes propriedades:

(B1) $\mathcal{B}$ é não vazio.

(B2) Se B_1, B_2 são membros de $\mathcal{B}$, e $x \in B_1 - B_2$, então existe um elemento $y \in B_2 - B_1$, tal que $(B_1 - x) \cup \{y\} \in \mathcal{B}$.

Demonstração: (B1) Como $\emptyset \in \mathcal{I}$ por (I1) então $\mathcal{I}$ contém pelo menos um conjunto independente maximal, i.e., $\mathcal{B}$ é não vazio.

(B2) Sejam $B_1, B_2 \in \mathcal{B}$. De fato, $|B_1| = |B_2|$. Seja $x \in B_1 - B_2$. Como $B_1 - x \subseteq B_1 \in \mathcal{I}$, segue-se por (I2) que $B_1 - x \in \mathcal{I}$. Observe que tanto $B_1 - x$ e B_2 são conjuntos independentes tal que $|B_1 - x| < |B_2|$. Por (I3), existe $y \in B_2 - (B_1 - x)$ tal que $B_3 = (B_1 - x) \cup \{y\} \in \mathcal{I}$. Como B_3 é independente, existe um conjunto independente maximal B^* que contém B_3 . Como $|(B_1 - x) \cup \{y\}| = |B_1|$, de fato, sabemos que $|B_3| = |B^*|$. Portanto, $(B_1 - x) \cup \{y\} \in \mathcal{B}$ é uma base de M . $\square$

A condição (B2) é uma das *propriedades de troca de bases* satisfeitas por Matróides.

Agora, demonstraremos que (B1) e (B2) caracterizam as bases de um Matróide.

Considere o Lema seguinte, o qual será usado e o Lema do próximo Teorema.

Teorema. 2.2.2 Seja E um conjunto finito e $\mathcal{B}$ é uma coleção de subconjuntos de E que satisfaz (B1) e (B2). Seja $\mathcal{I}$ uma coleção de subconjuntos de E que está contido em algum membro de $\mathcal{B}$. Então $(E, \mathcal{I})$ é um Matróide tendo $\mathcal{B}$ como sua coleção de bases.

Demonstração: Mostremos primeiramente que $(E, \mathcal{I})$ é um Matróide. Como a família $\mathcal{B}$ é não vazia, existe $B \in \mathcal{B}$ e $\emptyset \subset B$. Assim, por hipótese, $\emptyset \in \mathcal{I}$, donde (I1) é satisfeita. Para provar (I2), seja $I \in \mathcal{I}$ então existe um elemento B de $\mathcal{B}$, tal que $I \subset B$. Assim $I' \subset I$, então $I' \subset B$, $I' \in \mathcal{I}$.

Para mostrar que $\mathcal{I}$ satisfaz (I3) provaremos a seguinte afirmação.

Lema 2.2.2 Os membros de $\mathcal{B}$ são equicardinais.

Demonstração: Suponha que B_1 e B_2 sejam membros distintos de $\mathcal{B}$ tais que $|B_1| > |B_2|$ e, além disso, $|B_1 - B_2|$ seja mínima. Claramente $B_1 - B_2 \neq \emptyset$. Assim escolha $x \in B_1 - B_2$; pela propriedade (B2) existe $y \in B_2 - B_1$ tal que $(B_1 - x) \cup \{y\} \in \mathcal{B}$. Como $|(B_1 - x) \cup y| = |B_1| > |B_2|$ e $|(B_1 - x) \cup \{y\} - B_2| < |B_1 - B_2|$, obtém-se uma contradição, pois $|B_1 - B_2|$ é mínima.

Agora, demonstraremos que (I3) não valha para $\mathcal{I}$. Então existem I_1 e I_2 membros de $\mathcal{I}$ tais que, $|I_1| < |I_2|$ e para todo $e \in I_2 - I_1$, o conjunto $I_1 \cup \{e\}$ não pertence a $\mathcal{I}$. Por hipótese, existem elementos B_1 e B_2 de $\mathcal{B}$ tais que $I_1 \subseteq B_1$ e $I_2 \subseteq B_2$.

Suponha que o conjunto B_2 seja escolhido de forma que $|B_2 - (I_2 \cup B_1)|$ seja mínima. Note que, pela escolha de I_1 e I_2 temos.

$$I_2 - B_1 = I_2 - I_1 \quad (1)$$

De fato, a inclusão $I_2 - B_1 \subseteq I_2 - I_1$ é óbvia e se existe $x \in (I_2 - I_1) \cap B_1$ então $I_1 \cup x \subseteq B_1$ e assim $I_1 \cup x \in \mathcal{I}$, isso contradiz a escolha de I_1 e I_2 . Agora suponha que $B_2 - (I_2 \cup B_1)$ é não vazio e escolha um elemento x neste conjunto. Pela propriedade (B2) existe um elemento $y \in B_1 - B_2$ tal que $(B_2 - x) \cup \{y\} \in \mathcal{B}$. Mas $|(B_2 - x) \cup \{y\} - (I_2 \cup B_1)| < |B_2 - (I_2 \cup B_1)|$ e isto contradiz a escolha de B_2 .

Logo, temos que $B_2 - (I_2 \cup B_1)$ é vazia e assim $B_2 - B_1 = I_2 - B_1$. Como de (1) $I_2 - B_1 = I_2 - I_1$ temos que:

$$B_2 - B_1 = I_2 - I_1 \quad (1.1)$$

Vamos mostrar que $B_1 - (I_1 \cup B_1)$ é vazia. Se não, então existe elemento x neste conjunto e um elemento y em $B_2 - B_1$ tal que $(B_1 - x) \cup \{y\} \in \mathcal{B}$. Agora $I_1 \cup y \subseteq (B_1 - x) \cup y$ e, portanto $I_1 \cup \{y\} \in \mathcal{I}$. Como $B_2 - B_1 = I_2 - I_1$ de (1.1) temos que $y \in I_2 - I_1$ e chegamos a uma contradição com a hipótese que (I3) é falsa. Concluímos

que $B_1 - (I_1 \cup B_1)$ é vazia e disto temos que $B_1 - B_2 = I_1 - B_2$. Note que o conjunto $I_1 - B_2$ está contido em $I_1 - I_2$ e da última igualdade segue que

$$B_1 - B_2 \subseteq I_1 - I_2 \quad (1.2)$$

Sabemos que $|B_1| = |B_2|$, disto temos que $|B_1 - B_2| = |B_2 - B_1|$. Portanto por (1.1) e (1.2), $|I_1 - I_2| \geq |I_2 - I_1|$, assim $|I_1| \geq |I_2|$. Chegamos a uma contradição. Concluimos que $(E, \mathcal{I})$ é um Matróide e por construção a família $\mathcal{B}$ é conjunto de bases deste Matróide. $\square$

Ao combinar o Teorema 2.2.1 com o Lema 2.2.1 e as observações que o procedem, obtemos o seguinte resultado.

Corolário 2.2.1 Seja $\mathcal{B}$ uma família de subconjuntos de E . Então $\mathcal{B}$ é a coleção de bases de um Matróide sobre E se, e somente se, satisfaz as seguintes condições:

(B1) $\mathcal{B}$ é não vazio.

(B2) Se B_1, B_2 são membros de $\mathcal{B}$, e $x \in B_1 - B_2$, então existe um elemento y de $B_2 - B_1$, tal que $(B_1 - x) \cup \{y\} \in \mathcal{B}$.

Corolário 2.2.2 Seja $\mathcal{B}$ uma base de um Matróide M . Se $e \in E - B$ então $B \cup e$ contém um único circuito $C(e, B)$. Além disso, $e \in C(e, B)$.

Demonstração: Seja $e \in E - B$, então $B \cup e$ é dependente, assim existe um circuito $C \subset B \cup e$. Devemos ter $e \in C$, pois se $e \notin C$ teríamos $C \subset B$, um absurdo.

Unicidade: Supor $C_1 \neq C_2 \subset B \cup e$; sabemos que $e \in C_1 \cap C_2$. Por (C3), existe $C_3 \subset (C_1 \cup C_2 - e) \subset B$, um absurdo. $\square$

Exemplo 2.2.1 Sejam m e n inteiros não negativos tal que $m \leq n$. Seja E um conjunto com n elementos e $\mathcal{B}$ uma coleção de subconjuntos de E contendo m elementos. Então $\mathcal{B}$ é o conjunto de bases de um Matróide E . Denotaremos este Matróide por $U_{m, n}$ e o denominaremos Matróide *uniforme* de *rank* m em um conjunto de n elementos. Claramente.

$$\mathcal{I}(U_{m, n}) = \{X \subseteq E: |X| \leq m\}.$$

$$\mathcal{C}(U_{m, n}) = \{\emptyset \text{ se } m = n \text{ e } X \subset E: |X| \leq m\}.$$

Demonstração: Seja $E = \{a_1, \dots, a_n\}$ e $m \leq n$, $\mathcal{B} = \{B \subset E: |B| = m\}$ e $\mathcal{I} = \{I \subset E: |I| \leq m\}$.

(I1) $\emptyset \in \mathcal{I}$, pois $|\emptyset| \leq m$.

(I2) Supor $I \in \mathcal{I}$, isto é, $|I| \leq m$. Se $I' \subset I$, então $|I'| \leq |I| \leq m$. Então $|I'| \leq m$, ou seja, $I' \in \mathcal{I}$.

(I3) Supor $I_1, I_2 \in \mathcal{I}$ com $|I_1| < |I_2|$. Então $|I_1| \leq m-1$ e $|I_2| \leq m$. Então $\exists e \in I_2 - I_1$ tal que $|I_1 \cup \{e\}| \leq m$, donde $I_1 \cup \{e\} \in \mathcal{I}$. $\square$

Em um Matróide M , como foi especificado a relação entre $\mathcal{B}(M)$ e $\mathcal{I}(M)$ e entre $\mathcal{I}(M)$ e $\mathcal{C}(M)$ diretamente notamos que $\mathcal{B}(M)$ é a coleção maximal de subconjuntos de E que não contém nenhum membro de $\mathcal{C}(M)$, enquanto $\mathcal{C}(M)$ é a coleção de conjuntos minimais que não estão contidos em nenhum membro de $\mathcal{B}(M)$.

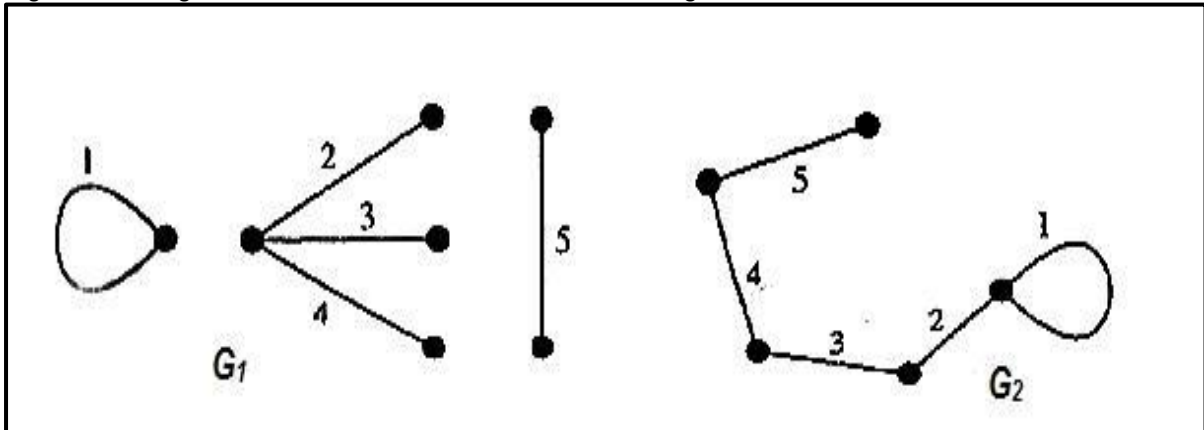
A seguir, investigaremos quais as bases de um Matróide gráfico e de Matróides representáveis.

Seja G um grafo. Observamos anteriormente que um subconjunto X de $E(G)$ é independente em $M(G)$ exatamente quando $G[X]$, que é o subgrafo de G induzido por X , é uma floresta. Assim, X é uma base de $M(G)$ precisamente quando $G[X]$ é uma floresta e, para todo $e \notin X$, $G[X \cup e]$ contém um ciclo. Segue que quando o grafo G é conexo, X é uma base de $M(G)$ se, e somente se, para cada componente H de G tendo pelo menos uma aresta não ciclo, $H[X \cap E(H)]$ é uma árvore geradora de H (uma árvore T , é denominada árvore geradora de um grafo G , se T é um subgrafo de G possui todos os vértices de G).

Definição 2.2.2 Seja X um subconjunto não vazio de $E(G)$, então $G[X]$ é o subgrafo de G induzido por X , que contém X como conjunto de arestas e o conjunto de pontos finais de arestas em X , como o conjunto de vértices.

Exemplo 2.2.2 Considere o grafo G_1 e G_2 na Fig (3), cada um tem um conjunto de arestas $\{1,2,3,4,5\}$. Além disso, cada um dos $M(G_1)$ e $M(G_2)$ têm $\{2,3,4,5\}$ como sua única base. Assim $M(G_1) = M(G_2)$, embora claramente $(G_1) \not\cong (G_2)$.

Figura 3. Dois grafos não isomorfos com ciclo Matróides iguais.



FONTE: Oxley, 1992.

Proposição 2.2.1. Seja M um Matróide gráfico. Então $M \cong M(G)$ para algum grafo conexo G .

Demonstração: Como M é gráfico, $M \cong M(H)$ para algum grafo H . Se H é conexo o resultado está demonstrado. Se não, suponhamos que $H_1, H_2, \dots, H_n$ são as componentes conexas de H . Para cada i em $\{1, 2, \dots, n\}$ escolhemos um vértice v_i em H_i e formamos um novo grafo G por $v_1, v_2, \dots, v_n$ identificando como um vértice único. Claramente $E(H) = E(G)$ e G é conexo. $\square$

Destacamos que na demonstração anterior, o ponto chave sobre a forma como $H_1, H_2, \dots, H_n$ foram unidos é que nenhum ciclo novo foi formado. De fato, esta condição garante que o ciclo Matróide do grafo resultante é isomorfo a $M(H)$.

Com relação aos Matróides representáveis, seja A uma matriz $m \times n$ sobre um corpo F . As colunas de A geram subespaços W de $V(m, F)$ de dimensão $r \leq 2$. Se H é um subconjunto das colunas rotuladas de A , então B é uma base de $M[A]$ se, e somente se, $|B| = r$ e as colunas rotuladas por B formam uma base para W .

Exemplo 2.2.3 Considere A a matriz.

$$A = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 1 & 0 & 1 & 1 \\ 0 & 1 & 1 & -1 \end{pmatrix},$$

sobre $GF(3)$, o corpo com 3 elementos $(-1, 0, 1)$. Neste caso $\dim W = 2$ e as bases de $M[A]$ são todos os subconjuntos de 2 elementos de $\{1, 2, 3, 4\}$. Assim $M[A] \cong U_{2,4}$.

Portanto, $U_{2,4}$ é ternário, isto é, $U_{2,4}$ é representável sobre $GF(3)$. Ou seja, $GF(q)$, é o corpo de Galois de q elementos.

Definição 2.2.3 Um Matróide binário é um Matróide que é representável sobre $GF(2)$.

Definição 2.2.4 Um Matróide regular é um Matróide que pode ser representável por uma matriz totalmente unimodular, i.e, uma matriz sobre $\mathbb{R}$ para qual cada submatriz quadrada tem determinante 0, 1 ou -1.

O Matróide $U_{2,4}$ que tem como conjunto de base $E = \{1,2,3,4\}$ e cuja coleção de conjuntos independentes é formado pelos subconjuntos de E que tem 2 ou menos elementos ou seja, $U_{2,4} = \{\emptyset, \{1\}, \{2\}, \{3\}, \{4\}, \{1,2\}, \{1,3\}, \{1,4\}, \{2,3\}, \{2,4\}, \{3,4\}\}$. Mostremos que $U_{2,4}$ não é binário.

Demonstração: Suponha que $U_{2,4}$ seja representado por uma matriz A e que as colunas do espaço vetorial de A seja de dimensão 2. Então, o maior conjunto de independentes, são os vetores $v_1 = (0,1)$ e $v_2 = (1,0)$. Como são quatro colunas, os outros dois vetores são conjuntos dependentes, já que, não existe quatro vetores diferentes com os elementos 0 e 1 e que seja conjuntos independentes. Concluimos que não existe uma matriz A que seja representada por $U_{2,4}$ sobre $GF(2)$. Portanto, $U_{2,4}$ não é binário. $\square$

2.3 FUNÇÃO RANK

Um Matróide pode ser definido de diversas maneiras diferentes, todas equivalentes. Como já vimos anteriormente, um dos modos de se definir Matróides é por meio de conjuntos independentes. Outra maneira interessante de se definir Matróides é mediante a função *rank*.

A principal ideia da função *rank* é generalizar a noção de dimensão e de conjuntos geradores de um dado espaço vetorial definido num corpo. Nesta seção, vamos definir o conceito de Matróide baseado na função *rank*.

Definição 2.3.1 Seja M o Matróide $(E, \mathcal{I})$ e suponha que $X \subseteq E$. Seja $\mathcal{I} \upharpoonright X = \{I \subseteq X : I \in \mathcal{I}\}$. Então é fácil verificar que o par $(X, \mathcal{I} \upharpoonright X)$ é um Matróide. Denotaremos por $M \upharpoonright X$ ou $M|(E - X)$ esse Matróide restrição de M para X ou inclusão de $E - X$ de M .

Note que $\mathcal{C}(M \upharpoonright X) = \{C \subseteq X : C \in \mathcal{C}(M)\}$, pois todos os subconjuntos de $\mathcal{C}(M \upharpoonright X)$ satisfazem claramente as condições (C1), (C2) e (C3) do Corolário 2.1.1.

Como $M \upharpoonright X$ é um Matróide, e de fato, $|B_1| = |B_2|$ implica que todas as suas bases são equicardinais.

Definição 2.3.2 Seja $(X, \mathcal{I} \upharpoonright X)$ um Matróide. Para cada $X \in 2^E$, definimos o *rank* de X , denotado por $r(X)$, como sendo o número de elementos de uma base B de $M \upharpoonright X$ (B é dita uma *base* de X). A função *rank* r é uma função $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ que, a cada $X \in 2^E$ associa um número inteiro não negativo $r(X)$.

Lema 2.3.1 Seja $(X, \mathcal{I} \upharpoonright X)$ um Matróide. A função *rank* $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ satisfaz as seguintes propriedades:

(R1) Se $X \subseteq E$, então $0 \leq r(X) \leq |X|$.

(R2) Se $X \subseteq Y \subseteq E$, então $r(X) \leq r(Y)$.

(R3) Se X e Y são subconjuntos de E , então.

$$r(X \cup Y) + r(X \cap Y) \leq r(X) + r(Y).$$

Demonstração: As demonstrações das propriedades (R1) e (R2) são diretas. Para provar (R3), seja $B_{X \cap Y}$ uma base para $X \cap Y$. Então $B_{X \cap Y}$ é um conjunto independente em $M \upharpoonright (X \cup Y)$, donde está contido numa base $B_{X \cup Y}$ desse Matróide. Sabemos que $B_{X \cup Y} \cap X$ e $B_{X \cup Y} \cap Y$ são independentes em $M \upharpoonright X$ e $M \upharpoonright Y$, respectivamente. Portanto, $|B_{X \cup Y} \cap X| \leq r(X)$ e $|B_{X \cup Y} \cap Y| \leq r(Y)$.

Assim, $r(X) + r(Y) \geq |B_{X \cup Y} \cap X| + |B_{X \cup Y} \cap Y|$. Utilizando a propriedade $n(A \cup B) + n(A \cap B) = n(A) + n(B)$, segue-se que $|B_{X \cup Y} \cap X| + |B_{X \cup Y} \cap Y| = |(B_{X \cup Y} \cap X) \cup (B_{X \cup Y} \cap Y)| + |(B_{X \cup Y} \cap X) \cap (B_{X \cup Y} \cap Y)| = |B_{X \cup Y} \cap (X \cup Y)| + |(B_{X \cup Y} \cap X \cap Y)|$, em que, na última equação utilizamos o fato de que $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$. Mas $B_{X \cup Y} \cap (X \cup Y) = B_{X \cup Y}$ e $B_{X \cup Y} \cap X \cap Y = B_{X \cap Y}$, donde

$$r(X) + r(Y) \geq |B_{X \cup Y}| + |B_{X \cap Y}|$$

$$r(X) + r(Y) = r(X \cup Y) + r(X \cap Y), \text{ como requerido.}$$

□

Observação 2.3.1 Note que (R3) é correspondente à identidade $\dim(V + W) + \dim(V \cap W) = \dim V + \dim W$ que vale para subespaços V e W de um dado espaço vetorial U , em que $V + W$ consiste em todos os vetores $v + w$, tal que $v \in V$ e $w \in W$. A desigualdade (R3) é conhecida como semimodularidade ou submodularidade.

O teorema apresentado a seguir, é uma espécie de recíproca do Lema 2.3.1. Mais precisamente, o resultado estabelece que se uma função $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ satisfaz as condições (R1), (R2) e (R3), então r origina um Matróide cuja função *rank* r_m é a própria r , ou seja, $r_m = r$.

Teorema 2.3.1 Seja $E \neq \emptyset$ um conjunto e r uma função que mapeia 2^E em um conjunto de inteiros não negativos e que satisfaz (R1) – (R3). Seja $\mathcal{I}$ uma coleção de subconjuntos X de E para qual $r(X) = |X|$. Então $(E, \mathcal{I})$ é um Matróide tendo $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ como sua função *rank*.

A demonstração do teorema utiliza o seguinte resultado.

Lema 2.3.2 Seja E um conjunto, r uma função sobre 2^E satisfazendo (R2) e (R3). Se X e Y são subconjuntos de E tal que para todo y em $Y - X$, $r(X \cup y) = r(X)$ então $r(X \cup Y) = r(X)$.

Demonstração: Seja $Y - X = \{y_1, y_2, \dots, y_k\}$. Iremos demonstrar por indução em k . Se $k = 1$, o resultado é imediato, ou seja, $Y - X = \{y_1\}$ e $r(X \cup Y) = r(X)$.

Suponha que o resultado seja válido para $k = n$ (HI). Devemos provar que $P(n)$ implica $P(n+1)$. Então, assumimos por hipótese de indução e (R3).

$$r(X) + r(X) = r(X \cup \{y_1, y_2, \dots, y_n\}) + r(X \cup y_{n+1}) \geq r((X \cup \{y_1, y_2, \dots, y_n\}) \cup r(X \cup y_{n+1})) + r((X \cup \{y_1, y_2, \dots, y_n\}) \cap r(X \cup y_{n+1}))$$

$$= r(X \cup \{y_1, y_2, \dots, y_{n+1}\}) + r(X).$$

$\geq r(X) + r(X)$, em que a última desigualdade segue por (R2). Como a primeira e a última linha são iguais, segue-se que $r(X \cup \{y_1, y_2, \dots, y_{n+1}\}) = r(X)$. $\square$

Demonstração do Teorema 2.3.1. Por (R1), $0 \leq r(\emptyset) \leq r|\emptyset|$, então $r(\emptyset) = |\emptyset|$ e $\emptyset \in \mathcal{J}$. Assim, $\mathcal{J}$ satisfaz (I1). Agora suponha que $I \in \mathcal{J}$ e $I' \subseteq I$. Então $r(I) = |I|$. Por (R3),

$$r(I' \cup (I - I')) + r(I' \cap (I - I')) \leq r(I') + r(I - I') \text{ isto é,}$$

$$r(I) + r(\emptyset) \leq r(I') + r(I - I') \quad (1)$$

Mas $r(I) = |I|$ e $r(\emptyset) = 0$. Além disso, por (R2), $r(I') \leq |I'|$ e $r(I - I') \leq |I - I'|$. Portanto, por (1) obtemos $|I| \leq r(I') + r(I - I') \leq |I'| + |I - I'| = |I|$.

Como a primeira e a última equações são iguais, segue-se que $r(I') = |I'|$ donde $I' \in \mathcal{J}$.

Para provar que $\mathcal{J}$ satisfaz (I3), iremos assumir o contrário: sejam I_1 e I_2 membros de $\mathcal{J}$ que $|I_1| < |I_2|$ e suponha que para todo e em $I_2 - I_1$ implique que $I_1 \cup \{e\} \notin \mathcal{J}$. Então, para todo e , tem-se que $r(I_1 \cup \{e\}) \neq |I_1 \cup \{e\}|$. Assim, por (R1), (R2) e o fato de que $I_1 \in \mathcal{J}$ segue-se que para todo e , $|I_1| + 1 > r(I_1 \cup \{e\}) \geq r(I_1) = |I_1|$, donde, $r(I_1 \cup \{e\}) = |I_1|$.

Agora, aplicando o Lema 2.3.2 com $X = I_1$ e $Y = I_2$, temos que $r(I_1) = r(I_1 \cup I_2)$. Mas, $|I_1| = r(I_1)$ e $r(I_1 \cup I_2) \geq r(I_2) = |I_2|$, logo $|I_1| \geq |I_2|$, uma contradição. Portanto, $\mathcal{J}$ satisfaz (I3), e assim $(E, \mathcal{J})$ é um Matróide.

Falta mostrar que r é a função *rank* r_m de M . Suponha que $X \subseteq E$. Se $X \in \mathcal{J}$ então $r(X) = |X|$ e, como X é uma base de $M|X$, segue-se que $r_m(X) = |X|$. Assim, se $X \in \mathcal{J}$ e seja B uma base para $M|X$. Então $r_m(X) = |B|$. Além disso, $B \cup \{x\} \notin \mathcal{J}$ para todo x em $X - B$. Assim, $|B| = r(B) \leq r(B \cup x) < |B \cup x|$, daí $r(B \cup x) = r(B)$. Pelo Lema 2.3.2 segue que $r(B \cup x) = r(B)$, que é $r(X) = r(B) = |B|$. Assim se $X \notin \mathcal{J}$, então $r(X) = r_m(X)$. Podemos concluir que $r = r_m$. $\square$

O seguinte corolário é a combinação do Teorema 2.3.1 e do Lema 2.3.1.

Corolário 2.3.1 Seja E um conjunto finito. Uma função $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ é uma função *rank* de um Matróide E se, e somente se, r satisfaz as seguintes condições:

(R1) Se $X \subseteq E$, então $0 \leq r(X) \leq |X|$.

(R2) Se $X \subseteq Y \subseteq E$, então $r(X) \leq r(Y)$

(R3) Se X e Y são subconjuntos de E , então $r(X \cup Y) + r(X \cap Y) \leq r(X) + r(Y)$.

Conjuntos independentes, bases e circuitos são caracterizados em termos da função *rank* como mostra a proposição a seguir.

Proposição 2.3.1 Seja M um Matróide com função *rank* r e suponha que $X \subseteq E$. Então as seguintes afirmações são válidas.

- (i) X é independente se, e somente se, $|X| = r(X)$
- (ii) X é uma base se, e somente se, $|X| = r(X) = r(M)$
- (iii) X é um circuito se, e somente se, X é não vazio e para todo $x \in X$ $r(X - x) = |X| - 1 = r(X)$.

Agora especificaremos a função *rank* para a classe de Matróides uniformes, representável e gráfico. Em cada caso suponhamos que $X \subseteq E(M)$. Se $M = U_{m,n}$ então claramente

$$r(X) = \begin{cases} |X|, & \text{se } |X| < m \\ m, & \text{se } |X| \geq m. \end{cases}$$

Se $M = M[A]$ onde A é uma matriz $m \times n$ sobre o corpo F , então, claramente, $r(X)$ é o *rank* de A_X , e o $m \times |X|$ submatrix de A constituída pelas colunas de A que são rotuladas por membros de X . Equivalentemente, $r(X)$ é igual à dimensão do subespaço de $V(m, F)$ que é abrangido pelas colunas de A_X .

Seja $M = M(G)$, onde G é um grafo e $V(T)$ número de vértices. Primeiramente determinaremos $r(M)$. Suponhamos inicialmente que G seja conexo; então uma base de G é um conjunto de arestas de uma árvore geradora em G . Utilizando-se indução matemática, pode-se facilmente demonstrar que, para uma árvore T vale $|V(T)| = |E(T)| + 1$. Portanto, se G é conexo, então

$$2.3.1. \quad r(M) = |V(G)| - 1$$

Estendendo-se 2.3.1, tem-se que se G tem $w(G)$ componentes conexas, então

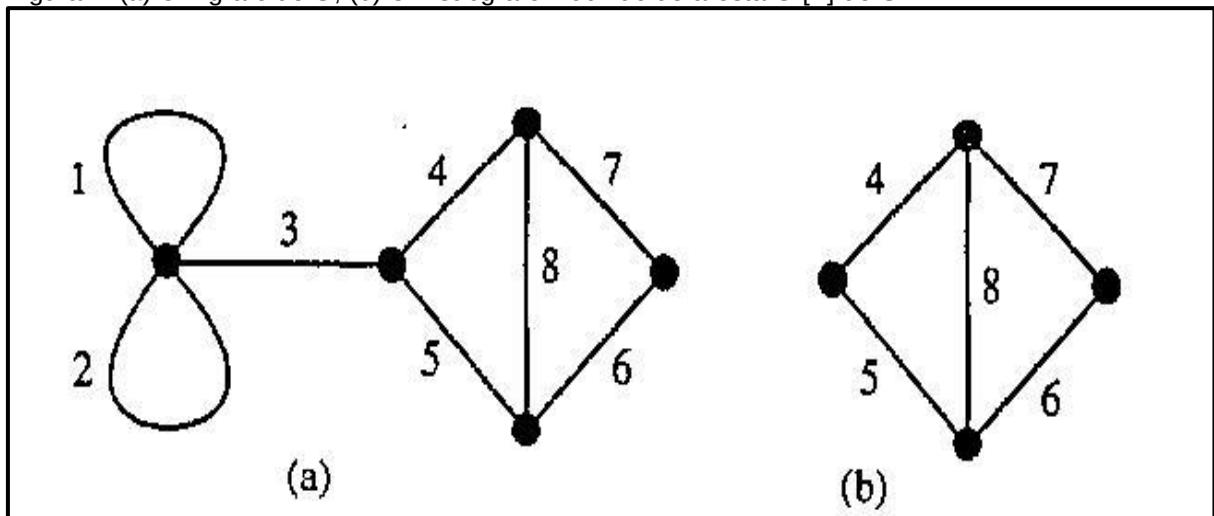
$$2.3.2. \quad r(M) = |V(G)| - w(G).$$

Segue que se $X \subseteq E(G)$, então

$$2.3.3. \quad r(X) = |V(G|X)| - w(G|X).$$

Exemplo 2.3.1 Seja $M = M(G)$, em que G é o grafo da Fig. 4 (a). Então, como G é conexo, $r(M) = |V(G)| - 1 = 4$. Se $X = \{4,5,6,7,8\}$, então uma base para $M(G)$ é $\{4,5,6\}$, assim $r(\{4,5,6,7,8\}) = 3$. Equivalentemente, podemos ver na Figura 4 (b), que $|V(G[X])| = 4$ e $w(G[X]) = 1$. Assim por 2.3.3, conclui-se que $r(X) = 3$. Se $Y = \{1,3,7\}$ então $|V(G[Y])| = 4$ e $w(G[Y]) = 2$; assim $r(Y) = 2$. Finalmente, se $Z = \{1,2\}$ então $r(Z) = 0$ porque $\emptyset$ é uma base para $M|Z$. $\square$

Figura 4. (a) Um grafo de G , (b) Um subgrafo induzido de aresta $G[X]$ de G .



FONTE: Oxley, 1992.

Proposição 2.3.2 Um Matróide M é uniforme, se, e somente se, M não tem circuitos de tamanho inferior a $r(M) + 1$.

Demonstração: Segue diretamente da definição de Matróide uniforme.

Definição 2.3.3 Diz-se que um Matróide M é *paving* se o mesmo não possui circuito de tamanho menor que $r(M)$.

Assim, a classe dos Matróides uniformes está contida na classe de Matróides *paving*. Entretanto, existem Matróides *paving* que não são uniformes. Um exemplo é o Matróide $M(K_4)$.

A proposição seguinte caracteriza Matróides *paving* em termos de sua coleção de circuitos.

Proposição 2.3.3 Seja D uma coleção de subconjuntos não vazios de um conjunto E . Então D é o conjunto de circuitos de um Matróide *paving* de *rank*- r em E se, e somente se, para algum inteiro positivo k , há uma partição (D', D'') de D tal

- (i) Cada membro de D' possui k elementos e, se dois membros distintos D_1 e D_2 de D' têm $k - 1$ elementos comuns, então cada subconjunto de k elementos de $D_1 \cup D_2$ está em D' ; e
- (ii) D'' consiste em todos os subconjuntos de $(k + 1)$ elementos de E que não contêm nenhum membro de D' .

2.4 OPERADOR FECHO

Um dos principais resultados desta seção caracterizará os Matróides em termos de seus respectivos operadores fecho.

Em um espaço vetorial V , o vetor v pertence ao espaço gerado por $\{v_1, v_2, \dots, v_m\}$ se o subespaço gerado por $\{v_1, v_2, \dots, v_m\}$ e por $\{v_1, v_2, \dots, v_m, v\}$ possuem a mesma dimensão. Isso motiva a seguinte definição de operador fecho de um Matróide.

Definição 2.4.1 Seja M um Matróide arbitrário tendo E como seu conjunto fundamental e r a função *rank* de M . Seja cl uma aplicação $cl : 2^E \rightarrow 2^E$ definida para todo $X \subseteq E$ como sendo $cl(X) = \{x \in E : r(X \cup x) = r(X)\}$. Esta aplicação é denominada *operador fecho* de M .

Exemplo 2.4.1 Considerando o Matróide construído no Exemplo 2.3.1, note que $cl(\emptyset) = \{1, 2\}$, $cl(\{1, 3, 5\}) = \{1, 2, 3, 5\}$ e $cl(\{4, 5, 6\}) = \{1, 2, 4, 5, 6, 7, 8\}$.

Veremos, a seguir, algumas das principais propriedades básicas do operador fecho.

Lema 2.4.1 O operador fecho de um Matróide sobre o conjunto E tem as seguintes propriedades:

(CL1) Se $X \subseteq E$, então $X \subseteq cl(X)$.

(CL2) Se $X \subseteq Y \subseteq E$, então $cl(X) \subseteq cl(Y)$.

(CL3) Se $X \subseteq E$, então $cl(cl(X)) = cl(X)$.

(CL4) Se $X \subseteq E$, $x \in E$ e $y \in cl(X \cup x) - cl(X)$, então $x \in cl(X \cup y)$.

Demonstração: O fato de que o operador fecho satisfaz (CL1) é imediato pela definição, pois, se $x \in X$ então $r(X \cup x) = r(X)$, donde $X \subseteq cl(X)$. Para provar (CL2), supor que $X \subseteq Y$ e $x \in cl(X) - X$. Então $r(X \cup x) = r(x)$. Assim, se B_x é à base de X , então B_x é à base de $X \cup x$. Portanto $Y \cup x$ tem uma base $B_{Y \cup x}$ que contém B_x mas, não contém x . Como $B_{Y \cup x}$ também deve ser uma base de Y , segue-se que $r(Y \cup x) = |B_{Y \cup x}| = r(Y)$ e, por conseguinte, $x \in cl(Y)$.

Para provar (CL3), note que por (CL1), $cl(X) \subseteq cl(cl(X))$. Agora, para estabelecer a inclusão inversa, escolhamos x em $cl(cl(X))$. Então $r(cl(X) \cup x) = r(cl(X))$. Mas, para todo y em $cl(X) - X$, $r(X \cup y) = r(X)$. Assim, pelo Lema 2.3.3, $r(X) = r(X \cup (cl(X) - X)) = r(cl(X))$. Daqui $r(cl(X) \cup x) = r(X)$. Mas por (R2), $r(cl(X) \cup x) \geq r(X \cup x) \geq r(x)$. Portanto, $cl(x)$. Logo $cl(cl(X)) \subseteq cl(X)$ e (CL3) está demonstrado.

Para provar (CL4), utilizaremos o resultado a seguir.

Lema 2.4.2 Se $X \subseteq E$ e $x \in E$, então $r(X) \leq r(X \cup x) \leq r(X) + 1$.

Demonstração: Seja B_x uma base de X . Também B_x ou $B_x \cup x$ é uma base de $X \cup x$ e assim $r(X \cup x)$ é igual a $r(X)$ ou $r(X) + 1$. $\square$

Agora, suponha que $y \in cl(X \cup x) - cl(X)$. Então $r(X \cup x \cup y) = r(X \cup x)$ e $r(X \cup y) \neq r(X)$. Partindo da última inequação e pelo Lema 2.4.2, podemos deduzir que $r(X \cup y) = r(X) + 1$. Assim, $r(X) + 1 = r(X \cup y) \leq r(X \cup y \cup x) = r(X \cup x) \leq r(X) + 1$, e então $r(X \cup y \cup x) = r(X \cup y)$, donde $x \in cl(X \cup y)$, como requerido. $\square$

As condições (CL1) - (CL4) caracterizam as aplicações que podem ser operadores fecho de Matróides. Às vezes (CL4) é referido como o *Mac Lane-Steinitz exchange property* propriedade de troca de MLS.

Teorema 2.4.1 Seja E um conjunto e cl uma aplicação $cl : 2^E \rightarrow 2^E$ satisfazendo (CL1) - (CL4). Seja $\mathcal{I} = \{X \subseteq E : x \notin cl(X - x) \text{ para todo } x \text{ em } X\}$. Então $(E, \mathcal{I})$ é um Matróide em que cl é sua respectiva aplicação fecho.

Demonstração: Evidentemente $\emptyset \in \mathcal{I}$, assim $\mathcal{I}$ satisfaz (I1). Suponha que $I \in \mathcal{I}$ e $I' \subseteq I$. Se $x \in I'$, então $x \in I$ assim $x \notin cl(I - x)$. Por (CL2), $cl(I - x)$ contém $cl(I' - x)$, então $x \notin cl(I' - x)$ e, por conseguinte $I' \in \mathcal{I}$. Assim $\mathcal{I}$ satisfaz (I2).

Para o restante da demonstração utilizaremos o seguinte resultado.

Lema 2.4.3 Suponha $X \subseteq E$ e $x \in E$. Se X está em $\mathcal{I}$, mas $X \cup x$ não está, então $x \in cl(X)$.

Demonstração: Como $X \cup x \notin \mathcal{I}$, existe um elemento y de $X \cup x$ tal que $y \in cl((X \cup x) - y)$. Se $y = x$, então o lema é válido. Se $y \neq x$, então $(X \cup x) - y = (X - y) \cup x$ e $y' \in cl((X - y) \cup x) - cl(X - y)$. Portanto, por (CL4) $x \in cl((X - y) \cup y) = cl(X)$. $\square$

Vamos provar que $\mathcal{I}$ satisfaz (I3). Suponha que I_1 e I_2 sejam membros de $\mathcal{I}$ tal que $|I_1| < |I_2|$ e (I3) falha para o par (I_1, I_2) . Assuma, além disso, que entre todos esses pares, $|I_1 \cap I_2|$ seja maximal. Escolher y em $I_2 - I_1$ e considerar $I_2 - y$. Assumir que $I_1 \subseteq cl(I_2 - y)$. Então, por (CL2) e (CL3), $cl(I_1) \subseteq cl(I_2 - y)$. Como $y \notin cl(I_2 - y)$, $y \notin cl(I_1)$. Então pelo Lema 2.4.3, $I_1 \cup y \in \mathcal{I}$ e assim (I3) vale para o par (I_1, I_2) , uma contradição. Concluimos então que $I_1 \not\subseteq cl(I_2 - y)$ e assim para algum elemento t de I_1 não está em $cl(I_2 - y)$. Evidentemente, $t \in I_1 - I_2$. Além disso, pelo Lema 2.4.3 novamente, $(I_2 - y) \cup t \in \mathcal{I}$. Como $|I_1 \cap ((I_2 - y) \cup t)| > |I_1 \cap I_2|$, (I3) vale para o par $(I_1, (I_2 - y) \cup t)$, isto é, para algum x em $((I_2 - y) \cup t) - I_1$, o conjunto $I_1 \cup x \in \mathcal{I}$. Mas $x \in I_2 - I_1$, assim (I3) é satisfeito para o par (I_1, I_2) . Isso completa a demonstração que $(E, \mathcal{I})$ é um Matróide M .

Agora temos de verificar que cl e o operador fecho cl_M de M coincidem. Suponha primeiro que $x \in cl_M(X) - X$. Então $r_M(X \cup x) = r_M(X)$. Seja B uma base de X . Então $B \in \mathcal{I}$ e $B \cup x \notin \mathcal{I}$; assim, pelo Lema 2.4.3, $x \in cl(B)$. Mas, por (CL2), $cl(B) \subseteq cl(X)$. Portanto, $x \in cl(X)$ e assim, $cl_M(X) \subseteq cl(X)$. Para demonstrar a desigualdade inversa, suponha que x está em $cl(X) - (X)$. Seja B uma base para X . Então $B \cup y \notin \mathcal{I}$ para algum y em $X - B$. Assim, pelo Lema 2.4.3, $X \subseteq cl(B)$, donde $cl(X) \subseteq cl(B)$. Como $X \subseteq cl(X)$, segue-se que $x \in cl(B)$. Então $B \cup x \notin \mathcal{I}$, donde conclui-se que B é

uma base para $X \cup x$ e $r_M(X \cup x) = |B| = r_M(X)$. Daqui $x \in cl_M(X)$ e concluímos que $cl(X) \subseteq cl_M(X)$ e, por conseguinte que $cl(X) = cl_M(X)$. $\square$

Podemos concluir que se, uma aplicação que satisfaz (CL1)-(CL4) se, e somente se, é uma aplicação fecho é considerado um Matróide.

Corolário 2.4.1 Seja E um conjunto finito. A aplicação $cl: 2^E \rightarrow 2^E$ é operador fecho de um Matróide sobre E se, e somente se, satisfaz as seguintes condições:

(CL1) Se $X \subseteq E$, então $X \subseteq cl(X)$.

(CL2) Se $X \subseteq Y \subseteq E$, então $cl(X) \subseteq cl(Y)$.

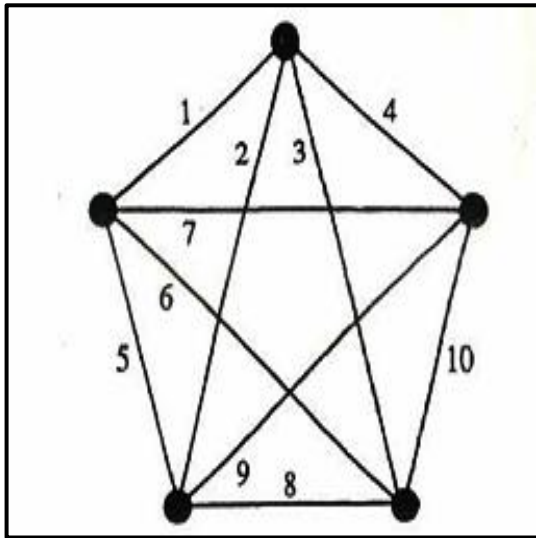
(CL3) Se $X \subseteq E$, então $cl(cl(X)) = cl(X)$.

(CL4) Se $X \subseteq E$, $x \in E$, e $y \in cl(X \cup x) - cl(X)$, então $x \in cl(X \cup y)$.

Se M é um Matróide e $X \subseteq E(M)$, denotamos $cl(X)$ como sendo o fecho ou gerador de X em M ; às vezes escreve-se como $cl_M(X)$ se consideramos mais de um Matróide. Alguns autores denotam $cl(X)$ por $\bar{X}$.

Definição 2.4.2 Se $X = cl(X)$, então X é dito um conjunto fechado (*flat*) de M . Um hiperplano de M é um conjunto fechado de $rank\ r(M) - 1$. Um subconjunto de X de $E(M)$ é um conjunto gerador de M se $cl(X) = E(M)$.

Exemplo 2.4.2 Seja $M = M(K_5)$ onde as arestas de K_5 são rotuladas na Figura 5. Então M tem um único conjunto fechado de $rank\ 0$, ou seja $\emptyset$; M tem dez conjuntos de fechados $rank-1$, as dez arestas de K_5 . Os planos de $rank-2$ de M são de dois tipos: conjuntos de triângulos de arestas, dos quais há dez; e pares de arestas não adjacentes, das quais há quinze. Os *flats* ou hiperplanos de $rank-3$ também são de dois tipos: conjuntos de arestas de subgrafos de K_4 , dos quais existem cinco; e conjunto de arestas de subgrafos isomorfos à união disjunta de um K_2 e K_3 , dos quais há dez. Finalmente, o único *flat* de $rank-4$ é $E(K_5)$. $\square$

Figura 5. K_5 

FONTE: Oxley, 1992.

A seguinte proposição mostra que as bases e hiperplanos são conjuntos geradores e não geradores, respectivamente. Para sua demonstração, utilizaremos o seguinte lema.

Lema 2.4.4 Suponha que M seja um Matróide e $X \subseteq E(M)$. Se $x \in cl(X)$, então $cl(X \cup x) = cl(X)$.

Demonstração: Como $X \subset X \cup x$ então $cl(X) \subset cl(X \cup x)$. Falta mostrar a inclusão inversa. Supor que y pertença a $cl(X \cup x)$. Por definição, $r(X \cup x \cup y) = r(X \cup x)$. Mas como $x \in cl(X)$, segue-se que $r(X \cup x) = r(X)$, ou seja, $r(X \cup x \cup y) = r(X)$. Como $X \cup y \subset X \cup x \cup y$, tem-se que $r(X) \leq r(X \cup y) \leq r(X \cup x \cup y) = r(X)$, isto é, $r(X \cup y) = r(X)$, donde $y \in cl(X)$. Assim, $cl(X \cup x) = cl(X)$. $\square$

Agora, estamos preparados para demonstrar a proposição dada a seguir.

Proposição 2.4.1 Seja M um Matróide e X um subconjunto de $E(M)$. Então

- (i) X é um conjunto gerador se, e somente se, $r(X) = r(M)$.
- (ii) X é uma base se, e somente se, esta é um conjunto gerador minimal; e
- (iii) X é um hiperplano se, e somente se, é um conjunto maximal não gerador.

Demonstração: Suponha que X seja um conjunto gerador de M . Então $cl(X) = E$. Seja B uma base de X . Então para todo x em $X - B$, o conjunto $B \cup x \notin \mathcal{I}(M)$. Assim, pelo Lema 2.4.3, $x \in cl(B)$. Daqui $X \subseteq cl(B)$ e assim $E = cl(X) \subseteq cl(B) \subseteq$

E . Portanto, para todo y em $E - B$, $y \in cl(B)$, assim, $B \cup y \notin \mathcal{I}$. Portanto, B é uma base de M , donde $r(X) = |B| = r(M)$.

(ii) Suponha que X seja um conjunto gerador minimal de M . Então por (i), $r(X) = r(M)$. Além disso, para todo x em X , o conjunto $X - x$ não é gerador e, pelo Lema 2.4.4, $x \notin cl(X - x)$. Portanto, $X \in \mathcal{I}(M)$. Como $r(X) = r(M)$, segue-se que X é uma base de M .

Mostraremos que $r(E(M) - f) = r(E(M)) - 1$. Suponha que exista uma base B , tal que $f \notin B$. Observe que $B \subseteq E(M) - f$, logo $cl(B) \subseteq E(M) - f$, mas $cl(B) = E(M)$ e assim $cl(E(M) - f) = E(M)$, absurdo, pois supusemos que $E(M) - f$ é fechado e tal base B não existe. Como f está contido em todas as bases, segue que $r(E(M) - f) = r(E(M)) - 1$. Desta forma concluímos que $r(M) - f$ é um hiperplano. $\square$

Deste resultado, deduzimos que se G é um grafo conexo, um subconjunto X de $E(G)$ é gerado em $M(G)$ se, e somente se, X contém o conjunto de arestas de uma árvore geradora G .

A seguir, relacionaremos o conceito de operador fecho com conjuntos independentes, bases, conjuntos geradores e hiperplanos.

Proposição 2.4.2 Seja M um Matróide e X um subconjunto de $E(M)$. Então:

(i) X é um circuito de M se, e somente se, X é um conjunto minimal com a propriedade que, para todo x em X , $x \in cl(X - x)$.

(ii) $cl(X) = X \cup \{x: M \text{ tem um circuito } C \text{ tal que } x \in C \subseteq X \cup x\}$

Demonstração: O item (i) apenas reafirma o fato de que um circuito é um conjunto dependente minimal. Para provar (ii), suponha primeiro que $x \in cl(X) - X$. Então $r(X \cup x) = r(X)$. Portanto, se B é uma base de X , então $B \cup x$ é dependente. Assim, pelo Corolário 2.2.2, há um circuito C tal que $x \in C \subseteq B \cup x$. Daqui $x \in C \subseteq X \cup x$. Portanto, se existir tal circuito, pelo Item (i) e por (CL2), $x \in cl(C - x) \subseteq cl(X)$, donde (ii) é válida. $\square$

A seguir, provaremos que a condição (C3) pode ser fortalecida do seguinte modo:

Proposição 2.4.3 Seja $\mathcal{C}$ um conjunto de circuitos de um Matróide M . Então $\mathcal{C}$ satisfaz a seguinte condição:

(C3)' Se C_1 e C_2 são membros de $\mathcal{C}$, $e \in C_1 \cap C_2$, e $f \in C_1 - C_2$, então existe um membro C_3 de $\mathcal{C}$ tal que $C_3 \subseteq (C_1 \cup C_2) - \{e\}$.

Demonstração: Como $e \in cl(C_2 - \{e\})$ e $C_2 - e \subseteq (C_1 \cup C_2) - \{e, f\}$, o elemento e está em $cl((C_1 \cup C_2) - \{e, f\})$. Então, pelo Lema 2.4.4, $cl((C_1 \cup C_2) - \{e, f\}) = cl((C_1 \cup C_2) - f)$. Mas $f \in cl(C_1 - f) \subseteq cl((C_1 \cup C_2) - f)$, e então $f \in cl((C_1 \cup C_2) - \{e, f\})$. Pela Proposição 2.4.2 (ii), M tem um circuito C tal que $f \in C \subseteq (C_1 \cup C_2) - e$, isto é, $\mathcal{C}$ satisfaz (C3)'. $\square$

Corolário 2.4.2 Seja $\mathcal{C}$ uma coleção de subconjuntos do conjunto E . Então $\mathcal{C}$ é o conjunto de circuitos de um Matróide em E se, e somente se, $\mathcal{C}$ satisfaz (C1), (C2) e (C3)'.

Vamos ver quais são os *flats* e hiperplanos em alguns tipos especiais de Matróides.

Em $U_{m,n}$ os *flats* consistem de todos os conjuntos com menos de m elementos junto com $E(U_{m,n})$. Os hiperplanos são os conjuntos com exatamente $m - 1$ elemento. Se $M = M[A]$, em que A é uma matriz $m \times n$ sobre o corpo F , então X é *flat* de M se, e somente se, existe um subespaço W de $V(m, F)$ tal que X é o conjunto dos rótulos do multiconjunto consistindo por essas colunas de A que são membros de W . Além disso, embora o subespaço W possa ser escolhido para ter dimensão $r(X)$, é claro que, em geral, a interseção de um subespaço de $V(m, F)$ com o multiconjunto de colunas de A não precisa gerar o subespaço.

Se $M = M[G]$ para algum grafo G , então, pela Proposição 2.4.2 (ii), um subconjunto X de $E(G)$ é *flat* em M se, e somente se, G não contém nenhum ciclo C , tal que X contenha todos exceto exatamente uma aresta C . Esta última é válida se, e somente se, para algum número inteiro positivo n , existe uma partição $\{V_1, V_2, \dots, V_n\}$ de $V(G)$ tal que X é a união de conjunto de arestas de $G[V_1], G[V_2], \dots, G[V_n]$.

Proposição 2.4.4 Seja G um grafo. Então H é um hiperplano em $M(G)$ se, e somente se, $E(G) - H$ é um conjunto minimal de arestas cuja remoção de G aumenta o número de componentes conexas.

Demonstração: Seja $E(G) - H$ um conjunto mínimo de arestas cuja remoção de G aumenta o número de componentes conexas. Então uma aresta e de G que

não está em H junta duas componentes distintas G_1 e G_2 de $G[H]$. Além disso, cada aresta de $E(G) - H$ deve ligar G_1 e G_2 . Daqui $H \cup e$ gera $M(G)$. Assim, H é um conjunto gerador não maximal de $M(G)$. Portanto, pela Proposição 2.4.3(ii), H é um hiperplano. $\square$

Teorema 2.4.2 Seja E um conjunto finito. A função $r : 2^E \rightarrow \mathbb{Z}^+ \cup \{0\}$ é a função *rank* de um Matróide sobre E se, e somente se, satisfaz as seguintes condições:

$$(R1)' \quad r(\emptyset) = 0.$$

$$(R2)' \quad \text{Se } X \subseteq E \text{ e } x \in E \text{ então, } r(X) \leq r(X \cup x) \leq r(X) + 1.$$

$$(R3)' \quad \text{Se } X \subseteq E \text{ e } x \text{ e } y \text{ são membros de } E \text{ tal que } r(X \cup x) = r(X \cup y) = r(X), \text{ então } r(X \cup x \cup y) = r(X).$$

No próximo capítulo, apresentaremos que um dual também é um Matróide sobre a mesma base.

3. DUALIDADE

Nesta seção, definiremos o dual de um Matróide e mostraremos que este dual também é um Matróide sobre o mesmo conjunto base.

Uma das características mais atraentes da teoria de Matróides é a existência da teoria de dualidade. Esta teoria, que foi introduzida por Whitney (1935), se estende de forma natural para noção de ortogonalidade em espaços vetoriais. A teoria de Matróides infinitos fornece ampla evidência do papel vital que a dualidade desempenha na teoria de Matróide. No contexto infinito, não há nenhuma teoria de dualidade tendo a riqueza da teoria no contexto finito. Esta parece ser uma das principais razões pelas quais a atenção é geralmente restrita ao caso finito, na maioria dos desenvolvimentos de Matróides. No caso de Matróides infinitos foi introduzido por Wesh (1976) e Oxley (1992).

3.1. AS DEFINIÇÕES E PROPRIEDADES BÁSICAS.

Nesta seção, definimos o Matróide dual, demonstrando que também é um Matróide e estabelecemos algumas conexões fundamentais entre Matróides e seus duais.

Teorema 3.1.1 Seja M um Matróide e $B^*(M) = \{E(M) - B : B \in \mathcal{B}(M)\}$. Então $B^*(M)$ é o conjunto de bases de um Matróide em $E(M)$.

A demonstração deste teorema usará o seguinte resultado.

Lema 3.1.1 O conjunto $\mathcal{B}$ de bases de um Matróide M satisfaz a seguinte condição:

(B2)* Se B_1 e B_2 estão em $\mathcal{B}$ e $x \in B_2 - B_1$, então existe um elemento y de $B_1 - B_2$ tal que $(B_1 - y) \cup \{x\} \in \mathcal{B}$.

Note que existe uma diferença real de (B2)* e (B2), que não podem ser conseguida simplesmente por mudança de estilo.

Demonstração do Lema 3.1.1. Pelo Corolário 2.2.2, $B_1 \cup x$ contém um único circuito $C(x, B_1)$. Note que $C(x, B_1) - B_2$ é não vazio, pois $C(x, B_1)$ é dependente e B_2 é independente. Seja y um elemento deste conjunto. Evidentemente, $y \in B_1 - B_2$. Além disso, como $(B_1 - y) \cup x$ não contém $C(x, B_1)$, $(B_1 - y) \cup x$ é independente.

Como $(B_1 - y) \cup x$ possui a mesma cardinalidade de B_1 e é independente, este é uma base. $\square$

Demonstração do Teorema 3.1.1. Com $\mathcal{B}(M)$ é não vazio, $\mathcal{B}^*(M)$ também é não vazio. Assim, $\mathcal{B}^*(M)$ satisfaz (B1). Suponha que B_1^* e B_2^* sejam membros de $\mathcal{B}^*(M)$ e $x \in B_1^* - B_2^*$. Para $i = 1, 2$, tomamos $B_i = E(M) - B_i^*$. Então $B_i \in \mathcal{B}(M)$ e $B_1^* - B_2^* = B_2 - B_1$. Por $(B_2)^*$, como $x \in B_2 - B_1$, existe um elemento $y \in B_1 - B_2$, tal que $(B_1 - y) \cup x \in \mathcal{B}(M)$. Disto segue que $y \in B_2^* - B_1^*$ e que $E(M) - ((B_1 - y) \cup x) \in \mathcal{B}^*(M)$. Mas $E(M) - ((B_1 - y) \cup x) = ((E(M) - B_1) - x) \cup y = (B_1^* - x) \cup y$. Concluimos que $\mathcal{B}^*(M)$ satisfaz (B2). Então $\mathcal{B}^*(M)$ é, de fato, o conjunto de bases de um Matróide em $E(M)$. $\square$

O Matróide construído no teorema acima, cujo conjunto base é $E(M)$ e cuja família de bases é $\mathcal{B}^*(M)$, é denominado dual de M e é denotado por M^* . Assim, $\mathcal{B}(M^*) = \mathcal{B}^*(M)$. Além disso, é evidente que 3.1.1. $(M^*)^* = M$.

Exemplo 3.1.1 Considere o Matróide uniforme $U_{m, n}$. As bases deste Matróide são todos subconjuntos de E com m elementos, em que $|E| = n$. Assim $\mathcal{B}^*(U_{m, n})$ consiste de todos os subconjuntos de E com $n - m$ elementos e então $U_{m, n}^* = U_{n-m, n}$.

O corolário a seguir mostra uma caracterização de um Matróide em termos de sua coleção de bases.

Corolário 3.1.1 Seja $\mathcal{B}$ uma coleção de subconjuntos do conjunto finito E . Então $\mathcal{B}$ é a coleção de bases de um Matróide em E se, e somente se, $\mathcal{B}$ satisfaz (B₁) e $(B_2)^*$.

Demonstração: Se $\mathcal{B}$ é a coleção de bases de um Matróide, então, pelo Corolário 2.2.1 e pelo Lema 3.1.1, $\mathcal{B}$ satisfaz (B₁) e $(B_2)^*$. Agora suponha que $\mathcal{B}$ satisfaz (B₁) e $(B_2)^*$ e considere $\mathcal{B}' = \{E - B : B \in \mathcal{B}\}$. Claramente $\mathcal{B}'$ satisfaz (B₁) e (B₂) assim, pelo Corolário 2.2.1, $\mathcal{B}'$ é o conjunto de bases de um Matróide M em E . Portanto, $\mathcal{B}$ é o conjunto de bases de um Matróide em E , a saber, M^* . $\square$

Definição 3.1.1 Uma base de M^* é denominada *cobase* de M . A mesma convenção se aplica também aos subconjuntos de $E(M^*)$. Por exemplo, circuitos, hiperplanos, conjuntos independentes e conjuntos geradores de M^* são denominados cocircuitos, cohiperplanos, conjuntos coindependentes e conjuntos cogeradores de M , respectivamente.

O próximo resultado descreve algumas relações elementares entre esses conjuntos.

Proposição 3.1.1 Seja M um Matróide em um conjunto finito E e suponha $X \subseteq E$. Então

- (i) X é independente se, e somente se, $E - X$ é cogerador.
- (ii) X é gerador se, e somente se, $E - X$ é coindependente.
- (iii) X é um hiperplano se, e somente se, $E - X$ é cocircuito e
- (iv) X é um circuito se, e somente se, $E - X$ é um cohiperplano.

Demonstração: Seja X um conjunto independente de M . Então $X \subset B$ para algum elemento da base B de M . Portanto, $E - B \subset E - X$ e $cl^*(E - B) = E \subset cl^*(E - X)$ em que cl^* é o fecho em M^* . Se $E - X$ é cohiperplano, então $E - X \supset B^*$ para algum elemento da base B^* de M^* . Isto significa que $X \subset E - B^*$ e X é um conjunto independente. (ii) é deduzido pela aplicação (i) para M^* . (iii) é obtido pelas seguintes afirmações equivalentes; a) X é um hiperplano de M . b) X é um conjunto não gerador de M , mas $X \cup y$ é gerador para todo $y \notin X$. c) $E - X$ é dependente de M^* mas, $(E - X) - y$ é independente em M^* para todo $y \in E - X$. d) $E - X$ é um cocircuito de M . (iv) é obtido pela aplicação (iii) para M^* . $\square$

Uma consequência do último resultado é que se X é um circuito hiperplano de um Matróide M , então $E(M) - X$ é um circuito hiperplano de M^* . Além disso, a operação relaxamento tem a seguinte propriedade.

Definição 3.1.2 O Matróide $M(E)$ é um relaxamento do Matróide $N(E)$, se para algum circuito-hiperplano H de N , o conjunto de bases de M é o conjunto de bases de N união com H . Neste caso, dizemos que M é obtido de relaxamento de H .

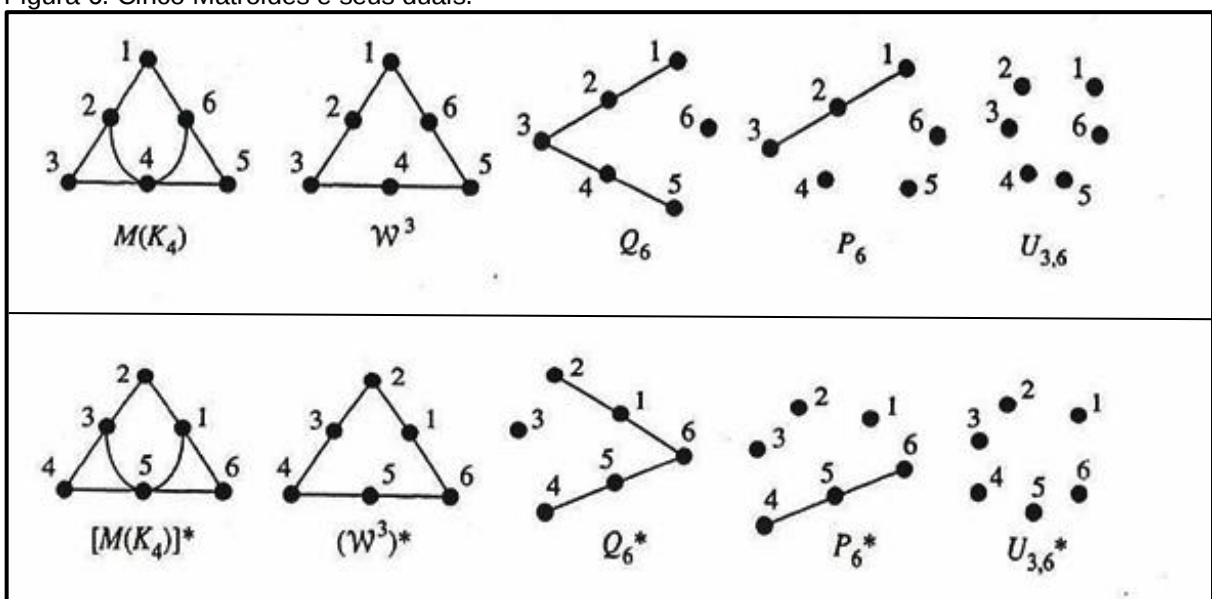
Exemplo 3.1.2. Seja C um circuito hiperplano (ou seja, um circuito e um hiperplano) de M , então há também um Matróide M' cujos conjuntos independentes são apenas os conjuntos independentes de M mais o conjunto C . Descreveremos que M' é obtido a partir de M , com relaxamento de C . Por exemplo, para o Matróide (binário) no conjunto $\{a, b, c, d\}$ com circuitos $\{a, b\}$, $\{a, c, d\}$, $\{b, c, d\}$, então $U_{2,4}$ é obtido de M , com relaxamento do circuito $\{a, b\}$. Este é um exemplo do relaxamento M' de um Matróide binário M .

Proposição 3.1.2 Seja M' obtido de M relaxamento o circuito hiperplano X de M , então $(M')^*$ pode ser obtido a partir de M^* relaxamento o circuito hiperplano $E(M) - X$ de M^* .

Demonstração: Como $B(M') = B(M) \cup \{x\}$, temos $B((M')^*) = B(M^*) \cup \{E(M) - X\}$, e o resultado segue imediato. $\square$

A proposição 3.1.2 é ilustrada na figura 6. Os Matróides na segunda fila são duais daquela da primeira fila. Em ambas as linhas, cada uma dos quatros Matróides é obtido a partir do seu antecessor, um relaxamento circuito hiperplano. Claramente cada uma $M(K_4)$, W^3 , Q_6 e P_6 são isomorfos a seu dual, enquanto $U_{3,6}$ é igual a seu dual. Nos seguintes Bondy e Welsh (1971) em chamar M , auto- dual se $M \cong M^*$ e identificar auto- dual se $M = M^*$.

Figura 6. Cinco Matróides e seus duais.



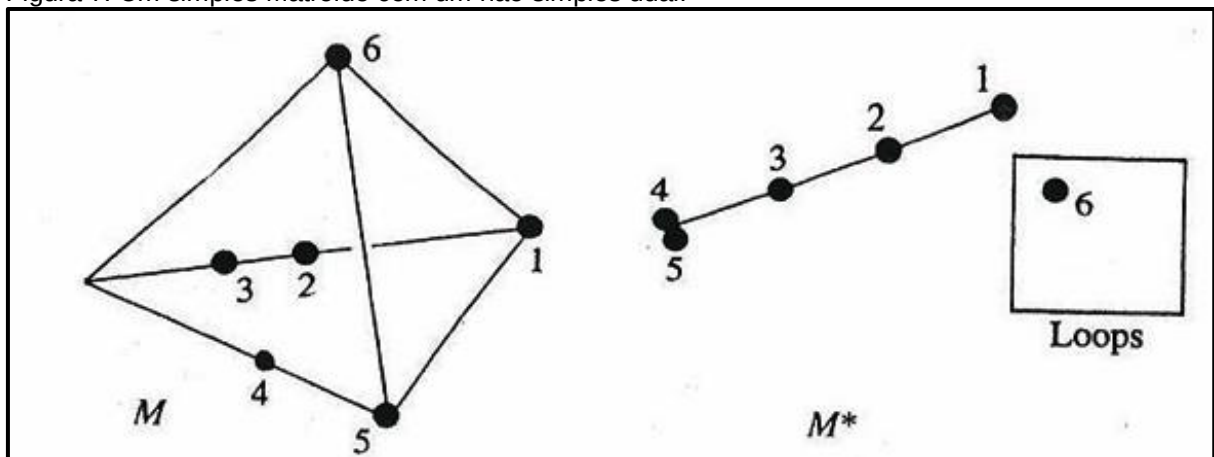
FONTE: Oxley, 1992

Note, no entanto, que alguns autores reservam o termo "auto-dual" para Matróides M , de modo que $M = M^*$.

O Matróide M na Figura 7 é simples. Porém, ele não é dual de M^* . Isto é principalmente por causa da classe de Matróides simples não é fechada sob dualidade então, o Matróide não é para os teóricos, em geral, uma limitação de Matróide simples; a dualidade é uma parte importantíssima da teoria dos Matróides. Isso fornece um contraste entre a teoria do Matróide e a teoria de grafos, pois, neste último, muitos trabalhos se concentram exclusivamente em grafos simples.

Na Figura 7, o elemento 6 que é um laço de M^* , é um colaço de M . Como 6 não é uma base de M^* , está em cada cobases de M^* , isso é, o elemento 6 está em todas as base de M .

Figura 7. Um simples Matróide com um não simples dual.



FONTE: Oxley, 1992.

Em geral, usamos asterisco para denotar a associação com o dual de um dado Matróide. Assim, por exemplo, r^* é a função *rank* de M^* e $\mathcal{C}^*$ denota o conjunto de circuitos de M^* . Evidentemente, tem-se que $r(M) + r^*(M) = |E(M)|$.

Proposição 3.1.3 Para todo subconjunto $X \subset E$ de um Matróide M , tem-se que

$$r^*(X) = |X| - r(M) + r(E - X).$$

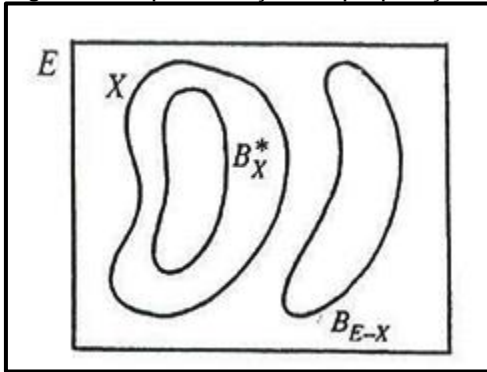
A demonstração disso seguirá facilmente a partir do seguinte resultado.

Lema 3.1.2 Sejam I e I^* subconjuntos disjuntos de $E(M)$ tais que I é independente e I^* é coindependente. Então M possui uma base B e uma cobase B^* tais que $I \subseteq B$, $I^* \subseteq B^*$, e B e B^* são disjuntas.

Demonstração: Em $M|(E-I^*)$, o conjunto I é independente. Portanto, contém uma base desta restrição. Assim $r(B) = r(E-I^*)$. Mas pela Proposição 3.1.1 (ii), $E-I^*$ é gerador em M . Consequentemente, $r(B) = r(M)$ e assim B é uma base de M . Como $I \subseteq B$ e $I^* \subseteq E-B$, segue-se que B e $E-B$ são as requeridas base e cobase de M . $\square$

Demonstração da Proposição 3.1.3. Seja B_X^* uma base para X em M^* e B_{E-X} uma base para $E-X$ em M (ver a Figura 8). Então $r^*(X) = |B_X^*|$ e $r(E-X) = |B_{E-X}|$. Sabemos que B_{E-X} e B_X^* são independentes, respectivamente. Assim, pelo Lema 3.1.2, conclui-se que M tem uma base B e uma cobase B^* de modo que B contenha B_{E-X} , B^* contenha B_X^* , e $B \cap B_X^* = \emptyset$. Como B_{E-X} é uma base para $M|(E-X)$, tem-se que $B \cap (E-X) = B_{E-X}$. Igualmente, $B^* \cap X = B_X^*$. Portanto, $B = B_{E-X} \cup (B \cap X)$, assim $|B \cap X| = |B| - |B_{E-X}|$. Segue que $|X| = |X \cap B| + |X \cap B^*| = |B| - |B_{E-X}| + |B_X^*| = r(M) - r(E-X) + r^*(X)$. $\square$

Figura 8. Representação da proposição 3.1.3



FONTE: Oxley, 1992.

A proposição seguinte expressa uma conexão fundamental entre circuitos e cocircuitos.

Proposição 3.1.4 Se C é um circuito e C^* é um cocircuito de um Matróide M , então $|C \cap C^*| \neq 1$.

Demonstração: Suponha que $|C \cap C^*| = 1$, digamos $C \cap C^* = \{e\}$; então, pelo Proposição 3.1.1, existe um hiperplano H tal que $H = E(M) - C^*$. Portanto $cl(H) = H$. Mas $C - \{e\} \subseteq H$ e $e \in C \subseteq H \cup e$. Assim, pela Proposição 2.4.2 segue-se que $e \in cl(H) = H$, uma absurdo, pois $e \in C^*$ e $H \cap C^* = \emptyset$. $\square$

Definição 3.1.3 Um *clutter* é uma coleção de conjuntos os quais não são subconjuntos próprios dos demais.

Existem diversos *clutters* que são associados naturalmente com um Matróide, como por exemplo, os conjuntos de bases, circuitos, hiperplanos, entre outros. Note que, caso $\mathcal{A}$ seja um *clutter* de subconjuntos de um dado conjunto S , então o conjunto $\mathcal{A}'$ dos complementares dos membros de $\mathcal{A}$ também é um clutter.

Definição 3.1.4 Supor que $\mathcal{A}$ seja um *clutter* de subconjuntos de um dado conjunto S . O *blocker* $b(\mathcal{A})$ de $\mathcal{A}$ consiste de subconjuntos minimais de S que possuem interseção não-vazia com cada membro de $\mathcal{A}$.

É fácil verificar que $b(\mathcal{A})$ é um *clutter*: de fato, se X, Y pertencem a $b(\mathcal{A})$ com $X \subsetneq Y$. Como X possui interseção não vazia com cada membro de $\mathcal{A}$ e está contido propriamente em Y , então Y não seria minimal. Assim, de fato $b(\mathcal{A})$ é um *clutter*. $\square$

Além disso, Edmonds e Fulkerson (1970) demonstraram que *clutters* possui a seguinte atraente propriedade.

Proposição 3.1.5 Se $\mathcal{A}$ é um *clutter* em S , então $b(b(\mathcal{A})) = \mathcal{A}$.

Seja M um Matróide e $\mathcal{B}(M)$, $\mathcal{C}(M)$, $\mathcal{H}(M)$, $\mathcal{B}^*(M)$, $\mathcal{C}^*(M)$ e $\mathcal{H}^*(M)$ *clutters* para bases, circuitos, hiperplanos, cobases, cocircuitos e cohiperplanos de M , respectivamente. Então as seguintes propriedades são válidas

$$3.1.2. \mathcal{B}(M)' = \mathcal{B}^*(M);$$

$$3.1.3. \mathcal{H}(M)' = \mathcal{C}^*(M) \text{ e}$$

3.1.4. $\mathcal{C}(M)' = \mathcal{H}^*(M)$. Além disso, o cocircuito de M é conjunto minimal, contendo interseção não vazia com muitas bases.

Proposição 3.1.6 $\mathcal{C}^*(M) = b(\mathcal{B}(M))$.

Demonstração: As seguintes afirmações são equivalentes.

- (i) $C^* \in \mathcal{C}^*(M)$
- (ii) $E(M) - C^* \in \mathcal{H}(M)$.
- (iii) $E(M) - C^*$ é um conjunto não gerador maximal de M .
- (iv) $E(M) - C^*$ não contém uma base de M , mas para todo x em C^* , $E(M) - (C^* - x)$ não contém uma base de M .
- (v) C^* atende a muitas bases de M , mas para todo x em C^* , $C^* - x$ evita alguma base de M .
- (vi) $C^* \in b(\mathcal{B}(M))$.

Concluiremos que $\mathcal{C}^*(M) = b(\mathcal{B}(M^*))$ e $b(\mathcal{C}^*(M)) = \mathcal{B}^*(M^*)$. Reescrevendo isso em termos de M , obtemos o seguinte.

Corolário 3.1.2 $\mathcal{C}(M) = b(\mathcal{B}^*(M))$ e $b(\mathcal{C}(M)) = \mathcal{B}^*(M)$.

O Corolário 3.1.2 é chamado o dual da Proposição 3.1.6. A primeira parte do corolário afirma que os circuitos de M são os conjuntos minimais que tem interseção não vazia com cada cobase. É fácil verificar que uma proposição é verdadeira se, e somente se, a proposição dual também é verdadeira. Se uma proposição coincide com sua dual então tal proposição é dita auto-dual. Um exemplo de proposição auto-dual é a Proposição 3.1.4

Como $\mathcal{H}(M)' = \mathcal{C}^*(M)$, podemos utilizar o Corolário 2.1.1 para derivar as seguintes características de Matróides em termos de seus hiperplanos.

Proposição 3.1.7 Seja $\mathcal{H}$ um conjunto de subconjuntos de um conjunto E . Então $\mathcal{H}$, é a coleção de hiperplanos de um Matróide E se, e somente se, satisfaz as seguintes condições.

- (H1) $E \notin \mathcal{H}$.

(H2) $\mathcal{H}$ é um *clutter*

(H3) Se H_1 e H_2 são membros distintos de $\mathcal{H}$ e $e \in E - (H_1 \cup H_2)$, então é um membro H_3 , de $\mathcal{H}$ tal que $H_3 \supseteq (H_1 \cap H_2) \cup e$.

A proposição dada a seguir estende a Proposição 3.1.4, caracterizando o conjunto de circuitos de um Matróide.

Lema 3.1.3 Seja $\mathcal{X}$ e $\mathcal{Y}$ uma coleção de subconjuntos de um conjunto finito E , tal que cada membro de $\mathcal{X}$ contém um membro de $\mathcal{Y}$ e cada membro de $\mathcal{Y}$ contém um membro de $\mathcal{X}$. Então os membros minimais de $\mathcal{X}$ são precisamente os membros minimais de $\mathcal{Y}$.

Proposição 3.1.8 Seja M um Matróide com conjunto base E . Então D é um circuito de M se, e somente se, D é um conjunto minimal não vazio de E tal que $|D \cap C^*| \neq 1$ para cada cocircuito C^* de M .

Demonstração: Usaremos o último lema tomando $\mathcal{X} = \mathcal{C}(M)$ e $\mathcal{Y}$ igual para o conjunto não vazio, subconjunto Y de E tal que $|Y \cap C^*| \neq 1$ para todo cocircuito C^* . Pela Proposição 3.1.4, se $x \in \mathcal{X}$, então $x \in \mathcal{Y}$. Agora suponha Y é um membro minimal de $\mathcal{Y}$. Se Y é independente então M contém uma base B contendo Y . Escolher y em Y e considerar $c/(B - y)$. Isto é um hiperplano de M . Seu complemento C_1^* é um cocircuito que cumpre Y em $\{y\}$ uma contradição. Consequentemente Y é dependente. Assim Y contém circuito C de M . Se $Y \neq C$, então, pela Proposição 3.1.4, obtemos uma contradição para a minimalidade de Y . Concluímos que Y é um circuito de M , em que $Y \in \mathcal{X}$. A proposição segue imediatamente do Lema 3.1.3. $\square$

Lembrando-se da Seção 2.3 que um Matróide pavimentação é um Matróide M que não tem circuitos de tamanho menor que $r(M)$. Assim, cada Matróide de *rank* inferior a dois é uma pavimentação.

Pra concluirmos esta parte, afirmamos uma caracterização em termos de seus hiperplanos, de Matróides de pavimentação de posição pelo menos dois (Hartmanis, 1959). Este resultado explica a razão para o nome Matróide de

pavimentação, que foi introduzido pelo Wesh (1976). Entendemos a definição de Wesh ligeiramente deixando cair sua exigência de que tais Matróides têm *rank* pelo menos dois. Seja K e m sendo inteiros com $K > 1$ e $m > 0$. Suponha que $\mathcal{T}$ é um conjunto $\{T_1, T_2, \dots, T_k\}$ de subconjuntos de um conjunto E tal que cada membro de $\mathcal{T}$ tem pelo menos m elementos e cada m – elementos subconjuntos de E está contido num único membro de $\mathcal{T}$ é chamado um m partição de E .

Com baseamento da Teoria de Matróides apresentado no capítulo anterior, em que foi demonstrado de uma forma mais clara, faz necessário um introdutório da Teoria da Computação Quântica que será apresentado a seguir, para fazer a conexões entre a Teoria de Matróides vetoriais e a Teoria de Codificação. Desta forma, poderá encontrar a relação da matriz de paridade de um código CSS e a matriz que gera um dado matróide vetorial.

4. CONSTRUÇÃO DE CÓDIGOS QUÂNTICOS

Na Teoria da Informação e Computação Quântica, uma das dificuldades encontradas é de proteger os estados quânticos originais, até seu destino, sem que ocorram alterações durante seu percurso.

Para proteger a informação quântica, surgiram os códigos corretores de erros quânticos, que tem como papel principal a confiabilidade de armazenamento e transmissão de informação quântica. Devido a esses fatos, várias construções de códigos quânticos estabilizadores (KETKAR, IEEE-IT, 2005; ASHIKHMIN, IEEE-IT, 2014; NIELSEN e CHUANG, 2005), incluindo o código de repetição (devido a Peter W. Shor), código de Hamming, códigos de Steane e a classe de códigos do tipo Calderbank-Shor-Steane (CSS).

O processamento confiável de informação quântica requer mecanismos para reduzir os efeitos do ruído operacional e ambiental. Assim, é possível minimizar os efeitos prejudiciais o erro, ao empregar códigos de correção de erros quânticos, de modo que se possa ter a comunicação quântica e computadores quânticos mais confiáveis.

Entretanto não se pode ter tanta certeza na fidelidade desta proteção, pois mesmo utilizando bons códigos quânticos podem ocorrer erros.

Os códigos CSS são atraentes para inúmeras aplicações, como na correção de erros na memória quântica, cálculos quantitativos tolerantes a falhas, criptografia quântica entre outros (ASHIKHMIN, IEEE-IT, 2014). Deste modo, o uso adequado do código CSS, é muito importante nestas aplicações.

Neste capítulo, definiremos códigos de bloco lineares bem como definiremos o que é um código quântico Calderbank-Shor-Steane (CSS).

4.1 CÓDIGOS LINEARES CLÁSSICOS

A classe dos códigos lineares clássicos é uma ferramenta fundamental para se gerar bons códigos quânticos mediante a construção CSS, que será abordada mais a frente. Na verdade, a construção CSS utiliza dois códigos clássicos lineares C_1 com parâmetros $[n, k_1, d_1]_q$, e C_2 com parâmetros $[n, k_2, d_2]_q$, de tal forma que C_1 será responsável por corrigir erros quânticos do tipo *qudit flip* (mudança de bit

quântico), e o dual euclidiano do código C_2 será responsável por corrigir erros quânticos do tipo *phase-shift* (mudança de fase). A construção precisa dos códigos quânticos CSS será apresentada no final deste capítulo.

Códigos clássicos lineares são, na verdade, espaços vetoriais definidos sobre corpos finitos. Assim, utilizamos a teoria de códigos lineares clássicos para a correção quântica de erros.

Primeiramente apresentamos uma breve revisão da teoria dos códigos lineares.

Definição 4.1.1 Um código linear C com parâmetros $[n, k, d]_q$, é um k -subespaço vetorial do espaço vetorial F_q^n , em que F_q é o corpo finito com q elementos, d é a distância mínima e F_q^n é o conjunto das n -uplas de elementos definidos no corpo F_q . O parâmetro n é denominado o *comprimento da palavra-código*.

Definição 4.1.2 Seja C um código linear com parâmetros $[n, k, d]_q$. Define-se o peso de Hamming de uma palavra-código $\mathbf{c} = (c_1, \dots, c_n)$ como sendo o número de coordenadas não nulas de $\mathbf{c}$. A distância de Hamming entre dois vetores $\mathbf{x}, \mathbf{y}$ em F_q^n é definido como sendo o número de coordenadas distintas entre o vetor diferença $\mathbf{v} = \mathbf{x} - \mathbf{y}$.

A distância mínima do código está diretamente relacionada com o poder de correção de erros do código, ou seja, quanto maior a distância mínima d , mais erros o código é capaz de corrigir. Mais precisamente, se um código linear C possui distância mínima de pelo menos $2t + 1$, para algum inteiro t , C é capaz de corrigir erros em até t bits, decodificado a mensagem y' como a única palavra y satisfeito $d(y, y') \leq t$.

Definição 4.1.3 Seja C um código linear com parâmetros $[n, k, d]_q$, sobre F_q . Se d é a distância mínima de C denotaremos os parâmetro do código por $[n, k, d]_q$. A taxa R do código é definida como sendo $R = k / n$.

Como os códigos são subespaços vetoriais sobre F_q , esses podem ser definidos mediante uma matriz, denominada *matriz geradora* $G_{k \times n}$, cujas linhas são vetores que formam uma base para C . Assim, a mensagem de comprimento k (que é um vetor de comprimento k , cujas coordenadas pertencem a F_q) é multiplicada pela matriz geradora $G_{k \times n}$, gerando uma palavra-código de comprimento n .

Definição 4.1.4 Seja C um código linear dado por $[n, k, d]_q$. Uma matriz geradora G de C é uma matriz $k \times n$, cujas linhas formam uma base para o código C , ou seja,

$$G_{k \times n} = \begin{pmatrix} g_0 \\ g_1 \\ \vdots \\ g_{k-1} \end{pmatrix} = \begin{pmatrix} g_{00} & g_{01} & g_{02} & \dots & g_{0,n-1} \\ g_{10} & g_{11} & g_{12} & \dots & g_{1,n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ g_{k-1,0} & g_{k-1,1} & g_{k-1,2} & \dots & g_{k-1,n-1} \end{pmatrix}$$

em que os $g_0, g_1, \dots, g_{k-1}$ são os vetores geradores.

A matriz geradora corresponde à conexão entre a mensagem que se deseja codificar e a palavra-código correspondente (mensagem codificada).

Na codificação, considerando $u = (u_0, u_1, \dots, u_{k-1})$ uma mensagem a ser codificada, a palavra-código resultante v pode ser expressa da seguinte forma.

$$v = u \cdot G = u_0 g_0 + u_1 g_1 + \dots + u_{k-1} g_{k-1}$$

Considere o código linear (7,4) com a seguinte matriz geradora.

$$G = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 & 0 & 1 \end{pmatrix}$$

A informação a ser codificada como sendo, $u = (1 \ 1 \ 0 \ 1)$, então v seria $v = 1 g_0 + 1 g_1 + 0 g_2 + 1 g_3 = (1000110) + (0100011) + (0001101) = (1101000)$.

Exemplo 4.1.1 Como um exemplo, apresentamos o código de repetição de comprimento 3, sobre F_2 , que mapeia um único bit em três bits repetidos; é especificado pela matriz geradora $G_{1 \times 3} = (1 \ 1 \ 1)_{1 \times 3}$. Assim, G mapeia as mensagens possíveis 0 e 1 para a forma codificada, $(0) \cdot G = (0,0,0)$ e $(1)G = (1,1,1)$. Assim, a taxa deste código é igual a $R = 1/3$. Note também que a distância do código é igual a três.

O código de repetição é forte em proteção e são utilizados nas senhas de cartões, ou seja, movimentações financeiras (BARROS, 2011)

Exemplo 4.1.2 Outro exemplo de código linear de dois bits de repetições de três bits para cada bit, ou seja, o código [6, 2]. A matriz geradora é:

$$G_{2 \times 6} = \begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 \\ 1 & 1 & 1 & 0 & 0 & 0 \end{pmatrix} \quad (4.1.1)$$

Portanto

$$(0,0)G = (0,0,0,0,0,0); G(0,1) = (0,0,0,1,1,1) \quad (4.1.2)$$

$$(1,0)G = (1,1,1,0,0,0); G(1,1) = (1,1,1,1,1,1) \quad (4.1.3)$$

Portanto, o conjunto de todas as palavras codificadas do código correspondente ao espaço vetorial é gerado pelas linhas de G .

A vantagem de se trabalhar com um código linear é que necessário apenas especificar os k vetores linearmente independentes (LI) da matriz geradora $G_{k \times n}$, o que consiste em uma economia de memória. A descrição compacta na capacidade de codificação e de decodificação eficiente são características importantes que os códigos lineares compartilham com seus “primos quânticos”, os códigos estabilizadores. Assim, para a codificação de um código linear clássico, multiplicamos a mensagem de k bits pela matriz geradora para se obter a uma mensagem codificada (palavra-código) com n bits.

Ao se trabalhar com a decodificação das palavras-código de um dado código linear, a matriz geradora não apresenta grandes vantagens. Assim, dado um código linear $C [n, k, d]_q$ definimos outra matriz associada a C , denominada *matriz verificadora de paridade*.

Definição 4.1.5 Seja $C [n, k, d]_q$ um código linear. Uma matriz verificadora de paridade para um código linear para C é uma matriz H de dimensão $(n - k) \times n$, cujas linhas são linearmente independentes e tais que $H x^T = 0 \forall x \in C$, em que

$$H = \begin{pmatrix} h_0 \\ h_1 \\ \vdots \\ h_{n-k-1} \end{pmatrix} = \begin{pmatrix} h_{00} & h_{01} & h_{02} & \dots & h_{0,n-1} \\ h_{10} & h_{11} & h_{12} & \dots & h_{1,n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ h_{n-k-1,0} & h_{n-k-1,1} & h_{n-k-1,2} & \dots & h_{n-k-1,n-1} \end{pmatrix}$$

A relação entre as matrizes geradora e verificadora de paridade de um dado código linear $C [n, k, d]_q$ é dada por $G.H^T = 0$, em que o subespaço gerado por H é dual do subespaço gerado por G , assim os vetores de G são ortogonais aos vetores de H .

Exemplo 4.1.3 Tomando como exemplo a matriz geradora do Exemplo 4.1.1 do código de repetição com parâmetros $[3, 1, 3]_2$. Para construir H , selecionamos $3 - 1 = 2$ vetores linearmente independentes ortogonais às linhas de G , dadas por $(1,1,0)$ e $(0,1,1)$. A matriz verificadora de paridade é

$$H \equiv \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix} \quad (4.1.4)$$

Logo, $H x^T = 0$ para as palavras código $x = (0,0,0)$ e $x = (1,1,1)$.

Nos códigos clássicos é possível encontrar uma matriz geradora na forma padrão. Em outras palavras, dada a matriz geradora G (ou a matriz verificação de paridade H) é possível encontrar uma matriz geradora G^* na forma padrão (H^* na forma padrão)

$$G^* = \begin{pmatrix} I_K \\ -A \end{pmatrix} \quad (4.1.5)$$

Em que I_k é a matriz identidade de ordem k e $A_{k \times (n-k)}$ é a matriz paridade de ordem $k \times (n - k)$.

$$H = [I_{(n-k) \times (n-k)} | A^T] = \left(\begin{array}{cccc|cccc} 1 & 0 & \dots & 0 & a_{00} & a_{10} & \dots & a_{k-1,0} \\ 0 & 1 & \dots & 0 & a_{01} & a_{11} & \dots & a_{k-1,1} \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 & a_{0,n-k-1} & a_{1,n-k-1} & \dots & a_{k-1,n-k-1} \end{array} \right)$$

$$G = [A_{k \times (n-k)} | I_{k \times k}] = \left(\begin{array}{cccc|cccc} a_{00} & a_{01} & \dots & a_{0,n-k-1} & 1 & 0 & \dots & 0 \\ a_{10} & a_{11} & \dots & a_{1,n-k-1} & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{k-1,0} & a_{k-1,1} & \dots & a_{k-1,n-k-1} & 0 & 0 & \dots & 1 \end{array} \right) \quad (4.1.6)$$

Suponha agora que se codifique a mensagem x como $y = xG$, mas um erro e devido ao ruído corrompa y , ao qual resulta a palavra modificada $y' = y + e$. Como

$Hy = 0$, para todas as palavras do código, temos, $Hy' = He$, em que Hy' é denotada a síndrome de erro. Uma função do estado corrompido y' , do mesmo modo que a síndrome de erro quântica, determinada por medidas do estado quântico corrompido. Como $Hy' = He$, a síndrome de erro contém informação sobre o erro cometido, o que deveria permitir a recuperação da palavra original y . Analisando essa possibilidade, suponha que nenhum, ou apenas um erro tenha ocorrido. Logo a síndrome de erro Hy' será igual a 0 no caso de nenhum erro, e igual a He_j no caso ter ocorrido erro no j -ésimo bit, consistindo e_j o vetor unitário com componentes 1 na j -ésima componente. Agora, se supusermos que ocorra apenas um erro, para corrigir este erro, calcula-se a síndrome de erro He_j para determinar o bit que precisa ser corrigido.

Abaixo, exibiremos uma importante desigualdade com relação aos parâmetros de um código linear $C [n, k, d]_q$, denominado limitante de Singleton.

Teorema 4.1.1 (Limite de Singleton) Se C é um código linear $[n, k, d]_q$, então $d \leq n - k + 1$. Se a igualdade é válida diz-se que o código é *maximum distance separable* (do inglês, *maximum distance separable* MDS), ou seja, possui a distância máxima de separação.

Observação 4.1.1 Códigos MDS são ótimos, ou seja, fixados os mesmos parâmetros n e k , não existem códigos capazes de corrigir mais erros que o correspondente MDS.

Demonstração do Teorema 4.1.1 Seja H uma matriz de paridade de um código linear $C [n, k, d]_q$. Como H é uma matriz de ordem $(n - k) \times n$, e que $(n - k)$ são linhas linearmente independentes. Então, cada coluna de H tem $n - k$ entradas, ou seja, de comprimento $n - k$. Agora, para quaisquer $d - 1$ coluna de H é linearmente independente. Então o conjunto de vetores $F^{(n - k)}$ tem no máximo $n - k$ vetores, portanto $d - 1 \leq n - k$. Conclui-se que $d \leq n - k + 1$. $\square$

Um exemplo da teoria acima. Seja o código linear $[6, 3, 3]$, então $n - k \geq d - 1 \Rightarrow 6 - 3 \geq 3 - 1 \Rightarrow 3 \geq 2$. Este teorema fornece um limitante superior para o código.

Exemplo 4.1.4 Um bom exemplo de códigos lineares de correção de erro são os códigos de Hamming, pois conseguem corrigir 1 erro por bit ou detectar 2 erros em bits diferentes. Seja $r \geq 2$ um inteiro, e H uma matriz cujas colunas são todas as $2^r - 1$ sequências de bits de comprimento r que não sejam identicamente iguais a zero. Essa matriz verificadora de paridade define um código linear $[2^r - 1, 2^r - r - 1, 3]_q$ conhecido como código de Hamming.

Um caso particular, para correção quântica de erro é de $r = 3$, o qual define um código $[7, 4]$ ou $C(7, 4, 3)$ em que as matrizes geradoras de verificação de paridade estão expressas abaixo.

$$G = \begin{pmatrix} 0 & 1 & 0 & 1 & 0 & 1 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 1 \\ 1 & 0 & 0 & 1 & 1 & 0 & 0 \end{pmatrix}$$

$$H = \begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{pmatrix} \quad (4.1.7)$$

Portanto, quaisquer duas colunas de H são diferentes e linearmente independentes. As três primeiras colunas são linearmente dependentes e a distância do código é igual a 3. Logo, este código é capaz de corrigir qualquer erro em qualquer q-bit isolado. Por (4.1.7) a síndrome He_j é uma representação binária para j , que nos diz qual bit deve ser corrigido.

4.2 CÓDIGOS QUÂNTICOS

Primeiramente, iremos introduzir os postulados da Mecânica Quântica. Para mais informações, referimos ao leitor a NIELSEN e CHUANG (2005). São estes:

Postulado 1: A qualquer sistema físico isolado existe, associado um espaço vetorial complexo com produto interno (ou seja, um espaço de Hilbert), conhecido como espaço de estados do sistema. O sistema é completamente descrito pelo seu vetor de estado, *um vetor unitário no espaço de estados*.

Definição 4.2.1 O estado quântico utilizado é o qubit (do inglês, *quantum bit*), ou seja, um vetor de estado arbitrário dado por $|\psi\rangle = a|0\rangle + b|1\rangle$, em que a e b são

números complexos e $|0\rangle$ e $|1\rangle$ (correspondentes aos estados clássicos 0 e 1), formam uma base ortogonal desse espaço. A condição de que $|\psi\rangle$ deve ser um vetor unitário $\langle\psi|\psi\rangle = 1$, é equivalente $|\alpha|^2 + |\beta|^2 = 1$, denotada como condição de normalização para o vetor de estado.

Postulado 2: A evolução do sistema quântico fechado é descrito por uma transformação unitária. Em outras palavras, o estado $|\psi\rangle$ de um sistema em um tempo t_1 está relacionado ao estado $|\psi'\rangle$ do sistema t_2 por um operador unitário U que depende somente de t_1 e t_2 : $|\psi'\rangle = U|\psi\rangle$.

Para o caso de um qubit qualquer operador unitário pode ser implementado em sistemas físicos reais. Na Computação e Informação Quântica são utilizados operadores que atuam em qubits separadamente, como por exemplo, as matrizes de Pauli e algumas portas quânticas.

O Postulado 2, somente é válido para sistemas isolados, i.e., sistemas que não interagem com outros sistemas. Ainda, descreve como os estados em dois instantes de tempo de um sistema quântico fechado estão relacionados por meio de transformações unitárias.

Postulado 3: As medidas quânticas são descritas por determinados operadores de medidas $\{M_m\}$. Esses operadores atuam sobre o espaço de estados do sistema. O índice m se refere aos possíveis resultados da medida. Se o estado de um sistema quântico for $|\psi\rangle$, imediatamente antes da medida, a probabilidade de um resultado m ocorrer é dada por: $p(m) = \langle\psi|M_m^\dagger M_m|\psi\rangle$, e o estado do sistema após a medida é:

$$\frac{M_m |\psi\rangle}{\sqrt{\langle\psi|M_m^\dagger M_m|\psi\rangle}}$$

Os operadores de medidas satisfazem a relação de completitude:

$$\sum_m M_m^\dagger M_m = I$$

A relação de completitude expressa o fato de que a soma das probabilidades deve ser igual a 1.

$$1 = \sum_m p(m) = \sum_m \langle \psi | M_m^\dagger M_m | \psi \rangle$$

Portanto, o Postulado 3 pode ser aplicado ao problema que se refere da distinção de estados quânticos. No mundo clássico, no caso de probabilidade de uma moeda sair cara ou coroa é fácil de ser resolvido. Já no mundo quântico a situação é mais complicada, pois estados quânticos que não são ortogonais não podem ser distinguidos, este Postulado 3 pode resolver de uma forma mais convincente.

Postulado 4: O espaço de estados de um sistema físico composto é o produto tensorial dos espaços de estados dos sistemas físicos individuais. Se os sistemas foram enumerados de 1 até n , e o sistema i for preparado no sistema $|\psi_i\rangle$ então o estado do sistema composto será $|\psi_1\rangle \otimes |\psi_2\rangle \otimes \dots \otimes |\psi_n\rangle$.

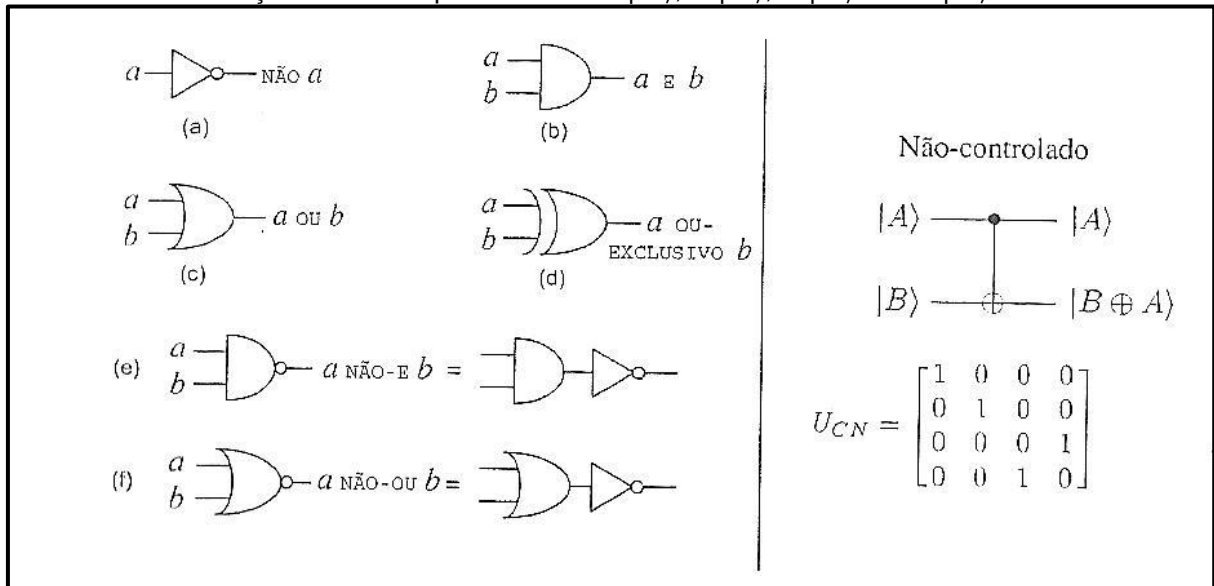
Este postulado, nos mostra que espaços de estados de sistemas quânticos devem combinar-se para formar um sistema quântico composto.

Definição 4.2.2 Uma porta quântica de um qubit ou múltiplos qubits são operadores unitários que podem atuar sobre um ou vários qubits. Tais operadores podem ser expressos de maneira matricial. Dessa forma, os qubits são representados por um vetor (ket) e a porta quântica por meio de um operador em forma de matriz de Pauli.

$$X \equiv \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, Y \equiv \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, Z \equiv \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

Uma porta lógica de muitos qubits é o NÃO- *controlado* ou porta CNOT (do inglês “*Controlled-NOT*”). Esta porta tem dois qubits de entrada, conhecidos como o qubit de controle e o qubit de alvo. Na figura 9 é mostrado o circuito que representa o CNOT.

Figura 9. Algumas portas clássicas de um ou mais bits (esquerda), e à direita o protótipo da porta quântica de muitos qubits, o NÃO-controlado. A representação matricial do NÃO-controlado, U_{CN} , é escrita com relação às amplitudes de $|00\rangle$, $|01\rangle$, $|10\rangle$ e $|11\rangle$ nessa ordem.



FONTE: Nielsen e Chuang, 2005.

A linha superior representa o qubit de controle e a linha inferior o qubit alvo. Deste modo, a ação da porta pode ser descrita da seguinte maneira: se o qubit de controle for colocado zero, nada acontece com o qubit alvo. Agora, para o caso de o qubit de controle for um, o qubit alvo troca seu estado. Como percebemos a seguir em símbolos:

$$C_{Not} |0\rangle \otimes |0\rangle = |00\rangle; C_{Not} |01\rangle = |01\rangle; C_{Not} |10\rangle = |11\rangle; C_{Not} |11\rangle = |10\rangle.$$

Como exemplo, as portas CNOT que tem a função de trocar os coeficientes α e β do qubit, para os estados quânticos ele leva o estado $|0\rangle$ para $|1\rangle$ e o estado $|1\rangle$ para $|0\rangle$, ou seja, inverte o qubit. A porta Hadamard é utilizada em qubits em uma superposição do estado de base $|0\rangle$ e $|1\rangle$ em que se destaca no uso do emaranhamento, pois permite criar os estados de superposição a partir dos estados quânticos.

Ao código quântico para proteção de inversão de bit e realizando essa comparação com código de repetição clássico, identificamos que as palavras do código de repetição são os rótulos para o código quântico. O rótulo de um estado quântico para representação binária, contida na simbologia de Dirac é $| \rangle_{bra}$. Para o mapeamento quântico, realizado pelo código de inversão de bit, dos rótulos de base padrão para o estado quântico de informação $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, seja mapeada

para estados quânticos de bases novas, em que sua representação é idêntica às palavras-código do código de repetição. Este procedimento está esquematizado pelas equações abaixo.

$$|u = 0\rangle \rightarrow |v = 0,0,0\rangle.$$

$$|u = 1\rangle \rightarrow |v = 1,1,1\rangle. \quad (4.2.1)$$

Sendo o erro de inversão de bit no caso quântico, um simples erro para o caso clássico no que refere aos rótulos dos estados quânticos, temos um modelo abaixo;

$$e = |100\rangle + v |000\rangle \rightarrow r |100\rangle \Leftrightarrow X_1 |101\rangle = |100\rangle.$$

$$e = |010\rangle + v |111\rangle \rightarrow r |101\rangle \Leftrightarrow X_2 |000\rangle = |101\rangle. \quad (4.2.2)$$

Desta forma, os operadores lineares de inversão de bit com representação binária na forma de vetores linhas, na matriz identidade colocando o vetor linha o bit zero e nas posições em que se encontra a matriz de inversão de bit, colocando se o bit um. Portanto, pode ser definido, pois a ação de inversão de bit no caso clássico é realizada quando à palavra- código é somado a um vetor linha. Exemplo.

$$e = [010] \Leftrightarrow E = IXI$$

$$e = [100] \Leftrightarrow E = XII \quad (4.2.3)$$

Sendo que qualquer código clássico passa a ser transportado, para apresentar uma aplicação quanto aos códigos quânticos, no que se refere a erros de inversão de bit. Assim, a capacidade de correção do código clássico, é transportada também para o código quântico, para erros de inversão de bit.

Concluimos que a correção clássica de erros, sobre a construção de códigos, é conhecido como construção dual.

Seja C um código $[n, k]$ com matriz geradora G e matriz verificadora de paridade H . É possível, definir outro código o dual de C , denotado por $C^\perp$, como o código cuja matriz geradora é H^T e uma matriz verificadora de paridade G^T . Portanto, o dual de C consiste em todas as palavras y tais que y é ortogonal a todas as palavras de C . Um código é fracamente auto- dual se $C \subseteq C^\perp$ e auto- dual se $C = C^\perp$.

$= C^\perp$. O fato de que a construção dual para códigos clássicos lineares aparece nos estudos das correções quântica de erro, em que esta classe é muito importante para a construção de códigos conhecido como CSS.

4.3 CÓDIGOS DE CALDERBANK–SHOR-STEANE (CSS)

Uma classe especial de códigos quânticos corretores de erros é a classe dos códigos Calderbank-Shor-Steane (CSS) em homenagem aos seus inventores, uma das mais importantes subclasses de uma classe mais geral de códigos estabilizadores (NIELSEN e CHUANG, 2005). Para os códigos CSS, daremos uma leve introdução da simbologia e a teoria de códigos clássicos que utilizaremos para a construção de códigos CSS.

Calderbank, Steane e Shor são conhecidos pelos códigos CSS, em que destacam a ortogonalidade entre espaços vetoriais que contém as palavras dos códigos.

Sejam C_1 e C_2 códigos clássicos lineares com parâmetros $[n, k_1, d_1]_2$ e $C_2[n, k_2, d_2]_2$, respectivamente, tais que $C_2 \subseteq C_1$ e C_1 e $C_2^\perp$ corrigem ambos t erros. Define-se um código quântico CSS (C_1, C_2) com parâmetros $[n, k_1 - k_2]$, capaz de corrigir erros em t qubits, por meio da seguinte construção.

Seja $x \in C_1$ uma codificação no código C_1 . Define-se o estado quântico $|x + C_2\rangle$ por $|x + C_2\rangle = \frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} |x + y\rangle$ (4.3.1)

Seja x' um elemento de C_1 tal que $x - x' \in C_2$. Portanto, é fácil ver que $|x + C_2\rangle = |x' + C_2\rangle$, donde o estado $|x + C_2\rangle$ depende somente do espaço- quociente de C_1 / C_2 que encontra x , explicando a notação de espaço-quociente para $|x + C_2\rangle$. Se x e x' pertencem a diferentes espaços-quociente de C_2 , resulta que, para nenhum $y, y' \in C_2$ ocorre $x + y = x' + y$, portanto $|x + C_2\rangle$ e $|x' + C_2\rangle$ são espaços ortogonais.

O código quântico CSS (C_1, C_2) é definido como sendo espaço vetorial gerado pelos estados $|x + C_2\rangle$ para $x \in C_1$. Assim, o número de espaços-quocientes de C_2 em C_1 é $|C_1| / |C_2|$. A dimensão de CSS (C_1, C_2) é $|C_1| / |C_2| = 2^{K_1 - K_2}$, assim CSS (C_1, C_2) é um código quântico $[[n, k_1 - k_2]]$. Logo, é possível

corrigir erros em até t qubits e erros de inversão de fase em CSS (C_1, C_2) usando as propriedades C_1 e $C_2^\perp$.

Supor que os erros de inversão de bit sejam descritos por um vetor de n bits e_1 com componentes iguais a 1 em que ocorrem as inversões, e 0 onde elas não ocorreram e que erros de inversão de fase sejam descritos por um vetor de n bits e_2 com componentes iguais a 1 onde ocorreram as inversões, e 0 onde tais não ocorreram. Se o estado original era $|x + C_2\rangle$, o estado corrompido é.

$$\frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} (-1)^{(x+y) \cdot e_2} |x + y + e_1\rangle \quad (4.3.2)$$

Para a introdução desse sistema quântico faz se uso da computação reversível, para aplicar a matriz de paridade H_1 para o código C_1 , e que seja iniciado pelo estado padrão da base computacional $|0\rangle$ para armazenar a síndrome para o código C_1 . Esse procedimento leva o estado $|x + y + e_1\rangle |0\rangle$ para $|x + y + e_1\rangle |H_1(x + y + e_1)\rangle = |x + y + e_1\rangle |H_1 e_1\rangle$ pois $(x + y) \in C_1$ é extinto pela matriz de paridade. Portanto, o efeito dessa operação é produzir o estado:

$$\frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} (-1)^{(x+y) \cdot e_2} |x + y + e_1\rangle |H_1 e_1\rangle \quad (4.3.3)$$

O erro de inversão de bit pode ser detectado realizando uma medida sobre o sistema quântico introduzido, já que neste sistema está contido o padrão de síndrome do erro que afetou o estado codificado. Essa medição não afeta o estado principal o qual permanece intacto como mostra a operação abaixo

$$\frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} (-1)^{(x+y) \cdot e_2} |x + y + e_1\rangle \quad (4.3.4)$$

Sabendo que a síndrome do erro $H_1 e_1$, podemos induzir o erro e_1 , pois C_1 pode corrigir até t erros completando a etapa de detecção. A recuperação é obtida aplicando-se as portas NÃO ao qubits cujas posições em e_1 detectaram uma inversão de bit. Consequentemente, todos os erros serão corrigidos e o estado fica.

$$\frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} (-1)^{(x+y) \cdot e_2} |x + y\rangle \quad (4.3.5)$$

Neste instante já foram corrigidos os erros de inversão de bit. Para a correção de erros de inversão de fase é requerido uma mudança de base computacional para a base de Hadamard. Aplica a cada qubit do estado codificado corrompido uma porta de Hadamard para que essa mudança de base seja efetuada levando o estado para

$$\frac{1}{\sqrt{|C_2|^{2n}}} \sum_z \sum_{y \in C_2} (-1)^{(x+y) \cdot (e_2+z)} |z\rangle \quad (4.3.6)$$

Portanto, o somatório introduzido é realizado sobre todos os valores possíveis de z de n bits. Faz necessária uma mudança de variável com $z' \equiv z + e_2$. Assim, o estado quântico pode ser reescrito como.

$$\frac{1}{\sqrt{|C_2|^{2n}}} \sum_{z'} \sum_{y \in C_2} (-1)^{(x+y) \cdot z'} |z' + e_2\rangle \quad (4.3.7)$$

Considerando z' pertencente ao espaço vetorial ortogonal à C_2 , supondo que $z' \in C_2^\perp$ é possível que $\sum_{y \in C_2} (-1)^{y \cdot z'} = |C_2|$, caso contrário se essa suposição não pudesse ser realizada ao passo que $z' \notin C_2^\perp$ então $\sum_{y \in C_2} (-1)^{y \cdot z'} = 0$. O estado quântico pode ser reescrito como

$$\frac{1}{\sqrt{2^n/|C_2|}} \sum_{z' \in C_2^\perp} (-1)^{x \cdot z'} |z' + e_2\rangle \quad (4.3.8)$$

Realizando a mudança de base, para uma base mais apropriada é que os erros de inversão de fase transformam erros de inversão de bit para essas bases. Introduzindo um sistema auxiliar e aplicando na matriz verificadora de paridade H_2 para $C_2^\perp$ para obter $H_2 e_2$ e programando a correção de inversão de bit para e_2 . Resulta no estado

$$\frac{1}{\sqrt{2^n/|C_2|}} \sum_{z' \in C_2^\perp} (-1)^{x \cdot z'} |z'\rangle \quad (4.3.9)$$

Para a última etapa da correção, aplicam-se portas Hadamard a cada um dos qubits para que se retorne ao estado quântico original. O resultado pode ser computado direto ou o efeito de aplicar portas Hadamard ao estado (4.3.8) com $e_2 = 0$. A operação leva de volta ao estado (4.3.5) com $e_2 = 0$, pois a porta de Hadamard é auto-inversa

$$\frac{1}{\sqrt{|C_2|}} \sum_{y \in C_2} |x + y\rangle, \quad (4.3.10)$$

que é o estado codificado original.

Assim, os códigos quânticos CSS são códigos quânticos construídos a partir de códigos clássicos lineares C_1 e C_2 com parâmetros $[n, k_1, d_1]$ e $[n, k_2, d_2]$, respectivamente, tal que $C_2 \subset C_1$ e tanto C_1 quanto $C_2^\perp$ devem corrigir t erros. Além disso, note que as etapas de detecção e correção dos erros requerem somente a aplicação de portas de Hadamard e CNOT, para cada caso em número linear no tamanho do código.

Agora, mostraremos alguns exemplos de códigos CSS. Um exemplo é o código de Steane. Esse código é construído tomando como base o código (clássico) de Hamming $C[7, 4, 3]$, cuja matriz verificadora de paridade é:

$$H = \begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{pmatrix} \quad (4.3.11)$$

Para definir um código CSS é preciso verificar se $C_2 \subset C_1$. Rotule esse código por C e defina $C_1 \equiv C$ e $C_2 \equiv C^\perp$. Por definição, a matriz verificadora de paridade de $C_2 \equiv C^\perp$ é igual à transposta da matriz geradora de códigos de $C_1 \equiv C$.

$$H[C_2] = G[C_1]^T = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 \end{pmatrix} \quad (4.3.12)$$

Analisando estas matrizes, e comparando com (4.3.11), observamos que o espaço gerado pelas linhas $H[C_2]$ contém o espaço gerado pelas linhas de $H[C_1]$, na medida em que os códigos são os kernels de $H[C_2]$ e $H[C_1]$, ou seja, conclui-se que $C_2 \subset C_1$. A segunda propriedade está relacionada à capacidade dos códigos, em que C_1 e $C^\perp$ corrija ambos t erros. O que ocorre é que $C_2^\perp = (C_1^\perp)^\perp = C$, onde se vê que C_1 e $C_2^\perp$ são códigos de distância 3 onde a propriedade é verificada. Como C_1 é um código $[7, 4]$ e C_2 um código $[7, 3]$ segue que CSS (C_1, C_2) é um código capaz de corrigir erros em único q-bit. Portanto, o código quântico $[[7, 1, 3]]$ é conhecido como código de Steane.

Note que, no código de Steane, o “zero lógico” $|0 + C_2\rangle$ é dado por:

$$|0_L\rangle = \frac{1}{\sqrt{8}} [|0000000\rangle + |1010101\rangle + |0110011\rangle + |1100110\rangle + |0001111\rangle + |1011010\rangle + |0111100\rangle + |1101001\rangle] \quad (4.3.13)$$

E o um lógico $|1_L\rangle$ é dado por

$$|1_L\rangle = \frac{1}{\sqrt{8}} [|1111111\rangle + |0101010\rangle + |1001100\rangle + |0011001\rangle + |1110000\rangle + |0100101\rangle + |1000011\rangle + |0010110\rangle] \quad (4.3.14)$$

Para os códigos CSS e código sete qubit. Os códigos CSS são um excelente exemplo de uma classe de códigos estabilizador, demonstrando o quão fácil o formalismo estabilizador faz com que ele compreenda a construção do código quântico. Suponha que C_1 e C_2 são $[n, k_1]$ e $[n, k_2]$ códigos lineares clássicos, tais que $C_2 \subset C_1$ e C_1 e $C_2^\perp$ corrigem t erros. Define uma matriz de verificação com a forma

$$\left[\begin{array}{c|c} H(C_2^\perp) & 0 \\ \hline 0 & H(C_1) \end{array} \right] \quad (4.3.15)$$

Mediante essa equação, poderão ser analisadas, como projeto futuro, as conexões entre a Teoria de Matróides e a Teoria de Codificação. A relação entre a matriz de paridade de um código CSS e a matriz que gera um dado matróide vetorial.

Para ver que isto define um código de estabilizador, precisamos da matriz de verificação para satisfazer a condição de comutatividade $H(C_2^\perp) H(C_1)^T = 0$. Mas nós temos $H(C_2^\perp) H(C_1)^T = [H(C_1) G(C_2)]^T = 0$

A construção dos operadores lógicos Z e X para um código de estabilizador é feito de maneira fácil de entender se colocarmos o código em forma padrão. Para entender o que é a fórmula padrão, considere a matriz de verificação para um código de estabilizador $C[n, k]$:

$$G = [G_1 | G_2] \quad (4.3.16)$$

Esta matriz possui $n - k$ linhas. As linhas de troca desta matriz correspondem a rotular de novo geradores, trocando colunas correspondentes em ambos os lados da matriz para rotular de novo qubits, e adicionar as duas linhas corresponde a geradores de multiplicação. É fácil ver que sempre podemos substituir um gerador g_i por gg_j quando $i \neq j$. Assim, existe um código equivalente com um conjunto diferente

de geradores cuja matriz de verificação corresponde à matriz G onde a eliminação gaussiana foi feita em G_1 , trocando qubits conforme necessário:

$$n - k - r \left\{ \begin{array}{cc|cc} \begin{matrix} r \\ \tilde{I} \\ 0 \end{matrix} & \begin{matrix} n-r \\ \tilde{A} \\ 0 \end{matrix} & \begin{matrix} r \\ \tilde{B} \\ D \end{matrix} & \begin{matrix} n-r \\ \tilde{C} \\ E \end{matrix} \end{array} \right. \quad (4.3.17)$$

Onde r é o *rank* d G_1 . Em seguida, trocando qubits conforme necessário, realizamos uma eliminação Gaussiana em E para obter

$$n - k - r - s \left\{ \begin{array}{ccc|ccc} \begin{matrix} r \\ \tilde{I} \\ 0 \\ 0 \end{matrix} & \begin{matrix} n-k-r-s \\ \tilde{A}_1 \\ 0 \\ 0 \end{matrix} & \begin{matrix} k+s \\ \tilde{A}_2 \\ 0 \\ 0 \end{matrix} & \begin{matrix} r \\ \tilde{B} \\ D_1 \\ D_2 \end{matrix} & \begin{matrix} n-k-r-s \\ \tilde{C}_1 \\ I \\ 0 \end{matrix} & \begin{matrix} k+s \\ \tilde{C}_2 \\ E_2 \\ 0 \end{matrix} \end{array} \right. \quad (4.3.18)$$

Os últimos geradores s não podem comutar com os primeiros geradores r , a menos que $D_2 = 0$ e, portanto, podemos assumir que $s = 0$. Além disso, podemos também definir $C_1 = 0$, tomando combinações lineares apropriadas de linhas, então nossa matriz de verificação tem a forma :

$$n - k - r \left\{ \begin{array}{ccc|cc} \begin{matrix} r \\ \tilde{I} \\ 0 \end{matrix} & \begin{matrix} n-k-r \\ \tilde{A}_1 \\ 0 \end{matrix} & \begin{matrix} k \\ \tilde{A}_2 \\ 0 \end{matrix} & \begin{matrix} r \\ \tilde{B} \\ D \end{matrix} & \begin{matrix} n-k-r \\ \tilde{C} \\ I \\ E \end{matrix} \end{array} \right. \quad (4.3.19)$$

Onde rotulamos E_2 como E e D_1 como D . Não é difícil ver que este procedimento não é único. No entanto, diremos que qualquer código com matriz de verificação na forma (4.3.19) está em forma padrão.

4.4 CÓDIGOS ESTABILIZADORES

Códigos estabilizadores forma uma importante classe de códigos quânticos, com construção análoga a dos códigos clássicos lineares. Para se compreender os códigos estabilizadores é necessário compreender o *formalismo estabilizador*.

O formalismo estabilizador é uma ferramenta de descrição de uma extensa classe de operações da Mecânica Quântica, em que, é possível exemplificar como as portas lógicas e as medidas podem ser descritas por meio desse formalismo, que quantifica as limitações das operações estabilizadoras.

4.4.1. O FORMALISMO DE ESTABILIZADORES

Considere o estado EPR (Estado de Bell ou par de Einstein Podolsky-Rosen) de dois qubits como um exemplo para o formalismo estabilizador.

$$|\psi\rangle = \frac{|00\rangle + |11\rangle}{\sqrt{2}} \quad (4.4.1)$$

Sobre este estado algumas características são bastante importantes, ou seja, esse estado satisfaz às identidades $X_1 X_2 |\psi\rangle = |\psi\rangle$ e $Z_1 Z_2 |\psi\rangle = |\psi\rangle$. O estado $|\psi\rangle$ é estabilizado pelos operadores $X_1 X_2$ e $Z_1 Z_2$ de fato, esse estado quântico é o único que é estabilizado por tais operadores. Em muitos códigos quânticos utilizam-se os estabilizadores em vez dos vetores estados, pelo fato de que os estabilizadores são mais facilmente manipulados. A principal ideia do formalismo é identificar os estados quânticos por meio dos operadores que os estabilizam.

Os códigos CSS e o código de Steane que são descritos em vetores de estados, podem ser descritos de forma mais compacta utilizando-se o formalismo dos estabilizadores. Assim, os erros de qubit e operações tais como a porta de Hadamard, porta de fase, porta de CNOT e medidas computacionais, são facilmente descritas pelo formalismo de estabilizadores.

No formalismo de estabilizadores a ideia principal está relacionada à teoria de grupos. Para as aplicações relacionadas aos códigos quânticos o grupo fundamental é o grupo das matrizes de Pauli, G_n que atua sobre n qubits, com os fatores multiplicativos ± 1 e $\pm i$.

$$G_1 = \{\pm I, \pm iI, \pm X, \pm iX, \pm Y, \pm iY, \pm Z, \pm iZ\} \quad (4.4.2)$$

Na operação de multiplicação matricial, não podemos omitir os fatores ± 1 e $\pm i$, pois a razão de inclusão é que estes garantem que G_1 seja fechado com respeito à multiplicação, e também formem um grupo. O grupo de Pauli, de n qubits, é definido como sendo o conjunto de todas as matrizes formadas pelos produtos tensoriais de ordem n das matrizes de Pauli com os fatores ± 1 e $\pm i$.

Considere S um subgrupo de G_n e V_S o conjunto de estados de n qubits tal que são estabilizados por cada elemento de S . É fácil ver que V_S é o *espaço vetorial estabilizado por S* ; analogamente, S estabiliza V_S , na medida em que cada elemento de V_S não é modificado quando sobre ele age um elemento S .

Vamos introduzir alguns exemplos de formalismo estabilizadores. Para o caso $n = 3$ qubits e $S \equiv \{I, Z_1 Z_2, Z_2 Z_3, Z_1 Z_3\}$. O subespaço definido por $Z_1 Z_2$ é gerado por $|000\rangle, |001\rangle, |110\rangle$ e $|111\rangle$, o subespaço definido por $Z_1 Z_3$, é gerado por $|000\rangle, |100\rangle, |011\rangle$ e $|111\rangle$. Observe que os elementos $|000\rangle$ e $|111\rangle$ são os únicos estados quânticos que fazem parte dos conjuntos estabilizadores. Assim, conclui-se que o conjunto vetorial V_S é estabilizado pelos operadores $Z_1 Z_2$ e $Z_1 Z_3$, V_S é o subespaço gerado pelos estados $|000\rangle$ e $|111\rangle$. Percebe que o espaço vetorial de estados estáveis foi determinado através de dois operadores de S , pois o conjunto de vetores pode ser determinado através dos operadores geradores. Sendo visível a descrição de um grupo por meio de seus geradores.

Considere um conjunto de elementos, $g_1, \dots, g_l$ em um grupo G . Estes elementos são geradores do grupo G , se todo e qualquer elemento do grupo pode ser representado como um produto destes elementos, $g_1, \dots, g_l$. Usamos a notação usual $G = \langle g_1, \dots, g_l \rangle$ para denotar conjunto gerador de um dado grupo. No exemplo acima, o conjunto de geradores é $S = \langle Z_1 Z_2, Z_2 Z_3 \rangle$, pois $Z_1 Z_3 = (Z_1 Z_2)(Z_2 Z_3)$ e $I = (Z_1 Z_2)^2$. A vantagem de usar geradores para descrever grupos é que eles são uma representação compacta. Um grupo G de dimensão $|G|$ possui um conjunto de no máximo $\log(|G|)$ geradores. Agora, para ver se um vetor particular é estabilizado por um grupo S , é preciso verificar se o vetor é estabilizado por todos os seus geradores, pois, automaticamente, o mesmo é estabilizado pelos produtos destes.

Nem todo subgrupo S de um grupo de Pauli pode ser utilizado como estabilizador para um espaço vetorial não trivial. Como um exemplo, considere o subgrupo de G_1 , formado por $\{\pm I, \pm X\}$. A única solução para $(-I)|\psi\rangle = |\psi\rangle$ é $|\psi\rangle = 0$ e, portanto $\{\pm I, \pm X\}$ é o estabilizador para um espaço vetorial trivial.

Quais condições devem ser satisfeitas por S a fim de estabilizar um espaço vetorial não trivial V_S ? Para isso, é necessário obedecer as seguintes condições:

- a) Os elementos de S devem apresentar a propriedade comutativa e;
- b) $-I$ não pode ser um elemento de S .

Considere que um espaço vetorial não trivial que contenha um vetor $|\psi\rangle$ não nulo. Sejam M e N elementos S . Portanto, M e N são operadores que são produtos tensoriais das matrizes de Pauli. Uma das características das matrizes de Pauli é que elas comutam ou anticomutam entre si. Para estabelecer a condição (a), a de que essas matrizes comutam, vamos supor que M e N anticomutam e mostraremos que isso leva a uma contradição. Por hipótese $-NM = MN$, e, portanto $-|\psi\rangle = -NM|\psi\rangle = MN|\psi\rangle = |\psi\rangle$. Assim, da primeira e última igualdade segue que M e N , estabilizam $|\psi\rangle$. Assim, $|\psi\rangle = -|\psi\rangle$ ou seja, $|\psi\rangle=0$, que é uma contradição. Para provar a condição (b) observe que se $-I$ é um elemento de S , resulta que $-I|\psi\rangle = |\psi\rangle$, o que também é uma contradição.

O código de Steane é um bom exemplo de formalismo de estabilizadores. Na Tabela 1 estão destacados os geradores do grupo estabilizadores de operadores para o código de Steane de sete bits.

Tabela 1. Geradores estabilizadores para o código de Steane de sete qubits

NOME	OPERADORES
g_1	$IIIXXXX$
g_2	$IXXIIIX$
g_3	$XIXIXIX$
g_4	$IIIZZZZ$
g_5	$IZZIIZZ$
g_6	$ZIZIZIZ$

FONTE: Nielsen e Chuang, 2005.

Os elementos representam produtos tensoriais atuando no qubit respectivo. Por exemplo: $ZIZIZIZ = Z \otimes I \otimes Z \otimes I \otimes Z \otimes I \otimes Z = Z_1 Z_3 Z_5 Z_7$.

Percebemos a semelhança da estrutura entre os geradores da Tabela 1 e as matrizes verificadora de paridade usados na construção do código de Steane, em que os três primeiros geradores do estabilizador têm operadores X nos lugares correspondentes aos “1s” da matriz de verificadora de paridade de C_1 , os três últimos geradores de g_4 até g_6 possuem operadores Z nos lugares correspondentes aos “1s” da matriz verificadora de paridade $C_2^\perp$.

Para que os geradores $g_1, \dots, g_l$ sejam independentes, a remoção de qualquer gerador g_i torna o grupo menor:

$$\langle g_1, \dots, g_{i-1}, g_{i+1}, \dots, g_l \rangle \neq \langle g_1, \dots, g_l \rangle \quad (4.4.3)$$

Seja $S = \langle g_1, \dots, g_l \rangle$. Uma forma útil de se representar os geradores de S são por meio da matriz de paridade. Na verdade, é uma matriz $l \times 2n$ cujas linhas correspondem aos geradores de g_1 a g_l . O lado esquerdo da matriz contém “1s” para indicar quais geradores contém X , e o lado direito contém “1s” que indicam quais geradores contém Z . Quando aparece um “1” nos dois lados, indica uma matriz Y no gerador. Através da Tabela 1, pode ser construída a matriz verificadora de paridade do código de Steane de sete qubits.

$$\begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{pmatrix} \quad (4.4.4)$$

Assim, a matriz verificadora não contém informação sobre os fatores multiplicativos na frente dos geradores. Se usarmos $r(g)$ para representar o vetor linha $2n$ - dimensional de um elemento g do grupo de Pauli se define uma matriz Λ , $2n \times 2n$, dada por:

$$\Lambda = \begin{pmatrix} 0 & I \\ I & 0 \end{pmatrix} \quad (4.4.5)$$

As linhas I fora da diagonal são matrizes identidade $n \times n$. Se verifica que os elementos g e g' do grupo de Pauli comutam se, e somente se, $r(g) \Lambda (g')^T = 0$. A formula $x \Lambda y^T$ define de produto interno “torcido” entre as matrizes linha x e y que expressa os elementos do grupo de Pauli a x e y comutam ou não.

Neste momento, proferido algumas proposições que estabelecem uma conexão entre a independência dos geradores e a matriz verificadora.

Proposição 4.4.1 Seja $S = \langle g_1, \dots, g_l \rangle$ tal que $-I$ não é elemento de S . Os geradores I de g_1 até g_l são independentes se, e somente se, as linhas correspondentes da matriz verificadora forem linearmente independentes.

Demonstração: É fácil ver que g_i^2 deve ser igual a I para todo i . Note que $r(g) + r(g) = r(gg)$, logo a adição na representação de vetores linha corresponde à multiplicação dos elementos do grupo. Assim, as linhas da matriz verificadora serão linearmente dependentes, $\sum_i a_i r(g_i) = 0$, com $a_j \neq 0$ para algum j , se, e somente se $\prod_i g_i^{a_i}$ for igual à identidade, a menos de um fator multiplicativo. Mas $-I \notin S$, e, portanto o fator multiplicativo deve ser igual a 1, e a última condição será $g_j = g_j^{-1} = \prod_{i \neq j} g_i^{a_i}$ portanto $g_1 \dots g_l$ não são geradores independentes. $\square$

A próxima proposição transporta a dimensão V_S é 2^k quando S , for gerado por $l = n - k$ geradores independentes comutativos, e $-I \notin S$.

Proposição 4.4.2 Seja $S = \langle g_1, \dots, g_l \rangle$ gerado por l geradores independentes, tal que $-I \notin S$. Fixe i no intervalo $1, \dots, l$. Resulta que existe $g \in G_n$ tal que $g g_i g^\dagger = -g_i$ e $g g_j g^\dagger = g_j$ para todo $j \neq i$.

Demonstração: Seja G a matriz verificadora associada a $g_1, \dots, g_l$. Com a Proposição 4.4.1, as linhas de G são linearmente independentes, logo existe um vetor x , $2n$ -dimensional, tal que $G \wedge x = e_i$, em que e_i é um vetor l -dimensional contendo "1" na i -ésima posição e 0 em todas as outras. Seja g tal que $r(g) = x^T$. Por definição de x , teremos $r(g_i) \wedge r(g)^T = 0$, para todo $j \neq i$ e $r(g_i) \wedge r(g)^T = 1$. Portanto, $g g_i g^\dagger = -g_i$ e $g g_j g^\dagger = g_j$ para todo $j \neq i$. $\square$

Além disso, existe ainda uma proposição que relaciona o fato espaço vetorial estabilizado que será não trivial de S se for gerado por geradores independentes comutativos e $-I \notin S$. Ou seja, se existem $l = n - k$, geradores ao qual será demonstrado que V_S seja 2^k dimensional em que cada gerador adicional para o estabilizador divide a dimensão de V_S por um fator 2, uma vez que os auto espaços $+1$ e -1 do produto tensorial das matrizes de Pauli dividem o espaço de Hilbert total em dois subespaços de igual dimensão como será visto na proposição a seguir.

Proposição 4.4.3 Seja $S = \langle g_1, \dots, g_{n-k} \rangle$ gerado por $n - k$ elementos independentes e que comutam de G_n , e suponha que $-I \notin S$. Resulta que V_S é um espaço vetorial 2^k - dimensional.

Nas descrições posteriores sobre o formalismo de estabilizadores, será usado a convenção de que estabilizadores são sempre descritos em termos de geradores independentes comutativos tais que $-I \notin S$.

Demonstração da Proposição 4.4.3: Seja $x = (x_1, \dots, x_{n-k})$ um vetor com $n - k$ elementos de Z_2 . Defina

$$P_S^x \equiv \frac{\prod_{j=1}^{n-k} (I + (-1)^{x_j} g_j)}{2^{n-k}} \quad (4.4.6)$$

Como $(I + g_j) / 2$ é o projetor sobre o espaço $+1$ de g_j , é fácil ver que $P_S^{(0, \dots, 0)}$ deve ser o projetor sobre V_S . De acordo com a proposição 4.4.2, para cada x existe g_x em G_n , tal que $g_x P_S^{(0, \dots, 0)} (g_x)^\dagger = P_S^x$, de onde se observa que a dimensão P_S^x é a mesma que a V_S . Pode-se ver facilmente que para diferentes x os P_S^x são ortogonais. A demonstração é concluída com ressalva de que:

$$I = \sum_x P_S^x \quad (4.4.7)$$

O lado esquerdo é um projetor sobre um espaço 2^n -dimensional, e o lado direito é a soma sobre 2^{n-k} projetores ortogonais da mesma dimensão que V_S , e portanto, a dimensão dos V_S tem que ser 2^k . $\square$

4.5. PORTAS UNITÁRIAS E O FORMALISMO DE ESTABILIZADORES

Foi analisado anteriormente o formalismo de estabilizadores para descrever espaços vetoriais. O processo de dinâmica desses espaços vetoriais diante das interações de outros operadores é útil na descrição dos efeitos dos ruídos sobre o espaço de estados o qual pertence o código corretor de erros. Considere que um operador linear unitário U seja posto a interagir com um estado quântico que é estabilizado por um grupo S . Seja $|\psi\rangle$ um elemento qualquer de V_S . Então, para qualquer elemento g de S , tem-se que

$$U |\psi\rangle = U g |\psi\rangle = U g U^\dagger U |\psi\rangle \quad (4.5.1)$$

O estado $U|\psi\rangle$ é estabilizado por $U g U^\dagger$. Ao qual pode ser estendido para todo espaço vetorial, já que não se realizou nenhuma observação que afetasse a generalidade das hipóteses, resultando em garantir que o espaço UV_S é estabilizado por $USU^\dagger \equiv \{U g U^\dagger | g \in S\}$. Essa extensão pode ser alcançada ao mesmo tempo aos geradores, pois que o conjunto $g_1, \dots, g_l$ gera o grupo S , resultando que $Ug_1U^\dagger, \dots, Ug_lU^\dagger$ gera $USU^\dagger$. Assim, para computar as mudanças ocorridas nos estabilizadores, faz necessário, computar a mudança sobre os geradores do estabilizador.

Para certas operações unitárias U , a transformação dos geradores adquire uma forma interessante.

Exemplo 4.2.1 Para uma porta Hadamard sobre um único qubit tem-se que

$$H X H^\dagger = Z; H Y H^\dagger = -Y; H Z H^\dagger = X \quad (4.5.2)$$

Em que, pode-se deduzir que após a porta de Hadamard ter sido aplicada ao estado estabilizador por Z (o estado $|0\rangle$), o estado resultante será estabilizado por X ou seja, (o estado $|+\rangle$).

Tabela 2. Propriedades de transformação dos elementos do grupo de Pauli sob conjugação por várias operações comuns. O NÃO CONTROLADO tem o q-bit 1 como controle e o q-bit 2 como alvo.

OPERAÇÃO	ENTRADA	SAÍDA
Não- controlado	X_1	$X_1 X_2$
	X_2	X_2
	Z_1	Z_1
	Z_2	$Z_1 Z_2$
H	X	Z
	Z	X
S	X	Y
	Z	Z
X	X	X
	Z	$-Z$
Y	X	$-X$
	Z	$-Z$
Z	X	$-X$
	Z	Z

FONTE: Nielsen e Chuang, 2005.

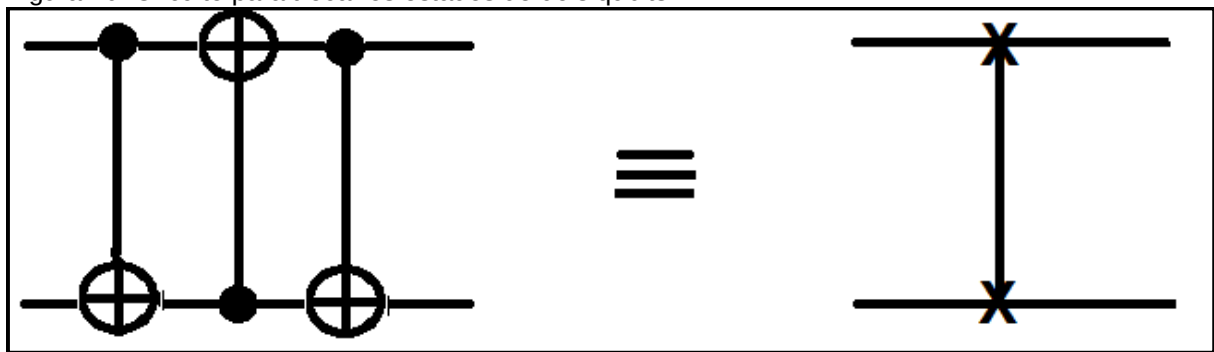
Agora, considere os operadores X_1 , X_2 , Z_1 , Z_2 se comportam sobre a operação conjugada pela porta CNOT. Denota-se essa porta por U , como o qubit 1 sendo o qubit de controle e o 2 como qubit alvo, temos:

$$U X_1 U^\dagger = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix} = X_1 X_2$$

Observamos que $U X_2 U^\dagger = X_2$, $U Z_1 U^\dagger = Z_1$ e $U Z_2 U^\dagger = Z_1 Z_2$. Como U conjuga os outros operadores do grupo de Pauli para dois qubits. Tomamos por exemplo, para calcular $U X_1 X_2 U^\dagger$ notamos que $U X_1 X_2 U^\dagger = U X_1 U^\dagger U X_2 U^\dagger = (X_1 X_2) X_2 = X_1$. A componente Y da matriz de Pauli tem $U Y_2 U^\dagger = i U X_2 Z_2 U^\dagger = i U X_2 U^\dagger U Z_2 U^\dagger = i X_2 (Z_1 Z_2) = Z_1 Y_2$.

Na dinâmica unitária, ainda para o formalismo de estabilizadores, considere um circuito para a porta TROCA, como mostra o circuito na Figura 10. A porta quântica TROCA transforma o operador Z_1 na sequência $Z_1 \rightarrow Z_1 \rightarrow Z_1 Z_2 \rightarrow Z_2$, e o operador Z_2 em $Z_2 \rightarrow Z_1 Z_2 \rightarrow Z_1 \rightarrow Z_1$. De forma análoga, tal porta opera em X_1 como, $X_1 \rightarrow X_2$ e $X_2 \rightarrow X_1$. Tomando U como operador TROCA conclui-se que $U Z_1 U^\dagger = Z_2$ e $U Z_2 U^\dagger = Z_1$, o mesmo para X_1 e X_2 como ocorre no circuito da Figura. 10.

Figura 10. Circuito para trocar os estados de dois qubits



FONTE: Nielsen e Chuang, 2005.

Um dos exemplos utilizados no formalismo de estabilizadores é o circuito da porta TROCA. Mas, portas juntas como a de Hadamard e CNOT podem criar emaranhamento.

Definição 4.5 O emaranhamento, aceita que dois ou mais subsistemas estejam correlacionados de tal forma que um subsistema não pode ser descrito corretamente sem que a sua contra-parte seja referenciada.

A próxima porta a de fase é mais importante desse conjunto.

Para a porta de fase, ou seja, uma porta de um único qubit é reescrito como.

$$S = \begin{pmatrix} 1 & 0 \\ 0 & i \end{pmatrix}$$

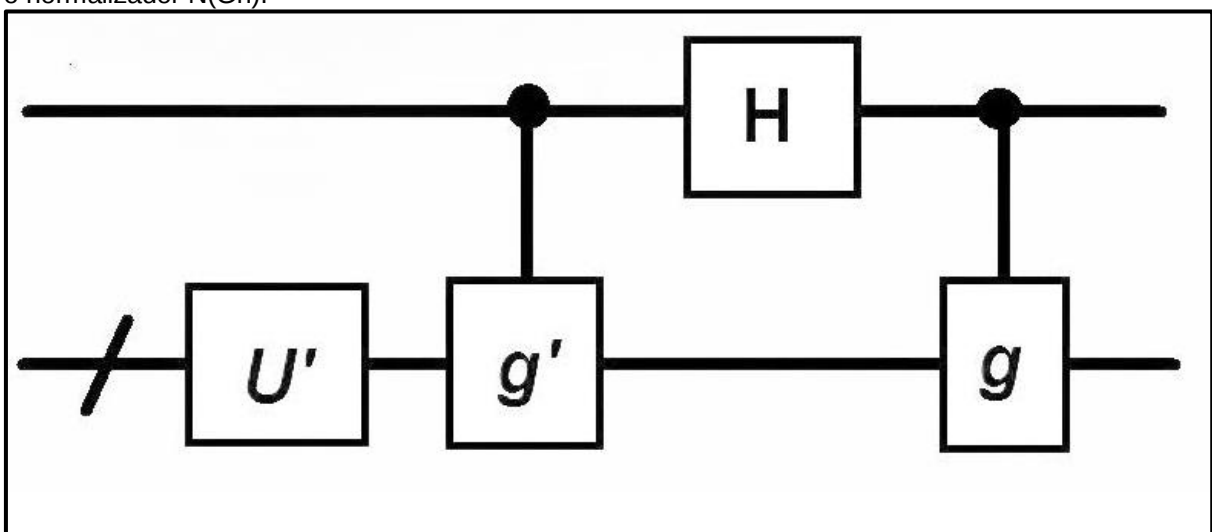
A porta de fase sobre a matriz de Pauli é computada por

$$S X S^\dagger = Y; \quad S Z S^\dagger = Z, \quad (4.5.5)$$

Por definição, diz-se que o conjunto de operadores U , tais que $U G_n U^\dagger = G_n$ é normalizador de G_n e o denotamos por $N(G_n)$. Assim, o normalizador de G_n é gerado pelas portas de Hadamard, porta de fase e porta de CNOT, em que essas portas também pode ser chamadas de portas normalizadoras.

Teorema 4.5.1 Considere U um operador unitário sobre n qubits, com a propriedade de que, se $g \in G_n$, então $U g U^\dagger \in G_n$. Resulta que, a menos de uma fase global U pode ser construído a partir de $O(n^2)$ portas Hadamard de fase e CNOT.

Figura 11. Construção usada na demonstração de que as portas Hadamard, de fase de CNOT geram o normalizador $N(G_n)$.



FONTE: Nielsen e Chuang, 2005.

4.6. MEDIDAS SEGUNDO O FORMALISMO DE ESTABILIZADORES

As medidas de base computacional também podem ser descritas nesse formalismo. Suponha que se realize uma medida de g e que este elemento pertença ao grupo de Pauli para n qubits. É possível, pois g é um operador linear e pode ser visto como observável ou seja, em que todas as grandezas reais que podem ser medidas são denotadas por operadores. Relacionando as matrizes de Pauli, não há problema em supor que este elemento é formado por uma multiplicação de Pauli, sem inserir qualquer fator multiplicativo, $-1 \pm i$. Considere ainda que o sistema quântico esteja no estado $|\psi\rangle$ que apresenta como grupo estabilizador $S = \langle g_1, \dots, g_n \rangle$. É possível descrever as transformações que ocorrem após a realização da medida. Existem duas possibilidades:

- g comuta com todos os geradores do estabilizador.
- g anticomuta com um ou mais geradores do estabilizador.

O estabilizador apresenta como geradores $g_1, \dots, g_n$, e que g anticomuta com g_1 . Não oferece dano à teoria e supor também que g comuta com os geradores $g_2, \dots, g_n$. Como isso não ocorre a anticomutatividade com um dos geradores, por exemplo, g_2 comutará $g_1 g_2$ em que pode ser substituído por g_2 na lista de geradores do estabilizador.

Analisando o primeiro caso, pode-se garantir que g ou $-g$ faz parte do conjunto de geradores, pois $g_j g|\psi\rangle = g g_j |\psi\rangle = g|\psi\rangle$. Isso implica que o estado $g|\psi\rangle$ faz parte do espaço vetorial V_S formado por estados quânticos que são estabilizados por S . Portanto, o estado $g|\psi\rangle$ é um múltiplo de $|\psi\rangle$, mas como $g \in G_n$ tem-se que $g^2 = I$. Então, $g|\psi\rangle = \pm |\psi\rangle$, implicando no fato que g ou $-g$ é um estabilizador de $g|\psi\rangle$. Nesse caso, $g|\psi\rangle = |\psi\rangle$, então, uma medida como g resulta sempre, no valor +1 com probabilidade igual a 1. Fato $g|\psi\rangle = -|\psi\rangle$, uma medida como g resulta sempre no valor -1 com probabilidade 1.

No segundo caso, o operador g anticomuta com g_1 e comuta com todos os outros geradores estabilizadores? Como $g \in G_n$, os seus possíveis autovalores de g são ± 1 , e seus projetores para as medidas são $(I + g) / 2$. Agora, as medidas podem ter os seguintes resultados de probabilidades.

$$p(+1) = \text{tr} \left(\frac{I+g}{2} |\psi\rangle \langle\psi| \right) \quad (4.6.1)$$

$$p(-1) = \text{tr} \left(\frac{I-g}{2} |\psi\rangle \langle\psi| \right) \quad (4.6.2)$$

Sabendo que $g_1 |\psi\rangle = |\psi\rangle$ e $gg_1 = -g_1g$, é possível expressar a probabilidade de resultado 1 como:

$$p(+1) = \text{tr} \left(\frac{I+g}{2} g_1 |\psi\rangle \langle\psi| \right) \quad (4.6.3)$$

$$p(+1) = \text{tr} \left(g_1 \frac{I+g}{2} |\psi\rangle \langle\psi| \right) \quad (4.6.4)$$

Aplicando a propriedade cíclica do traço, pode-se levar g_1 ao lado direito do produto e absorver em $\langle\psi|$ usando $g_1 = g_1^\dagger$, temos.

$$p(+1) = \text{tr} \left(\frac{I-g}{2} |\psi\rangle \langle\psi| \right) = p(-1) \quad (4.6.5)$$

Portanto, $p(+1) + p(-1) = 1$ então $p(+1) = p(-1) = 1/2$. No caso do resultado da medida ser de valor 1, o estado quântico é modificado com a medida e o conjunto estabilizador é modificado de modo acrescentar novo estado quântico. $\square$

5. CONCLUSÃO

As principais contribuições desta dissertação de mestrado foram as seguintes: Demonstrar alguns resultados que não aparecem no livro texto, com relação à Teoria de Matróides. Descrever de forma mais clara parte da Teoria de Matróides bem como parte da Teoria de Codificação Quântica (códigos estabilizadores). Utilização da Teoria de Codificação Clássica para o entendimento de uma classe particular de Códigos Quânticos, a saber, a classe dos códigos quânticos estabilizadores Calderbank-Shor-Steane (CSS). E Início da tentativa de relacionar a Teoria de Matróides com a Teoria de Codificação Quântica, mediante a possível identificação das matrizes de um código CSS com as matrizes que geram os respectivos Matróides vetoriais. Além disso, conclui-se que estas teorias são bastante abrangentes tanto na Matemática quanto na Física, acreditamos ter possibilitado uma interação eficaz entre tais áreas de conhecimento, o que pode contribuir para despertar o interesse de futuros pesquisadores nessas áreas.

6. SUGESTÕES PARA TRABALHOS FUTUROS

Com o advento do computador quântico, faz-se necessário ampliar e aprimorar a construção de códigos quânticos existentes na literatura. Como contribuições futuras, esperamos continuar o trabalho no doutorado gerando-se, possivelmente, códigos quânticos eficientes. Além disso, uma tentativa natural é a criação de uma classe de códigos quânticos derivados de Matróides vetoriais, ou mesmo derivados de módulos vetoriais (que é uma generalização do conceito de espaços vetoriais, em que os escalares são definidos sobre um anel, e não sobre um corpo).

Outro tópico para trabalhos futuros é a criação de códigos quânticos topológicos provenientes de Matróides (vetoriais ou outras classes de Matróides). Evidentemente, este último tópico requer bastante estudo e investigação, pois ambas as teorias apresentam as dificuldades de aprendizado, inerentes ao processo intrínseco de abstração.

7. REFERÊNCIAS

ASHIKHMIN, A.; MEMBER, S.; IEEE. **Fidelity Lower Bounds for Stabilizer and CSS Quantum Codes**. IEEE Transaction on information. Theory, vol.60. No June, 2014.

BARROS, C. M. F. **Um estudo sobre a decodificação de códigos corretores de erros quânticos**. Recife, julho , 2011. Disponível em <<http://www.de.ufpe.br/~hmo/dissertacaoCAIO.pdf> > Acesso 05/ 05/ 2016.

BIRKHOFF,; G. **Abstract Linear Dependence and Lattices**. Amer. J. Math. 57 (1935).

BONDY, J. A e WESH, D. J. A. **Some results ou transversal matróids and construction for identically self-dual matroids**.(1971).

COSTALONGA, J. P. **Cocircuitos Não-separados que Evitam um Elemento e Graficidade em Matróides Binárias**. Universidade Federal de Pernambuco. 2011. Disponível em: <http://repositorio.ufpe.br:8080/xmlui/bitstream/handle/123456789/7100/arquivo648_1.pdf?sequence=1&isAllowed=y>. Acesso 22/03/2017

EDMONDS, J. e FULKERSON, D. R. **Combinatorial Theory**. Bottleneck extema ,J. (1970).

LA GUARDIA, G.G.; PALAZZO, J. R.; LAVOR, C. **Relações entre Matróides e Códigos de Blocos Lineares**. 2007. Disponível em: <http://www.sbmec.org.br/eventos/cnmac/xxx_cnmac/PDF/784.pdf > Acesso em 12/01/2016.

HARTMANIS, J. **Lattice Theory of Generalized Partitions**. Canadian Journal of Mathematics 11, 97–106. 1959).

KETKAR, A.; KLAPPENECKER, A.; KUNAR, S.; SARVEPALLI, P. M. **Nonbinary Stabilizer Codes over Finite Fields**. University, Department of Computer Science, College Station, TX77843-3112. No 17 Aug, 2005.

MACLANE, S.: **Some interpretations of abstract linear dependence in terms of projective geometry**, American Journal of Mathematics 58 (1936).

MACWILLIAMS, F.J.: N.J.A SLOANE. **The theory of error correcting codes**. Amterdam (1977).

NIELSEN, M. A.; CHUANG. I.L. **Computação Quântica e Informação Quântica**. 3º ed. Porto Alegre, 2005.

OXLEY, J.G. **Matroid Theory**. Oxford university Press, 1992.

PALAZZO, R, J.; ALMEIDA, A. C. **Códigos Concatenados Corretores de Erros Quânticos**. Revista da Sociedade Brasileira de Telecomunicações. Volume 21. Agosto, 2006. Disponível em: <
https://www.researchgate.net/profile/Reginaldo_Palazzo/publications >. Acesso em 12/ 06/ 2016.

SIMÃO, J. B. **Minimização de funções submodulares**. Dissertação de Mestrado. São Paulo, 2009.

VAN DER WAERDEN, B. L.: **Moderne Algebra** Vol. 1. Berlin, Springer Verlag, 1937.

WELSH, D.J. **Matroid Theory**. Academic Press, London, 1976.

WHITNEY, H. **“On the abstract proprieties of linear dependence”**. Amer. J. Math. (1935).