
Universidade Estadual de Campinas

Instituto de Matemática, Estatística e Computação Científica

Departamento de Matemática

Dissertação de Mestrado

UM ESTUDO SOBRE FATORAÇÕES DE MATRIZES E A RESOLUÇÃO DE SISTEMAS LINEARES

por

Ludio Edson da Silva Campos

Mestrado Profissional em Matemática - Campinas - SP

Orientadora: Profa. Dra. Maria Zoraide Martins Costa Soares

UM ESTUDO SOBRE FATORAÇÕES DE MATRIZES E A RESOLUÇÃO DE SISTEMAS LINEARES

Este exemplar corresponde à redação final da tese devidamente corrigida e defendida por **Ludio Edson da Silva Campos** e aprovada pela Comissão Julgadora.

Campinas, 07 de Março de 2008.



Profa. Dra. Maria Zoraide M. C. Soares
Orientadora

Banca Examinadora:

Profa. Dra. Maria Zoraide M. C. Soares

Prof. Dr. Leônidas de Oliveira Brandão

Profa. Dra. Maria Cristina de Castro Cunha

Dissertação apresentada ao Instituto de Matemática, Estatística e Computação Científica, UNICAMP, como requisito parcial para a obtenção do título de **Mestre em Matemática**.

**FICHA CATALOGRÁFICA ELABORADA PELA
BIBLIOTECA DO IMECC DA UNICAMP**

Bibliotecária: Crislene Queiroz Custódio – CRB8a 162/2005

Campos, Ludio Edson da Silva

C157e Um estudo sobre fatorações de matrizes e a resolução de sistemas lineares / Ludio Edson da Silva Campos -- Campinas, [S.P. : s.n.], 2008.

Orientador : Maria Zoraide Martins Costa Soares

Trabalho final (Mestrado Profissional) - Universidade Estadual de Campinas, Instituto de Matemática, Estatística e Computação Científica.

1. Matrizes (Matemática). 2. Fatoração (Matemática). 3. Sistemas lineares. 4. Algoritmos. I. Soares, Maria Zoraide Martins Costa. II. Universidade Estadual de Campinas. Instituto de Matemática, Estatística e Computação Científica. III. Título.

Titulo em inglês: A study on matrix factorizations and the resolution of linear systems.

Palavras-chave em inglês (Keywords): 1. Matrices (Mathematics). 2. Factorization (Mathematics). 3. Linear systems. 4. Algorithms.

Área de concentração: Álgebra linear

Titulação: Mestre Profissional em Matemática

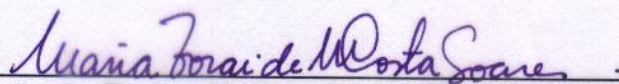
Banca examinadora: Profa. Dra. Maria Zoraide Martins Costa Soares (IMECC/UNICAMP)
Prof. Dr. Leônidas de Oliveira Brandão (USP)
Profa. Dra. Maria Cristina de Castro Cunha (IMECC/UNICAMP)

Data da defesa: 07/03/2008

Programa de Pós-Graduação: Mestrado Profissional em Matemática

Dissertação de Mestrado defendida em 07 de março de 2008 e aprovada

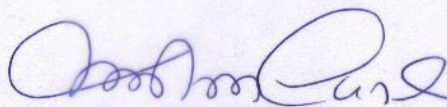
Pela Banca Examinadora composta pelos Profs. Drs.



Prof. (a). Dr (a). MARIA ZORAIDE MARTINS COSTA SOARES



Prof. (a). Dr (a). LEÔNIDAS DE OLIVEIRA BRANDÃO



Prof. (a). Dr (a). MARIA CRISTINA DE CASTRO CUNHA

Àqueles que enxergam nas ciências a possibilidade de construir um mundo melhor.

“Os ideais que iluminaram meu caminho e sempre me deram coragem para enfrentar a vida com alegria foram a Verdade, a Bondade e a Beleza.”

(Albert Einstein)

Agradecimentos

- A Deus pela experiência da vida e por poder desfrutar do amor dos meus pais, irmãos e amigos. Agradeço “por este mundo, nosso grande lar; por sua vastidão e riqueza, e pela vida multiforme que nele estua e de que todos fazemos parte. (...) pelo esplendor do céu azul e pela brisa da tarde, e pelas nuvens rápidas e pelas constelações nas alturas.(...) pelos oceanos imensos, pela água corrente, pelas montanhas eternas, pelas árvores frondosas e pela relva macia em que os nossos pés repousam. (...) [pelos] múltiplos encantos com que podemos sentir, em nossa alma, as belezas da Vida e do Amor!”
- Aos meus pais pelo amor que me devotam, pelos sutis ensinamentos que constituem meu caráter e minha humanidade. Por estarem sempre comigo, mesmo à distância. Amo-os indescritivelmente.
- À professora Zoraide, minha orientadora e amiga. Sempre será para mim um exemplo e uma inspiração muito forte. Pela dedicação que tem me dispensado, pelos estímulos e cobranças.
- Ao Cristiano Torezzan, doutorando em Matemática Aplicada na UNICAMP, pela imensa colaboração durante todo este trabalho, em particular na etapa final. Pela amizade, incentivo e apoio.
- A todos meus amigos e parentes, muitos dos quais torceram bastante para que este momento se concretizasse. Em especial à minha irmã Zenaide.
- Aos idealizadores desse projeto, por terem sonhado este sonho, contribuindo com o crescimento de muita gente, atentos ao fato de que educação se multiplica. Nós havemos de conduzir outros ao caminho libertador do conhecimento.
- Aos amigos do mestrado, em especial Cleuber, Donizete, Filardes (Maranhão), Giseli, Luiz Antonio, Mara (MA) e as duas Veras. A amizade de cada um de vocês é muito importante.

- Aos amigos Astrogildo, Bete, Edna Laet, Elídia, Elza, Emerson, Fátima Ricci, Greici, Irene, Magna, Michele, Naná, Paulo, Rosilane, Rosirlene, dentre outros igualmente amados.

Resumo

CAMPOS, Ludio Edson da Silva. *Um Estudo Sobre Fatorações de Matrizes e a Resolução de Sistemas Lineares*. Campinas - SP: Universidade Estadual de Campinas, 2008. Dissertação apresentada como requisito parcial para obtenção do Título de Mestre em Matemática.

Neste trabalho abordamos algumas fatorações de matrizes, com vistas à resolução de sistemas lineares através de métodos diretos. Enfocamos particularmente as decomposições LU, Cholesky e QR, cujo uso tem sido largamente difundido em implementações computacionais. Nosso objetivo é apresentar um texto didático, acessível a alunos de graduação, que contemple a teoria básica de cada fatoração, incluindo a demonstração dos principais resultados, e que também forneça condições para uma primeira implementação de cada decomposição. Sugerimos alguns algoritmos, que foram implementados no software livre OCTAVE, através dos quais comparamos o tempo gasto para resolução de alguns sistemas lineares, utilizando as fatorações citadas.

Palavras-Chave: Matrizes, Fatorações, Sistemas Lineares, Algoritmos.

Abstract

CAMPOS, Ludio Edson da Silva. *Um Estudo Sobre Fatorações de Matrizes e a Resolução de Sistemas Lineares*. Campinas - SP: Universidade Estadual de Campinas, 2008. Dissertação apresentada como requisito parcial para obtenção do Título de Mestre em Matemática.

In this work we discuss some matrix factorizations, with a view to the resolution of linear systems through direct methods. We focus particularly the LU, Cholesky and QR decompositions, whose use has been widely spread in computer implementations. Our goal is to present a didactic text, accessible to undergraduate students, which contemplates the basic theory of each factorization, including the demonstration of the main result and that also provide conditions for a first implementation of each decomposition. We suggest some algorithms that were scheduled in the free software OCTAVE, through which we compare the time elapsed for the resolution of a few linear systems, using the factorizations cited.

Keywords: Matrices, Factorizations, Linear Systems, Algorithms.

Sumário

Agradecimentos	vi
Resumo	viii
Abstract	ix
Introdução	1
1 Matrizes	2
1.1 Matrizes Elementares	2
1.2 Matriz Escalonada	7
1.3 Matrizes de Transformação de Gauss	9
1.4 Matrizes Triangulares	14
2 Fatoração de Matrizes	19
2.1 Fatoração LU	19
2.1.1 Contagem do número de operações para obter a fatoração $A=LU$. .	25
2.1.2 Fatoração LU de uma Matriz Retangular	28
2.2 Fatoração de Cholesky	29
2.3 Fatoração QR	41
2.3.1 Projeções Ortogonais	45
3 Sistemas Lineares	54
3.1 Eliminação Gaussiana	55
3.2 Resolução de Sistemas Lineares Usando Fatorações	58
3.2.1 Resolvendo um sistema linear usando a fatoração LU	58
3.2.2 Resolvendo um sistema linear usando a fatoração de Cholesky	63
3.3 Pivoteamento Parcial	65

3.4	Mínimos Quadrados e a Solução de Sistemas Lineares	68
3.4.1	Projeções Ortogonais e Aproximações	68
3.4.2	Resolução de sistemas lineares por mínimos quadrados usando QR . .	75
3.5	Implementação de métodos para resolver sistemas lineares usando o OCTAVE	78
3.5.1	Algoritmos para Resolver Sistemas Lineares por Mínimos Quadrados .	83
Considerações Finais		86
Referências Bibliográficas		87

Introdução

Pretendemos apresentar neste trabalho um texto introdutório sobre métodos diretos de resolução de sistemas de equações lineares, mais particularmente, aqueles que fazem uso das fatorações LU, Cholesky e QR. Este material é direcionado a alunos que tenham tido um primeiro contato com os assuntos de um curso de Álgebra Linear, pois alguns assuntos podem exigir um pouco desse conhecimento para serem melhor compreendidos. Tivemos o cuidado de apresentar as demonstrações dos resultados, excetuando-se alguns que são facilmente encontrados em livros com publicações em português.

Apesar de todo o trabalho estar relacionado ao estudo de matrizes, não tratamos neste texto de matrizes esparsas, cuja aritmética é especial, com vistas ao aproveitamento de apenas os elementos não-nulos da matriz.

No primeiro capítulo procuramos apresentar as principais definições e resultados envolvendo matrizes e que usamos nos capítulos seguintes. Existe uma teoria bastante ampla referente ao estudo de matrizes, sendo boa parte relacionada às aplicações computacionais. Neste trabalho o enfoque é relativamente restrito, pois nos atemos às propriedades matriciais ligadas às fatorações.

No segundo capítulo apresentamos as fatorações, exemplificando passo a passo como obtê-las. Fizemos uma análise de complexidade do número de operações executadas na obtenção da fatoração LU, no intento de mostrar como se deve fazer esse cálculo.

Reservamos para o terceiro capítulo o estudo de sistemas lineares, algoritmos e algumas análises numéricas. Lá apresentamos uma implementação dos métodos de fatoração e de resolução no software livre Octave. Resolvemos alguns sistemas lineares usando implementações, com o intuito de comparar o tempo decorrido para obter a solução através dos métodos estudados. Assim, as últimas seções apresentam uma ligação entre Álgebra Linear e Cálculo Numérico.

Matrizes

Neste capítulo objetivamos cobrir as definições que reiteradas vezes são utilizadas neste trabalho, bem como alguns teoremas cujos resultados serão necessários para a demonstração de outros.

Desse modo, omitiremos definições que julgamos mais comuns ou menos relevantes aqui, como, por exemplo, matriz quadrada, matriz coluna, matriz linha, matriz identidade, etc. Consideraremos também de conhecimento do leitor as principais operações e notações envolvendo matrizes.

1.1 Matrizes Elementares

O estudo de matrizes elementares é particularmente importante para a compreensão de vários resultados, visto que são imprescindíveis no processo de diversas demonstrações.

Veremos a seguir que o método de escalonamento, realizado por meio de operações elementares aplicadas às linhas de uma matriz, é equivalente à multiplicação de matrizes elementares à esquerda da matriz a ser escalonada.

As operações elementares aplicadas sobre as linhas de uma matriz são:

- (I) Trocar duas linhas entre si.
- (II) Multiplicar uma linha por uma constante não-nula.
- (III) Somar a uma linha dada um múltiplo de outra linha.

Definição 1.1 Chamamos de matriz elementar à matriz obtida ao se aplicar uma operação elementar sobre as linhas da matriz identidade.

Vejam os adiante exemplos de matrizes elementares 4×4 . Denotando L_i como a i -ésima linha de uma dada matriz e I_n a matriz identidade de ordem n , temos que as matrizes E_1 , E_2 e E_3 a seguir foram obtidas da identidade I_4 ao fazermos, respectivamente, *troca entre L_2 e L_4* , kL_2 e $L_3 + kL_1$:

$$E_1 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix} \quad E_2 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & k & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \quad E_3 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ k & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

Observamos que se B é uma matriz obtida de A ao aplicarmos uma operação elementar sobre uma linha de A , então B pode ser expressa da forma $B = E.A$, onde E é a matriz elementar correspondente à operação sobre A . Ilustraremos esse fato para o caso de uma matriz $A_{4 \times 4}$, ressaltando que o mesmo vale para uma matriz $A_{m \times n}$.

Seja

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix}$$

(i) Denotando por B_1 a matriz obtida ao trocarmos as linhas L_2 e L_4 de A temos:

$$B_1 = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{41} & a_{42} & a_{43} & a_{44} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{bmatrix}$$

Fazendo agora o produto $E_1 A$, onde E_1 é obtida da identidade I_4 pela troca entre as linhas L_2 e L_4 , segue-se que:

$$E_1 A = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{41} & a_{42} & a_{43} & a_{44} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{bmatrix} = B_1$$

(ii) Chamando de B_2 a matriz obtida de A ao multiplicarmos a linha L_3 por uma constante não-nula k , temos:

$$B_2 = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ ka_{31} & ka_{32} & ka_{33} & ka_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix}$$

Se E_2 é obtida ao efetuarmos o produto da terceira linha da identidade I_4 por k , então:

$$E_2 A = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & k & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ ka_{31} & ka_{32} & ka_{33} & ka_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} = B_2$$

(iii) Chamando de B_3 a matriz resultante da substituição em A da linha L_2 pela soma $L_2 + kL_4$, temos:

$$B_3 = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} + ka_{41} & a_{22} + ka_{42} & a_{23} + ka_{43} & a_{24} + ka_{44} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix}$$

Determinando $E_3 A$, onde E_3 difere da identidade I_4 pela substituição da linha L_2 pela soma $L_2 + kL_4$:

$$E_3 A = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & k \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} + ka_{41} & a_{22} + ka_{42} & a_{23} + ka_{43} & a_{24} + ka_{44} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix} = B_3$$

Dada uma operação elementar, denotamos por sua inversa à operação elementar que desfaz a alteração provocada numa matriz pela aplicação da primeira. Desse modo, a inversa das operações elementares *troca das linhas L_i e L_j* , kL_i e $L_i + kL_j$ são, respectivamente, *troca das linhas L_j e L_i* , $\frac{1}{k}L_i$ e $L_i - kL_j$. De fato, se A é uma matriz sobre a qual executamos uma operação elementar, então somente a aplicação da operação elementar que indicamos como inversa nos retornará à matriz A .

Vejamos a seguir um exemplo usando a operação elementar (III), bem como sua inversa. Também colocamos ao lado a equivalência usando produto por matrizes elementares à esquerda.

Exemplo: Seja

$$A = \begin{bmatrix} 1 & 0 & 3 \\ -4 & 7 & 6 \\ 10 & 2 & 5 \end{bmatrix}$$

Então aplicando a operação elementar $L_2 + 4L_1$ em A obtemos:

$$B = \begin{bmatrix} 1 & 0 & 3 \\ 0 & 7 & 18 \\ 10 & 2 & 5 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 4 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & 3 \\ -4 & 7 & 6 \\ 10 & 2 & 5 \end{bmatrix} \quad (E_1 A = B)$$

Executando agora em B a operação inversa $L_2 - 4L_1$ segue-se

$$A = \begin{bmatrix} 1 & 0 & 3 \\ -4 & 7 & 6 \\ 10 & 2 & 5 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ -4 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & 3 \\ 0 & 7 & 18 \\ 10 & 2 & 5 \end{bmatrix} \quad (E_1^{-1} B = A)$$

Teorema 1.1 *Toda matriz elementar é invertível e sua inversa é também elementar.*

Prova: Considere a aplicação de uma operação elementar sobre alguma linha da matriz identidade I . Isso produzirá uma matriz elementar E .

Suponha agora que apliquemos em E a operação elementar inversa da que foi aplicada em I . Isso resultará novamente a I . Mas se chamarmos de E_1 à matriz elementar correspondente à operação elementar inversa podemos representar os fatos descritos como $E_1 E = I$. Notando

que um raciocínio inteiramente análogo nos conduz a $EE_1 = I$, concluímos que E tem inversa E_1 , que também é elementar. ■

Note que por este teorema e pelos comentários anteriores temos que a inversa de uma matriz elementar E , obtida da identidade I pela *troca de linhas*, é a própria E , visto que $EE = I$.

De um modo geral, se em I forem efetuadas k trocas de linhas, então a matriz resultante pode ser indicada por $E_k \dots E_1$, a qual denotaremos por P e, doravante, chamaremos de **matriz de permutação**.

Lema 1.1 *Se a matriz elementar E é obtida da identidade I pela troca de linhas, então $E = E^t$.*

Prova: Seja I_n a matriz identidade de ordem n e $E = [e_{ij}]$ a matriz obtida de I_n ao se trocar as linhas L_k e L_m . Denotando a matriz identidade na forma $I = [a_{ij}]$ onde $a_{ii} = 1$ e $a_{ij} = 0$ se $i \neq j$ temos que $e_{ij} = a_{ij}$ para $i \neq k, m$ e $j \neq k, m$ uma vez que $e_{km} = e_{mk} = 1$ e $e_{kk} = e_{mm} = 0$. Notando que a transposta da matriz identidade é ela mesma, isto é, $a_{ij} = a_{ji}$, e tendo em vista que $e_{ij} = a_{ij}$ com exceção de e_{km} , e_{mk} , e_{kk} e e_{mm} temos $E^t = [e_{ji}] = E$, pois podemos garantir que $e_{ij} = e_{ji}$ para $e_{ij} = a_{ij}$ quando $i \neq k, m$ e $j \neq k, m$, e $e_{ij} = e_{ji}$ para $i = k, m$ e $j = k, m$, pois $e_{km} = e_{mk} = 1$. ■

No próximo teorema usaremos uma propriedade relativa à transposta de um produto de matrizes. Por esta propriedade temos que se A e B são matrizes de ordens $m \times k$ e $k \times n$, respectivamente, então $(AB)^t = B^t A^t$. Uma demonstração desse fato pode ser encontrada em [1], p. 53.

Teorema 1.2 *Se P é uma matriz de permutação então P é invertível e, além disso, $P^{-1} = P^t$.*

Prova: Se P é a matriz de permutação obtida da identidade I por k trocas de linhas, então existem matrizes elementares $E_1, E_2, \dots, E_k$ tais que $P = E_k \dots E_2 E_1$ e, portanto,

$$E_1^{-1} E_2^{-1} \dots E_k^{-1} P = I \quad \text{e} \quad P E_1^{-1} E_2^{-1} \dots E_k^{-1} = I$$

isto é, P é invertível e sua inversa é dada por

$$P^{-1} = E_1^{-1} E_2^{-1} \dots E_k^{-1}$$

Mas como $E_i^{-1} = E_i$, pois E_i é uma matriz elementar de troca de linhas, então $P^{-1} = E_1 E_2 \dots E_k$. Assim, para provarmos que $P^{-1} = P^t$ basta mostrar que $P^t = E_1 E_2 \dots E_k$. Primeiramente, provaremos que $P^t = E_1^t E_2^t \dots E_k^t$. Por indução, esse resultado é claramente válido para $k = 1$, pois se $P = E_1$ então $P^t = E_1^t$. Supondo $k > 1$ e o resultado válido para $k - 1$, ou seja, para $P = E_{k-1} \dots E_2 E_1$ temos

$$P^t = (E_{k-1} \dots E_2 E_1)^t = E_1^t E_2^t \dots E_{k-1}^t$$

então se $P = E_k E_{k-1} \dots E_2 E_1$ segue que

$$P^t = (E_k E_{k-1} \dots E_2 E_1)^t = (E_{k-1} \dots E_2 E_1)^t E_k^t$$

e assim, pela hipótese de indução, chegamos que

$$P^t = E_1^t E_2^t \dots E_{k-1}^t E_k^t$$

Como pelo Lema 1.1 temos $E_i^t = E_i$, então

$$P^t = E_1 E_2 \dots E_k$$

■

1.2 Matriz Escalonada

Definição 1.2 Dizemos que uma matriz está na forma escalonada por linhas, ou simplesmente na forma escalonada, ou ainda, na forma escada, quando:

- (i) o primeiro elemento não-nulo de cada linha está à esquerda do primeiro elemento não-nulo das linhas seguintes.
- (ii) as linhas nulas, se existirem, encontram-se abaixo das não-nulas.

Assim, dizemos que as matrizes abaixo são escalonadas:

$$\begin{bmatrix} 3 & 5 & 2 & 4 \\ 0 & 0 & 1 & 4 \\ 0 & 0 & 0 & 7 \end{bmatrix} \quad \begin{bmatrix} 0 & 6 & 2 & 4 & 1 \\ 0 & 0 & 1 & 3 & 4 \\ 0 & 0 & 0 & 5 & 7 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$

Ao primeiro elemento não-nulo de cada linha de uma matriz escalonada chamamos de **pivô** ou **líder**. E ao processo de obtenção da forma escalonada chamamos de **escalonamento**. O escalonamento de uma matriz é realizado pela aplicação de operações elementares sobre suas linhas de modo a cumprir as duas condições expostas na definição acima.

Para escalonar uma matriz $A = [a_{ij}]$ devemos executar os seguintes passos:

- (1) Se $a_{11} \neq 0$, faça $L_i - \frac{a_{i1}}{a_{11}}L_1$, para $i \geq 2$. Isso anulará os elementos da primeira coluna abaixo de a_{11} (assim a_{11} é um pivô).
- (2) Se $a_{11} = 0$, faça a troca entre L_1 e outra linha L_i tal que $a_{i1} \neq 0$. E então é só prosseguir como recomendado em (1). Se, porém, todos os elementos da primeira coluna são nulos então segue-se para a segunda coluna e se repete o processo de modo análogo ao já indicado. Caso ainda a segunda coluna seja nula, então busca-se a primeira coluna com algum elemento não-nulo e segue-se semelhantemente ao descrito em (1). Do mesmo modo repete-se esse procedimento para as demais linhas, o que fornecerá a forma escalonada de A .

Doravante usaremos o símbolo $\leftrightarrow$ para indicar uma troca de linhas. Assim, para efetuar a troca entre a i -ésima e a j -ésima linhas de uma matriz usaremos $L_i \leftrightarrow L_j$.

Exemplo: Vejamos a seguir o escalonamento da matriz

$$A = \begin{bmatrix} 0 & 0 & 1 & 1 & 1 \\ 3 & -1 & 2 & -3 & 1 \\ 5 & 2 & -2 & 0 & -1 \\ 1 & -4 & 6 & -6 & 3 \end{bmatrix}$$

Como vemos é preciso primeiramente executar uma troca de linhas em A . Esse e outros procedimentos seguem abaixo com as respectivas operações elementares executadas:

$$\begin{aligned} \begin{bmatrix} 0 & 0 & 1 & 1 & 1 \\ 3 & -1 & 2 & -3 & 1 \\ 5 & 2 & -2 & 0 & -1 \\ 1 & -4 & 6 & -6 & 3 \end{bmatrix} & \xrightarrow{L_1 \leftrightarrow L_2} \begin{bmatrix} 3 & -1 & 2 & -3 & 1 \\ 0 & 0 & 1 & 1 & 1 \\ 5 & 2 & -2 & 0 & -1 \\ 1 & -4 & 6 & -6 & 3 \end{bmatrix} \xrightarrow{\begin{matrix} L_3 - \frac{5}{3}L_1 \\ L_4 - \frac{1}{3}L_1 \end{matrix}} \\ \begin{bmatrix} 3 & -1 & 2 & -3 & 1 \\ 0 & 0 & 1 & 1 & 1 \\ 0 & \frac{11}{3} & -\frac{16}{3} & 5 & -\frac{8}{3} \\ 0 & -\frac{11}{3} & \frac{16}{3} & -5 & \frac{8}{3} \end{bmatrix} & \xrightarrow{L_2 \leftrightarrow L_3} \begin{bmatrix} 3 & -1 & 2 & -3 & 1 \\ 0 & \frac{11}{3} & -\frac{16}{3} & 5 & -\frac{8}{3} \\ 0 & 0 & 1 & 1 & 1 \\ 0 & -\frac{11}{3} & \frac{16}{3} & -5 & \frac{8}{3} \end{bmatrix} \xrightarrow{L_4 + L_2} \end{aligned}$$

$$\begin{bmatrix} 3 & -1 & 2 & -3 & 1 \\ 0 & \frac{11}{3} & -\frac{16}{3} & 5 & -\frac{8}{3} \\ 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$

Uma maneira alternativa de proceder ao escalonamento de uma matriz é tornar cada pivô igual a 1, o que pode ser feito usando a operação elementar do tipo (II). Inclusive alguns autores definem matriz escalonada acrescentando essa propriedade.

Ressaltamos que a forma escalonada de uma matriz não é única nem na definição exposta, nem na definição com pivôs iguais a 1. Porém, há outra definição segundo a qual a matriz na forma escada é única. Nesta definição exige-se que, além de os pivôs serem iguais a 1, cada coluna que contém um pivô tenha todos os seus outros elementos nulos. Tal definição, bem como o teorema garantindo a unicidade podem ser encontrados em [2], p. 37 e 38. No entanto, neste trabalho nos referiremos à forma escada ou forma escalonada de uma matriz conforme definição adotada, salvo menção em contrário.

Como vimos no exemplo acima, o escalonamento da matriz A anulou a última linha. Ao número de linhas não-nulas de uma matriz escalonada denominamos **posto**. Portanto, no exemplo anterior o posto da matriz é 3.

1.3 Matrizes de Transformação de Gauss

Sabemos que o escalonamento de uma matriz $A = [a_{ij}]$, de ordem $m \times n$, consiste em aplicar sobre as linhas de A uma seqüência de operações elementares. Mas conforme observado na página 3 existe uma matriz elementar E tal que EA resulta na mesma matriz obtida ao aplicarmos diretamente uma operação elementar sobre A .

Definindo como **escalonamento do tipo (III)** ao escalonamento de uma matriz que requer apenas operações elementares do tipo (III), definido na p. 2, temos que se A é uma matriz com essa propriedade e, além disso, A não tem colunas nulas, então $a_{11} \neq 0$ e, portanto, um pivô.

Tomando como **primeira etapa do escalonamento do tipo (III)** a colocação de zeros abaixo de a_{11} , então podemos dizer que existem $m-1$ matrizes elementares $E_{21}, E_{31}, \dots, E_{m1}$ tais que $E_{m1} \dots E_{31} E_{21} A = A^{(1)}$, onde $A^{(1)}$ indica a matriz obtida de A após a primeira etapa do escalonamento do tipo (III). De um modo geral denotaremos $A^{(k)}$ a matriz resultante da

k-ésima etapa do escalonamento do tipo (III). Ilustraremos essas etapas para o caso 4×3 , onde $k_{ij} = \frac{a_{ij}^{(j-1)}}{a_{jj}^{(j-1)}}$ e $a_{ij}^{(j-1)}$ indica o novo elemento obtido de A após a $(j-1)$ -ésima etapa do escalonamento do tipo (III). Além disso, tomaremos $a_{ij}^0 = a_{ij}$. Assim,

$$\begin{aligned}
 E_{21}A &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ -k_{21} & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \\
 &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ -\frac{a_{21}}{a_{11}} & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} \\
 E_{31}(E_{21}A) &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ -k_{31} & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \\
 &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ -\frac{a_{31}}{a_{11}} & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} \\
 E_{41}(E_{31}E_{21}A) &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -k_{41} & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \\
 &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -\frac{a_{41}}{a_{11}} & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = A^{(1)}
 \end{aligned}$$

Para a segunda etapa do escalonamento do tipo (III) temos que existem $m - 2$ matrizes elementares $E_{32}, E_{42}, \dots, E_{m2}$ tais que $E_{m2} \dots E_{42} E_{32} A^{(1)} = A^{(2)}$. Vejamos isto no caso exemplificado:

$$E_{32}A^{(1)} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & -k_{32} & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} =$$

$$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & -\frac{a_{32}^{(1)}}{a_{22}^{(1)}} & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix}$$

$$E_{42}(E_{32}A^{(1)}) = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & -k_{42} & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} =$$

$$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & -\frac{a_{42}^{(1)}}{a_{22}^{(1)}} & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} = A^{(2)}$$

Assim, neste exemplo, a terceira e última etapa consiste na multiplicação à esquerda de $A^{(2)}$ por E_{43} , conforme segue:

$$E_{43}A^{(2)} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & -k_{43} & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} =$$

$$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & -\frac{a_{43}^{(2)}}{a_{33}^{(2)}} & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & 0 \end{bmatrix} = A^{(3)}$$

Note que o produto $E_{41}E_{31}E_{21}$ necessário à conclusão da primeira etapa do escalonamento do tipo (III) pode ser substituído por uma única matriz M_1 da forma:

$$M_1 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ -k_{21} & 1 & 0 & 0 \\ -k_{31} & 0 & 1 & 0 \\ -k_{41} & 0 & 0 & 1 \end{bmatrix}$$

A esta matriz chamamos de *matriz de primeira transformação de Gauss*. De modo análogo dizemos que $E_{42}E_{32} = M_2$ é a *matriz da segunda transformação de Gauss*. E assim por diante. Quanto ao número $-k_{ij}$, definido anteriormente, o denominamos **multiplicador**.

Observamos que, de um modo geral, podemos denotar a j -ésima etapa de um escalonamento do tipo (III) por $E_{mj} \dots E_{(j+2)j} E_{(j+1)j} A^{(j-1)} = A^{(j)}$, onde $A^{(0)} = A$ e E_{ij} difere da identidade I_m pela substituição do zero encontrado na i -ésima linha e j -ésima coluna pelo multiplicador $-k_{ij}$. Desse modo,

$$E_{ij} = \begin{bmatrix} 1 & 0 & \dots & 0 & \dots & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 & \dots & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 & \dots & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & -k_{ij} & \dots & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & \dots & 0 & \dots & 1 \end{bmatrix}$$

E portanto, a **matriz da j -ésima transformação de Gauss** pode ser expressa como:

$$M_j = \begin{bmatrix} 1 & 0 & \cdots & 0 & 0 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 & 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 & 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & -k_{(j+1)j} & 1 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & -k_{(j+2)j} & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & -k_{mj} & 0 & 0 & \cdots & 1 \end{bmatrix}$$

Com isso, podemos a seguir definir esse tipo especial de matriz.

Definição 1.3 *Seja a matriz $A = [a_{ij}]$ de ordem $m \times n$ uma matriz que requer apenas escalonamento do tipo (III) para sua redução à forma escada. Dizemos que uma matriz quadrada M_j de ordem m é a **matriz da j-ésima transformação de Gauss** se $M_j A^{(j-1)} = A^{(j)}$.*

Como exemplo consideramos novamente a matriz A de ordem 4×3 , tendo a propriedade de requerer apenas operações elementares do tipo (III) em seu escalonamento. Procederemos agora ao escalonamento usando apenas matrizes de transformação de Gauss. Vejamos:

$$\begin{aligned} M_1 A &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ -k_{21} & 1 & 0 & 0 \\ -k_{31} & 0 & 1 & 0 \\ -k_{41} & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \\ &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ -\frac{a_{21}}{a_{11}} & 1 & 0 & 0 \\ -\frac{a_{31}}{a_{11}} & 0 & 1 & 0 \\ -\frac{a_{41}}{a_{11}} & 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{bmatrix} = \\ &= \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = A^{(1)} \end{aligned}$$

A segunda etapa do escalonamento é então:

$$\begin{aligned}
M_2 A^{(1)} &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & -k_{32} & 1 & 0 \\ 0 & -k_{42} & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = \\
&= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & -\frac{a_{32}^{(1)}}{a_{22}^{(1)}} & 1 & 0 \\ 0 & -\frac{a_{42}^{(1)}}{a_{22}^{(1)}} & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & a_{32}^{(1)} & a_{33}^{(1)} \\ 0 & a_{42}^{(1)} & a_{43}^{(1)} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} = A^{(2)}
\end{aligned}$$

E, por fim, a última etapa é dada por:

$$\begin{aligned}
M_3 A^{(2)} &= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & -k_{43} & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} = \\
&= \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & -\frac{a_{43}^{(2)}}{a_{33}^{(2)}} & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & a_{43}^{(2)} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} \\ 0 & 0 & a_{33}^{(2)} \\ 0 & 0 & 0 \end{bmatrix} = A^{(3)}
\end{aligned}$$

Como vimos, o escalonamento de uma matriz que requer apenas operações elementares do tipo (III) pode ser dada pela multiplicação à esquerda de matrizes de transformação de Gauss. No caso da matriz A , exemplificado acima, temos que seu escalonamento pode ser dado por $M_3 M_2 M_1 A = A^{(3)}$.

1.4 Matrizes Triangulares

Definição 1.4 Uma matriz quadrada $A = [a_{ij}]$ é chamada triangular inferior se todos os elementos acima da diagonal principal são nulos, isto é, se $a_{ij} = 0$ para $i < j$. Analogamente, diz-se que uma matriz quadrada $A = [a_{ij}]$ é triangular superior se todos os elementos abaixo da diagonal principal são nulos, isto é, se $a_{ij} = 0$ para $i > j$.

Assim, nos exemplos seguintes T_3 é triangular inferior e T_4 é triangular superior:

$$T_3 = \begin{bmatrix} 1 & 0 & 0 \\ 2 & -3 & 0 \\ 6 & 10 & 3 \end{bmatrix} \quad T_4 = \begin{bmatrix} 1 & 5 & 3 & 1 \\ 0 & 1 & 4 & 0 \\ 0 & 0 & 1 & 7 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

Dizemos ainda que uma matriz triangular inferior (superior) tem diagonal *unitária* quando todos os elementos da sua diagonal principal são iguais a 1. Então, dos exemplos acima, temos que T_4 é triangular superior com diagonal unitária.

Um exemplo de matriz triangular inferior com diagonal unitária é qualquer matriz M_j de transformação de Gauss.

Teorema 1.3

- (a) *Sejam A e B duas matrizes triangulares inferiores (superiores), então o produto AB é uma matriz triangular inferior (superior). Além disso, se A e B têm diagonal unitária, então AB também tem diagonal unitária.*
- (b) *Uma matriz triangular T é invertível se, e somente se, seus elementos na diagonal principal são todos não-nulos.*
- (c) *A inversa de uma matriz triangular inferior (superior) invertível T é triangular inferior (superior). Se T tem diagonal unitária, então a sua inversa também terá diagonal unitária.*

Nas provas a seguir trataremos apenas dos casos de matrizes triangulares inferiores, pois a demonstração para matrizes triangulares superiores é totalmente análoga.

Prova de (a): Sejam $A = [a_{ij}]$ e $B = [b_{ij}]$ duas matrizes triangulares inferiores de ordem n e $C = [c_{ij}]$ o produto de A por B .

Para provar que C é triangular inferior basta mostrar que $c_{ij} = 0$ se $i < j$.

Pela definição de multiplicação matricial temos que

$$c_{ij} = a_{i1}b_{1j} + a_{i2}b_{2j} + \dots + a_{in}b_{nj}$$

Supondo $i < j$ podemos reescrever a expressão acima na forma

$$c_{ij} = \underbrace{a_{i1}b_{1j} + a_{i2}b_{2j} + \dots + a_{i(j-1)}b_{(j-1)j}}_{(I)} + \underbrace{a_{ij}b_{jj} + \dots + a_{in}b_{nj}}_{(II)}$$

onde em (I) aparecem todos os termos em que o número de linha de B é menor que o número de coluna de B , isto é, nesse grupo todos os elementos b_{kj} são nulos, visto que B é triangular inferior. Do mesmo modo, em (II) se encontram todos os termos em que o número de linha de A é menor que o número de coluna de A , ou seja, nesse grupo todos os elementos a_{ik} são nulos, pois A é triangular inferior. Assim, $c_{ij} = 0$ se $i < j$, e portanto, C é triangular inferior.

Para a prova da segunda parte do teorema basta mostrarmos que $c_{ii} = 1$, para $i = 1, 2, \dots, n$. Observando novamente que

$$c_{ii} = a_{i1}b_{1i} + a_{i2}b_{2i} + \dots + a_{in}b_{ni}$$

então reescrevendo essa expressão acima como

$$c_{ii} = \underbrace{a_{i1}b_{1i} + \dots + a_{i(i-1)}b_{(i-1)i}}_{(I)} + a_{ii}b_{ii} + \underbrace{a_{i(i+1)}b_{(i+1)i} + \dots + a_{in}b_{ni}}_{(II)}$$

temos pela análise feita na prova da primeira parte que os elementos em (I) e (II) são nulos. Deste modo,

$$c_{ii} = a_{ii}b_{ii}$$

Se A e B têm diagonal unitária, então $a_{ii} = b_{ii} = 1$ e, portanto, $c_{ii} = 1$.

Prova de (b): Para esta prova relembremos a definição de determinante. Define-se o determinante de uma matriz quadrada A_n da forma

$$\det(A) = \sum \pm a_{1j_1} a_{2j_2} \dots a_{nj_n}$$

onde o sinal positivo ou negativo é atribuído a uma parcela conforme o número de inversões da permutação $j_1 j_2 \dots j_n$ seja par (positivo) ou ímpar (negativo). Esta soma é estendida às $n!$ permutações da seqüência $(1, 2, \dots, n)$.

Observando, na definição de determinante, que em cada parcela do somatório existe um e somente um elemento de cada linha e cada coluna da matriz, temos que o determinante de uma matriz triangular inferior T terá em todas as parcelas do somatório pelo menos um elemento nulo, com exceção, possivelmente, da parcela constituída apenas com elementos da diagonal. Considerando que uma matriz quadrada é invertível se, e somente se, seu determinante é diferente de zero (para uma prova consulte [1], p. 87), temos T invertível se, e somente se, $\det(T) \neq 0$, o que acontece somente quando os elementos da diagonal de T são todos não-nulos.

Prova de (c): Faremos esta prova por indução e usaremos o produto de matrizes particionadas em blocos. O produto de duas matrizes A e B , particionadas em bloco, é feito semelhantemente ao produto matricial comum, porém, neste caso os “elementos” são submatrizes e o produto AB está definido somente se o produto das submatrizes estiver definido. Seja n a ordem de uma matriz triangular inferior invertível $T = [t_{ij}]$, então para $n = 1$ temos $T = [t_{11}]$. Se T é invertível $t_{11} \neq 0$ e, portanto, a inversa T^{-1} existe e é triangular inferior. Além disso, se $t_{11} = 1$ (isto é, T tem diagonal unitária) teremos também T^{-1} com diagonal unitária.

Considere uma matriz triangular inferior invertível $\mathcal{T}$ de ordem $n + 1$ escrita na forma

$$\mathcal{T}_{(n+1) \times (n+1)} = \left[\begin{array}{c|c} T & 0 \\ \hline v^t & d_{n+1} \end{array} \right]$$

onde $v \in \mathbb{R}^n$ (v^t é v transposto), d_{n+1} é um número real não-nulo e 0 é o vetor nulo de $\mathbb{R}^n$. Do fato de $\mathcal{T}$ ser invertível segue, pela parte (b) deste teorema, que $\mathcal{T}$ tem somente elementos não-nulos na sua diagonal principal, e como os elementos diagonais de T estão na diagonal de $\mathcal{T}$ então T também é invertível. Suponha que a inversa de T é uma matriz triangular inferior T^{-1} , ou seja, a proposição é válida para n . Precisamos mostrar que ela continua sendo válida para $n + 1$. Seja w um vetor de $\mathbb{R}^n$ tal que

$$w^t = -\frac{1}{d_{n+1}} v^t T^{-1}$$

Definindo uma matriz triangular inferior de ordem $n + 1$ da forma

$$\mathcal{W}_{(n+1) \times (n+1)} = \left[\begin{array}{c|c} T^{-1} & 0 \\ \hline w^t & \frac{1}{d_{n+1}} \end{array} \right]$$

temos que $\mathcal{W}$ é a inversa de $\mathcal{T}$ pois

$$\mathcal{W}\mathcal{T} = \left[\begin{array}{c|c} T^{-1}T + 0v^t & T^{-1}0 + 0d_{n+1} \\ \hline w^tT + \frac{1}{d_{n+1}}v^t & w^t0 + \frac{1}{d_{n+1}}d_{n+1} \end{array} \right] = \left[\begin{array}{c|c} I & 0 \\ \hline w^tT + \frac{1}{d_{n+1}}v^t & 1 \end{array} \right]$$

onde I é a matriz identidade de ordem n .

Como

$$w^tT + \frac{1}{d_{n+1}}v^t = -\frac{1}{d_{n+1}}v^tT^{-1}T + \frac{1}{d_{n+1}}v^t = -\frac{1}{d_{n+1}}v^t + \frac{1}{d_{n+1}}v^t = 0^t$$

então

$$\mathcal{W}\mathcal{T} = \left[\begin{array}{c|c} I & 0 \\ \hline 0^t & 1 \end{array} \right]$$

isto é, $\mathcal{T}_{n+1}$ é invertível e sua inversa $\mathcal{W}$ é triangular inferior. Por fim, se T e T^{-1} têm diagonal unitária, então basta tomar $d_{n+1} = 1$ para termos $\mathcal{W} = \mathcal{T}^{-1}$ triangular inferior unitária. Isto conclui a prova. ■

Fatoração de Matrizes

2.1 Fatoração LU

Definição 2.1 Se A é uma matriz quadrada que pode ser fatorada na forma $A = LU$, onde L é triangular inferior e U triangular superior, então dizemos que $A = LU$ é uma fatoração LU de A .

Alternativamente chamamos a fatoração LU por *decomposição LU de A* , ou ainda, por *decomposição triangular de A* . As letras L e U são abreviaturas das palavras *lower* e *upper*, que advêm da língua inglesa e significam *inferior* e *superior*, respectivamente.

Na verdade, alguns autores concebem ainda a decomposição LU de matrizes retangulares $A \in \mathbb{R}^{m \times n}$. Ao final desta seção dedicaremos mais alguns comentários a esse respeito. No entanto, até lá consideraremos apenas a fatoração LU conforme exposto na definição, isto é, quando A é quadrada.

São exemplos de fatoração LU as decomposições abaixo:

$$A = \begin{bmatrix} 3 & -1 & 2 \\ 4 & 3 & 7 \\ -2 & 1 & 5 \end{bmatrix} = \begin{bmatrix} 3 & 0 & 0 \\ 4 & \frac{13}{3} & 0 \\ -2 & \frac{1}{3} & 6 \end{bmatrix} \cdot \begin{bmatrix} 1 & -\frac{1}{3} & \frac{2}{3} \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix}$$

$$B = \begin{bmatrix} 2 & 6 & 4 \\ 4 & 4 & -1 \\ -2 & 2 & 5 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ -1 & -1 & 1 \end{bmatrix} \cdot \begin{bmatrix} 2 & 6 & 4 \\ 0 & -8 & -9 \\ 0 & 0 & 0 \end{bmatrix}$$

Agora é conveniente dizer quando a fatoração $A = LU$ é possível. Uma matriz quadrada só admitirá a decomposição LU se no seu escalonamento não for necessário troca de linhas. De fato, como a redução de uma matriz quadrada A à forma escada consiste em aplicar sobre A uma seqüência de operações elementares, as quais podem ser substituídas pela aplicação de matrizes elementares à esquerda de A , então existem matrizes elementares $E_1, E_2, \dots, E_k$ tais que a forma escada de A é obtida por $E_k \dots E_2 E_1 A$. Considerando a hipótese de que A é quadrada e o escalonamento de A não requer troca de linhas, então sua forma escada é triangular superior. Logo, existe uma matriz triangular superior U resultante da redução à forma escada de A , obtida ao se fazer

$$E_k \dots E_2 E_1 A = U \quad (I)$$

Mas se as matrizes $E_1, E_2, \dots, E_k$ são triangulares inferiores, então as suas inversas são também triangulares inferiores. Assim, existe uma matriz triangular inferior L tal que

$$E_1^{-1} E_2^{-1} \dots E_k^{-1} = L \quad (II)$$

Tomando o primeiro membro de (I) e de (II) podemos expressar A como

$$A = (E_1^{-1} E_2^{-1} \dots E_k^{-1})(E_k \dots E_2 E_1 A)$$

ou, equivalentemente,

$$A = LU$$

Observação: Na prática, pode ser inviável trabalhar com aplicações de matrizes elementares sobre A . Esse processo é mais conveniente no trato teórico do assunto. Assim, se A não requer troca de linhas no seu escalonamento, então o que se faz é escalonar A até a obtenção de U , observar as operações elementares utilizadas e então construir L a partir da identidade, com as operações inversas que as aplicadas no escalonamento de A .

Tomando novamente as fatorações

$$A = \begin{bmatrix} 3 & -1 & 2 \\ 4 & 3 & 7 \\ -2 & 1 & 5 \end{bmatrix} = \begin{bmatrix} 3 & 0 & 0 \\ 4 & \frac{13}{3} & 0 \\ -2 & \frac{1}{3} & 6 \end{bmatrix} \cdot \begin{bmatrix} 1 & -\frac{1}{3} & \frac{2}{3} \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix}$$

$$B = \begin{bmatrix} 2 & 6 & 4 \\ 4 & 4 & -1 \\ -2 & 2 & 5 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ -1 & -1 & 1 \end{bmatrix} \cdot \begin{bmatrix} 2 & 6 & 4 \\ 0 & -8 & -9 \\ 0 & 0 & 0 \end{bmatrix}$$

observamos que na decomposição de A temos U triangular superior com diagonal unitária, enquanto que na de B temos L com diagonal unitária. Essas diferentes fatorações são obtidas por uma pequena modificação no escalonamento, onde num método cada pivô é igual a 1, e noutro, não necessariamente.

O método de fatoração $A = LU$, onde L é triangular inferior e U triangular superior com diagonal unitária é chamado de **método de Crout**, ao passo que se tivermos uma fatoração $A = LU$ em que L é triangular inferior com diagonal unitária e U triangular superior, dizemos que essa fatoração foi feita pelo **método de Doolittle**. Numa ilustração da fatoração LU pelo método de Doolittle pode-se usar matrizes de transformações de Gauss ao invés das matrizes elementares, como fizemos há pouco. Pois como as matrizes de transformações de Gauss anulariam apenas os elementos abaixo da diagonal de $A^{(k)}$, então U não teria necessariamente diagonal unitária ao passo que L teria, visto que esta seria produto das inversas M_k^{-1} que são triangulares inferiores com diagonal unitária. Antes de enunciar o teorema que estabelece as condições de existência e unicidade da fatoração LU vejamos a definição de submatriz principal, pois a usaremos na demonstração desse teorema.

Definição 2.2 Se $A = [a_{ij}]$ é uma matriz de ordem $m \times n$, então chamamos de **submatriz principal** de ordem k da matriz A a matriz formada pelas primeiras k linhas e k colunas de A e a denotamos por $A_{(k,k)}$, onde $k \leq \min\{m, n\}$.

No próximo teorema consideraremos o método de Doolittle na decomposição LU .

Teorema 2.1 (Fatoração LU) Uma matriz $A \in \mathbb{R}^{n \times n}$ tem uma fatoração LU se $\det(A_{(k,k)}) \neq 0$ para $k = 1, 2, \dots, n-1$. Se a fatoração LU existe e A é não-singular, então a fatoração LU é única e $\det(A) = u_{11}u_{22} \dots u_{nn}$, onde $u_{11}, u_{22}, \dots, u_{nn}$ são os elementos da diagonal de U .

Prova: Seja

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\ a_{31} & a_{32} & a_{33} & \dots & a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & a_{n3} & \dots & a_{nn} \end{bmatrix}$$

Temos por hipótese que $\det(A_{(k,k)}) \neq 0$ para $k = 1, 2, \dots, n-1$. Então a submatriz $A_{(1,1)} = [a_{11}]$ tem determinante não-nulo. Consequentemente, $a_{11} \neq 0$ e portanto é um pivô.

Temos também que $\det(A_{(2,2)}) = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} \neq 0$ e como a_{11} é um pivô então a primeira etapa do escalonamento de A pode ser escrita como $M_1 A = A^{(1)}$. Logo

$$A_{(2,2)}^{(1)} = M_{1(2,2)} A_{(2,2)} = \begin{bmatrix} a_{11} & a_{12} \\ 0 & a_{22}^{(1)} \end{bmatrix}$$

Agora usando o fato de que o determinante de um produto de matrizes é o produto dos determinantes dessas matrizes temos que se $\det(A_{(2,2)}) \neq 0$ então $\det(A_{(2,2)}^{(1)}) \neq 0$ o que obriga $a_{22}^{(1)} \neq 0$ e, portanto, o segundo pivô.

De um modo geral, supondo que $k-1$ etapas do escalonamento do tipo (III) foram realizadas, então teremos $M_{k-1} \dots M_1 A = A^{(k-1)}$, onde $a_{kk}^{(k-1)}$ é o k -ésimo pivô. Considerando que

$$\det(M_{k-1} \dots M_1 A) = \det(A^{(k-1)})$$

$$\Leftrightarrow \det(M_{k-1}) \dots \det(M_1) \det(A) = \det(A^{(k-1)})$$

$$\Leftrightarrow 1 \dots 1 \det(A) = \det(A^{(k-1)})$$

$$\Leftrightarrow \det(A) = \det(A^{(k-1)})$$

e portanto $\det(A_{(k,k)}) = \det(A_{(k,k)}^{(k-1)})$, então se $\det(A_{(k,k)}) \neq 0$ segue-se que $\det(A_{(k,k)}^{(k-1)}) \neq 0$ e, por conseguinte, $A_{(k,k)}^{(k-1)}$ é triangular superior com k pivôs. Assim a hipótese $\det(A_{(k,k)}) \neq 0$ para $k = 1, 2, \dots, n-1$ garante que a matriz $A_{(n-1,n-1)}^{(n-2)}$ é triangular superior. Observamos que o fato de k variar apenas até $n-1$ é o suficiente para que a forma escalonada de A seja triangular superior porque após a $(n-1)$ -ésima etapa do escalonamento do tipo (III) temos que $A^{(n-1)}$ possui na n -ésima linha somente elementos zeros, com exceção, possivelmente, de $a_{nn}^{(n-1)}$. Isto é, a forma escalonada de A é uma matriz triangular superior U . Por

fim, existem $M_1^{-1}, \dots, M_{n-1}^{-1}$ (que também são matrizes de transformação de Gauss, tais) tais que $M_1^{-1} \dots M_{n-1}^{-1} = L$ e, portanto, $A = LU$. Logo, a hipótese $\det(A_{(k,k)}) \neq 0$ para $k = 1, 2, \dots, n-1$ é suficiente para que uma matriz A tenha a decomposição LU .

Para provar a segunda parte do teorema observamos inicialmente que se existe a decomposição $A = LU$, então podemos escrever $M_{n-1} \dots M_2 M_1 A = U$. Acrescentando a hipótese de que A é não singular concluímos que U deve ser também não singular. Portanto, supondo que $A = L_1 U_1$ e $A = L_2 U_2$ sejam duas decomposições LU de uma matriz não singular A , então podemos escrever

$$\begin{aligned} L_1 U_1 &= L_2 U_2 \\ \Rightarrow L_1 U_1 U_1^{-1} &= L_2 U_2 U_1^{-1} \\ \Rightarrow L_1 &= L_2 U_2 U_1^{-1} \\ \Rightarrow L_2^{-1} L_1 &= L_2^{-1} L_2 U_2 U_1^{-1} \\ \Rightarrow L_2^{-1} L_1 &= U_2 U_1^{-1} \end{aligned}$$

Como $L_2^{-1} L_1$ é triangular inferior com diagonal unitária e $U_2 U_1^{-1}$ é triangular superior, então existe uma matriz diagonal D tal que

$$L_2^{-1} L_1 = D = U_2 U_1^{-1}$$

Da observação de que $L_2^{-1} L_1$ é triangular inferior com diagonal unitária temos que D é igual à identidade I , isto é,

$$L_2^{-1} L_1 = I = U_2 U_1^{-1}$$

Mas isso significa que L_2^{-1} é a inversa de L_1 e U_1^{-1} a inversa de U_2 , ou seja, $L_2 = L_1$ e $U_2 = U_1$, ou ainda, a fatoração LU é única, respeitadas as hipóteses.

Por fim, se $A = LU$, então $\det(A) = \det(LU) = \det(L)\det(U) = \det(U) = u_{11}u_{22} \dots u_{nn}$. Isto conclui a prova. ■

No caso em que no teorema acima é tomado o método de Crout para a decomposição LU temos que substituir no enunciado $\det(A) = u_{11}u_{22} \dots u_{nn}$ por $\det(A) = \ell_{11}\ell_{22} \dots \ell_{nn}$ e na demonstração as matrizes de transformações de Gauss por matrizes elementares.

Obtendo uma Fatoração LU

Considere a matriz A abaixo, cujos determinantes das submatrizes são não-nulos:

$$A = \begin{bmatrix} 2 & 3 & -1 \\ 1 & -2 & 2 \\ 3 & 1 & -1 \end{bmatrix}$$

Primeiramente, aplicamos o escalonamento em A e obtemos U :

$$A = \begin{bmatrix} 2 & 3 & -1 \\ \boxed{1} & -2 & 2 \\ \boxed{3} & 1 & -1 \end{bmatrix} \begin{array}{l} L_2 - L_1/2 \\ L_3 - 3L_1/2 \end{array} \rightarrow$$

$$\begin{bmatrix} 2 & 3 & -1 \\ 0 & -\frac{7}{2} & \frac{5}{2} \\ 0 & \boxed{-\frac{7}{2}} & \frac{1}{2} \end{bmatrix} L_3 - L_2 \rightarrow \begin{bmatrix} 2 & 3 & -1 \\ 0 & -\frac{7}{2} & \frac{5}{2} \\ 0 & 0 & -2 \end{bmatrix} = U$$

Sabemos que a matriz $L = E_1^{-1}E_2^{-1} \dots E_k^{-1}$, e como vimos na observação anterior, para construir L podemos simplesmente tomar as operações elementares inversas das que utilizamos no escalonamento de A . Desse modo, segue que

$$L = \begin{bmatrix} 1 & 0 & 0 \\ \frac{1}{2} & 1 & 0 \\ \frac{3}{2} & 1 & 1 \end{bmatrix}$$

Ou seja, A pode ser fatorada na forma $A = LU$:

$$\begin{bmatrix} 2 & 3 & -1 \\ 1 & -2 & 2 \\ 3 & 1 & -1 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ \frac{1}{2} & 1 & 0 \\ \frac{3}{2} & 1 & 1 \end{bmatrix} \cdot \begin{bmatrix} 2 & 3 & -1 \\ 0 & -\frac{7}{2} & \frac{5}{2} \\ 0 & 0 & -2 \end{bmatrix}$$

Dada uma matriz quadrada A pode ser inviável usar o Teorema 2.1 para determinar se A tem decomposição LU . No entanto, se no escalonamento de A for necessário trocas de linhas, podemos tomar uma **matriz de permutação P** obtida da identidade pela aplicação das trocas de linhas necessárias no escalonamento de A . Desse modo, o produto PA é uma

matriz que não requer permutação de linhas em seu escalonamento e, então, poderemos escrever $PA = LU$. Como pelo Teorema 1.2 temos que P é invertível e que sua inversa é dada por $P^{-1} = P^t$, então podemos escrever $A = P^{-1}LU$, ou ainda, $A = P^tLU$.

2.1.1 Contagem do número de operações para obter a fatoração $A=LU$

Detalharemos a seguir a contagem das operações requeridas na obtenção da fatoração $A = LU$ pelo método de Doolittle. O número de operações requeridas usando o método de Crout é o mesmo.

Seja

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\ a_{31} & a_{32} & a_{33} & \dots & a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & a_{n3} & \dots & a_{nn} \end{bmatrix}$$

Para facilitar a nossa notação durante as etapas de contagem das operações colocaremos o símbolo $\times$ representando quais elementos sofreram alteração na última operação realizada. Observamos que a colocação de um elemento zero abaixo de um pivô não é contado como uma operação, pois esse procedimento não exige cálculo.

Escalonando $A_{n \times n}$ para obter U :

Passo 1: Colocação de zeros abaixo do pivô a_{11} .

Operação realizada: $L_k - \frac{a_{k1}}{a_{11}}L_1$, para $k = 2, 3, \dots, n$.

$$A^{(1)} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ 0 & \times & \times & \dots & \times \\ 0 & \times & \times & \dots & \times \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & \times & \times & \dots & \times \\ 0 & \times & \times & \dots & \times \end{bmatrix}$$

Número de operações realizadas:

Adições/Subtrações: $n - 1$ por linha em $n - 1$ linhas. Total: $(n - 1)(n - 1) = (n - 1)^2$.

Multiplicações/Divisões: $2(n-1)$ por linha em $n-1$ linhas. Total: $2(n-1)(n-1) = 2(n-1)^2$.

Passo 2: Colocação de zeros abaixo do pivô $a_{22}^{(1)}$.

Operação realizada: $L_k - \frac{a_{k2}^{(1)}}{a_{22}^{(1)}} L_2$, para $k = 3, 4, \dots, n$.

$$A^{(2)} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & \dots & a_{2n}^{(1)} \\ 0 & 0 & \times & \dots & \times \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \times & \dots & \times \\ 0 & 0 & \times & \dots & \times \end{bmatrix}$$

Número de operações realizadas:

Adições/Subtrações: $n-2$ por linha em $n-2$ linhas. Total: $(n-2)(n-2) = (n-2)^2$.

Multiplicações/Divisões: $2(n-2)$ por linha em $n-2$ linhas. Total: $2(n-2)(n-2) = 2(n-2)^2$.

Passo n-2: Colocação de zeros abaixo do pivô $a_{(n-2)(n-2)}^{(n-3)}$.

Operação realizada: $L_k - \frac{a_{k(n-2)}^{(n-3)}}{a_{(n-2)(n-2)}^{(n-3)}} L_{(n-2)}$, para $k = n-1, n$.

$$A^{(n-2)} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1(n-2)} & a_{1(n-1)} & a_{1n} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & \dots & a_{2(n-2)}^{(1)} & a_{2(n-1)}^{(1)} & a_{2n}^{(1)} \\ 0 & 0 & a_{33}^{(2)} & \dots & a_{3(n-2)}^{(2)} & a_{3(n-1)}^{(2)} & a_{3n}^{(2)} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & a_{(n-2)(n-2)}^{(n-3)} & a_{(n-2)(n-1)}^{(n-3)} & a_{(n-2)n}^{(n-3)} \\ 0 & 0 & 0 & \dots & 0 & \times & \times \\ 0 & 0 & 0 & \dots & 0 & \times & \times \end{bmatrix}$$

Número de operações realizadas:

Adições/Subtrações: 2 por linha em 2 linhas. Total: 2^2 .

Multiplicações/Divisões: $2 \cdot 2$ por linha em 2 linhas. Total: $2(2)^2$

Passo n-1: Colocação de zeros abaixo do pivô $a_{(n-1)(n-1)}^{(n-2)}$.

Operação realizada: $L_n - \frac{a_{n(n-1)}^{(n-2)}}{a_{(n-1)(n-1)}^{(n-2)}} L_{(n-1)}$.

$$A^{(n-1)} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1(n-2)} & a_{1(n-1)} & a_{1n} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & \dots & a_{2(n-2)}^{(1)} & a_{2(n-1)}^{(1)} & a_{2n}^{(1)} \\ 0 & 0 & a_{33}^{(2)} & \dots & a_{3(n-2)}^{(2)} & a_{3(n-1)}^{(2)} & a_{3n}^{(2)} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & a_{(n-2)(n-2)}^{(n-3)} & a_{(n-2)(n-1)}^{(n-3)} & a_{(n-2)n}^{(n-3)} \\ 0 & 0 & 0 & \dots & 0 & a_{(n-1)(n-1)}^{(n-2)} & a_{(n-1)n}^{(n-2)} \\ 0 & 0 & 0 & \dots & 0 & 0 & \times \end{bmatrix}$$

Número de operações realizadas:

Adições/Subtrações: 1 por linha em 1 linha. Total: 1^2 .

Multiplicações/Divisões: $2 \cdot 1$ por linha em 1 linha. Total: $2(1)^2$

Obtemos então a forma escalonada de A após a execução de $n - 1$ passos:

$$U = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1(n-2)} & a_{1(n-1)} & a_{1n} \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & \dots & a_{2(n-2)}^{(1)} & a_{2(n-1)}^{(1)} & a_{2n}^{(1)} \\ 0 & 0 & a_{33}^{(2)} & \dots & a_{3(n-2)}^{(2)} & a_{3(n-1)}^{(2)} & a_{3n}^{(2)} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & a_{(n-2)(n-2)}^{(n-3)} & a_{(n-2)(n-1)}^{(n-3)} & a_{(n-2)n}^{(n-3)} \\ 0 & 0 & 0 & \dots & 0 & a_{(n-1)(n-1)}^{(n-2)} & a_{(n-1)n}^{(n-2)} \\ 0 & 0 & 0 & \dots & 0 & 0 & a_{nn}^{(n-1)} \end{bmatrix}$$

Total de operações para reduzir A a U :

Adições/Subtrações:

$$(n-1)^2 + (n-2)^2 + \dots + 2^2 + 1^2 = \frac{(n-1)n(2(n-1)+1)}{6} = \frac{(n-1)n(2n-1)}{6}$$

Multiplicações/Divisões:

$$2(n-1)^2 + 2(n-2)^2 + \dots + 2(2)^2 + 2(1)^2 = 2 \frac{(n-1)n(2(n-1)+1)}{6} = \frac{(n-1)n(2n-1)}{3}$$

Para uma dedução das igualdades acima consulte [1], p. 326, exercícios 7 e 8, onde são apresentados os passos necessários à obtenção da fórmula $1^2 + 2^2 + 3^2 + \dots + (n-1)^2 = \frac{(n-1)n(2n-1)}{6}$.

A matriz triangular inferior com diagonal unitária L é determinada diretamente, apenas por armazenamento das operações inversas às aplicadas a A para obter U . Assim, a contagem indicada acima já resume o total de operações executadas na fatoração $A = LU$.

A matriz L pode ser escrita genericamente na forma:

$$L = \begin{bmatrix} 1 & 0 & 0 & \dots & 0 & 0 \\ \frac{a_{21}}{a_{11}} & 1 & 0 & \dots & 0 & 0 \\ \frac{a_{31}}{a_{11}} & \frac{a_{32}^{(1)}}{a_{22}^{(1)}} & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \frac{a_{(n-1)1}}{a_{11}} & \frac{a_{(n-1)2}^{(1)}}{a_{22}^{(1)}} & \frac{a_{(n-1)3}^{(2)}}{a_{33}^{(2)}} & \dots & 1 & 0 \\ \frac{a_{n1}}{a_{11}} & \frac{a_{n2}^{(1)}}{a_{22}^{(1)}} & \frac{a_{n3}^{(2)}}{a_{33}^{(2)}} & \dots & \frac{a_{n(n-1)}^{(n-2)}}{a_{(n-1)(n-1)}^{(n-2)}} & 1 \end{bmatrix}$$

2.1.2 Fatoração LU de uma Matriz Retangular

No início deste capítulo falamos da existência de uma fatoração LU para matrizes retangulares $A \in \mathbb{R}^{m \times n}$. É claro que L e U não podem ser quadradas ao mesmo tempo, pois se $A_{m \times n} = L_{m \times k} U_{k \times n}$ e $m \neq n$ o máximo que podemos ter é uma das matrizes L ou U quadrada. Se $k = m$, então L é triangular inferior (portanto quadrada); se $k = n$ então U é triangular superior (portanto quadrada). Podemos ainda ter L e U não-quadradas, quando $k \neq m, n$. No entanto, se definirmos como **matriz com forma triangular** inferior (superior) à matriz $A = [a_{ij}]$ tal que $a_{ij} = 0$ se $i < j$ ($i > j$), então L e U têm forma triangular, pois são obtidas, respectivamente, do escalonamento de A e das operações elementares inversas das aplicadas no escalonamento.

Vejamos a seguir uma fatoração LU de duas matrizes retangulares, onde os casos $m > n$ e $m < n$ estão ilustrados.

$$A_{3 \times 2} = \begin{bmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 2 & 1 \\ 3 & 2 \end{bmatrix} \cdot \begin{bmatrix} 1 & 4 \\ 0 & -3 \end{bmatrix} = L_{3 \times 2} \cdot U_{2 \times 2}$$

$$B_{2 \times 3} = \begin{bmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 2 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 3 & 5 \\ 0 & -2 & -4 \end{bmatrix} = L_{2 \times 2} \cdot U_{2 \times 3}$$

A condição para que exista uma fatoração LU de uma matriz $A \in \mathbb{R}^{m \times n}$ é que $A_{(k,k)}$ seja não-singular para $k = 1, 2, \dots, \min\{m, n\}$. Nós omitimos a demonstração desse resultado uma vez que ela é totalmente análoga à encontrada na prova da primeira parte do Teorema 2.1.

A seguir apresentamos outra forma de fatorar as matrizes exemplificadas $A_{3 \times 2}$ e $B_{2 \times 3}$, a qual é menos econômica, pois nela acrescentamos colunas e linhas só de zeros.

$$A_{3 \times 2} = \begin{bmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 3 & 2 & 0 \end{bmatrix} \cdot \begin{bmatrix} 1 & 4 \\ 0 & -3 \\ 0 & 0 \end{bmatrix} = L_{3 \times 3}^1 \cdot U_{3 \times 2}^1$$

$$B_{2 \times 3} = \begin{bmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \end{bmatrix} \cdot \begin{bmatrix} 1 & 3 & 5 \\ 0 & -2 & -4 \\ 0 & 0 & 0 \end{bmatrix} = L_{2 \times 3}^2 \cdot U_{3 \times 3}^2$$

Desse modo, queremos dizer que a unicidade não é garantida para a decomposição LU de matrizes retangulares.

2.2 Fatoração de Cholesky

Definição 2.3 Chamamos de **fatoração de Cholesky** de uma matriz A à fatoração $A = GG^t$, onde G é uma matriz triangular inferior com elementos positivos na diagonal principal.

Exemplificamos a seguir dois casos de fatoração de Cholesky:

$$A = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix} = \begin{bmatrix} \sqrt{2} & 0 & 0 \\ -\frac{1}{\sqrt{2}} & \sqrt{\frac{3}{2}} & 0 \\ 0 & -\sqrt{\frac{2}{3}} & \frac{2}{\sqrt{3}} \end{bmatrix} \cdot \begin{bmatrix} \sqrt{2} & -\frac{1}{\sqrt{2}} & 0 \\ 0 & \sqrt{\frac{3}{2}} & -\sqrt{\frac{2}{3}} \\ 0 & 0 & \frac{2}{\sqrt{3}} \end{bmatrix} = GG^t$$

$$B = \begin{bmatrix} 4 & -2 & 4 & 10 \\ -2 & 10 & 1 & -2 \\ 4 & 1 & 6 & 13 \\ 10 & -2 & 13 & 31 \end{bmatrix} = \begin{bmatrix} 2 & 0 & 0 & 0 \\ -1 & 3 & 0 & 0 \\ 2 & 1 & 1 & 0 \\ 5 & 1 & 2 & 1 \end{bmatrix} \cdot \begin{bmatrix} 2 & -1 & 2 & 5 \\ 0 & 3 & 1 & 1 \\ 0 & 0 & 1 & 2 \\ 0 & 0 & 0 & 1 \end{bmatrix} = G_1 G_1^t$$

Desde já pode ser observado que as matrizes A e B são simétricas e como veremos adiante isso é uma das condições necessárias para que haja a decomposição de Cholesky. No entanto, o caminho que tomamos nos mostrará a existência de duas outras fatorações antes da de Cholesky. A primeira delas é a fatoração LDM^t , onde L e M são matrizes triangulares inferiores com diagonal unitária e D é uma matriz diagonal. Tal fatoração pode ser obtida da fatoração LU de uma matriz não-singular A . Assim, se no Teorema 2.1, p. 21, tomarmos $k = 1, 2, \dots, n$ temos garantida a decomposição LDM^t . A segunda fatoração, LDL^t , será um caso particular da fatoração LDM^t para A simétrica. Desse modo, dada uma decomposição LU de uma matriz A simétrica não-singular, é possível construir a fatoração LDL^t .

Teorema 2.2 *Se todas as submatrizes principais de $A \in \mathbb{R}^{n \times n}$ são não-singulares então existem matrizes triangulares inferiores com diagonal unitária L e M únicas e uma única matriz diagonal D tal que $A = LDM^t$.*

Prova: Pelo Teorema 2.1 sabemos que A tem uma fatoração LU . Considerando a matriz diagonal D tal que $d_{ii} = u_{ii}$ para $i = 1, \dots, n$ temos que D é invertível e então podemos escrever $A = LU = LDD^{-1}U$. Fazendo $M^t = D^{-1}U$ temos a decomposição $A = LDM^t$, onde $M^t = D^{-1}U$ é uma matriz triangular superior com diagonal unitária, pois a inversa de uma matriz diagonal D cuja diagonal é u_{ii} é outra matriz diagonal D^{-1} com elementos diagonais $\frac{1}{u_{ii}}$. A unicidade segue da fatoração LU já demonstrada no Teorema 2.1. ■

Teorema 2.3 *Se todas as submatrizes principais de uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ são não-singulares então existe uma única matriz triangular inferior com diagonal unitária L e uma única matriz diagonal D tal que $A = LDL^t$.*

Prova: Pelo Teorema 2.2 a matriz A admite a decomposição LDM^t . Considerando que M é triangular inferior com diagonal unitária, então existe M^{-1} e assim valem as igualdades a seguir:

$$\begin{aligned}
A &= LDM^t \\
AM^{-t} &= LD \\
M^{-1}AM^{-t} &= M^{-1}LD
\end{aligned}$$

Do primeiro membro da última igualdade temos uma matriz simétrica, pois

$$(M^{-1}AM^{-t})^t = (AM^{-t})^t(M^{-1})^t = (M^{-t})^t A^t M^{-t} = M^{-1}AM^{-t}$$

Do segundo membro de $M^{-1}AM^{-t} = M^{-1}LD$ temos uma matriz triangular inferior, visto que o produto de triangulares inferiores é triangular inferior (D também pode ser interpretada como uma matriz triangular inferior). Mas uma matriz ao mesmo tempo simétrica e triangular só pode ser diagonal. Como D é não-singular temos que $M^{-1}L$ deve ser também diagonal. Por outro lado, $M^{-1}L$ é triangular inferior com diagonal unitária e, portanto, $M^{-1}L = I$, isto é, $M = L$, o que estabelece a decomposição $A = LDL^t$. A unicidade segue da fatoração LDM^t . ■

Agora passamos à definição de matriz definida positiva, visto que outra condição necessária à existência da fatoração de Cholesky de uma matriz é que esta seja definida positiva. Vários resultados semelhantes aos anteriores serão enunciados e provados para o caso de uma matriz definida positiva, culminando com o teorema que estabelece a fatoração de Cholesky.

Definição 2.4 Chamamos de matriz **definida positiva** a uma matriz A tal que $x^t Ax > 0$ para todo $x \neq 0$.

Ressaltamos que alguns autores assumem na definição de matriz definida positiva a hipótese de A ser simétrica, apesar de esta não ser uma condição necessária para que $x^t Ax > 0$. Porém, muitos resultados importantes envolvendo matrizes definidas positivas exigem simetria.

Existe ainda uma outra maneira equivalente de apresentar a definição de uma matriz definida positiva. Nessa definição dizemos que uma matriz $A = [a_{ij}] \in \mathbb{R}^{n \times n}$ é dita definida positiva se para qualquer vetor não-nulo $x \in \mathbb{R}^n$ tivermos $\sum_{i,j=1}^n a_{ij}x_i x_j > 0$.

Antes de enunciarmos o lema do próximo teorema, veremos um resultado que relaciona a solução de um sistema linear homogêneo à invertibilidade de uma matriz. Dedicaremos maior destaque aos sistemas lineares no Capítulo 3.

Teorema 2.4 *Uma matriz A é invertível se, e somente se, o sistema linear homogêneo $Ax = 0$ tem somente a solução trivial.*

Prova: Suponha que A seja uma matriz de ordem $n \times n$ invertível e que $\bar{x}$ seja uma solução qualquer de $Ax = 0$, isto é, $A\bar{x} = 0$. Multiplicando ambos os membros de $Ax = 0$ por A^{-1} resulta $A^{-1}(A\bar{x}) = A^{-1}0$, ou seja, $(A^{-1}A)\bar{x} = 0$, ou ainda, $\bar{x} = 0$, implicando que se A é invertível então $Ax = 0$ tem somente a solução trivial.

Suponha agora que $Ax = 0$ tem somente a solução trivial. Escrevendo $Ax = 0$ como uma combinação linear dos vetores-coluna de A temos $x_1A_1 + \dots + x_nA_n = 0$, onde x_i indica uma componente do vetor x e A_i um vetor-coluna de A , para qualquer $i = 1, 2, \dots, n$. Como a solução $x_i = 0$ é a única solução para a equação $x_1A_1 + \dots + x_nA_n = 0$ então os vetores-colunas de A são linearmente independentes. Denotando a forma escada de A como $Esc(A)$ temos que existem matrizes elementares $E_1, \dots, E_k$ tais que $E_k \dots E_1 A = Esc(A)$ e, dessa forma,

$$\begin{aligned} E_1^{-1} \dots E_k^{-1} E_k \dots E_1 A &= E_1^{-1} \dots E_k^{-1} Esc(A) \Leftrightarrow \\ A &= E_1^{-1} \dots E_k^{-1} Esc(A) \Leftrightarrow \\ \det(A) &= \det(E_1^{-1} \dots E_k^{-1} Esc(A)) \Leftrightarrow \\ \det(A) &= \det(E_1^{-1}) \dots \det(E_k^{-1}) \det(Esc(A)) \end{aligned}$$

Conforme teorema 1.1 toda matriz elementar é invertível e como A é uma matriz quadrada com vetores-coluna linearmente independentes então a forma escada de A é uma matriz triangular superior invertível, implicando que no segundo membro da última equação temos todos os determinantes diferentes de zero, de modo que devemos ter $\det(A) \neq 0$ e, portanto, A invertível. ■

Lema 2.1 *Se $A \in \mathbb{R}^{n \times n}$ é uma matriz definida positiva então A é invertível.*

Prova: Supondo que A seja não-invertível, então existe um vetor não-nulo $x \in \mathbb{R}^n$ tal que $Ax = 0$ e, portanto, $x^t Ax = 0$, o que é uma contradição. ■

Teorema 2.5 *Se uma matriz $A \in \mathbb{R}^{n \times n}$ é definida positiva, então todas as suas submatrizes principais $A_{(k,k)}$, $1 \leq k \leq n$, são também definidas positivas.*

Prova: Considere o vetor não-nulo $x = (x_1, \dots, x_k)$. Tomando um vetor $y = (y_1, \dots, y_n)$ tal que $y_i = x_i$ para $i = 1, \dots, k$ e $y_i = 0$ para $i = k+1, \dots, n$ temos que y é um vetor não-nulo de $\mathbb{R}^n$. Portanto,

$$\sum_{i,j=1}^k a_{ij}x_ix_j = \sum_{i,j=1}^n a_{ij}y_iy_j > 0.$$

■

Corolário 2.1 Se A é uma matriz definida positiva, então A admite a decomposição LU .

Prova: Como todas as submatrizes principais de A são também definidas positivas, então pelo Lema 2.1 seus determinantes são não-nulos, o que garante a decomposição LU , conforme Teorema 2.1. ■

Teorema 2.6 Se $A \in \mathbb{R}^{n \times n}$ é definida positiva, então todos os seus elementos diagonais são positivos.

Prova: Para todo $i \in \{1, 2, \dots, n\}$ seja $x = (x_1, \dots, x_n)$, definido da forma $x_i = 1$ e $x_j = 0$ para $j \neq i$, isto é, x é um vetor não-nulo. Assim,

$$a_{ii} = \sum_{i,j=1}^n a_{ij}x_ix_j = x^t Ax > 0.$$

■

Teorema 2.7 Se $A \in \mathbb{R}^{n \times n}$ é definida positiva e $X \in \mathbb{R}^{n \times k}$ tem posto k , com $k \leq n$, então $B = X^t AX \in \mathbb{R}^{k \times k}$ é também definida positiva.

Prova: Seja $z \in \mathbb{R}^k$ tal que $0 \geq z^t Bz = z^t (X^t AX)z = (Xz)^t A(Xz)$. Suponha que o vetor $(Xz) \in \mathbb{R}^n$ seja nulo, isto é, $Xz = 0$. Como X tem posto e número de incógnitas iguais a k , então o sistema homogêneo $Xz = 0$ tem solução única, a saber $z = 0$. Assim, não existe um vetor não-nulo $z \in \mathbb{R}^k$ tal que $z^t Bz = 0$.

Para provar que não existe $z \in \mathbb{R}^k$ tal que $z^t Bz < 0$ suponhamos que z seja não-nulo, pois para z nulo $z^t Bz = 0$. Definindo $v = Xz$ temos $z^t Bz = (Xz)^t A(Xz) = v^t Av > 0$, pois A é definida positiva. Desse modo, $z^t Bz \leq 0$ não ocorre se z é não-nulo, e, portanto, B é definida positiva. ■

Teorema 2.8 Se A é uma matriz definida positiva, então A admite a decomposição LDM^t e D tem elementos diagonais positivos.

Prova: Pelo Teorema 2.5 e Lema 2.1 temos que as submatrizes $A_{(k,k)}$ são não-singulares para $k = 1, 2, \dots, n$, então, pelo Teorema 2.2, temos garantida a existência da fatoração $A =$

LDM^t . Tomando no Teorema 2.7 $X = L^{-t}$ segue que $B = X^tAX = (L^{-t})^tAL^{-t} = L^{-1}AL^{-t}$ é definida positiva. Mas, por outro lado, se $A = LDM^t$, então $B = L^{-1}LDM^tL^{-t} = DM^tL^{-t}$. Considerando que M^tL^{-t} é triangular superior com diagonal unitária, e como $B = DM^tL^{-t}$ então B e D têm a mesma diagonal. Mas B é definida positiva e, conforme Teorema 2.6, tem elementos diagonais positivos. Logo, se B e D têm diagonais iguais, então D tem elementos diagonais positivos. ■

Corolário 2.2 *Se A é uma matriz simétrica definida positiva, então A admite a decomposição LDL^t e D tem elementos diagonais positivos.*

Prova: Pelo Teorema 2.5 temos que A tem todas as submatrizes principais não-singulares, então pelo Teorema 2.3 a fatoração $A = LDL^t$ existe. A garantia de que os elementos diagonais de D são positivos segue do Corolário 2.8, substituindo-se apenas M por L . ■

Teorema 2.9 (Fatoração de Cholesky) *Uma matriz $A \in \mathbb{R}^{n \times n}$ é simétrica definida positiva se, e somente se, existe uma única matriz triangular inferior $G \in \mathbb{R}^{n \times n}$ com elementos diagonais positivos tal que $A = GG^t$.*

Prova: Para provar que se $A \in \mathbb{R}^{n \times n}$ é simétrica definida positiva então existe uma única matriz triangular inferior $G \in \mathbb{R}^{n \times n}$ com elementos diagonais positivos tal que $A = GG^t$ considere o Corolário 2.2 que garante a existência da fatoração $A = LDL^t$, onde L é triangular inferior com diagonal unitária e D é diagonal com elementos diagonais positivos. Denotando os elementos diagonais de D por $d_1, \dots, d_n$ e definindo uma matriz diagonal $S = [s_{ij}]$ de modo que $s_{ii} = \sqrt{d_i}$, para $i = 1, \dots, n$, podemos escrever $D = SS$. Como toda matriz diagonal é simétrica, então $S = S^t$ e, portanto, $D = SS^t$. Desse modo, $A = LDL^t = LSS^tL^t = LS(LS)^t$. Definindo a matriz $G = LS$ temos $A = GG^t$, onde G é uma matriz triangular inferior com elementos diagonais positivos. A unicidade segue da fatoração LDL^t .

Por outro lado, se para $A \in \mathbb{R}^{n \times n}$ existe uma única matriz G com elementos diagonais positivos tal que $A = GG^t$, então $A^t = (GG^t)^t = (G^t)^tG^t = GG^t = A$, isto é, A é simétrica. Tomando $x \neq 0$ temos $x^tAx = x^t(GG^t)x = (x^tG)(G^tx) = (G^tx)^t(G^tx)$. Fazendo $G^tx = y$ segue que $x^tAx = y^ty$. Como $x \neq 0$ e G é não-singular então $y \neq 0$ e, portanto, $x^tAx > 0$, ou seja, A é definida positiva. ■

Os dois teoremas seguintes fornecem, em particular, uma maneira de determinarmos se uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ é ou não definida positiva.

Teorema 2.10 *Uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ é definida positiva se, e somente se, seus pivôs são todos positivos.*

Prova: Supondo $A \in \mathbb{R}^{n \times n}$ simétrica definida positiva então pelo Corolário 2.1 e pelo Teorema 2.9 segue que $A = LU = GG^t$. Definindo uma matriz diagonal D tal que $d_{ii} = g_{ii}$ temos $A = GD^{-1}DG$, onde GD^{-1} é triangular inferior com diagonal unitária e DG^t é triangular superior. Mas pela unicidade da decomposição LU devemos ter $L = GD^{-1}$ e $U = DG^t$. Como D e G possuem elementos diagonais positivos ($d_{ii} = g_{ii}$ e $g_{ii} > 0$), então U tem elementos diagonais positivos. Mas U é obtida de A por escalonamento, resultando que os pivôs u_{ii} de A são todos positivos.

Por outro lado, se $A \in \mathbb{R}^{n \times n}$ é simétrica com pivôs positivos as submatrizes principais de A são não-singulares e, então, A pode ser fatorada na forma $A = LDL^t$. Denotando $D = [d_{ij}]$ e $S = [s_{ij}]$ como a matriz diagonal tal que $s_{ii} = \sqrt{d_{ii}}$ e observando que uma matriz diagonal é simétrica, então $S = S^t$ e assim

$$x^t Ax = x^t LDL^t x = (L^t x)^t S S (L^t x) = (L^t x)^t S^t S (L^t x) = (S L^t x)^t (S L^t x).$$

Definindo ainda um vetor-coluna $y = (y_1, y_2, \dots, y_n) = S L^t x$ temos que

$$x^t Ax = y^t y = \sum_{i=1}^n y_i^2 \geq 0$$

Precisamos mostrar que essa soma não é zero. Para tanto, consideremos que $S L^t$ é triangular inferior com diagonal positiva, portanto invertível, então o sistema $S L^t x = y$ só terá solução $y = 0$ se $x = 0$. Mas x é por hipótese um vetor não-nulo, obrigando y também a ser não-nulo, garantindo que A é definida positiva. ■

Teorema 2.11 *Uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ é definida positiva se, e somente se, os determinantes das submatrizes principais são todos positivos.*

Prova: Para provar que se uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ é definida positiva então os determinantes de suas submatrizes principais são todos positivos, consideremos o Teorema 2.10 que garante serem positivos os pivôs de A . Como A admite a decomposição LU temos $\det(A) = \det(LU)$ e, por sua vez, $\det(A_{(k,k)}) = \det(L_{(k,k)}) \cdot \det(U_{(k,k)}) = u_{11} u_{22} \dots u_{kk}$. Mas u_{ii} são os pivôs de A e, como sabemos, são positivos, implicando que $\det(A_{(k,k)}) > 0$ para $k = 1, 2, \dots, n$.

Supondo agora que os determinantes das submatrizes principais de uma matriz simétrica $A \in \mathbb{R}^{n \times n}$ são positivos temos que os pivôs são positivos e pelo Teorema 2.10 segue que A é definida positiva. ■

Note que num sistema $Ax = b$, onde A é simétrica definida positiva temos A invertível ($\det(A) \neq 0$), garantindo que $Ax = b$ é um sistema possível e determinado, isto é, tem solução e ela é única.

A seguir usaremos o teste dos determinantes para descobrir se a matriz simétrica abaixo é definida positiva ou não.

Seja

$$A = \begin{bmatrix} 1 & 2 & 4 \\ 2 & 20 & 16 \\ 4 & 16 & 29 \end{bmatrix}$$

Vejam os:

$$\det(A_{(1,1)}) = |1| = 1 > 0$$

$$\det(A_{(2,2)}) = \begin{vmatrix} 1 & 2 \\ 2 & 20 \end{vmatrix} = 20 - 4 = 16 > 0$$

$$\det(A_{(3,3)}) = \begin{vmatrix} 1 & 2 & 4 \\ 2 & 20 & 16 \\ 4 & 16 & 29 \end{vmatrix} = 580 + 128 + 128 - 320 - 256 - 116 = 144 > 0$$

Portanto, a matriz A é definida positiva.

É bastante comum determinar se uma matriz é ou não definida positiva avaliando se a mesma tem ou não a fatoração de Cholesky.

Métodos práticos para obter a fatoração de Cholesky

Agora vamos apresentar dois métodos para obter a decomposição de Cholesky de uma matriz simétrica definida positiva $A \in \mathbb{R}^{n \times n}$.

O primeiro método segue por escalonamento e serve tanto para determinar se uma matriz simétrica é definida positiva (observando os pivôs) quanto para, em caso afirmativo, determinar a sua decomposição de Cholesky. Supondo que $A \in \mathbb{R}^{n \times n}$ é simétrica definida positiva

temos, pelos resultados anteriores que $A = LU = LDL^t = LSS^tL^t = (LS)(LS)^t = GG^t$, onde $D = [d_{ij}]$ e $S = [s_{ij}]$ são diagonais tais que $d_{ii} = u_{ii}$ e $s_{ij} = \sqrt{u_{ij}}$, e por fim, $G = LS$. Assim, na prática, escalonamos A até obter U e L , construímos a matriz diagonal $S = [s_{ij}]$, onde $s_{ij} = \sqrt{u_{ij}}$ e obtemos G fazendo o produto LS . Vejamos isso para a matriz A exemplificada anteriormente, a qual sabemos ser simétrica definida positiva.

(1) Escalonando A para obter U :

$$\begin{bmatrix} 1 & 2 & 4 \\ 2 & 20 & 16 \\ 4 & 16 & 29 \end{bmatrix} \xrightarrow[L_3 - 4L_1]{L_2 - 2L_1} \begin{bmatrix} 1 & 2 & 4 \\ 0 & 16 & 8 \\ 0 & 8 & 13 \end{bmatrix} \xrightarrow{L_3 - \frac{1}{2}L_2} \begin{bmatrix} 1 & 2 & 4 \\ 0 & 16 & 8 \\ 0 & 0 & 9 \end{bmatrix} = U$$

(2) Obtendo L e também a matriz diagonal $S = [s_{ij}]$ tal que $s_{ij} = \sqrt{u_{ij}}$:

$$L = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 4 & \frac{1}{2} & 1 \end{bmatrix} \quad S = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 4 & 0 \\ 0 & 0 & 3 \end{bmatrix}$$

(3) Obtendo G pelo produto LS :

$$G = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 4 & \frac{1}{2} & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & 0 \\ 0 & 4 & 0 \\ 0 & 0 & 3 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 4 & 0 \\ 4 & 2 & 3 \end{bmatrix}$$

(4) Finalmente escrevemos a decomposição de Cholesky de A :

$$A = \begin{bmatrix} 1 & 2 & 4 \\ 2 & 20 & 16 \\ 4 & 16 & 29 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 4 & 0 \\ 4 & 2 & 3 \end{bmatrix} \cdot \begin{bmatrix} 1 & 2 & 4 \\ 0 & 4 & 2 \\ 0 & 0 & 3 \end{bmatrix} = GG^t$$

O outro método consiste no fato de que se $A = [a_{ij}]_{n \times n}$ é uma matriz simétrica definida positiva então a decomposição de Cholesky de A pode ser expressa na forma

$$\begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{k1} & a_{k2} & \dots & a_{kn} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix} =$$

$$\begin{bmatrix} g_{11} & 0 & \dots & 0 & \dots & 0 \\ g_{21} & g_{22} & \dots & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ g_{k1} & g_{k2} & \dots & g_{kk} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ g_{n1} & g_{n2} & \dots & g_{nk} & \dots & g_{nn} \end{bmatrix} \cdot \begin{bmatrix} g_{11} & g_{21} & \dots & g_{k1} & \dots & g_{n1} \\ 0 & g_{22} & \dots & g_{k2} & \dots & g_{n2} \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & g_{kk} & \dots & g_{nk} \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & \dots & g_{nn} \end{bmatrix}$$

Então a obtenção de $G = [g_{ij}]$ consiste em determinar inicialmente os elementos da primeira coluna, em seguida os da segunda coluna e assim por diante.

Note que se efetuarmos o produto GG^t teremos válidas as igualdades entre as colunas correspondentes de A e GG^t , onde cada elemento da k -ésima coluna de G pode ser determinada como segue:

PRIMEIRA COLUMNA

$$\begin{bmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{k1} \\ \vdots \\ a_{n1} \end{bmatrix} = \begin{bmatrix} g_{11}^2 \\ g_{21}g_{11} \\ \vdots \\ g_{k1}g_{11} \\ \vdots \\ g_{n1}g_{11} \end{bmatrix}$$

Então

$$g_{11} = \sqrt{a_{11}}$$

$$g_{i1} = \frac{a_{i1}}{g_{11}}, \text{ para } i = 2, 3, \dots, n$$

SEGUNDA COLUMNA

$$\begin{bmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{k2} \\ \vdots \\ a_{n2} \end{bmatrix} = \begin{bmatrix} g_{11}g_{21} \\ g_{21}^2 + g_{22}^2 \\ \vdots \\ g_{k1}g_{21} + g_{k2}g_{22} \\ \vdots \\ g_{n1}g_{21} + g_{n2}g_{22} \end{bmatrix}$$

Já são conhecidos os elementos g_{i1} . Logo,

$$g_{22} = \sqrt{a_{22} - g_{21}^2}$$

$$g_{i2} = \frac{a_{i2} - g_{i1}g_{21}}{g_{22}}, \text{ para } i = 3, 4, \dots, n$$

K-ÉSIMA COLUNA

$$\begin{bmatrix} a_{1k} \\ a_{2k} \\ \vdots \\ a_{kk} \\ \vdots \\ a_{ik} \\ \vdots \\ a_{nk} \end{bmatrix} = \begin{bmatrix} g_{11}g_{k1} \\ g_{21}g_{k1} + g_{22}g_{k2} \\ \vdots \\ g_{k1}^2 + g_{k2}^2 + \dots + g_{kk}^2 \\ \vdots \\ g_{i1}g_{k1} + g_{i2}g_{k2} + \dots + g_{ik}g_{kk} \\ \vdots \\ g_{n1}g_{k1} + g_{n2}g_{k2} + \dots + g_{nk}g_{kk} \end{bmatrix}$$

Os elementos g_{ij} , para $i = 1, 2, \dots, n$ e $j = 1, 2, \dots, k-1$, já são conhecidos. Então

$$g_{kk} = \sqrt{a_{kk} - (g_{k1}^2 + g_{k2}^2 + \dots + g_{k(k-1)}^2)}$$

$$g_{ik} = \frac{a_{ik} - g_{i1}g_{k1} - g_{i2}g_{k2} - \dots - g_{i(k-1)}g_{k(k-1)}}{g_{kk}}, \text{ para } i = k+1, k+2, \dots, n$$

Exemplo: Seja

$$B = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix}$$

Verificaremos que B é definida positiva usando o teste dos determinantes:

$$\det(B_{(1,1)}) = |2| = 2 > 0$$

$$\det(B_{(2,2)}) = \begin{vmatrix} 2 & -1 \\ -1 & 2 \end{vmatrix} = 4 - 1 = 3 > 0$$

$$\det(B_{(3,3)}) = \begin{vmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{vmatrix} = 8 - 2 - 2 = 4 > 0$$

Logo, B é definida positiva e, portanto, admite a decomposição de Cholesky.

Vamos determiná-la fazendo $B = GG^t$

$$\begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix} = \begin{bmatrix} g_{11} & 0 & 0 \\ g_{21} & g_{22} & 0 \\ g_{31} & g_{32} & g_{33} \end{bmatrix} \cdot \begin{bmatrix} g_{11} & g_{21} & g_{31} \\ 0 & g_{22} & g_{32} \\ 0 & 0 & g_{33} \end{bmatrix}$$

$$\begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix} = \begin{bmatrix} g_{11}^2 & g_{21}g_{11} & g_{31}g_{11} \\ g_{21}g_{11} & g_{21}^2 + g_{22}^2 & g_{21}g_{31} + g_{22}g_{32} \\ g_{31}g_{11} & g_{31}g_{21} + g_{32}g_{22} & g_{31}^2 + g_{32}^2 + g_{33}^2 \end{bmatrix}$$

Igualando a primeira coluna de B e de GG^t para calcular g_{11} , g_{21} e g_{31} :

$$\begin{bmatrix} 2 \\ -1 \\ 0 \end{bmatrix} = \begin{bmatrix} g_{11}^2 \\ g_{21}g_{11} \\ g_{31}g_{11} \end{bmatrix}$$

Então

$$\boxed{g_{11} = \sqrt{2}} \quad g_{21} = -\frac{1}{g_{11}} \quad \boxed{g_{31} = 0}$$

$$\boxed{g_{21} = -\frac{1}{\sqrt{2}}}$$

Igualando a segunda coluna de B e de GG^t para encontrar g_{22} e g_{32} :

$$\begin{bmatrix} -1 \\ 2 \\ -1 \end{bmatrix} = \begin{bmatrix} g_{21}g_{11} \\ g_{21}^2 + g_{22}^2 \\ g_{31}g_{21} + g_{32}g_{22} \end{bmatrix}$$

Então

$$g_{21}^2 + g_{22}^2 = 2$$

$$\left(-\frac{1}{\sqrt{2}}\right)^2 + g_{22}^2 = 2$$

$$g_{22}^2 = 2 - \frac{1}{2}$$

$$\boxed{g_{22} = \sqrt{\frac{3}{2}}}$$

$$g_{31}g_{21} + g_{32}g_{22} = -1$$

$$0 \cdot \left(-\frac{1}{\sqrt{2}}\right) + g_{32} \cdot \sqrt{\frac{3}{2}} = -1$$

$$\boxed{g_{32} = -\sqrt{\frac{2}{3}}}$$

Igualando a terceira coluna de B e de GG^t para determinar g_{33} :

$$\begin{bmatrix} 0 \\ -1 \\ 2 \end{bmatrix} = \begin{bmatrix} g_{31}g_{11} \\ g_{21}g_{31} + g_{22}g_{32} \\ g_{31}^2 + g_{32}^2 + g_{33}^2 \end{bmatrix}$$

Então

$$\begin{aligned} g_{31}^2 + g_{32}^2 + g_{33}^2 &= 2 \\ g_{33}^2 &= 2 - 0^2 - \left(-\sqrt{\frac{2}{3}}\right)^2 \\ g_{33}^2 &= 2 - \frac{2}{3} \\ \boxed{g_{33} &= \frac{2}{\sqrt{3}}} \end{aligned}$$

Logo a fatoração de Cholesky de B é:

$$B = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix} = \begin{bmatrix} \sqrt{2} & 0 & 0 \\ -\frac{1}{\sqrt{2}} & \sqrt{\frac{3}{2}} & 0 \\ 0 & -\sqrt{\frac{2}{3}} & \frac{2}{\sqrt{3}} \end{bmatrix} \cdot \begin{bmatrix} \sqrt{2} & -\frac{1}{\sqrt{2}} & 0 \\ 0 & \sqrt{\frac{3}{2}} & -\sqrt{\frac{2}{3}} \\ 0 & 0 & \frac{2}{\sqrt{3}} \end{bmatrix}$$

2.3 Fatoração QR

Nesta seção consideraremos de conhecimento do leitor algumas definições e resultados relativos a espaços vetoriais, subespaços, combinação linear, dependência e independência linear, bases e dimensão. Assim, passamos aos conceitos mais diretamente usados neste trabalho.

Como não consideraremos espaços vetoriais com escalares complexos poderemos chamar um espaço vetorial real simplesmente de espaço vetorial. O **espaço vetorial euclidiano de dimensão n** é o conhecido $\mathbb{R}^n$.

Produto Interno

Definição 2.5 *Seja V um espaço vetorial. Chamamos de produto interno em V à função, de $V \times V$ em $\mathbb{R}$, que associa a cada par de vetores $u, v \in V$ um número real $\langle u, v \rangle$, chamado produto interno de u por v , tal que para quaisquer vetores $u, v, w \in V$ e qualquer escalar $k \in \mathbb{R}$ se verifiquem as seguintes propriedades:*

$$(1) \langle u, v \rangle = \langle v, u \rangle$$

(simetria)

$$(2) \langle u + v, w \rangle = \langle u, w \rangle + \langle v, w \rangle \quad (\text{aditividade})$$

$$(3) \langle ku, v \rangle = k \langle u, v \rangle \quad (\text{homogeneidade})$$

$$(4) \langle v, v \rangle \geq 0 \text{ e } \langle v, v \rangle = 0 \text{ se, e somente se, } v = 0 \quad (\text{positividade})$$

Exemplo: Se $u = (x_1, y_1, z_1)$ e $v = (x_2, y_2, z_2)$ são vetores de $V = \mathbb{R}^3$, então definimos por produto interno euclidiano de $\mathbb{R}^3$ ao número

$$\langle u, v \rangle = x_1x_2 + y_1y_2 + z_1z_2.$$

De um modo geral, se $u = (u_1, \dots, u_n)$ e $v = (v_1, \dots, v_n)$ são dois vetores do espaço euclidiano $\mathbb{R}^n$ então o produto interno euclidiano de u por v é dado por

$$\langle u, v \rangle = u_1v_1 + \dots + u_nv_n.$$

Norma de um Vetor

Definição 2.6 Se V é um espaço vetorial com produto interno e u é um vetor de V , então a **norma** ou o **comprimento** de u é denotado por $\|u\|$ e é definido como

$$\|u\| = \sqrt{\langle u, u \rangle}.$$

Observe que a expressão acima satisfaz às fórmulas usadas para calcular o comprimento de um vetor em $\mathbb{R}$, $\mathbb{R}^2$ e $\mathbb{R}^3$ se o produto interno definido é o euclidiano. Por exemplo, se estivermos usando o produto interno euclidiano em $\mathbb{R}^2$ e se $u = (x, y)$ é um vetor de $\mathbb{R}^2$, então seu comprimento é $\sqrt{x^2 + y^2} = \sqrt{\langle u, u \rangle}$.

Definição 2.7 Seja V um espaço vetorial com produto interno. Dizemos que dois vetores $u, v \in V$ são **ortogonais** ou **perpendiculares** se $\langle u, v \rangle = 0$.

Por definição o vetor nulo 0 é ortogonal a qualquer outro vetor do espaço, pois $\langle 0, v \rangle = \langle 0.v, v \rangle = 0$. $\langle v, v \rangle = 0$.

O teorema seguinte é uma generalização do Teorema de Pitágoras para qualquer espaço vetorial com produto interno.

Teorema 2.12 (Teorema de Pitágoras) Se u e v são dois vetores ortogonais de um espaço com produto interno, então

$$\|u + v\|^2 = \|u\|^2 + \|v\|^2$$

Prova: Usando as propriedades (1) e (2) do produto interno e a definição de norma de um vetor, podemos escrever

$$\|u + v\|^2 = \langle u + v, u + v \rangle = \langle u, u + v \rangle + \langle v, u + v \rangle = (\langle u, u \rangle + \langle u, v \rangle) + (\langle v, u \rangle + \langle v, v \rangle) = \|u\|^2 + 2\langle u, v \rangle + \|v\|^2$$

Por hipótese u e v são ortogonais, implicando que $\langle u, v \rangle = 0$ e, portanto,

$$\|u + v\|^2 = \|u\|^2 + \|v\|^2$$

■

Em várias aplicações a ortogonalidade entre vetores é fundamental para simplificar os cálculos. É também freqüente a utilização de **vetores com norma 1**, chamados *unitários* ou *normalizados*. Dado um vetor não-nulo v de um espaço vetorial V munido de um produto interno, então o vetor $\frac{v}{\|v\|}$ tem norma 1, pois

$$\left\| \frac{v}{\|v\|} \right\| = \left\langle \frac{v}{\|v\|}, \frac{v}{\|v\|} \right\rangle^{\frac{1}{2}} = \frac{1}{\|v\|^{\frac{1}{2}}} \left\langle v, \frac{v}{\|v\|} \right\rangle^{\frac{1}{2}} = \frac{1}{\|v\|} \langle v, v \rangle^{\frac{1}{2}} = \frac{\|v\|}{\|v\|} = 1$$

Chamamos de **normalização** o processo de obtenção de um vetor unitário a partir de um vetor não-nulo v de um espaço vetorial com produto interno.

Definição 2.8 *Seja V um espaço vetorial munido de um produto interno. Dizemos que um conjunto $W \subset V$ é ortogonal se para quaisquer dois vetores $u, v \in W$ tivermos $\langle u, v \rangle = 0$. Se, além disso, os vetores de W são unitários, então dizemos que W é um conjunto ortonormal.*

O próximo teorema mostra como é simples escrever as coordenadas de um vetor arbitrário usando uma base ortonormal.

Teorema 2.13 *Se $W = \{v_1, v_2, \dots, v_n\}$ é uma base ortonormal de um espaço vetorial V com produto interno, então para qualquer vetor $u \in V$ temos*

$$u = \langle u, v_1 \rangle v_1 + \langle u, v_2 \rangle v_2 + \dots + \langle u, v_n \rangle v_n$$

Prova: Por hipótese W é uma base de V , então podemos escrever u como combinação linear dos vetores de W , isto é, existem escalares $k_1, k_2, \dots, k_n$ tais que

$$u = k_1 v_1 + k_2 v_2 + \dots + k_n v_n$$

Considere agora o produto interno

$$\langle u, v_i \rangle = \langle (k_1 v_1 + k_2 v_2 + \dots + k_n v_n), v_i \rangle$$

Usando as propriedades (2) e (3) do produto interno temos:

$$\begin{aligned} \langle u, v_i \rangle &= \langle k_1 v_1, v_i \rangle + \langle k_2 v_2, v_i \rangle + \dots + \langle k_n v_n, v_i \rangle \\ \langle u, v_i \rangle &= k_1 \langle v_1, v_i \rangle + k_2 \langle v_2, v_i \rangle + \dots + k_n \langle v_n, v_i \rangle \end{aligned}$$

Mas como W é ortonormal temos $\langle v_i, v_j \rangle = 0$ para $i \neq j$ e $\langle v_i, v_i \rangle = 1$, então

$$\langle u, v_i \rangle = k_i$$

Logo

$$u = \langle u, v_1 \rangle v_1 + \langle u, v_2 \rangle v_2 + \dots + \langle u, v_n \rangle v_n$$

■

Teorema 2.14 *Se $W = \{v_1, v_2, \dots, v_n\}$ é um conjunto ortogonal de vetores não-nulos de um espaço com produto interno, então W é linearmente independente.*

Prova: Supondo que $k_1 v_1 + k_2 v_2 + \dots + k_n v_n = 0$, para mostrar que W é linearmente independente basta provar que $k_1 = k_2 = \dots = k_n = 0$. Como para $v_i \in W$ temos

$$\langle k_1 v_1 + k_2 v_2 + \dots + k_n v_n, v_i \rangle = \langle 0, v_i \rangle = 0$$

e

$$\langle k_1 v_1 + k_2 v_2 + \dots + k_n v_n, v_i \rangle = k_1 \langle v_1, v_i \rangle + k_2 \langle v_2, v_i \rangle + \dots + k_n \langle v_n, v_i \rangle$$

então do fato de $\langle v_i, v_j \rangle = 0$ para $i \neq j$ temos

$$k_1 \langle v_1, v_i \rangle + k_2 \langle v_2, v_i \rangle + \dots + k_n \langle v_n, v_i \rangle = k_i \langle v_i, v_i \rangle$$

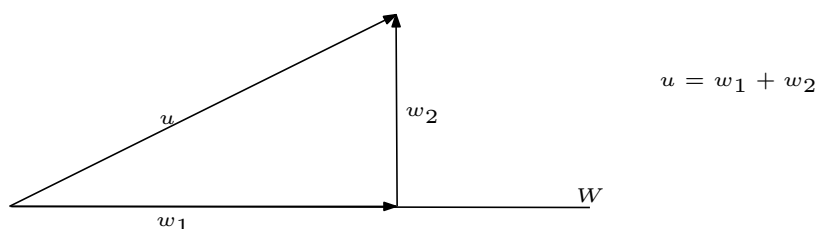
Como v_i é não-nulo segue que $\langle v_i, v_i \rangle \neq 0$, obrigando $k_i = 0$. Da arbitrariedade de i segue que $k_1 = k_2 = \dots = k_n = 0$ e, portanto, W é linearmente independente. ■

2.3.1 Projeções Ortogonais

Considere o espaço $\mathbb{R}^2$ munido do produto interno euclidiano. Podemos afirmar que um vetor $u \in \mathbb{R}^2$ pode ser expresso como uma soma de dois vetores ortogonais w_1 e w_2 . De fato, pois se tomarmos o subespaço W determinado por uma reta que passa pela origem e, perpendicularmente a W , levantarmos o vetor w_2 até a extremidade da flecha que representa o vetor u , então chamando de w_1 o vetor com extremidade no pé de w_2 temos

$$u = w_1 + w_2$$

Denotando por W o subespaço gerado pela reta que contém w_1 e por $W^\perp$ o subespaço ortogonal a W e que contém w_2 , podemos escrever $u = w_1 + w_2$, onde $w_1 \in W$ e $w_2 \in W^\perp$. Veja ilustração abaixo:



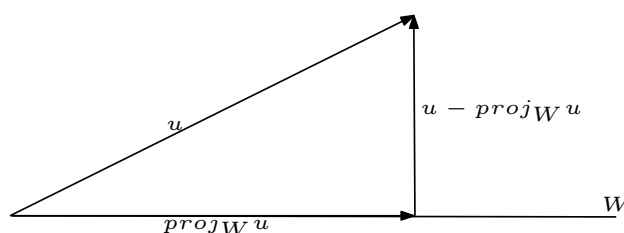
Chamamos o vetor w_1 de *projeção ortogonal de u em W* e o denotamos por $proj_W u$. Ao vetor w_2 chamamos de *componente de u ortogonal a W* e o denotamos por $proj_{W^\perp} u$. Desse modo, podemos escrever

$$u = proj_W u + proj_{W^\perp} u$$

Como $w_2 = u - w_1$ pode ser expresso na forma $proj_{W^\perp} u = u - proj_W u$, então

$$u = proj_W u + (u - proj_W u)$$

Assim, na ilustração teríamos:



De modo análogo ao que fizemos para $\mathbb{R}^2$ podemos fazer em $\mathbb{R}^3$ (munido do produto interno euclidiano) para expressar um vetor como soma de dois vetores ortogonais. Dado um vetor $u \in \mathbb{R}^3$ basta que tomemos w_1 em um subespaço W formado por um plano (ou reta) que contém a origem e w_2 num subespaço $W^\perp$ contendo a extremidade de u .

O próximo teorema estende esse resultado para qualquer espaço vetorial com produto interno definido e que tenha um subespaço de dimensão finita.

Teorema 2.15 *Se W é um subespaço de dimensão finita de um espaço V com produto interno, então cada vetor u de V pode ser expresso de maneira única como*

$$u = w_1 + w_2$$

onde $w_1 \in W$ e $w_2 \in W^\perp$

Para uma prova consulte a *Prova Adicional* de [1], p. 220 e 221.

Teorema 2.16 *Seja W um subespaço de dimensão finita de um espaço V com produto interno.*

(a) *Se $\{v_1, v_2, \dots, v_m\}$ é uma base ortonormal de W e u é um vetor qualquer de V , então*

$$\text{proj}_W u = \langle u, v_1 \rangle v_1 + \langle u, v_2 \rangle v_2 + \dots + \langle u, v_m \rangle v_m$$

(b) *Se $\{v_1, v_2, \dots, v_m\}$ é uma base ortogonal de W e u é um vetor qualquer de V , então*

$$\text{proj}_W u = \frac{\langle u, v_1 \rangle}{\|v_1\|^2} v_1 + \frac{\langle u, v_2 \rangle}{\|v_2\|^2} v_2 + \dots + \frac{\langle u, v_m \rangle}{\|v_m\|^2} v_m$$

Prova de (a): Analogamente ao feito na prova do Teorema 2.13 escrevemos o vetor $\text{proj}_W u$ como combinação linear dos vetores da base ortonormal de W :

$$\text{proj}_W u = k_1 v_1 + k_2 v_2 + \dots + k_m v_m$$

Substituindo na demonstração do Teorema 2.13 o vetor u por $\text{proj}_W u$ chegamos com a mesma argumentação que

$$\langle \text{proj}_W u, v_i \rangle = k_i \text{ para } i = 1, 2, \dots, m$$

Logo

$$\text{proj}_W u = \langle \text{proj}_W u, v_1 \rangle v_1 + \langle \text{proj}_W u, v_2 \rangle v_2 + \dots + \langle \text{proj}_W u, v_m \rangle v_m$$

Pelo Teorema 2.15 podemos escrever $u = w_1 + w_2$, onde $w_1 \in W$ e $w_2 \in W^\perp$, isto é,

$$u = \text{proj}_W u + w_2 \quad \Longleftrightarrow \quad \text{proj}_W u = u - w_2$$

e, por conseguinte, podemos expressar $\text{proj}_W u$ como

$$\begin{aligned} \text{proj}_W u &= \langle u - w_2, v_1 \rangle v_1 + \langle u - w_2, v_2 \rangle v_2 + \dots + \langle u - w_2, v_m \rangle v_m \\ \text{proj}_W u &= (\langle u, v_1 \rangle - \langle w_2, v_1 \rangle) v_1 + \dots + (\langle u, v_m \rangle - \langle w_2, v_m \rangle) v_m \end{aligned}$$

Como $w_2 \in W^\perp$ e $v_i \in W$, então $\langle w_2, v_i \rangle = 0$ resultando que

$$\text{proj}_W u = \langle u, v_1 \rangle v_1 + \langle u, v_2 \rangle v_2 + \dots + \langle u, v_m \rangle v_m.$$

Prova de (b): Observe que se $\{v_1, v_2, \dots, v_m\}$ é uma base ortogonal de W , então o conjunto $\left\{ \frac{v_1}{\|v_1\|}, \frac{v_2}{\|v_2\|}, \dots, \frac{v_m}{\|v_m\|} \right\}$ constitui uma base ortonormal de W .

Logo, pela parte (a) do teorema podemos escrever

$$\text{proj}_W u = \left\langle u, \frac{v_1}{\|v_1\|} \right\rangle \frac{v_1}{\|v_1\|} + \left\langle u, \frac{v_2}{\|v_2\|} \right\rangle \frac{v_2}{\|v_2\|} + \dots + \left\langle u, \frac{v_m}{\|v_m\|} \right\rangle \frac{v_m}{\|v_m\|}$$

Como

$$\left\langle u, \frac{v_i}{\|v_i\|} \right\rangle \frac{v_i}{\|v_i\|} = \frac{\langle u, v_i \rangle v_i}{\|v_i\|^2}$$

segue que

$$\text{proj}_W u = \frac{\langle u, v_1 \rangle}{\|v_1\|^2} v_1 + \frac{\langle u, v_2 \rangle}{\|v_2\|^2} v_2 + \dots + \frac{\langle u, v_m \rangle}{\|v_m\|^2} v_m.$$

■

Processo de Gram-Schmidt

O processo de Gram-Schmidt consiste em se obter uma base ortogonal a partir de uma base qualquer. Para compreender esse processo veja os passos da demonstração do teorema seguinte, pois foram feitos usando o processo de Gram-Schmidt.

Teorema 2.17 *Se V é um espaço vetorial não-nulo de dimensão finita n , com produto interno, então V possui uma base ortonormal.*

Prova: Seja V um espaço vetorial não-nulo de dimensão n , sobre o qual está definido um produto interno. Seja ainda $\{u_1, u_2, \dots, u_n\}$ uma base de V . Tendo em vista que todo vetor não-nulo de V pode ser normalizado, basta que provemos que V tem uma base ortogonal $\{v_1, v_2, \dots, v_n\}$.

Nos passos seguintes, onde determinamos essa base, tomaremos W_i como o subespaço de V gerado por $v_1, v_2, \dots, v_i$ e o denotaremos por $W_i = [v_1, v_2, \dots, v_i]$.

Passo 1: Tome $v_1 = u_1$ (1)

Passo 2: Sejam $W_1 = [v_1]$ e v_2 o componente de u_2 ortogonal a W_1 , então

$$v_2 = u_2 - \text{proj}_{W_1} u_2 = u_2 - \frac{\langle u_2, v_1 \rangle}{\|v_1\|^2} v_1 \quad (2)$$

Como o vetor nulo é por definição ortogonal a qualquer vetor, então precisamos garantir que v_2 não é nulo, pois do contrário não poderá pertencer à base ortogonal de V que queremos construir. Suponha que tivéssemos $v_2 = 0$, então de (1) e (2) teríamos

$$u_2 = \frac{\langle u_2, v_1 \rangle}{\|v_1\|^2} v_1 = \frac{\langle u_2, u_1 \rangle}{\|u_1\|^2} u_1$$

implicando que u_2 seria um múltiplo de u_1 , o que é impossível, pois u_1 e u_2 são elementos de uma base de V e, portanto, linearmente independentes.

Passo 3: Sejam $W_2 = [v_1, v_2]$ e v_3 o componente de u_3 ortogonal a W_2 (portanto ortogonal a v_1 e v_2), então

$$v_3 = u_3 - \text{proj}_{W_2} u_3 = u_3 - \frac{\langle u_3, v_1 \rangle}{\|v_1\|^2} v_1 - \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} v_2 \quad (3)$$

Novamente justificaremos porque $v_3 \neq 0$. Supondo que $v_3 = 0$ então de (3) segue que

$$u_3 = \frac{\langle u_3, v_1 \rangle}{\|v_1\|^2} v_1 + \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} v_2$$

Substituindo (1) e (2) em (3) temos

$$u_3 = \frac{\langle u_3, u_1 \rangle}{\|u_1\|^2} u_1 + \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} \cdot \left(u_2 - \frac{\langle u_2, u_1 \rangle}{\|u_1\|^2} u_1 \right)$$

$$u_3 = \left(\frac{\langle u_3, u_1 \rangle}{\|u_1\|^2} - \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} \cdot \frac{\langle u_2, u_1 \rangle}{\|u_1\|^2} \right) u_1 + \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} u_2$$

ou seja, u_3 é combinação linear de u_2 e u_1 , o que é uma contradição, uma vez que u_1 , u_2 e u_3 são linearmente independentes. Prosseguindo desse modo podemos escrever o n -ésimo passo como segue.

Passo n : Sejam $W_{n-1} = [v_1, v_2, \dots, v_{n-1}]$ e v_n o componente de u_n ortogonal a W_{n-1} , então

$$v_n = u_n - \text{proj}_{W_{n-1}} u_n = u_n - \sum_{i=1}^{n-1} \frac{\langle u_n, v_i \rangle}{\|v_i\|^2} v_i \quad (n)$$

Assim, obtemos uma base ortogonal $\{v_1, v_2, \dots, v_n\}$ de V , visto que V tem dimensão n e vetores ortogonais são linearmente independentes. ■

Corolário 2.3 *Se $q_1, q_2, \dots, q_n$ é uma base ortonormal obtida pelo processo de Gram-Schmidt, com normalização subsequente, dos vetores da base $u_1, u_2, \dots, u_n$ de V , então q_k é ortogonal a $u_1, u_2, \dots, u_{k-1}$ para $k \geq 2$.*

Prova: Seja q_k o vetor normalizado de v_k , obtido no teorema acima por Gram-Schmidt aplicado a $\{u_1, u_2, \dots, u_k\}$. Como $q_k = \frac{v_k}{\|v_k\|}$ então para provar a ortogonalidade entre q_k e u_i , para $1 \leq i \leq k-1$ é suficiente mostrar que $\langle v_k, u_i \rangle = 0$, uma vez que

$$\langle q_k, u_i \rangle = 0 \Leftrightarrow \left\langle \frac{v_k}{\|v_k\|}, u_i \right\rangle = 0 \Leftrightarrow \frac{1}{\|v_k\|} \langle v_k, u_i \rangle = 0 \Leftrightarrow \langle v_k, u_i \rangle = 0$$

pois $\|v_k\| \neq 0$ visto que v_k é elemento de base.

Pelo teorema acima temos

$$v_k = u_k - \sum_{i=1}^{k-1} \frac{\langle u_k, v_i \rangle}{\|v_i\|^2} v_i$$

ou seja,

$$u_k = \sum_{i=1}^{k-1} \frac{\langle u_k, v_i \rangle}{\|v_i\|^2} v_i + v_k = \frac{\langle u_k, v_1 \rangle}{\|v_1\|^2} v_1 + \frac{\langle u_k, v_2 \rangle}{\|v_2\|^2} v_2 + \dots + \frac{\langle u_k, v_{k-1} \rangle}{\|v_{k-1}\|^2} v_{k-1} + v_k$$

Logo u_k é combinação linear de $W_k = [v_1, v_2, \dots, v_k]$, isto é, $u_k \in W_k$. Para mostrar que v_k é ortogonal a $u_1, u_2, \dots, u_{k-1}$ observemos que $v_k \perp W_i$ para $1 \leq i \leq k-1$ e, portanto, v_k

é ortogonal aos vetores $\{u_1, u_2, \dots, u_{k-1}\}$, pois estes são combinações lineares de W_j . Em outras palavras $\langle v_k, u_i \rangle = 0$ para $1 \leq i \leq k-1$. ■

Corolário 2.4 *Se $q_1, q_2, \dots, q_n$ é uma base ortonormal obtida pelo processo de Gram-Schmidt, com normalização subsequente, dos vetores da base $u_1, u_2, \dots, u_n$ de V , então $\langle u_i, q_i \rangle \neq 0$ para $i = 1, 2, \dots, n$.*

Prova: O produto interno entre dois vetores só é nulo se eles são ortogonais, então precisamos mostrar que u_i e q_i não são ortogonais. Considerando que $q_i = \frac{v_i}{\|v_i\|}$, isto é, v_i é múltiplo de q_i , então se u_i é ortogonal a q_i também o será a v_i . Assim, é suficiente provar que $\langle u_i, v_i \rangle \neq 0$. Como

$$\langle u_i, v_i \rangle = \langle u_i, (u_i - \text{proj}_{W_{i-1}} u_i) \rangle = \langle u_i, u_i \rangle - \langle u_i, \text{proj}_{W_{i-1}} u_i \rangle$$

e $\langle u_i, u_i \rangle \neq 0$, então $\langle u_i, v_i \rangle = 0$ só se $\langle u_i, u_i \rangle = \langle u_i, \text{proj}_{W_{i-1}} u_i \rangle$, ou ainda, se $\text{proj}_{W_{i-1}} u_i = u_i$. Tendo em vista que $\text{proj}_{W_{i-1}} u_i \in W_{i-1}$ e que os conjuntos $\{v_1, v_2, \dots, v_{i-1}\}$ e $\{u_1, u_2, \dots, u_{i-1}\}$ são bases de W_{i-1} , pois são linearmente independentes e têm dimensão $i-1$, então u_i é linearmente independente com $\{u_1, u_2, \dots, u_{i-1}\}$. Logo $\langle u_i, v_i \rangle \neq 0$. ■

Agora vamos ver uma propriedade relativa à multiplicação matricial, pois ela será usada no próximo teorema. Considere as matrizes A e x de ordens $m \times n$ e $n \times 1$, respectivamente. Afirmamos que o produto Ax pode ser expresso como uma combinação linear das colunas de A com os coeficientes de x . De fato, pois se

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} \quad \text{e} \quad x = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}$$

então

$$Ax = \begin{bmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{bmatrix} = x_1 \begin{bmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{bmatrix} + x_2 \begin{bmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{m2} \end{bmatrix} + \dots + x_n \begin{bmatrix} a_{1n} \\ a_{2n} \\ \vdots \\ a_{mn} \end{bmatrix} \quad (2.1)$$

Teorema 2.18 (Decomposição QR) Se A é uma matriz $m \times n$ com vetores-coluna linearmente independentes ($m \geq n$), então A pode ser fatorada como

$$A = QR$$

onde Q é uma matriz $m \times n$ com vetores-coluna ortonormais e R é uma matriz $n \times n$ triangular superior invertível.

Prova: Sejam $u_1, u_2, \dots, u_n$ os vetores-coluna de A e $q_1, q_2, \dots, q_n$ os vetores-coluna obtidos ao aplicarmos o processo de Gram-Schmidt e normalizarmos u_i , para $i = 1, 2, \dots, n$.

Pelo Teorema 2.13 podemos escrever

$$\begin{aligned} u_1 &= \langle u_1, q_1 \rangle q_1 + \langle u_1, q_2 \rangle q_2 + \dots + \langle u_1, q_n \rangle q_n \\ u_2 &= \langle u_2, q_1 \rangle q_1 + \langle u_2, q_2 \rangle q_2 + \dots + \langle u_2, q_n \rangle q_n \\ &\vdots \\ u_n &= \langle u_n, q_1 \rangle q_1 + \langle u_n, q_2 \rangle q_2 + \dots + \langle u_n, q_n \rangle q_n \end{aligned}$$

Considerando a fórmula 2.1 podemos escrever essas igualdades na forma matricial como

$$\begin{bmatrix} u_1 & u_2 & \dots & u_n \end{bmatrix} = \begin{bmatrix} q_1 & q_2 & \dots & q_n \end{bmatrix} \cdot \begin{bmatrix} \langle u_1, q_1 \rangle & \langle u_2, q_1 \rangle & \dots & \langle u_n, q_1 \rangle \\ \langle u_1, q_2 \rangle & \langle u_2, q_2 \rangle & \dots & \langle u_n, q_2 \rangle \\ \vdots & \vdots & \ddots & \vdots \\ \langle u_1, q_n \rangle & \langle u_2, q_n \rangle & \dots & \langle u_n, q_n \rangle \end{bmatrix}$$

Pelo Corolário 2.3 o vetor q_j é ortogonal a $u_1, u_2, \dots, u_{j-1}$, implicando que todos os elementos abaixo da diagonal principal da última matriz do segundo membro são nulos. Então a igualdade acima pode ser expressa na forma

$$\begin{bmatrix} u_1 & u_2 & \dots & u_n \end{bmatrix} = \begin{bmatrix} q_1 & q_2 & \dots & q_n \end{bmatrix} \cdot \begin{bmatrix} \langle u_1, q_1 \rangle & \langle u_2, q_1 \rangle & \dots & \langle u_n, q_1 \rangle \\ 0 & \langle u_2, q_2 \rangle & \dots & \langle u_n, q_2 \rangle \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \langle u_n, q_n \rangle \end{bmatrix}$$

Pelo Corolário 2.4 temos que $\langle u_i, q_i \rangle \neq 0$ e pelo Teorema 1.3, parte (b), uma matriz triangular com elementos diagonais não-nulos é invertível; logo a matriz triangular acima é invertível. Então, denotando as matrizes do segundo membro da igualdade acima por Q e R , respectivamente, temos a fatoração $A = QR$. ■

Exemplo: Obtenha a fatoração QR da matriz

$$A = \begin{bmatrix} 1 & 1 & 1 \\ 1 & -1 & 2 \\ -1 & 1 & 0 \\ 1 & 5 & 1 \end{bmatrix}$$

Solução: Escrevendo os vetores-coluna de A :

$$u_1 = \begin{bmatrix} 1 \\ 1 \\ -1 \\ 1 \end{bmatrix} \quad u_2 = \begin{bmatrix} 1 \\ -1 \\ 1 \\ 5 \end{bmatrix} \quad u_3 = \begin{bmatrix} 1 \\ 2 \\ 0 \\ 1 \end{bmatrix}$$

Tomando $v_1 = u_1$ passamos à construção de v_2 e v_3 , ortogonais com v_1 e entre si:

$$v_2 = u_2 - \frac{\langle u_2, v_1 \rangle}{\|v_1\|^2} v_1 = \begin{bmatrix} 1 \\ -1 \\ 1 \\ 5 \end{bmatrix} - \frac{4}{4} \cdot \begin{bmatrix} 1 \\ 1 \\ -1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ -2 \\ 2 \\ 4 \end{bmatrix}$$

$$v_3 = u_3 - \frac{\langle u_3, v_1 \rangle}{\|v_1\|^2} v_1 - \frac{\langle u_3, v_2 \rangle}{\|v_2\|^2} v_2 = \begin{bmatrix} 1 \\ 1 \\ -1 \\ 1 \end{bmatrix} - \frac{4}{4} \cdot \begin{bmatrix} 1 \\ 1 \\ -1 \\ 1 \end{bmatrix} - \frac{0}{24} \cdot \begin{bmatrix} 0 \\ -2 \\ 2 \\ 4 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 1 \\ 0 \end{bmatrix}$$

Fazemos agora a normalização de v_1 , v_2 e v_3 obtendo os vetores ortonormais q_1 , q_2 e q_3 , respectivamente.

$$q_1 = \frac{v_1}{\|v_1\|} = \frac{v_1}{2} = \begin{bmatrix} \frac{1}{2} \\ \frac{1}{2} \\ -\frac{1}{2} \\ \frac{1}{2} \end{bmatrix} \quad q_2 = \frac{v_2}{\|v_2\|} = \frac{v_2}{2\sqrt{6}} = \begin{bmatrix} 0 \\ -\frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{6}} \\ \frac{2}{\sqrt{6}} \end{bmatrix} \quad q_3 = \frac{v_3}{\|v_3\|} = \frac{v_3}{\sqrt{2}} = \begin{bmatrix} 0 \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \\ 0 \end{bmatrix}$$

Logo

$$Q = \begin{bmatrix} q_1 & q_2 & q_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & 0 & 0 \\ \frac{1}{2} & -\frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{2} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{2}{\sqrt{6}} & 0 \end{bmatrix}$$

e

$$R = \begin{bmatrix} \langle u_1, q_1 \rangle & \langle u_2, q_1 \rangle & \langle u_3, q_1 \rangle \\ 0 & \langle u_2, q_2 \rangle & \langle u_3, q_2 \rangle \\ 0 & 0 & \langle u_3, q_3 \rangle \end{bmatrix} = \begin{bmatrix} 2 & 2 & 2 \\ 0 & 2\sqrt{6} & 0 \\ 0 & 0 & \sqrt{2} \end{bmatrix}$$

Finalmente escrevemos

$$A = \begin{bmatrix} 1 & 1 & 1 \\ 1 & -1 & 2 \\ -1 & 1 & 0 \\ 1 & 5 & 1 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & 0 & 0 \\ \frac{1}{2} & -\frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{2} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{2}{\sqrt{6}} & 0 \end{bmatrix} \cdot \begin{bmatrix} 2 & 2 & 2 \\ 0 & 2\sqrt{6} & 0 \\ 0 & 0 & \sqrt{2} \end{bmatrix} = QR$$

Observamos que a matriz Q obtida acima é tal que $Q^t Q = I$, onde I é a matriz identidade. De fato, pois se Q é uma matriz de ordem $m \times n$ formada por vetores-coluna ortonormais, então

$$Q^t Q = \begin{bmatrix} q_1^t \\ q_2^t \\ \vdots \\ q_n^t \end{bmatrix} \cdot \begin{bmatrix} q_1 & q_2 & \dots & q_n \end{bmatrix} = \begin{bmatrix} \langle q_1, q_1 \rangle & \langle q_1, q_2 \rangle & \dots & \langle q_1, q_n \rangle \\ \langle q_2, q_1 \rangle & \langle q_2, q_2 \rangle & \dots & \langle q_2, q_n \rangle \\ \vdots & \vdots & \ddots & \vdots \\ \langle q_n, q_1 \rangle & \langle q_n, q_2 \rangle & \dots & \langle q_n, q_n \rangle \end{bmatrix}$$

Mas como $\langle q_i, q_j \rangle = 0$ para $i \neq j$ e $\langle q_i, q_i \rangle = 1$ para $i = j$ segue que

$$Q^t Q = \begin{bmatrix} \langle q_1, q_1 \rangle & \langle q_1, q_2 \rangle & \dots & \langle q_1, q_n \rangle \\ \langle q_2, q_1 \rangle & \langle q_2, q_2 \rangle & \dots & \langle q_2, q_n \rangle \\ \vdots & \vdots & \ddots & \vdots \\ \langle q_n, q_1 \rangle & \langle q_n, q_2 \rangle & \dots & \langle q_n, q_n \rangle \end{bmatrix} = \begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{bmatrix} = I$$

Sistemas Lineares

Em diversas ciências existem problemas que recaem na resolução de sistemas lineares. Geralmente esses problemas são descritos e resolvidos computacionalmente, pois suas dimensões e complexidade podem tornar impraticável uma solução manual. É comum, por exemplo, um sistema linear ser gerado por um banco de dados de um experimento, fornecendo um número expressivo de equações e incógnitas. Por isso, muitos pesquisadores têm se dedicado a estudar maneiras eficientes de resolver sistemas lineares, e nesse sentido o estudo das fatorações colaborou no desenvolvimento de novas ferramentas computacionais úteis.

Freqüentemente denotamos um sistema linear de uma forma bem compacta $Ax = b$, onde A é a matriz dos coeficientes, x o vetor das incógnitas e b o vetor dos termos independentes. De um modo geral a matriz A e os vetores x e b são, respectivamente, de ordens $m \times n$, $n \times 1$ e $m \times 1$.

Existem duas classes de métodos para resolução de sistemas de equações lineares: os **métodos diretos** e os **métodos iterativos**. Teoricamente podemos definir *métodos diretos* como técnicas que fornecem a solução exata, se existir, em um número finito de operações, enquanto que os *métodos iterativos* partem de uma aproximação inicial da solução e, caso haja convergência, fornecem uma aproximação para solução exata. Na prática, como a resolução de muitos sistemas envolvem números racionais e irracionais cujas representações numéricas podem ter muitas ou infinitas casas decimais, uma resolução computacional de um sistema normalmente apresenta erros de arredondamento mesmo se usado um método direto. A minimização desses erros pode significar muito na resolução de um problema. Assim, o

tipo de sistema e do método empregado para resolvê-lo implicam em diferenças relevantes na aproximação da solução exata. Métodos que produzem soluções “ruins” são ditos *instáveis*.

Outro fator primordial quando tratamos da elaboração de algoritmos para implementação computacional é o tempo gasto para realizar os cálculos, o que está diretamente relacionado ao número de operações envolvidas no algoritmo. Por exemplo, para um sistema linear $Ax = b$ onde A é uma matriz invertível, poderíamos calcular a inversa A^{-1} e multiplicar em ambos os membros de $Ax = b$, obtendo a solução $x = A^{-1}b$. Apesar de ser teoricamente correto, esse procedimento é computacionalmente desvantajoso. Além de erros numéricos que são introduzidos naturalmente no decorrer do processo, o tempo gasto na determinação da inversa é relativamente grande, se comparado com diversos outros métodos. Por esse motivo, não é usual a utilização da inversa na resolução de sistemas lineares em implementações computacionais.

Neste trabalho tratamos de métodos diretos de resolução de sistemas lineares, mais especificamente daqueles baseados nas fatorações LU, Cholesky e QR. Por estes métodos, a solução de um sistema linear $Ax = b$ pode ser obtida resolvendo-se um ou dois sistemas triangulares, cuja definição que adotamos segue abaixo:

Definição 3.1 *Seja $Ax = b$ um sistema linear cuja matriz dos coeficientes $A = [a_{ij}]$ é de ordem $m \times n$. Dizemos que $Ax = b$ é um sistema triangular inferior (superior) se $a_{ij} = 0$ para $i < j$ ($i > j$).*

Os sistemas triangulares são facilmente resolvidos por substituições sucessivas. Se o sistema é triangular inferior obtemos sua solução de cima para baixo (substituição direta), enquanto que num sistema triangular superior devemos efetuar as substituições de baixo para cima (retro-substituição).

Antes de apresentarmos a resolução de sistemas lineares usando as fatorações estudadas no capítulo anterior vamos falar da eliminação gaussiana, um dos métodos diretos mais usados e que é baseado na triangularização do sistema.

3.1 Eliminação Gaussiana

Considere um sistema linear de m equações e n incógnitas como segue:

$$\left[\begin{array}{ccc|c} 1 & 1 & -1 & 6 \\ 0 & -6 & -1 & -11 \\ 0 & 0 & \frac{31}{6} & -\frac{25}{6} \end{array} \right]$$

Escrevendo o sistema equivalente ao inicial e já obtendo a sua solução:

$$\begin{cases} x_1 + x_2 - x_3 = 6 \\ -6x_2 - x_3 = -11 \\ \frac{31}{6}x_3 = -\frac{25}{6} \end{cases}$$

$$x_3 = -\frac{25}{31}$$

$$-6x_2 - x_3 = -11$$

$$-6x_2 - \left(-\frac{25}{31}\right) = -11$$

$$-6x_2 = -11 - \frac{25}{31}$$

$$x_2 = \frac{61}{31}$$

$$x_1 + x_2 - x_3 = 6$$

$$x_1 + \frac{61}{31} - \left(-\frac{25}{31}\right) = 6$$

$$x_1 = 6 - \frac{86}{31}$$

$$x_1 = \frac{100}{31}$$

Nem todo sistema linear tem uma única solução como no caso exemplificado. Existem sistemas lineares que admitem infinitas soluções ou nenhuma solução. As condições que determinam o número de soluções de um sistema linear são geralmente estudadas num curso de álgebra linear e podem ser encontrados em [2], p. 45.

Agora achamos conveniente justificar a validade do método de eliminação gaussiana, ou seja, porque o sistema linear $Ax = b$ é equivalente a $\bar{A}x = \bar{b}$, construído após operarmos com a matriz ampliada $[A|b]$. A justificativa está no fato de que aplicar operações elementares sobre as linhas da matriz ampliada $[A|b]$ de um sistema $Ax = b$ é equivalente a multiplicar à esquerda de $[A|b]$ matrizes elementares. Assim, se E_i , para $i = 1, 2, \dots, k$, é uma matriz elementar, então a forma escalonada de $[A|b]$ pode ser expressa como

$$E_k \dots E_1 [A|b] = [E_k \dots E_1 A | E_k \dots E_1 b] \quad (3.1)$$

Denotando $E_k \dots E_1 A$ por $\bar{A}$ e $E_k \dots E_1 b$ por $\bar{b}$ podemos escrever 3.1 como $[\bar{A}|\bar{b}]$. Mas note que se aplicássemos as mesmas matrizes elementares $E_1, \dots, E_k$ diretamente a $Ax = b$ teríamos

$$E_k \dots E_1 Ax = E_k \dots E_1 b \quad \Longleftrightarrow \quad \bar{A}x = \bar{b}$$

Ou seja, $Ax = b$ e $\bar{A}x = \bar{b}$ são sistemas equivalentes (possuem o mesmo conjunto solução) e, portanto, a solução de $Ax = b$ pode ser obtida de $[\bar{A}|\bar{b}]$, forma escada da matriz ampliada de $Ax = b$, resolvendo $\bar{A}x = \bar{b}$.

3.2 Resolução de Sistemas Lineares Usando Fatorações

No capítulo anterior vimos como obter as decomposições LU, Cholesky e QR de uma matriz A , respeitadas as condições de existência de cada fatoração. Nesta seção apresentaremos como resolver um sistema linear $Ax = b$ usando as duas primeiras fatorações. Contaremos o número de operações envolvidas na resolução de um sistema usando LU, lembrando que já fizemos a contagem de operações da fatoração em si (Capítulo 2, Seção 2.1.1). O número de operações executadas na resolução de um sistema usando as fatorações de Cholesky e QR serão apenas indicadas.

Resolveremos um sistema linear usando a fatoração QR numa nova seção, onde relacionamos a resolução de sistemas lineares a mínimos quadrados. Lá também consideraremos a fatoração de Cholesky.

3.2.1 Resolvendo um sistema linear usando a fatoração LU

Se uma matriz A de ordem $n \times n$ admite a fatoração $A = LU$, então o sistema linear $Ax = b$ pode ser resolvido seguindo os passos abaixo:

- (1) Reescreva $Ax = b$ na forma $LUx = b$.
- (2) Resolva $Ly = b$, onde y é um vetor de incógnitas de ordem $n \times 1$. (Note que $y = Ux$)
- (3) Resolva $Ux = y$, para finalmente encontrar x , pois y já está determinado.

Exemplo: Use a decomposição LU para resolver o sistema abaixo:

$$\begin{cases} 2x_1 + 3x_2 - x_3 = -4 \\ x_1 - 2x_2 + 2x_3 = 14 \\ 3x_1 + x_2 - x_3 = 0 \end{cases}$$

Note que tomamos como matriz dos coeficientes do sistema acima a mesma matriz A cuja fatoração foi descrita passo a passo no capítulo anterior. Portanto, podemos reescrever esse sistema em forma matricial e resolvê-lo seguindo os passos indicados acima:

- (1) Reescrevendo $Ax = b$ na forma $LUx = b$:

$$\begin{bmatrix} 1 & 0 & 0 \\ \frac{1}{2} & 1 & 0 \\ \frac{3}{2} & 1 & 1 \end{bmatrix} \cdot \begin{bmatrix} 2 & 3 & -1 \\ 0 & -\frac{7}{2} & \frac{5}{2} \\ 0 & 0 & -2 \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} -4 \\ 14 \\ 0 \end{bmatrix}$$

(2) Resolvendo $Ly = b$:

$$\begin{bmatrix} 1 & 0 & 0 \\ \frac{1}{2} & 1 & 0 \\ \frac{3}{2} & 1 & 1 \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix} = \begin{bmatrix} -4 \\ 14 \\ 0 \end{bmatrix}$$

obtemos, resolvendo de cima para baixo, a solução $y_1 = -4$, $y_2 = 16$ e $y_3 = -10$.

(3) Resolvendo de baixo para cima $Ux = y$ obtemos a solução do sistema inicial: $x_1 = 2$, $x_2 = -1$ e $x_3 = 5$.

Contagem do número de operações necessárias para resolver $LUx=b$:

Como vimos, depois de obtidas as matrizes L e U , o sistema $Ax = b$ pode ser escrito na forma $LUx = b$. Então podemos construir uma matriz-coluna de incógnitas y tal que $Ly = b$, cuja matriz ampliada pode ser expressa da forma:

$$\begin{bmatrix} 1 & 0 & 0 & \dots & 0 & 0 & b_1 \\ \frac{a_{21}}{a_{11}} & 1 & 0 & \dots & 0 & 0 & b_2 \\ \frac{a_{31}}{a_{11}} & \frac{a_{32}^{(1)}}{a_{22}^{(1)}} & 1 & \dots & 0 & 0 & b_3 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ \frac{a_{(n-1)1}}{a_{11}} & \frac{a_{(n-1)2}^{(1)}}{a_{22}^{(1)}} & \frac{a_{(n-1)3}^{(2)}}{a_{33}^{(2)}} & \dots & 1 & 0 & b_{n-1} \\ \frac{a_{n1}}{a_{11}} & \frac{a_{n2}^{(1)}}{a_{22}^{(1)}} & \frac{a_{n3}^{(2)}}{a_{33}^{(2)}} & \dots & \frac{a_{n(n-1)}^{(n-2)}}{a_{(n-1)(n-1)}^{(n-2)}} & 1 & b_n \end{bmatrix}$$

Resolvendo por substituição direta $Ly=b$ para determinar y :

Para as etapas seguintes denotaremos o vetor y por $y = (y_1, y_2, \dots, y_n)$.

Passo 1: Determinação de y_1 .

Operação realizada: $y_1 = b_1$.

Número de operações realizadas:

Adições/Subtrações: 0

Multiplicações/Divisões: 0

Passo 2: Determinação de y_2 .

Operação realizada: $y_2 = b_2 - \frac{a_{21}}{a_{11}}x_1$.

Número de operações realizadas:

Adições/Subtrações: 1

Multiplicações/Divisões: 1

Passo 3: Determinação de y_3 .

Operação realizada: $y_3 = b_3 - \frac{a_{31}}{a_{11}}x_1 - \frac{a_{32}^{(1)}}{a_{22}^{(1)}}x_2$.

Número de operações realizadas:

Adições/Subtrações: 2

Multiplicações/Divisões: 2

Passo n-1: Determinação de y_{n-1} .

Operação realizada: $y_{n-1} = b_{n-1} - \frac{a_{(n-1)1}}{a_{11}}x_1 - \frac{a_{(n-1)2}^{(1)}}{a_{22}^{(1)}}x_2 - \dots - \frac{a_{(n-1)(n-2)}^{(n-3)}}{a_{(n-2)(n-2)}^{(n-3)}}x_{n-2}$.

Número de operações realizadas:

Adições/Subtrações: $n - 2$

Multiplicações/Divisões: $n - 2$

Passo n: Determinação de y_n .

Operação realizada: $y_n = b_n - \frac{a_{n1}}{a_{11}}x_1 - \frac{a_{n2}^{(1)}}{a_{22}^{(1)}}x_2 - \dots - \frac{a_{n(n-1)}^{(n-2)}}{a_{(n-1)(n-1)}^{(n-2)}}x_{n-1}$.

Número de operações realizadas:

Adições/Subtrações: $n - 1$

Multiplicações/Divisões: $n - 1$

Total de operações para determinar y:

Adições/Subtrações:

$$1 + 2 + \dots + (n-2) + (n-1) = \frac{(n-1)n}{2}$$

Multiplicações/Divisões:

$$1 + 2 + \dots + (n-2) + (n-1) = \frac{(n-1)n}{2}$$

Assim, temos o vetor y determinado, bastando para a determinação do vetor x fazermos $Ux = y$, cuja matriz ampliada está descrita a seguir:

$$\left[\begin{array}{ccccccc} a_{11} & a_{12} & a_{13} & \dots & a_{1(n-1)} & a_{1n} & y_1 \\ 0 & a_{22}^{(1)} & a_{23}^{(1)} & \dots & a_{2(n-1)}^{(1)} & a_{2n}^{(1)} & y_2 \\ 0 & 0 & a_{33}^{(2)} & \dots & a_{3(n-1)}^{(2)} & a_{3n}^{(2)} & y_3 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & a_{(n-1)(n-1)}^{(n-2)} & a_{(n-1)n}^{(n-2)} & y_{n-1} \\ 0 & 0 & 0 & \dots & 0 & a_{nn}^{(n-1)} & y_n \end{array} \right]$$

Resolvendo por retro-substituição $Ux=y$ para determinar x :

Nas etapas a seguir o vetor x será denotado por $x = (x_1, x_2, \dots, x_n)$.

Passo 1: Determinação de x_n .

Operação realizada: $x_n = \frac{y_n}{a_{nn}^{(n-1)}}$.

Número de operações realizadas:

Adições/Subtrações: 0.

Multiplicações/Divisões: 1.

Passo 2: Determinação de x_{n-1} .

Operação realizada: $x_{n-1} = \frac{y_{n-1} - a_{(n-1)n}^{(n-2)} x_n}{a_{(n-1)(n-1)}^{(n-2)}}$.

Número de operações realizadas:

Adições/Subtrações: 1.

Multiplicações/Divisões: 2.

Passo n-1: Determinação de x_2 .

Operação realizada: $x_2 = \frac{y_2 - a_{23}^{(1)}x_3 - \dots - a_{2(n-1)}^{(1)}x_{n-1} - a_{2n}^{(1)}x_n}{a_{22}^{(1)}}.$

Número de operações realizadas:

Adições/Subtrações: $n - 2.$

Multiplicações/Divisões: $n - 1.$

Passo n: Determinação de $x_1.$

Operação realizada: $x_1 = \frac{y_1 - a_{12}x_2 - \dots - a_{1(n-1)}x_{n-1} - a_{1n}x_n}{a_{11}}.$

Número de operações realizadas:

Adições/Subtrações: $n - 1.$

Multiplicações/Divisões: $n.$

Total de operações para determinar x em $Ux=y$:

Adições/Subtrações:

$$1 + 2 + \dots + (n - 2) + (n - 1) = \frac{(n - 1)n}{2}$$

Multiplicações/Divisões:

$$1 + 2 + \dots + (n - 2) + (n - 1) + n = \frac{(n + 1)n}{2}$$

Total de operações utilizadas na resolução de $Ax=b$ usando a fatoração LU:

Adições/Subtrações:

$$\frac{(n - 1)n(2n - 1)}{6} + \frac{(n - 1)n}{2} + \frac{(n - 1)n}{2} = \frac{n^3}{3} + \frac{n^2}{2} - \frac{5n}{6}$$

Multiplicações/Divisões:

$$\frac{(n - 1)n(n + 1)}{3} + \frac{(n - 1)n}{2} + \frac{(n + 1)n}{2} = \frac{n^3}{3} + n^2 - \frac{n}{3}$$

Observamos que no total de operações acima já consideramos, inclusive, a contagem das operações utilizadas para fatorar A, cujos detalhes estão na seção do capítulo anterior

dedicada à fatoração LU . O total de operações envolvidas nesse processo é $\frac{2n^3}{3} + \frac{3n^2}{2} - \frac{7n}{6}$. Na prática, consideramos apenas o termo de maior expoente e assim podemos dizer que este é um método tem o custo assintótico de $2n^3/3$ operações, ou ainda, é da ordem de $O(n^3)$.

Note que se tivermos vários sistemas $Ax = b_1, Ax = b_2, \dots, Ax = b_k$, então a decomposição $A = LU$ pode ser muito vantajosa computacionalmente, visto que não necessitaremos fazer mais de uma vez o escalonamento. Uma aplicação dessa propriedade pode ser a determinação da inversa de uma matriz usando a fatoração LU . Uma vez fatorada a matriz $A = LU$ podemos determinar se ela é ou não invertível apenas observando se a última linha de U é não-nula, pois uma matriz invertível tem determinante não-nulo. Assim, se $A = LU$ então $\det(A) = \det(L)\det(U)$ implicando que se U tem última linha não-nula então $\det(A) \neq 0$ (veja demonstração do Teorema 2.1). Existindo a inversa A^{-1} temos que $AA^{-1} = I$, então se denotarmos os vetores-colunas de A^{-1} por vetores de incógnitas $x_1, x_2, \dots, x_n$ e os vetores-coluna da identidade I por $e_1, e_2, \dots, e_n$ podemos determinar A^{-1} resolvendo os sistemas $Ax = e_1, Ax = e_2, \dots, Ax = e_n$. Essa abordagem é muito utilizada na prática, principalmente quando se trabalha com matrizes de médio a grande porte.

3.2.2 Resolvendo um sistema linear usando a fatoração de Cholesky

Considere o sistema

$$\begin{cases} x_1 + 2x_2 + 4x_3 = 7 \\ 2x_1 + 20x_2 + 16x_3 = 4 \\ 4x_1 + 16x_2 + 29x_3 = 1 \end{cases} \iff \begin{bmatrix} 1 & 2 & 4 \\ 2 & 20 & 16 \\ 4 & 16 & 29 \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 7 \\ 4 \\ 1 \end{bmatrix}$$

Veja que a matriz dos coeficientes desse sistema é a mesma para a qual obtivemos a decomposição de Cholesky usando o método do escalonamento. A sua solução pode ser obtida analogamente ao indicado para a decomposição LU . Vejamos:

(1) Reescrevendo $Ax = b$ da forma $GG^t x = b$:

$$Ax = b \iff \begin{bmatrix} 1 & 0 & 0 \\ 2 & 4 & 0 \\ 4 & 2 & 3 \end{bmatrix} \cdot \begin{bmatrix} 1 & 2 & 4 \\ 0 & 4 & 2 \\ 0 & 0 & 3 \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 7 \\ 4 \\ 1 \end{bmatrix}$$

(1) Resolvendo $Gy = b$ de cima para baixo:

$$\begin{bmatrix} 1 & 0 & 0 \\ 2 & 4 & 0 \\ 4 & 2 & 3 \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix} = \begin{bmatrix} 7 \\ 4 \\ 1 \end{bmatrix}$$

Segue-se que

$$\begin{array}{ll} 2y_1 + 4y_2 = 4 & 4y_1 + 2y_2 + 3y_3 = 1 \\ \boxed{y_1 = 7} & 4 \cdot (7) + 4y_2 = 4 \\ & 4y_2 = -10 \\ & \boxed{y_2 = -\frac{5}{2}} \end{array} \quad \begin{array}{l} 4y_1 + 2y_2 + 3y_3 = 1 \\ 4 \cdot (7) + 2 \cdot \left(-\frac{5}{2}\right) + 3y_3 = 1 \\ 3y_3 = 1 - 28 + 5 \\ \boxed{y_3 = -\frac{22}{3}} \end{array}$$

(2) Resolvendo $G^t x = y$ por retro-substituição:

$$\begin{bmatrix} 1 & 2 & 4 \\ 0 & 4 & 2 \\ 0 & 0 & 3 \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 7 \\ -\frac{5}{2} \\ -\frac{22}{3} \end{bmatrix}$$

Donde obtemos

$$\begin{array}{ll} 4x_2 + 2x_3 = -\frac{5}{2} & x_1 + 2x_2 + 4x_3 = 7 \\ 4x_2 + 2 \cdot \left(-\frac{22}{9}\right) = -\frac{5}{2} & x_1 + 2 \cdot \left(-\frac{43}{72}\right) + 4 \cdot \left(-\frac{22}{9}\right) = 7 \\ 3x_3 = -\frac{22}{3} & x_1 = 7 + \frac{43}{36} + \frac{88}{9} \\ \boxed{x_3 = -\frac{22}{9}} & x_1 = \frac{252+43+352}{36} \\ & \boxed{x_1 = \frac{647}{36}} \\ 4x_2 = -\frac{5}{2} + \frac{44}{9} & \\ 4x_2 = \frac{-45+88}{18} & \\ 4x_2 = -\frac{43}{18} & \\ \boxed{x_2 = -\frac{43}{72}} & \end{array}$$

Logo, a solução do sistema linear dado é $x_1 = \frac{647}{36}$, $x_2 = -\frac{43}{72}$ e $x_3 = -\frac{22}{9}$.

Assim como fizemos para o método LU é possível contar o número de operações envolvidas na resolução de um sistema linear utilizando a fatoração de Cholesky. Procedendo de forma análoga, concluímos que, neste caso, o número de operações envolvidas é $\frac{n^3}{6} + n^2 - \frac{7n}{6}$ Adições/Subtrações e $\frac{n^3}{6} + \frac{3n^2}{2} - \frac{2n}{3}$ Multiplicações/Divisões. Dizemos que esse número é assintoticamente de ordem $n^3/3$ e, portanto, aproximadamente metade do custo computacional da resolução utilizando a fatoração LU. Cumpre ressaltar, que ambos procedimentos são de ordem $O(n^3)$.

3.3 Pivoteamento Parcial

Nesta seção estudaremos um pouco sobre como um erro de arredondamento pode distorcer a solução de um sistema linear, e veremos que uma forma de minorar essa distorção é usar o **pivoteamento parcial**, muito importante quando na eliminação gaussiana os pivôs são pequenos se comparados com outros elementos da mesma coluna da matriz ampliada. Resolver sistemas lineares desse tipo usando arredondamento pode produzir soluções bastante ruins.

O pivoteamento parcial consiste em se trocar as linhas da matriz ampliada, escolhendo como pivô o maior elemento da coluna. Em outras palavras, compara-se todos os elementos da coluna que contém o pivô e toma-se o maior, efetuando uma troca de linhas. Isso contribui bastante para diminuir os erros de arredondamento, produzindo, portanto, soluções melhores.

Antes de exemplificarmos o uso do pivoteamento parcial na resolução de um sistema linear, falaremos sobre a representação de um número em **ponto fixo** e em **ponto flutuante**. Considere as representações $-40, 19$ e $-0,4019 \times 10^2$. Essas duas maneiras de escrever um mesmo número são chamadas, respectivamente, de *representação em ponto fixo* e *representação em ponto flutuante*. O armazenamento dos números em calculadoras e computadores é feito usando a representação em ponto flutuante. Numa representação em ponto flutuante constam o *sinal*, a *parte fracionária* (chamada *mantissa* ou *significando*), a *base* e o *expoente*. Desse modo, em $-0,4019 \times 10^2$ temos sinal negativo, mantissa 0,4019, base 10 e expoente 2. A maneira geral de representar um número em ponto flutuante é

$$\pm \cdot d_1 d_2 \dots d_n \times \beta^k$$

onde d_i , para $i = 1, 2, \dots, n$, é um dígito entre 0 e $\beta - 1$ e $d_1 \neq 0$, β é a base tomada e k o expoente. Assim, se tomarmos a base 10 e chamarmos a mantissa de M (ou seja, $M = 0, d_1 d_2 \dots d_n$), então $0,1 \leq M < 1$. Geralmente as bases consideradas em computadores são 2, 10 ou 16. Quando n é o número máximo de dígitos armazenados dizemos que existem n *dígitos significativos*.

O número de dígitos n é limitado e daí surgem os erros de arredondamento. Por exemplo, algumas calculadoras armazenam 8 ou 12 dígitos, desprezando os dígitos seguintes (*truncamento*) ou fazendo um *arredondamento*. Os computadores têm maior capacidade de armazenamento, porém ainda limitada. Para saber mais sobre operações com números representados em ponto flutuante consulte [4], p. 21 a 23.

No exemplo seguinte veremos o efeito de um erro de arredondamento na solução de um sistema linear resolvido por eliminação gaussiana sem e com pivoteamento parcial.

Exemplo: Considere o sistema linear a seguir:

$$\begin{cases} 0,001x + 0,995y = 1,00 \\ -10,2x + 1,00y = -50,0 \end{cases}$$

Usando a notação em ponto flutuante com 3 dígitos significativos, resolva-o, por eliminação gaussiana,

(a) sem fazer uso do pivoteamento parcial.

(b) com o pivoteamento parcial.

Resolução de (a): Construindo a matriz ampliada do sistema com representação em ponto flutuante e procedendo ao escalonamento.

$$\left[\begin{array}{cc|c} 0,100 \times 10^{-2} & 0,995 \times 10^0 & 0,100 \times 10^1 \\ -0,102 \times 10^2 & 0,100 \times 10^1 & -0,500 \times 10^2 \end{array} \right] L_2 + \frac{0,102 \times 10^2}{0,100 \times 10^{-2}} L_1$$

Fazendo os cálculos:

$$\frac{0,102 \times 10^2}{0,100 \times 10^{-2}} = 1,02 \times 10^4 = 0,102 \times 10^5 \text{ (multiplicador)}$$

$$0,100 \times 10^1 + (0,102 \times 10^5)(0,995 \times 10^0) \approx$$

$$0,100 \times 10^1 + 0,101 \times 10^5 \approx$$

$$0,000001 \times 10^5 + 0,101 \times 10^5 \approx$$

$$0,101 \times 10^5$$

$$-0,500 \times 10^2 + (0,102 \times 10^5)(0,100 \times 10^1) \approx$$

$$-0,500 \times 10^2 + 0,0102 \times 10^6 \approx$$

$$-0,500 \times 10^2 + 0,102 \times 10^5 \approx$$

$$-0,0005 \times 10^5 + 0,102 \times 10^5 \approx$$

$$0,102 \times 10^5$$

Assim, a matriz escalonada do sistema é

$$\left[\begin{array}{cc|c} 0,100 \times 10^{-2} & 0,995 \times 10^0 & 0,100 \times 10^1 \\ 0 & 0,101 \times 10^5 & 0,102 \times 10^5 \end{array} \right]$$

Resultando que

$$y \approx \frac{0,102 \times 10^5}{0,101 \times 10^5} \approx 0,101 \times 10^1$$

e

$$x \approx \frac{0,100 \times 10^1 - (0,995 \times 10^0)(0,101 \times 10^1)}{0,100 \times 10^{-2}} \approx \frac{0,100 \times 10^1 - 0,100 \times 10^1}{0,100 \times 10^{-2}} \approx 0$$

Obtemos, portanto, como solução $x = 0$ e $y = 1,01$.

Resolução de (b): Escrevendo a matriz ampliada do sistema usando representação em ponto flutuante, com pivoteamento parcial executado, e procedendo ao escalonamento.

$$\left[\begin{array}{cc|c} -0,102 \times 10^2 & 0,100 \times 10^1 & -0,500 \times 10^2 \\ 0,100 \times 10^{-2} & 0,995 \times 10^0 & 0,100 \times 10^1 \end{array} \right] L_2 - \frac{0,100 \times 10^{-2}}{(-0,102 \times 10^2)} L_1$$

Fazendo os cálculos:

$$\frac{0,100 \times 10^{-2}}{0,102 \times 10^2} = 0,980 \times 10^{-4} \text{ (multiplicador)}$$

$$0,995 \times 10^0 + (0,980 \times 10^{-4})(0,100 \times 10^1) \approx$$

$$0,995 \times 10^0 + 0,098 \times 10^{-3} \approx$$

$$0,995 \times 10^0 + 0,980 \times 10^{-4} \approx$$

$$0,995 \times 10^0 + 0,000098 \times 10^0 \approx 0,995 \times 10^0$$

$$0,100 \times 10^1 + (0,980 \times 10^{-4})(-0,500 \times 10^2) \approx$$

$$0,100 \times 10^1 - 0,490 \times 10^{-2} \approx$$

$$0,100 \times 10^1 - 0,00049 \times 10^1 \approx$$

$$0,100 \times 10^1$$

Então a matriz escalonada do sistema, com pivoteamento parcial, é:

$$\left[\begin{array}{cc|c} -0,102 \times 10^2 & 0,100 \times 10^1 & -0,500 \times 10^2 \\ 0 & 0,980 \times 10^{-4} & 0,100 \times 10^1 \end{array} \right]$$

Assim, temos que:

$$y \approx \frac{0,100 \times 10^1}{0,995 \times 10^0} \approx 0,100 \times 10^1$$

e

$$x \approx \frac{-0,500 \times 10^2 - (0,100 \times 10^1)(0,100 \times 10^1)}{-0,102 \times 10^2} \approx \frac{-0,510 \times 10^2}{-0,102 \times 10^2} \approx 5$$

Portanto, obtemos como solução $x = 5$ e $y = 1$, que é a solução exata do sistema.

Ressaltamos que nem sempre a solução obtida por pivoteamento parcial é exatamente a solução do sistema, como no caso exemplificado. Muitas vezes ela oferece apenas uma aproximação melhor do que se desprezássemos o pivoteamento.

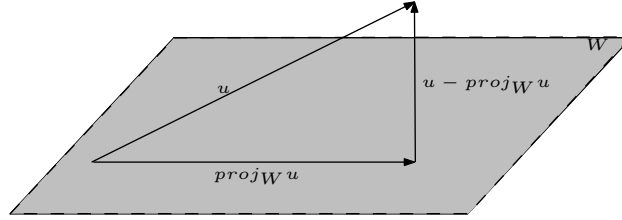
Existe um resultado garantindo que se os elementos da diagonal principal, em módulo, de uma matriz quadrada A são maiores que outros da coluna em que se encontram, então não é preciso fazer o pivoteamento parcial na eliminação gaussiana do sistema $Ax = b$, onde x e b são, respectivamente, os vetores-coluna das incógnitas e dos termos independentes. Esse resultado, bem como sua prova podem ser encontrados em [3], p. 346 e 347.

3.4 Mínimos Quadrados e a Solução de Sistemas Lineares

3.4.1 Projeções Ortogonais e Aproximações

No capítulo anterior tratamos da projeção ortogonal com o objetivo de, usando o processo de Gram-Schmidt, construir bases ortonormais. Agora o tratamento que daremos às projeções ortogonais está diretamente relacionado à minimização de distância. E mais adiante veremos o quanto isso pode ser importante para encontrar a melhor aproximação da solução de um sistema impossível, ou mesmo, a solução exata, para sistemas possíveis e determinados.

Novamente consideremos um espaço vetorial V munido de produto interno e um subespaço de dimensão finita W de V . Vimos que se $u \in V$ podemos escrever $u = \text{proj}_W u + (u - \text{proj}_W u)$. Ilustrando esse fato para $V = \mathbb{R}^3$ e W um plano que passa pela origem temos



Intuitivamente enxergamos o vetor $proj_W u$ como o vetor de W que está à menor distância de u , isto é, a menor distância entre u e o subespaço W pode ser dada por $\|u - proj_W u\|$. Essa idéia intuitiva é confirmada pelo teorema a seguir, onde o vetor $proj_W u$ é chamado de *melhor aproximação de u em W* .

Teorema 3.1 (Teorema da Melhor Aproximação) *Se W é um subespaço de dimensão finita de um espaço V com produto interno e u é um vetor de V , então $proj_W u$ é a melhor aproximação para u em W .*

Prova: Para provar esse teorema basta mostrarmos que para qualquer vetor $w \in W$ tal que $w \neq proj_W u$ então $\|u - proj_W u\| < \|u - w\|$, isto é, a distância entre os vetores u e sua projeção ortogonal $proj_W u$ é menor que a distância entre u e qualquer outro vetor w de W . Então tomemos $w \in W$ tal que $w \neq proj_W u$. Podemos escrever o vetor diferença $u - w$ da forma

$$u - w = (u - proj_W u) + (proj_W u - w)$$

Como o vetor $(proj_W u - w) \in W$, pois é a diferença entre dois vetores de W , e $(u - proj_W u)$ é ortogonal a W , então $(proj_W u - w)$ e $(u - proj_W u)$ são ortogonais entre si. Desse modo, segue pelo Teorema 2.12 (Teorema de Pitágoras) que

$$\|u - w\|^2 = \|u - proj_W u\|^2 + \|proj_W u - w\|^2$$

Por hipótese $w \neq proj_W u$, então $\|proj_W u - w\| > 0$, implicando que

$$\|u - w\|^2 > \|u - proj_W u\|^2$$

ou, equivalentemente,

$$\|u - w\| > \|u - proj_W u\|$$

■

Agora vamos relacionar a idéia da projeção ortogonal à solução de um sistema linear $Ax = b$. O resultado que faz essa relação é especialmente importante para determinar a melhor aproximação da solução de um sistema impossível $Ax = b$. No caso de um sistema possível $Ax = b$ temos que $Ax - b = 0$, ou ainda, $\|Ax - b\| = 0$.

O que queremos, no caso de um sistema impossível, é determinar qual vetor $x \in \mathbb{R}^n$ que minimiza $\|Ax - b\|$. Tal solução é conhecida como *solução por mínimos quadrados* do sistema $Ax = b$, e deriva do fato de que resolver $\min_{x \in \mathbb{R}^n} \|Ax - b\|$ é equivalente a $\min_{x \in \mathbb{R}^n} \|Ax - b\|^2$.

Vamos verificar que este problema de otimização é equivalente a um sistema linear denominado sistema de equações normais.

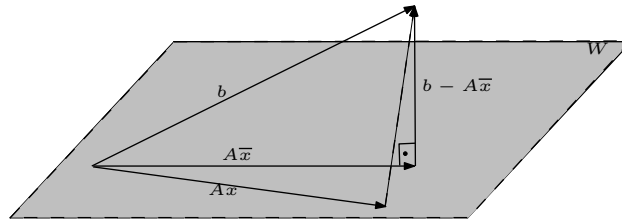
Definição 3.2 *Seja $Ax = b$ um sistema linear com m equações e n incógnitas. Definimos sistema de equações normais ao sistema dado por $A^t Ax = A^t b$.*

Teorema 3.2 *Sejam as matrizes $A \in \mathbb{R}^{m \times n}$ e $b \in \mathbb{R}^m$, então o sistema $Ax = b$ tem pelo menos uma solução $\bar{x}$ por mínimos quadrados e, além disso, o vetor $\bar{x}$ é uma solução por mínimos quadrados de $Ax = b$ se, e somente se, $\bar{x}$ é uma solução do sistema de equações normais $A^t Ax = A^t b$.*

Prova: Seja W o espaço-coluna de $A \in \mathbb{R}^{m \times n}$. Como para cada vetor $x \in \mathbb{R}^n$ o vetor Ax é uma combinação linear dos vetores-coluna de A , então ao tomarmos x variando sobre $\mathbb{R}^n$ estamos variando as combinações lineares Ax dos vetores-coluna de A . Desse modo, se x percorre todo $\mathbb{R}^n$ então Ax percorre todo o espaço-coluna W de A . Com isso, encontrar a solução por mínimos quadrados de $Ax = b$ é equivalente a encontrar $\bar{x} \in \mathbb{R}^n$ tal que $A\bar{x}$ é o vetor de W que melhor se aproxima de b , ou ainda, se $\bar{x} \in \mathbb{R}^n$ é a solução por mínimos quadrados de $Ax = b$, então

$$\|b - A\bar{x}\| \leq \|b - Ax\|$$

para todo $x \in \mathbb{R}^n$. Veja a interpretação geométrica dessa equivalência.



Conforme Teorema da Melhor Aproximação o vetor $A\bar{x}$ é dado pela projeção ortogonal de b sobre W , isto é,

$$A\bar{x} = \text{proj}_W b$$

Como $\text{proj}_W b$ sempre existe então $Ax = b$ tem pelo menos uma solução por mínimos quadrados (provando a existência da solução por mínimos quadrados).

Considerando que

$$b - A\bar{x} = b - \text{proj}_W b$$

e que $b - \text{proj}_W b$ é ortogonal a W e, por conseguinte, é ortogonal a todos os vetores-coluna de A , que geram W , então se a_i é um vetor-coluna de A segue que

$$a_i^t(b - A\bar{x}) = \langle a_i^t, b - A\bar{x} \rangle = 0$$

o que é equivalente a

$$\begin{aligned} A^t(b - A\bar{x}) &= \begin{bmatrix} a_1 & \dots & a_n \end{bmatrix}^t (b - A\bar{x}) \\ &= \begin{bmatrix} a_1^t \\ \vdots \\ a_n^t \end{bmatrix} (b - A\bar{x}) = \begin{bmatrix} a_1^t(b - A\bar{x}) \\ \vdots \\ a_n^t(b - A\bar{x}) \end{bmatrix} = \begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix} \end{aligned}$$

ou, resumidamente,

$$A^t(b - A\bar{x}) = 0$$

$$A^t b - A^t A\bar{x} = 0$$

$$A^t A\bar{x} = A^t b$$

■

Para o próximo teorema precisamos da definição de matriz semi-definida positiva, a qual fazemos abaixo.

Definição 3.3 Chamamos de matriz **semi-definida positiva** a uma matriz A tal que $x^t A x \geq 0$ para todo x .

Teorema 3.3 Se $A \in \mathbb{R}^{m \times n}$ então $A^t A$ é uma matriz simétrica semi-definida positiva e $A^t A$ é definida positiva se, e somente se, A tem colunas linearmente independentes.

Prova: A prova de que $A^t A$ é simétrica segue do fato de

$$(A^t A)^t = A^t (A^t)^t = A^t A$$

Agora provaremos que $A^t A$ é semi-definida positiva. Seja $x \in \mathbb{R}^n$ um vetor não-nulo, denotando $Ax = y$ temos

$$x^t A^t A x = (Ax)^t (Ax) = y^t y = \|y\|^2 \geq 0$$

provando que $A^t A$ é semi-definida positiva.

Considerando que

$$\|y\| = 0 \quad \Longleftrightarrow \quad y = 0 \quad \Longleftrightarrow \quad Ax = 0$$

então, denotando x_i como o i -ésimo elemento de x e a_j como o j -ésimo vetor-coluna de A segue que

$$Ax = 0 \quad \Longleftrightarrow \quad x_1 a_1 + x_2 a_2 + \dots + x_n a_n = 0$$

Mas a combinação linear de vetores linearmente independentes só resulta no vetor nulo se, e somente se, os coeficientes x_i são todos nulos. Como por hipótese $x \neq 0$ então $x^t A^t A x = 0$ se, e somente se, os vetores-coluna de A são linearmente dependentes, ou, equivalentemente, $x^t A^t A x > 0$ se, e somente se, os vetores-coluna de A são linearmente independentes. ■

Como vimos encontrar uma solução por mínimos quadrados de um sistema linear $Ax = b$ é o mesmo que encontrar a solução do sistema de equações normais $A^t A x = A^t b$. A seguir vamos mostrar como obter uma solução por mínimos quadrados do sistema linear $Ax = b$ usando as fatorações de Cholesky e QR .

Resolução de sistemas lineares por mínimos quadrados usando Cholesky

Conforme Teorema 3.3 a matriz $A^t A$ é simétrica definida positiva se os vetores-coluna de A são linearmente independentes. Então, supondo A com vetores-coluna linearmente independentes, podemos obter a fatoração de Cholesky $A^t A = GG^t$. Assim, o sistema que fornece solução por mínimos quadrados para $Ax = b$ pode ser dado por

$$A^t A x = A^t b \quad \Longleftrightarrow \quad GG^t x = A^t b$$

cuja solução pode ser obtida resolvendo-se dois sistemas triangulares: $Gy = A^t b$ (onde y é uma matriz-coluna de incógnitas) e, em seguida, $G^t x = y$.

Exemplo 1: Obtenha a solução por mínimos quadrados do sistema linear $Ax = b$ indicado abaixo:

$$\begin{cases} 3x_1 + 2x_2 = 13 \\ -4x_1 + 9x_2 = 76 \end{cases}$$

Solução: A matriz dos coeficientes de $Ax = b$ e a matriz dos coeficientes do sistema de equações normais $A^t Ax = A^t b$ são dadas, respectivamente, por

$$A = \begin{bmatrix} 3 & 2 \\ -4 & 9 \end{bmatrix} \quad A^t A = \begin{bmatrix} 25 & -30 \\ -30 & 85 \end{bmatrix}$$

Como $A^t A$ é simétrica definida positiva, então $A^t A$ tem decomposição de Cholesky GG^t :

$$A^t A = \begin{bmatrix} 25 & 42 \\ 42 & 85 \end{bmatrix} = \begin{bmatrix} 5 & 0 \\ -6 & 7 \end{bmatrix} \cdot \begin{bmatrix} 5 & -6 \\ 0 & 7 \end{bmatrix} = GG^t$$

Assim, escrevemos $A^t Ax = A^t b$ da forma $GG^t x = A^t b$. Então, construindo uma matriz-coluna de incógnitas y tal que $G^t x = y$, resolvemos o sistema triangular $Gy = A^t b$ para obter y :

$$\begin{bmatrix} 5 & 0 \\ -6 & 7 \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \begin{bmatrix} 3 & -4 \\ 2 & 9 \end{bmatrix} \cdot \begin{bmatrix} 13 \\ 76 \end{bmatrix} \iff \begin{bmatrix} 5 & 0 \\ -6 & 7 \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \begin{bmatrix} -265 \\ 710 \end{bmatrix}$$

cuja resolução fornece $y_1 = -53$ e $y_2 = 56$.

Agora o sistema $G^t x = y$ pode ser expresso da forma

$$\begin{bmatrix} 5 & -6 \\ 0 & 7 \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} -53 \\ 56 \end{bmatrix}$$

cuja solução é $x_1 = -1$ e $x_2 = 8$.

É fácil verificar que a solução por mínimos quadrados do sistema $Ax = b$ acima é a mesma que a obtida se tivéssemos resolvido $Ax = b$, por exemplo, por eliminação gaussiana. Isso acontece porque esse sistema é possível. No próximo exemplo tomamos um sistema impossível

e encontraremos sua solução por mínimos quadrados, ou seja, o vetor do espaço-coluna da matriz dos coeficientes que mais se aproxima do vetor dado pelos termos independentes.

Exemplo 2: Obtenha a solução por mínimos quadrados do sistema linear $Bx = b_1$, dado por

$$\begin{cases} 2x_1 + 3x_2 = 5 \\ 9x_1 + 4x_2 = 51 \\ 6x_1 - x_2 = 40 \end{cases}$$

Solução: As matrizes dos coeficientes de $Bx = b_1$ e do sistema de equações normais $B^t Bx = B^t b_1$ são, respectivamente,

$$B = \begin{bmatrix} 2 & 3 \\ 9 & 4 \\ 6 & -1 \end{bmatrix} \quad B^t B = \begin{bmatrix} 121 & 36 \\ 36 & 26 \end{bmatrix}$$

A decomposição de Cholesky da matriz simétrica definida positiva $B^t B$ é

$$B^t B = \begin{bmatrix} 121 & 36 \\ 36 & 26 \end{bmatrix} = \begin{bmatrix} 11 & 0 \\ \frac{36}{11} & \frac{\sqrt{1850}}{11} \end{bmatrix} \cdot \begin{bmatrix} 11 & \frac{36}{11} \\ 0 & \frac{\sqrt{1850}}{11} \end{bmatrix} = G_1 G_1^t$$

Escrevendo $B^t Bx = B^t b_1$ da forma $G_1 G_1^t x = B^t b_1$, podemos construir uma matriz-coluna de incógnitas y tal que $G_1^t x = y$ e, assim, resolvemos o sistema triangular $G_1 y = B^t b_1$ para obter y :

$$\begin{bmatrix} 11 & 0 \\ \frac{36}{11} & \frac{\sqrt{1850}}{11} \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \begin{bmatrix} 2 & 9 & 6 \\ 3 & 4 & -1 \end{bmatrix} \cdot \begin{bmatrix} 5 \\ 51 \\ 40 \end{bmatrix} \iff \begin{bmatrix} 11 & 0 \\ \frac{36}{11} & \frac{\sqrt{1850}}{11} \end{bmatrix} \cdot \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \begin{bmatrix} 709 \\ 179 \end{bmatrix}$$

donde obtemos $y_1 = \frac{709}{11}$ e $y_2 = -\frac{3865}{11\sqrt{1850}}$.

Agora o sistema $G_1^t x = y$ pode ser expresso da forma

$$\begin{bmatrix} 11 & \frac{36}{11} \\ 0 & \frac{\sqrt{1850}}{11} \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} \frac{709}{11} \\ -\frac{3865}{11\sqrt{1850}} \end{bmatrix}$$

cuja solução é $x_1 = \frac{1199}{185} \approx 6,4811$ e $x_2 = -\frac{773}{370} \approx -2,0892$.

Para se ter uma noção do erro cometido na aproximação da solução de um sistema, costumamos calcular $A\bar{x} - b$, onde $\bar{x}$ é a solução por mínimos quadrados de $Ax = b$. À diferença $A\bar{x} - b$ chamamos **resíduo** ou **erro** da solução por mínimos quadrados. Assim, o resíduo da solução por mínimos quadrados do sistema $Bx = b_1$ dado acima é obtida ao se fazer

$$B\bar{x} - b = \begin{bmatrix} 2 & 3 \\ 9 & 4 \\ 6 & -1 \end{bmatrix} \cdot \begin{bmatrix} \frac{1199}{185} \\ -\frac{773}{370} \end{bmatrix} - \begin{bmatrix} 5 \\ 51 \\ 40 \end{bmatrix} \approx \begin{bmatrix} 1,69459 \\ -1,02703 \\ 0,97568 \end{bmatrix}$$

Se calcularmos a norma $\|B\bar{x} - b\|$ veremos a que distância do vetor b está o vetor-solução $\bar{x}$, encontrado por mínimos quadrados. Assim, se a norma é não-nula significa que o sistema é impossível, ou que erros de arredondamento o perturbaram gerando apenas uma aproximação da solução. No caso de sistema impossível dizemos ainda que o vetor b não está no subespaço gerado pelos vetores-coluna da matriz dos coeficientes.

No Exemplo 1 temos um sistema possível e determinado, então o vetor que indica o erro cometido é nulo, pois o vetor b pertence ao subespaço gerado pelas colunas de A . Assim, a norma $\|A\bar{x} - b\| = 0$.

Observamos que o produto $A^t A$ pode gerar uma matriz muito mal-condicionada, ou seja, embora tenhamos garantia teórica de que $A^t A$ seja simétrica definida positiva, na prática, essa matriz pode trazer alguns problemas numéricos.

3.4.2 Resolução de sistemas lineares por mínimos quadrados usando QR

Já sabemos que encontrar uma solução por mínimos quadrados de $Ax = b$ é o mesmo que encontrar a solução de $A^t Ax = A^t b$. Supondo que A tem colunas linearmente independentes, então A tem fatoração QR . Assim,

$$\begin{aligned} A^t Ax &= A^t b \\ (QR)^t QRx &= (QR)^t b \\ R^t Q^t QRx &= R^t Q^t b \\ R^t Rx &= R^t Q^t b \end{aligned}$$

$$R^{-t}R^tRx = R^{-t}R^tQ^tb$$

$$Rx = Q^tb$$

isto é, a solução por mínimos quadrados de $Ax = b$ pode ser obtida resolvendo um sistema triangular. Uma observação importante é que resolver um sistema linear usando a fatoração QR nos fornece necessariamente a solução por mínimos quadrados, pois se A tem colunas linearmente independentes, então A pode ser fatorada como $A = QR$ e, portanto,

$$Ax = b \Leftrightarrow (QR)x = b \Leftrightarrow Q^tQRx = Q^tb \Leftrightarrow Rx = Q^tb$$

isto é, resolver $Ax = b$ simplesmente substituindo A por QR já garante que a solução obtida é a mesma de $A^tAx = A^tb$ e, portanto, a solução por mínimos quadrados. Isso significa que um sistema $Ax = b$, em que A tem colunas linearmente independentes, sempre terá solução se usada a fatoração QR , a qual será a solução por mínimos quadrados.

Exemplo: Considere o sistema linear $Ax = b$ seguinte:

$$\begin{cases} x_1 + x_2 + x_3 = 8 \\ x_1 - x_2 + 2x_3 = 17 \\ -x_1 + x_2 = -5 \\ x_1 + 5x_2 + x_3 = 6 \end{cases}$$

Solução: A matriz dos coeficientes A tem colunas linearmente independentes, então pode ser expressa pela fatoração QR . No capítulo anterior, na seção que trata da decomposição QR obtemos a fatoração QR dessa matriz, então vamos apenas reescrevê-la:

$$A = \begin{bmatrix} 1 & 1 & 1 \\ 1 & -1 & 2 \\ -1 & 1 & 0 \\ 1 & 5 & 1 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & 0 & 0 \\ \frac{1}{2} & -\frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{2} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{2}{\sqrt{6}} & 0 \end{bmatrix} \cdot \begin{bmatrix} 2 & 2 & 2 \\ 0 & 2\sqrt{6} & 0 \\ 0 & 0 & \sqrt{2} \end{bmatrix} = QR$$

Assim $Rx = Q^tb$ pode ser escrito como

$$\begin{bmatrix} 2 & 2 & 2 \\ 0 & 2\sqrt{6} & 0 \\ 0 & 0 & \sqrt{2} \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & \frac{1}{2} & -\frac{1}{2} & \frac{1}{2} \\ 0 & -\frac{1}{\sqrt{6}} & \frac{1}{\sqrt{6}} & \frac{2}{\sqrt{6}} \\ 0 & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 \end{bmatrix} \cdot \begin{bmatrix} 8 \\ 17 \\ -5 \\ 6 \end{bmatrix}$$

ou ainda, fazendo o produto do segundo membro, como

$$\begin{bmatrix} 2 & 2 & 2 \\ 0 & 2\sqrt{6} & 0 \\ 0 & 0 & \sqrt{2} \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 18 \\ -\frac{5}{3}\sqrt{6} \\ 6\sqrt{2} \end{bmatrix}$$

resultando que $x_1 = \frac{23}{6}$, $x_2 = -\frac{5}{6}$ e $x_3 = 6$.

Agora vamos calcular o vetor erro $A\bar{x} - b$, cometido na aproximação da solução desse sistema:

$$\begin{bmatrix} 1 & 1 & 1 \\ 1 & -1 & 2 \\ -1 & 1 & 0 \\ 1 & 5 & 1 \end{bmatrix} \cdot \begin{bmatrix} \frac{23}{6} \\ -\frac{5}{6} \\ 6 \end{bmatrix} - \begin{bmatrix} 8 \\ 17 \\ -5 \\ 6 \end{bmatrix} \approx \begin{bmatrix} 9 \\ 16,666 \\ -4,666 \\ 5,666 \end{bmatrix} - \begin{bmatrix} 8 \\ 17 \\ -5 \\ 6 \end{bmatrix} \approx \begin{bmatrix} 1 \\ -0,333 \\ 0,333 \\ -0,333 \end{bmatrix}$$

O sistema que tomamos no exemplo acima é impossível, pois não efetuamos arredondamentos na sua resolução e ainda assim obtivemos vetor erro não-nulo. Indicando por $\bar{x}$ o vetor solução obtido, temos que a distância do vetor b do espaço-coluna de A é $\|A\bar{x} - b\| \approx 1,1547$, isto é, o vetor b dista da solução por mínimos quadrados em aproximadamente 1,1547.

Agora faremos alguns comentários relacionando a existência e a unicidade da solução por mínimos quadrados de um sistema linear $Ax = b$ com a dependência e a independência linear dos vetores-coluna de A . Vimos que uma solução por mínimos quadrados sempre existe. Ela será única quando os vetores-coluna de A forem linearmente independentes, pois nesse caso, se b não pertence ao espaço-coluna de A então existe um único vetor $A\bar{x}$ que é projeção ortogonal de b sobre o espaço-coluna de A , e como $A\bar{x}$ é combinação linear dos vetores-coluna de A , que são linearmente independentes, então, denotando o i -ésimo vetor-coluna de A por A_i , isto é, $A = [A_1 \ A_2 \ \dots \ A_n]$ para A de ordem $m \times n$, e $\bar{x} = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n)$ temos pela fórmula 2.1, p. 50, que $A\bar{x} = \bar{x}_1 A_1 + \dots + \bar{x}_n A_n$. Podemos afirmar que se um vetor v pode ser representado como combinação linear de um conjunto de vetores linearmente independentes $v_1, \dots, v_n$, então essa representação é única. De fato, pois se existem escalares c_i e k_i tais que

$$v = c_1 v_1 + \dots + c_n v_n \quad (I)$$

e

$$v = k_1 v_1 + \dots + k_n v_n \quad (II)$$

então fazendo $(I) - (II)$ temos

$$0 = (c_1 - k_1) v_1 + \dots + (c_n - k_n) v_n.$$

Como $v_1, \dots, v_n$ são linearmente independentes segue que $c_1 - k_1 = 0, \dots, c_n - k_n = 0$, ou seja, $c_1 = k_1, \dots, c_n = k_n$. Daí concluímos que a solução $\bar{x} = (\bar{x}_1, \dots, \bar{x}_n)$ é única, pois a representação de $A\bar{x}$ como combinação linear de vetores-coluna linearmente independentes $A_1, \dots, A_n$ é única.

No caso em que os vetores-coluna de A são linearmente dependentes e b não pertence ao espaço-coluna de A não podemos usar o recurso das fatorações de Cholesky ou QR , como apresentado neste texto, para obter uma solução por mínimos quadrados. Neste caso o sistema de equações normais $A^t A x = A^t b$ possui infinitas soluções, pois a matriz $A^t A$ é não-invertível ([1], p. 225), ou seja, o sistema $A^t A x = A^t b$ é possível e indeterminado.

3.5 Implementação de métodos para resolver sistemas lineares usando o OCTAVE

Nesta seção apresentamos alguns algoritmos para resolução de sistemas lineares $Ax = b$ usando as fatorações LU , Cholesky e QR , que foram implementados no software livre Octave.

O Octave permite programação numa linguagem de alto nível, interpretada, voltada para computações numéricas. Ele é um software com distribuição GNU e apresenta grande compatibilidade com o Matlab. Também oferece uma interface adequada para utilização da notação matricial para resolver problemas numéricos.

A seguir apresentaremos a descrição de alguns comandos e notações que aparecerão nos algoritmos. Para uma referência mais completa sobre o Octave consulte a página <http://www.gnu.org/software/octave/>.

- $[m, n] = \text{size}(A)$ significa tamanho e nos algoritmos se refere às dimensões da matriz. Assim, $[m, n] = \text{size}(A)$ significa que a matriz A tem m linhas e n colunas.
- $A(i, j)$ denotará o elemento da matriz A que se encontra na i -ésima linha e j -ésima coluna.

- “:” indicará que há uma lista ordenada contendo todos os inteiros entre o número que está antes dos *dois pontos* e o que está depois. Assim, $i=j:n$ deve ser interpretado como $i = j, j+1, \dots, n$, para $j < n$, e $i=n:-1:j$ deve ser interpretado como $i = n, n-1, \dots, j$.
- A' indicará a transposta de A , que no decorrer dos capítulos anteriores foi denotada como A^t .
- `max` significa máximo.
- `abs` significa valor absoluto.

O leitor poderá notar que alguns algoritmos que apresentaremos diferem ligeiramente de muitos dos seus equivalentes apresentados nos livros de análise numérica. Isso se deve ao fato de que, na prática, quando implementamos um algoritmo, procuramos aproveitar ao máximo os recursos da linguagem de programação utilizada. Só para exemplificar, o uso do comando `for`, no Octave, produz um custo de processamento considerável, enquanto que na linguagem Fortran é natural o uso deste comando sem prejuízos computacionais. Assim, sempre que possível evitamos o uso de laços do tipo `for` em nossos algoritmos.

Iniciamos apresentando dois algoritmos que serão utilizados para resolver sistemas triangulares.

Algoritmo para resolver um sistema triangular inferior

```
function [ x ] = trianinf(A,b)
    [m,n] = size(A);
    x = zeros(n,1);
    x(1) = b(1)/A(1,1);
    for i = 2 :n
        x(i)=(b(i)-A(i,[1:i-1])*x(1:i-1))/A(i,i);
    end
endfunction
```

Note que o comando

```
x(i)=(b(i)-A(i,[1:i-1])*x(1:i-1))/A(i,i)
```

3.5 Implementação de métodos para resolver sistemas lineares usando o OCTAVE 80

usado para obter os elementos x_i do vetor solução x , para $i=2:n$, é um produto interno. Uma alternativa seria usar

```
s=0
for j=1:i-1
    s=s+A(i,j)*x(j);
end
x(i)=(b(i)-s)/A(i,i);
```

o que seria computacionalmente mais caro, no Octave.

Algoritmo para resolver um sistema triangular superior

```
function [ x ] = triansup(A,b)
    [m,n] = size(A);
    x = zeros(n,1);
    x(n) = b(n)/A(n,n);
    for i = n-1: -1 :1
        x(i)=(b(i)-A(i,[i+1:n])*x(i+1:n))/A(i,i);
    end
endfunction
```

Algoritmo para fatorar $A = LU$

```
function [L,U,p] = novolu(A)
    [n,n] = size(A);
    p = (1:n)'; %constroi o vetor de permutacoes p
    for k = 1:n-1
        [r,m] = max(abs(A(k:n,k))); %escolhe como pivo o maior elemento,
        %em modulo, da coluna (pivoteamento)
        m = m+k-1;
        if A(m,k) ~=0
            if (m ~= k)
                A([k m],:) = A([m k],:); %troca linhas
                p([k m]) = p([m k]); %armazena a posicao das trocas
```

```

        end
        i = k+1:n;
        A(i,k) = A(i,k)/A(k,k); %constroi os multiplicadores
        A(i,i) = A(i,i) - A(i,k)*A(k,i);
    end
end
L = tril(A,-1) + eye(n,n); %constroi a matriz L colocando na identidade
%os elementos abaixo da diagonal de A
U = triu(A); %constroi uma matriz triangular superior com os respecti-
%vos elementos de A

```

Chamamos à atenção para o fato de que as matrizes L e U foram obtidas com as permutações necessárias nas linhas da matriz A . Cada vez que uma permutação é executada a nova matriz sobrepõe-se à matriz A antiga e o vetor p armazenará as trocas efetuadas. Assim, o que obtemos após a execução do programa é a fatoração $A(p, :)=LU$. Como veremos adiante o vetor p será importante na resolução do sistema para permutar adequadamente o vetor b de $Ax = b$.

Outra observação relevante sobre o algoritmo é que praticamente todo o tempo gasto na fatoração usando o programa `novolu` deve-se ao comando

```
A(i,i) = A(i,i) - A(i,k)*A(k,i)
```

que calcula os elementos da matriz U . O comando $A(i,k)*A(k,i)$ cria uma matriz quadrada de ordem $n - k$ que será subtraída da submatriz $A_{(k,k)}$. Esse procedimento possibilita a eliminação de um duplo `for`, resultando numa significativa economia de tempo.

Algoritmo para fatorar $A = GG^t$

```

function [ G ] = novocholesky(A)
[m,n] = size(A);
A(:,1)=A(:,1)/sqrt(A(1,1));
for j=2:n
    A(j:n,j)=A(j:n,j)-A(j:n,1:j-1)*A(j,1:j-1)';
    A(j:n,j)=A(j:n,j)/sqrt(A(j,j));
end
G = tril(A);

```

3.5 Implementação de métodos para resolver sistemas lineares usando o OCTAVE 82

Assim, como no algoritmo para a fatoração LU tomamos o cuidado de eliminar laços desnecessários com o comando *for*.

Algoritmo para fatorar $A = QR$

```
function [Q, R]=novoqr(A)
[m,n]=size(A);
Q = zeros(m,n);
Q(:,1)=(1/norm(A(:,1)))*A(:,1); %cria a primeira coluna da matriz Q
for j = 2:n;
    uj=A(:,j);
    i = 1:j-1;
    aa=(uj'*Q(:,i)); %calcula os produtos internos usados no
    %processo de Gram-Schmidt
    bb=Q(:,i);
    qj = uj - bb*aa'; %determina um vetor ortogonal as colunas ja
    %existentes de Q
    Q(:,j)=qj/norm(qj); %constroi a j-ésima coluna de Q
end
R = Q'*A;
```

A fatoração QR por este algoritmo realiza um total de $mn^2 - n^2 + 2mn + \frac{n}{2} - 3m$ operações. Assim, para $n \approx m$ o custo de realizar uma fatoração QR é ligeiramente superior às fatorações LU e Cholesky.

Algoritmo para resolver um sistema $LUx = b$

```
function [x]=resolvelu(A,b)
[L,U,p]=novolu(A);
y = trianinf(L,b(p));
x = triansup(U,y)';
```

Algoritmo para resolver um sistema $GG^t x = b$

```
function [x]=resolvechol(A,b)
[G]=novocholesky(A);
y = trianinf(G,b);
x = triansup(G',y)';
```

Algoritmo para resolver um sistema $QRx = b$

```
function [x]=resolveqr(A,b)
[Q, R]=novoqr(A);
x = triansup(R,Q'*b);
```

3.5.1 Algoritmos para Resolver Sistemas Lineares por Mínimos Quadrados

Algoritmo para resolver um sistema $QRx = b$

```
function [x]=resolveqr(A,b)
[Q, R]=novoqr(A);
x = triansup(R,Q'*b);
```

Algoritmo para resolver um sistema $GG^t x = A^t b$

```
function [x]=resolvecchol(A,b)
[G]=novocholesky(A'*A);
y = trianinf(G,A'*b);
x = triansup(G',y)';
```

Resultados Numéricos

A tabela 1 mostra, a título de exemplo, o tempo gasto na resolução de 3 sistemas lineares usando as fatorações LU, Cholesky e QR. As matrizes dos coeficientes, cujas dimensões são indicadas na primeira coluna, foram geradas aleatoriamente no Octave em duas etapas. Para cada matriz de coeficientes A , geramos uma matriz aleatória B com elementos inteiros entre -20 e 20 e posto completo. Em seguida construímos $A = B^t * B$ garantindo assim que A é simétrica definida positiva. Os tempos apresentados na tabela foram obtidos executando as implementações num computador Intel(R) Core(TM)2 CPU 6600 @2.40GHz 2 GB de RAM.

3.5 Implementação de métodos para resolver sistemas lineares usando o OCTAVE 84

$n \times n$	LU	Cholesky	QR
10×10	0,224s	0,047s	0,045s
100×100	0,286s	0,091s	0,085s
1000×1000	19,954s	4,674s	27,138s

Tabela 3.1: tempo decorrido, em segundos, na resolução de sistemas lineares $Ax=b$ usando programas implementados no Octave

Observamos que dentre os métodos utilizados, o que apresentou melhor desempenho é o que faz uso da fatoração de Cholesky, confirmando a expectativa gerada pela contagem de operações, e justificando porque tal método é melhor na prática, se pudermos usá-lo.

É possível notar que o tempo decorrido no processamento dos cálculos utilizando a fatoração LU está acima do esperado para as matrizes 10×10 e 100×100 se considerarmos o número de operações requeridas por esse método e pela fatoração de Cholesky. Investigamos possíveis motivos para essa inconsistência mas não obtivemos uma resposta concisa. O mais provável é que tal oscilação esteja relacionada à estrutura interna do próprio Octave. De qualquer forma, isso é uma evidência prática de que a escolha da linguagem e a forma de implementar o algoritmo influenciam bastante no custo de execução.

Note também que para pequenas dimensões o método usando QR se mostrou bastante eficiente. Contudo, para dimensões maiores, esse método não obteve bom desempenho, o que é compatível com o número de operações executadas pelo algoritmo.

Em muitos problemas práticos é importante determinar uma equação que ajuste um conjunto de dados obtidos experimentalmente. Problemas dessa natureza são convertidos em problemas de mínimos quadrados [7], p. 223, cuja matriz dos coeficientes tem a forma $m \times n$, onde $m > n$. Esse problema pode ser solucionado resolvendo um sistema de equações normais associado.

Na tabela abaixo apresentamos o tempo gasto na solução de alguns problemas de quadrados mínimos utilizando a fatoração QR e Cholesky.

Observamos que no caso de sistemas lineares com matrizes dos coeficientes $m \times n$ com $m \gg n$ a solução de quadrados mínimos via sistema de equações normais por fatoração de Cholesky é menos vantajosa que quando fazemos uso da fatoração QR . Porém, essa relação se inverte à medida que a matriz dos coeficientes vai se aproximando da forma quadrada. Neste caso, como já vimos, o método de Cholesky é mais eficiente. A escolha do melhor

$m \times n$	Cholesky	QR
10×2	0,181s	0,044s
100×20	0,187s	0,049s
100×90	0,216s	0,079s
1000×20	0,187s	0,057s
1000×500	1,069s	6,563s

Tabela 3.2: tempo decorrido, em segundos, na resolução por mínimos quadrados de sistemas lineares $Ax=b$ usando programas implementados no Octave

método para resolver um problema de quadrados mínimos é ainda um assunto em aberto e muitas variáveis estão envolvidas, conforme observado em [5], p. 224.

Ressaltamos que o tempo computado para a resolução via fatoração de Cholesky levou em conta as multiplicações necessárias para a construção do sistema de equações normais, enquanto que a solução por QR sempre fornece uma solução por mínimos quadrados, de modo que não foi preciso executar operações extras.

Considerações Finais

De uma maneira geral, procuramos apresentar alguns dos vários importantes resultados relacionados às matrizes, assim como as condições de existência e unicidade das fatorações estudadas, culminando na resolução de sistemas de equações lineares utilizando métodos diretos envolvendo essas fatorações e, mais ultimamente, na implementação desses métodos usando o software Octave. Nesse sentido, nos esforçamos por exibir a importância das fatorações na elaboração de métodos diretos para resolução de sistemas lineares, os quais se constituem como ferramentas úteis em implementações computacionais. Dentro dessa perspectiva, quisemos motivar a compreensão de como os algoritmos podem ser otimizados para a programação nos softwares.

Muito embora a análise dos algoritmos e o custo computacional não tenham sido o enfoque deste trabalho, a última seção do Capítulo 3 pode manifestar-se como bom incentivo à implementação computacional, à medida que oferece algoritmos prontos e programados num software livre.

Por fim, é natural algumas expectativas para estudos futuros. Por exemplo, a análise de sensibilidade em sistemas lineares e o estudo do comportamento do resíduo em função do condicionamento da matriz foram assuntos que, às vezes de forma periférica, permearam nossas pesquisas e que gostaríamos de investigar mais profundamente. Além disso, seria interessante comparar os métodos que expusemos com outros métodos diretos e também iterativos, como por exemplo, o método dos gradientes conjugados, explorando tanto a parte teórica quanto numérica.

Referências Bibliográficas

- [1] ANTON, H., RORRES, C. *Álgebra Linear com Aplicações*. 8. ed., Porto Alegre: Bookman, 2001.
- [2] BOLDRINI, J. L., COSTA, S. I. R., FIGUEIREDO, V. L., WETZLER, H. G. *Álgebra Linear*. 3. ed., São Paulo: Harbra, 1986.
- [3] BURDEN, R. L., FAIRES, J. D. *Análise Numérica*. São Paulo: Pioneira Thomson Learning, 2003.
- [4] CAMPOS, filho, F. F. *Algoritmos Numéricos*. Rio de Janeiro: LTC, 2001.
- [5] GOLUB, G. H., LOAN, C. F. V., *Matrix Computations*. 3. ed., The Johns Hopkins University Press, Baltimore, 1996.
- [6] LIMA, Elon Lages. *Álgebra Linear*. 7. ed., Rio de Janeiro: IMPA, 2004.
- [7] MEYER, C. D. *Matrix Analysis and Applied Linear Algebra*. SIAM, 2000.
- [8] NOBLE, B., DANIEL, J. W. *Applied Linear Algebra*. 3. ed., Prentice-Hall, 1988.
- [9] TREFETHEN, L. N., BAU, D. *Numerical Linear Algebra*. SIAM, Filadélfia, PA, 1997.
- [10] POOLE, D. *Álgebra Linear*. São Paulo: Pioneira Thomson Learning, 2004.
- [11] STEWART, G. W. *Afternotes on Numerical Analysis*. SIAM Publications, Filadélfia, PA, 1996.
- [12] WENDROFF, B. *Theoretical Numerical Analysis*. Academic Press, Nova Iorque, 1996.
- [13] OCTAVE, *Gnu Octave*. <http://www.gnu.org/software/octave/>. Acesso em 01/03/2008.