

# Aula 02 - Algoritmo de Thomas e Métodos Iterativos

Élton Fontana

Como visto na Aula 01, métodos de eliminação direta, como o método de Gauss, podem ser aplicados para a resolução de sistemas de equações algébricas lineares. Diversos outros métodos são muito utilizados, como os métodos de Gauss-Jordan, fatoração LU e da matriz inversa. Estes métodos estão baseados em processo de eliminação de determinados termos da matriz dos coeficientes através de operações algébricas que não alterem o sistema.

Apesar de não haver uma regra geral, os métodos de eliminação costumam ser aplicados quando as seguintes condições são satisfeitas: (a) o número de equações é pequeno (menos de 100), (b) a maioria dos coeficientes da matriz  $\mathbf{A}$  são não-nulos, (c) a matriz dos coeficientes não é predominantemente diagonal ou (d) o sistema de equações é mal condicionado (quando pequenas mudanças nos parâmetros de entrada causam grandes mudanças nos resultados obtidos).

A princípio, os métodos de eliminação poderiam ser aplicados para qualquer sistema. No entanto, existem alguns problemas que limitam a utilização destes métodos para certas classes de problemas, especialmente envolvendo um grande número de equações. Nestes casos, é necessário um número muito grande de operações para se obter a matriz escalonada, o que faz com que os erros de arredondamento sejam muito significativos.

De modo geral, a utilização de métodos iterativos é mais adequada para a resolução de sistemas com muitas equações onde não existe uma forma adequada de controlar o erro. No entanto, quando a matriz possui alguns formatos específicos, a utilização de métodos de eliminação próprios para cada formato pode permitir a resolução de forma simples, com um erro associado baixo e com um gasto mínimo de memória computacional. Em particular, são de especial interesse as matriz *tridiagonais*, pois estas surgem com frequência na resolução de EDP's e podem ser resolvidas facilmente com um método de eliminação chamado de algoritmo de Thomas, como será apresentado a seguir.

# 1 Matrizes Tridiagonais e o Algoritmo de Thomas

Uma matriz é dita tridiagonal quando possui uma largura de banda igual a 3, ou seja, somente a diagonal principal e os elementos vizinhos acima e abaixo são não-nulos. Este tipo de matriz surge naturalmente na resolução de EDP's através de métodos implícitos. De forma geral, um sistema linear tridiagonal  $n \times n$  pode ser expresso como:

$$\begin{bmatrix} a_{11} & a_{12} & 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ a_{21} & a_{22} & a_{23} & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & a_{32} & a_{33} & a_{34} & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & a_{43} & a_{44} & a_{4,5} & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \dots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & 0 & \dots & a_{n-1,n-2} & a_{n-1,n-1} & a_{n-1,n} \\ 0 & 0 & 0 & 0 & 0 & \dots & 0 & a_{n,n-1} & a_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ \vdots \\ x_{n-1} \\ x_n \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ b_3 \\ b_4 \\ \vdots \\ b_{n-1} \\ b_n \end{bmatrix}$$

Para resolver este tipo de sistema, pode-se utilizar uma versão simplificada do método de eliminação de Gauss conhecida como algoritmo de Thomas. Como todos os elementos da primeira coluna abaixo da segunda linha são nulos, o único elemento que precisa ser eliminado nesta coluna é  $a_{21}$ . Assim, pode-se fazer a seguinte operação elementar na segunda linha:

$$L2 \rightarrow L2 - (a_{21}/a_{11})L1$$

onde  $L2$  representa a linha 2 e  $L1$  a linha 1. Com isso, a linha 2 passa a ser escrita como:

$$\begin{bmatrix} 0 & a_{22} - (a_{21}/a_{11})a_{12} & a_{23} & 0 & 0 & \dots & 0 & 0 & 0 \end{bmatrix}$$

De forma similar, para deixar a matriz no formato triangular, na coluna 2 somente o termo  $a_{32}$  precisa ser eliminado, na coluna 3 somente o termo  $a_{43}$  e assim sucessivamente. Como visto no exemplo para a primeira coluna, somente o elemento da diagonal superior será alterado pelo processo de eliminação (o elemento acima não será afetado pois será descontado o valor de zero).

De maneira generalizada, os elementos da diagonal principal após a eliminação passam a ser avaliados como:

$$a'_{i,i} = a_{i,i} - (a_{i,i-1}/a_{i-1,i-1})a_{i-1,i} \quad (i = 2, 3, 4, \dots, n)$$

Como discutido na aula anterior, as operações elementares utilizadas no método de eliminação são aplicadas na matriz aumentada, levando em conta também a parte não-homogênea do sistema. Por isso, as operações que afetam a diagonal principal também irão afetar o vetor  $\mathbf{b}$ . Assim, os elementos deste vetor passam a ser avaliados como:

$$b'_i = b_i - (a_{i,i-1}/a_{i-1,i-1})b_{i-1} \quad (i = 2, 3, 4, \dots, n)$$

Com isso, pode-se construir uma matriz triangular  $\mathbf{A}'$  e o sistema linear pode ser reescrito como:

$$\begin{bmatrix} a_{11} & a_{12} & 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & a'_{22} & a_{23} & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & a'_{33} & a_{34} & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 0 & a'_{44} & a_{4,5} & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \dots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & 0 & \dots & 0 & a'_{n-1,n-1} & a_{n-1,n} \\ 0 & 0 & 0 & 0 & 0 & \dots & 0 & 0 & a'_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ \vdots \\ x_{n-1} \\ x_n \end{bmatrix} = \begin{bmatrix} b'_1 \\ b'_2 \\ b'_3 \\ b'_4 \\ \vdots \\ b'_{n-1} \\ b'_n \end{bmatrix}$$

A partir deste sistema, a solução  $\mathbf{x}$  pode ser avaliada, partindo do último elemento  $x_n$  e avançando até o primeiro:

$$x_n = b'_n/a'_{nn}$$

$$x_i = (b'_i - a'_{i,i+1}x_{i+1})/a'_{i,i} \quad i = n-1, n-2, n-3, \dots, n$$

Uma das principais vantagens do algoritmo de Thomas é sua facilidade de implementação. A estrutura do código pode ser dividida em duas partes:

(a) *Substituição:*

Para  $i = 2$  até  $n$  :

$$e_i = a_{i,i-1}/a_{i-1,i-1}$$

$$a_{i,i} = a_{i,i} - e_i a_{i-1,i}$$

$$b_i = b_i - e_i b_{i-1}$$

(b) *Resolução:*

$$x_n = b_n/a_{n,n}$$

Para  $i = n-1$  até  $1$  :

$$x_i = (b_i - a_{i,i+1}x_{i+1})/a_{i,i}$$

É importante observar que o contador na etapa (b) é *decrecente*, por isso, os valores  $x_{i+1}$  já são conhecidos quando  $x_i$  está sendo calculado.

**Exemplo 01:** Através da aplicação do método de diferenças finitas, obteve-se o seguinte conjunto de equações para avaliar a distribuição de temperatura  $T_i$  ao longo de um objeto:

$$\begin{aligned} T_1 &= 0 \\ T_{i-1} - (2 + \alpha^2 \Delta x^2) T_i + T_{i+1} &= 0 \quad i = 2, 3, 4 \\ T_5 &= 10 \end{aligned}$$

Considerando que  $\alpha = 4$  e  $\Delta x = 0.125$ , utilize o algoritmo de Thomas para encontrar os valores de  $T_i$ .<sup>1</sup>

$$R: T = (0 \quad 1.451 \quad 3.266 \quad 5.896 \quad 10)^T$$

## 2 Métodos Iterativos para Sistemas Lineares

O algoritmo de Thomas, apesar de muito simples, está restrito a matrizes estritamente tridiagonais. Quando o sistema não possui este formato e a utilização dos métodos de eliminação não é conveniente, costuma-se utilizar métodos iterativos para a resolução do sistema linear. Estes métodos são particularmente úteis para avaliar matrizes muito esparsas (com muitos elementos iguais a zero), pois neste caso os métodos tendem a convergir rapidamente.

Para iniciar os métodos iterativos, deve-se assumir uma solução inicial  $\mathbf{x}^0$  (chute inicial). Através de alguma estratégia específica de cada método, este vetor é então corrigido para um valor aprimorado  $\mathbf{x}^1$ . Este processo é então repetido (iterado) até que a diferença entre o vetor obtido e o anterior seja menor que um valor especificado. Este processo é convergente se cada iteração produz um valor que se aproxima cada vez mais da solução exata conforme o número de iterações aumenta. O número de iterações necessárias para se atingir um determinado critério de convergência estabelecido depende de vários fatores, podendo-se destacar:

- O formato da matriz dos coeficientes. Matrizes com superior dominância diagonal convergem mais rapidamente;

---

<sup>1</sup>As resoluções dos exemplos no Scilab são apresentadas como anexo, ao final do arquivo

- O método de iteração utilizado, ou seja, a forma como a solução inicial é corrigida;
- A solução inicial assumida. Quanto mais próximo da solução exata, mais rapidamente o método irá convergir;
- O critério de convergência estabelecido.

Dentre os métodos iterativos mais simples para a resolução de sistemas lineares, pode-se destacar os de Jacobi e Gauss-Siedel, que serão apresentados a seguir.

## 2.1 Método de Jacobi

O método de Jacobi é o método iterativo mais simples para a resolução de sistemas lineares. Apesar de apresentar uma convergência relativamente lenta, o método funciona bem para sistemas esparsos com grande dominância diagonal. No entanto, este método não funciona em todos os casos. Em particular, uma condição necessária é que todos os elementos da diagonal principal sejam não nulos. Quando algum elemento é nulo, pode-se eventualmente resolver o problema trocando-se a ordem das linhas (lembrando que esta é uma operação elementar que não altera o sistema).

Considere novamente o sistema linear da forma:

$$\begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix}$$

Assumindo que todos os termos  $a_{ii}$  são todos não-nulos, a solução deste sistema linear pode ser dada por:

$$\begin{aligned} x_1 &= \frac{1}{a_{11}}(b_1 - a_{12}x_2 - a_{13}x_3 - \dots - a_{1n}x_n) \\ x_2 &= \frac{1}{a_{22}}(b_2 - a_{21}x_1 - a_{23}x_3 - \dots - a_{2n}x_n) \\ &\vdots \\ x_n &= \frac{1}{a_{nn}}(b_n - a_{n1}x_1 - a_{n2}x_2 - \dots - a_{n,n-1}x_{n-1}) \end{aligned}$$

Como pode ser visto, para determinar o valor de algum termo  $x_i$ , é preciso conhecer todos os outros valores  $x_j$ ,  $j \neq i$ . O método de Jacobi consiste em assumir um valor

inicial para o vetor  $\mathbf{x}^0$ , substituir no lado direito das equações anteriores, encontrar um novo valor  $\mathbf{x}^1$  e continuar com este processo até que a diferença entre o valor  $\mathbf{x}^k$  e  $\mathbf{x}^{k+1}$  seja pequena o suficiente. Assim, pode-se estabelecer a seguinte relação para avançar as iterações:

$$\begin{aligned}x_1^{k+1} &= \frac{1}{a_{11}}(b_1 - a_{12}x_2^k - a_{13}x_3^k - \dots - a_{1n}x_n^k) \\x_2^{k+1} &= \frac{1}{a_{22}}(b_2 - a_{21}x_1^k - a_{23}x_3^k - \dots - a_{2n}x_n^k) \\&\vdots \\x_n^{k+1} &= \frac{1}{a_{nn}}(b_n - a_{n1}x_1^k - a_{n2}x_2^k - \dots - a_{nn-1}x_{n-1}^k)\end{aligned}$$

As equações acima podem ser expressas de forma geral como:

$$x_i^{k+1} = \frac{1}{a_{ii}} \left( b_i - \sum_{j=1}^{i-1} a_{ij}x_j^k - \sum_{j=i+1}^n a_{ij}x_j^k \right)$$

Uma maneira mais conveniente de expressar a relação anterior pode ser obtida somando-se e subtraindo-se (portanto, sem alterar a igualdade) o termo  $x_i^k$  do lado direito da equação. Com isso, temos:

$$x_i^{k+1} = x_i^k + \frac{1}{a_{ii}} \left( b_i - \sum_{j=1}^n a_{ij}x_j^k \right)$$

Separando o termo entre parêntesis e definido:

$$R_i = b_i - \sum_{j=1}^n a_{ij}x_j^k \quad \rightarrow \quad x_i^{k+1} = x_i^k + \frac{R_i}{a_{ii}}$$

Quando o método numérico atingir a convergência, o termo  $R_i$  será nulo, por isso o termo  $R_i$  é muitas vezes chamado de *resíduo* da equação. De fato, deve-se garantir que o resíduo diminua ao longo das iterações para garantir que o método está convergido para algum lugar.

### 2.1.1 Critério das Linhas

Se o método for convergente, conforme o valor de  $k$  aumenta, mais próximo se estará da solução verdadeira. Um critério que pode ser utilizado para determinar se o método será convergente é avaliar se a matriz dos coeficientes possui a diagonal principal dominante,

ou seja, se os termos da diagonal principal em cada linha são maiores (em módulo) que a soma dos demais termos da mesma linha:

$$|a_{ii}| > \sum_{\substack{j=1 \\ j \neq i}} |a_{ij}|$$

Este critério é muitas vezes chamado de *critério das linhas* e é uma condição suficiente, mas não necessária, para que o método de Jacobi seja convergente para qualquer chute inicial. Quando esta condição não é satisfeita, o método ainda pode convergir, mas irá depender de outros fatores, como por exemplo solução inicial adotada.

Anteriormente, foi especificado que o processo iterativo deve continuar até que a diferença entre o valor atual e o anterior seja pequena o suficiente. Para definir o que isto significa, é importante entender como o erro pode ser calculado.

## 2.2 Convergência de Métodos Iterativos

Todo sistema não-singular de equações algébricas lineares possui uma solução exata. A princípio, os métodos de resolução direta, como os de eliminação, são capazes de fornecer esta solução exata. No entanto, como os valores são avaliados computacionalmente com uma quantidade finita de algarismos significativos, sempre irão existir erros de arredondamento associados à solução, seja ela por métodos diretos ou iterativos.

De forma geral, os métodos iterativos costumam ser menos afetados pelos erros de arredondamento do que os diretos, devido principalmente a três fatores: (a) os sistemas de equações resolvidos iterativamente normalmente a diagonal dominante, (b) usualmente são esparsos e (c) cada iteração ao longo da resolução é independente dos erros de arredondamento do passo anterior, ou seja, a diferença causada pelo erro somente altera o valor inicial utilizado em cada iteração, mas não se acumula ao longo das iterações.

Quando um método numérico convergente é utilizado, a solução obtida se aproxima assintoticamente da solução exata conforme o número de iterações aumenta. Quando o número de iterações tende ao infinito, a diferença entre a solução obtida numericamente e a solução exata é da magnitude da precisão com que os valores são computados (8 algarismos significativos para precisão simples e 16 para precisão dupla). Normalmente não é necessário atingir uma precisão tão grande, sendo por isso estabelecido um critério de parada.

A precisão de um método numérico é medida em termos do **erro** associado ao método. Existem duas formas de especificar o erro, de forma absoluta ou relativa:

$$\text{Erro absoluto} = \text{Valor aproximado obtido} - \text{Valor exato}$$

$$\text{Erro relativo} = \frac{\text{Erro absoluto}}{\text{Valor exato}}$$

A utilização do erro absoluto como critério de convergência só faz sentido quando é conhecida a magnitude da solução exata. Por exemplo, um critério absoluto de  $10^{-3}$  costuma ser suficiente se a solução exata possuir valores da ordem de  $10^2$ , por exemplo. No entanto, se a solução é da ordem de  $10^{-4}$ , o critério é totalmente incorreto. Por isso, é sempre mais adequado especificar o erro relativo.

No entanto, durante a resolução do sistema linear, o valor da solução exata não é conhecido, portanto os erros definidos anteriormente não podem ser calculados. Por isso, durante a resolução, o erro é calculado baseado na variação dos valores ao longo das iterações. Ao longo da solução, o erro absoluto  $\Delta x_i = x_i^{k+1} - x_i^{exato}$  é aproximado como  $\Delta x_i = x_i^{k+1} - x_i^k$ .

Em cada iteração realizada o erro pode ser pequeno para um determinado valor de  $i$  e grande para outro valor, por isso deve-se ter cuidado quando é analisado se o critério de convergência foi atingido. Existem diferentes formas de fazer isto, sendo que as mais comuns são avaliar o valor máximo entre os erros para cada variável ou o somatório dos erros de cada variável. Considerando um critério de convergência  $\varepsilon$ , estes critérios podem ser respectivamente expressos como:

$$|(x_i^{k+1} - x_i^k)_{max}| \leq \varepsilon \qquad \sum_{i=1}^n |x_i^{k+1} - x_i^k| \leq \varepsilon$$

É importante destacar que o erro sempre deve ser avaliado em **módulo** e que a soma dos módulos é *diferente* do módulo das somas. De forma equivalente, os erros relativos podem ser expressos como:

$$\left| \frac{(x_i^{k+1} - x_i^k)_{max}}{x_i^{k+1}} \right| \leq \varepsilon \qquad \sum_{i=1}^n \left| \frac{x_i^{k+1} - x_i^k}{x_i^{k+1}} \right| \leq \varepsilon$$

**Exemplo 02** Resolva os Exemplo 01 utilizando o método de Jacobi, tendo como critério de convergência um erro relativo menor que  $10^{-3}$ .

### 3 Condicionamento de Sistemas Lineares

Uma dúvida comum que surge na resolução de sistemas lineares é em relação à precisão com que os coeficientes e a solução precisam ser avaliados, ou seja, quantos algarismos significativos devem ser utilizados nos cálculos. De forma geral, todo sistema linear  $\mathbf{Ax} = \mathbf{b}$  com  $\det \mathbf{A} \neq 0$  admite solução única. A princípio, esta solução pode sempre ser obtida utilizando o método de eliminação de Gauss ou algum método iterativo, como o método de Jacobi, com algarismos com precisão infinita. No entanto, todas os cálculos são realizados com valores com precisão limitada. Como consequência sempre existem erros de arredondamento associados e estes erros podem ou não alterar a solução do sistema.

Para alguns sistemas, pequenas variações nos coeficientes causam uma grande variação na solução obtida. Como muitas vezes os coeficientes são obtidos através de medidas físicas (que possuem erros associados) ou advêm de outros métodos matemáticos, deve-se avaliar a sensibilidade do sistema em relação aos coeficientes. Com base nisso, pode-se dividir os problemas em duas classes:

- **Problemas bem condicionados:** são aqueles onde uma pequena variação em qualquer um dos elementos do problema causa somente uma pequena variação na solução do problema;

- **Problemas mal condicionados:** são problemas onde uma pequena variação em algum dos elementos causa uma grande variação na solução obtida. Estes problemas tendem a ser muito sensíveis em relação a erros de arredondamento.

Considere o seguinte sistema:

$$\begin{bmatrix} x_1 + x_2 = 2 \\ x_1 + 1.0001x_2 = 2.0001 \end{bmatrix}$$

Utilizando o método de eliminação de Gauss, pode-se reescrever o sistema como:

$$\begin{bmatrix} x_1 + x_2 = 2 \\ 0.0001x_2 = 0.0001 \end{bmatrix}$$

Assim, a solução do sistema é  $x_1 = 1$ ,  $x_2 = 1$ . Considere agora que o coeficiente  $a_{22}$  seja alterado para 0.9999 (uma redução de menos de 0.02%):

$$\begin{bmatrix} x_1 + x_2 = 2 \\ x_1 + 0.9999x_2 = 2.0001 \end{bmatrix}$$

Utilizando o método de eliminação, podemos reescrever o sistema como:

$$\begin{bmatrix} x_1 + x_2 = 2 \\ -0.0001x_2 = 0.0001 \end{bmatrix}$$

Assim, a solução obtida neste caso é  $x_1 = 3$ ,  $x_2 = -1$ , o que se distancia muito da solução obtida anteriormente. Este exemplo mostra como um sistema mal condicionado pode ser sensível aos coeficientes.

Um indicativo de que um sistema pode ser mal condicionado é quando o determinante da matriz é muito próximo a zero, porém este critério não representa uma avaliação quantitativa do condicionamento. Uma maneira de avaliar o quanto um determinado sistema é mal condicionado é através da determinação do **número de condicionamento**, que é definido com base na norma da matriz dos coeficientes e sua inversa.

O número de condicionamento é uma medida da sensibilidade do sistema a pequenas variações em qualquer de seus elementos. A origem deste número não será apresentada aqui, mas pode ser encontrada em Hoffman (2001). Considerando um sistema linear da forma  $\mathbf{Ax} = \mathbf{b}$ , o número de condicionamento da matriz  $\mathbf{A}$  é definido como:

$$C(\mathbf{A}) = \|\mathbf{A}\| \|\mathbf{A}^{-1}\|$$

Pequenos valores de  $C(\mathbf{A})$ , da ordem de uma unidade, indicam uma pequena sensibilidade da solução em relação a variações nos coeficientes, ou seja, indicam problemas bem condicionados. Valores grandes de  $C(\mathbf{A})$  mostram que o sistema é mal condicionado.

Na definição do número de condicionamento, a norma utilizada é a norma Euclidiana (ou de Frobenius), definida como:

$$\|\mathbf{A}\| = \left( \sum_{i=1}^n \sum_{j=1}^n (a_{ij})^2 \right)^{1/2}$$

**Exemplo 03:** Determine o número de condicionamento da seguinte matriz:

$$\mathbf{A} = \begin{bmatrix} 1 & 1 \\ 1 & 1.0001 \end{bmatrix}$$

Primeiramente, pode-se determinar a inversa da matriz  $\mathbf{A}$  para na sequência calcular as normas. Lembrando da definição, a matriz inversa  $\mathbf{A}^{-1}$  é uma matriz tal que:

$$\mathbf{A} \cdot \mathbf{A}^{-1} = \mathbf{I}$$

Neste caso, como a matriz é pequena, pode-se formar um sistema linear para determinar a inversa:

$$\begin{bmatrix} 1 & 1 \\ 1 & 1.0001 \end{bmatrix} \begin{bmatrix} a & b \\ c & d \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

Assim:

$$a + c = 1$$

$$b + d = 0$$

$$a + 1.0001c = 0$$

$$b + 1.0001d = 1$$

Resolvendo o sistema, obtém-se:

$$\begin{bmatrix} 10001 & -10000 \\ -10000 & 10000 \end{bmatrix}$$

Com isso, as normas podem ser calculadas:

$$\| \mathbf{A} \| = (1^2 + 1^2 + 1^2 + 1.0001^2)^{1/2} = 2.00005$$

$$\| \mathbf{A}^{-1} \| = (10001^2 + (-10000)^2 + (-10000)^2 + 10000^2)^{1/2} = 20000.5$$

Dessa forma, o número de condicionamento será:

$$C(\mathbf{A}) = \| \mathbf{A} \| \| \mathbf{A}^{-1} \| = (2.00005)(20000.5) = 40002.0$$

Como visto, o número de condicionamento é muito elevado, indicando que a matriz é muito mal condicionada, como discutido anteriormente.

## Anexos

Resolução do Exemplo 01 no Scilab 5.5.2:

```
1 //Aula 02 -- Exemplo 01 -- Algoritmo de Thomas
2 alpha=4;
3 dx=0.125;
4 k=-(2+alpha^2*dx^2);
5
6 //Matriz dos coeficientes:
7 A=[1 0 0 0 0; -1 k 1 0 0; 0 1 k 1 0; 0 0 1 k 1; 0 0 0 0 1];
8
9 //Parte não-homogêneo
10 b=[0 0 0 0 10]
11
12 //Substituição
13 for i=2:5
14     e(i)=A(i,i-1)/A(i-1,i-1);
15     A(i,i)=A(i,i)-e(i)*A(i-1,i);
16     b(i)=b(i)-e(i)*b(i-1);
17 end
18
19 //Resolução
20 T(5)=b(5)/A(5,5);
21 for i=4:-1:1 //-1 para indicar que é decrescente
22     T(i)=(b(i)-A(i,i+1)*T(i+1))/A(i,i);
23 end
```

Resolução do Exemplo 02 no Scilab 5.5.2:

```
1 //Aula 02 -- Exemplo 02 -- Jacobi
2 clear;
3 alpha=4;
4 dx=0.125;
5 k=-(2+alpha^2*dx^2)
6
7 //Matriz dos coeficientes:
8 A=[1 0 0 0 0; 1 k 1 0 0; 0 1 k 1 0; 0 0 1 k 1; 0 0 0 0 1];
9
10 //Parte não-homogêneo
11 b=[0 0 0 0 10]
12
13 //Inicialização
14 k=1;
15 erro=1;
16 for i=1:5
17     T(i,k)=0;
18 end
19
20 //Iteração
21 while erro>0.001
22     for i=1:5
23         T(i,k+1)=T(i,k)+(1/A(i,i))*(b(i)-sum(A(i,:)*T(:,k)));
24     end
25     erro=sum(abs((T(:,k+1)-T(:,k))/T(:,k+1)));
26     k=k+1;
27 end
```

## References

- [1] Hoffman, J. D. Numerical Methods for Engineers and Scientists. 2nd ed., Marcel Dekker, New York: 2001.
- [2] Chapra, S. C.; Canale R. P. Numerical Methods for Engineers. 6th ed., McGraw Hill, New York: 2010.