

INSTITUTO BAIANO DE EDUCAÇÃO SUPERIOR – IBES
CURSO: CIÊNCIA DA COMPUTAÇÃO
DISCIPLINA: MÉTODOS NUMÉRICOS EM COMPUTAÇÃO
PROFESSOR: ATAUALPA MAGNO FERRAZ DE NOVAES
PRIMEIRO SEMESTRE DE 2006

Prezados Alunos, sejam bem-vindos ao nosso curso de Métodos Numéricos em Computação.

Desejo que vocês possam desfrutar destas notas de aula. O sabor coloquial com que procurei “temperar” estas anotações certamente facilitará a aprendizagem da matéria.

Um agradecimento muito especial aos autores de livros, às excelentes mestres Ruggiero e Lopes como também ao professor Barroso. Desde que desejamos aprimorar este trabalho ao longo do tempo, sugestões e críticas serão bem vindas.

Email: amferraznovaes@ig.com.br ou ataualpa@im.ufba.br.

Página na Internet: [http:// geocities.yahoo.com.br/magnoferraz/](http://geocities.yahoo.com.br/magnoferraz/)

Telefones: 3353-4784 ou 9179-1925

Primeiramente apresentarei a vocês a EMENTA do nosso curso e o CONTEÚDO PROGRAMÁTICO

I – EMENTA

Aritmética de ponto flutuante. Zeros de funções reais. Sistemas lineares.

II – OBJETIVOS GERAIS

Familiarizar o aluno com conceitos e aplicações numéricas de computação.

III – OBJETIVOS ESPECÍFICOS

Esta disciplina cobre os tópicos fundamentais de métodos numéricos, abordando métodos gerais e problemas numéricos como aplicações desses métodos.

IV – CONTEÚDO PROGRAMÁTICO

Erros

- Simplificação no modelo matemático.
- Erro de truncamento.
- Erro de arredondamento.
- Erro nos dados.
- Aritmética de ponto flutuante.

Zeros de funções

- Localização de raízes isoladas. Teorema de Bolzano.
- Processos iterativos.

- Método da Dicotomia.
- Método das aproximações sucessivas.
- Método de Newton-Raphson.

Sistemas Lineares

- Introdução: Esforço computacional.
- Método da eliminação de Gauss.
- Método da eliminação de Gauss com condensação pivotal. Matriz inversa.
- Refinamento da solução.
- Método iterativo de Gauss-Seidel. Critérios de convergência

V – ESTRATÉGIA DE TRABALHO

Aulas expositivas e recursos audiovisuais

VI – AVALIAÇÃO

Provas bimestrais e trabalhos práticos

VII – BIBLIOGRAFIA BÁSICA

HUMES, A. F. P. C.; MELO, I. S. H.; YOSHIDA, L. K.; MARTINS, W. T. **Noções de Cálculo numérico**. São Paulo: McGraw-Hill, 1984.

RUGGIERO, M. A. G.; LOPES, V. L. R. **Cálculo numérico: aspectos teóricos e computacionais**. São Paulo: Makron Books, 1996.

NOTAS DE AULA DE CÁLCULO NUMÉRICO

O que é o Cálculo Numérico?

O **Cálculo Numérico** corresponde a um conjunto de **métodos** usados para se obter a solução de problemas científicos na maioria das vezes de forma **aproximada**, usando técnicas matemáticas e o computador.

Porém, o profissional que se defrontar com o problema terá que tomar uma série de decisões antes de resolvê-lo. E para tomar essas decisões, é preciso ter conhecimento de métodos numéricos. O profissional terá que decidir, dentre outras coisas:

- ✓ Pela utilização ou não de um método numérico (se existirem métodos numéricos para se resolver o tal problema.);
- ✓ Escolher o método a ser utilizado, procurando aquele que é mais adequado para o seu problema. Analisar quais vantagens cada método oferece e as limitações que eles apresentam;
- ✓ Saber avaliar a qualidade da solução obtida. Para isso, é importante ele saber exatamente o que está sendo feito pelo computador ou calculadora, isto é, como determinado método é operacionalizado.

Os principais objetivos do nosso curso são:

- Apresentar diversos métodos numéricos para a resolução de diferentes problemas matemáticos. Pretende-se esclarecer a importância desses métodos, mostrando:
 - ✓ a essência de um método numérico;
 - ✓ a diferença em relação a soluções analíticas;
 - ✓ as situações em que eles devem ser aplicados;
 - ✓ as vantagens de se utilizar um método numérico;
 - ✓ as limitações na sua aplicação e confiabilidade na solução obtida.
- Melhorar a familiarização e “intimidade” do aluno com a matemática, mostrando seu lado prático e sua utilidade no dia-a-dia de um cientista. Rever conceitos, exercitá-los e utilizá-los de maneira prática.

Erros Numéricos

Vamos supor o seguinte problema: como calcular o valor de $\sqrt{2}$ (ou, por exemplo uma dízima como $\frac{2}{3}$)? Provavelmente, a primeira resposta que vem à mente de qualquer pessoa esclarecida será: utilizando uma calculadora ou um computador. Indiscutivelmente, essa é a resposta mais sensata e prática. Porém, um profissional que utilizará o resultado fornecido pela calculadora para projetar, construir ou manter pontes, edifícios, máquinas, sistemas, dispositivos eletrônicos, etc., não pode aceitar o valor obtido antes de fazer alguns questionamentos (pelo menos uma vez na sua vida profissional!). Será que esse resultado é confiável? (por exemplo, será que a ponte pode desabar?)

Essa pergunta faz sentido pois $\sqrt{2}$ é um número irracional, isto é, não existe uma forma de representá-lo com um número finito de algarismos. Portanto, o número apresentado pela calculadora é uma **aproximação do valor real** de $\sqrt{2}$, já que ela não pode mostrar infinitos algarismos. E quão próximo do valor real está o resultado mostrado?

O erro cometido ao se calcular o valor de $\sqrt{2}$, se refere à inevitável limitação na representação de números irracionais e é apenas um tipo de erro que pode surgir ao se resolver um problema real. Esse tipo de erro é chamado de **erro de arredondamento**.

Outros tipos de erros também podem aparecer devido a outros tipos de problemas ou limitações.

Tipos de Erros

A solução matemática de um determinado problema envolve diversas etapas. A solução do problema se inicia com a criação de um **modelo matemático** que melhor se ajuste ao problema em questão. Esse modelo sempre apresentará aproximações e limitações. Esse tipo de erro é chamado de **erro na simplificação do modelo matemático**. Além disso, na grande maioria das vezes, **dados experimentais** (alguns autores o chamam de **erro inerente ao modelo matemático utilizado**) serão utilizados para se obter a solução. **Como toda medida experimental apresenta uma incerteza**, a solução do problema será influenciada pela mesma. Esse tipo de erro é chamado de **erro de entrada de dados**. Portanto, logo de início, existem diversos fatores que introduzem erros na solução numérica do problema.

Vamos considerar um outro tipo de erro que pode surgir ao realizarmos determinadas operações. Digamos que precisamos calcular o valor de e^x . Mais uma vez, iremos utilizar uma máquina digital (calculadora ou computador). Como esse equipamento irá realizar essa operação?

Sabemos, do Cálculo Diferencial, que a *exponencial* é uma função que pode ser representada por uma série infinita dada por: $e^x = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots + \frac{x^n}{n!} + \dots$ e na prática é impossível calcular seu valor

exato. Portanto, mais uma vez, teremos que fazer uma aproximação, que levará a um erro no resultado final de e^x .

Por exemplo, se considerarmos $e^x = 1 + x + \frac{x^2}{2!}$ estaremos fazendo um **truncamento** dessa série, e o erro gerado no valor de e^x é chamado de **erro de truncamento** (é claro que estamos nos referindo à função e^x de um modo geral, pois para $x = 0$ não há erro algum).

Podemos criar algumas definições a fim de facilitar as discussões e trocas de informação sobre esse problema. Vamos definir o **módulo da diferença entre o valor real da grandeza que queremos calcular e o valor aproximado** que efetivamente calculamos como sendo o **erro absoluto**, ou seja:

$$\text{Erro Absoluto} = E_A = | \text{valor real} - \text{valor aproximado} |$$

Quanto menor for esse erro, mais **preciso** será o resultado da operação.

Porém, se estivermos trabalhando com números muito grandes, o erro pode ser grande em termos absolutos, mas o resultado ainda será preciso. E o caso inverso também pode ocorrer: um erro absoluto pequeno, mas um resultado impreciso. Por exemplo, digamos que o resultado de uma operação nos forneça o valor 2.123.542,7 enquanto o valor real que deveríamos obter é 2.123.544,5. O erro absoluto neste caso é 1,8. Comparada com o valor real, essa diferença (o erro absoluto) é bem pequena, portanto, podemos considerar o resultado preciso. Em um outro caso, digamos que o resultado da operação seja 0,234 e o resultado correto era 0,128. Desta vez o erro absoluto será igual a 0,106, porém o resultado é bastante impreciso.

A fim de evitar esse tipo de ambigüidade, podemos criar uma nova definição. Podemos definir o **erro relativo** como sendo o quociente entre o erro absoluto e o valor real da grandeza a ser calculada, ou seja:

$$\text{erro relativo} = \frac{|\text{valor real} - \text{valor aproximado}|}{\text{valor real}}$$

O erro relativo é uma forma mais interessante de se avaliar a **precisão** de um cálculo efetuado. No exemplo acima, teremos um erro relativo de 0,0000008 ou 0,00008% no primeiro caso e um erro relativo igual a 0,83 ou 83% no segundo caso.

Propagação e Condicionamento de Erros Numéricos

Vamos supor que queremos calcular o valor de $\sqrt{2} - e^3$. Como vimos anteriormente, ao calcularmos o valor de $\sqrt{2}$, teremos que realizar um arredondamento, que leva ao um resultado aproximado de $\sqrt{2}$, ou seja, existe um **erro de arredondamento** associado ao resultado. Para calcularmos o valor de e^3 teremos que fazer um truncamento, que também irá gerar um **erro de truncamento** (ao usarmos a função e^x) no resultado obtido. Portanto, o resultado da operação de

subtração entre $\sqrt{2}$ e e^3 apresentará um erro que é proveniente dos erros nos valores de $\sqrt{2}$ e e^3 separadamente. Em outras palavras, os erros nos valores de $\sqrt{2}$ e e^3 se **propagam** para o resultado de $\sqrt{2} - e^3$. Podemos concluir então que, ao se resolver um problema numericamente, a cada etapa e a cada operação realizada, devem surgir diferentes tipos de erros gerados das mais variadas maneiras, e estes erros se propagam e determinam o erro no resultado final obtido.

Representação Numérica

Introdução

A fim de realizarmos de maneira prática qualquer operação com números, nós precisamos representá-los em uma determinada base numérica. Podemos escrevê-lo na **base decimal**, por exemplo, que é a base mais usada atualmente pela humanidade (graças à nossa anatomia).

Voltemos ao número $\sqrt{2}$. O valor de $\sqrt{2}$ na base decimal pode ser escrito como 1,41 ou 1,4142 ou ainda 1,41421356237. Qual é a diferença entre essas várias formas de representar $\sqrt{2}$?

RESPOSTA: A diferença é a quantidade de **algarismos significativos** usados em cada representação.

Em uma máquina digital, como uma calculadora ou um computador, os números não são representados na base decimal. Eles são representados na base binária, ou seja, usam o número 2 como base ao invés do número 10.

EXERCÍCIOS

1) Converta os números da base 10 para a base 2:

- | | | | | | | |
|---------|-----------|--------|----------|--------|--------|--------|
| a) 2 | b) 5 | c) 10 | d) 15 | e) 25 | f) 37 | g) 347 |
| h) 2345 | i) 0,1875 | j) 0,6 | l) 13,25 | m) 3,5 | n) 0,1 | |

2) Converta os números da base 2 para a base 10:

- | | | | | |
|----------|----------|----------|---------|----------|
| a) 101 | b) 1000 | c) 1001 | d) 1011 | e) 1111 |
| f) 10100 | g) 10101 | h) 0,10 | i) 0,11 | j) 0,101 |
| l) 1,01 | m) 1,001 | n) 1,011 | | |

Ponto Fixo e Ponto Flutuante

A princípio, toda vez que escrevemos um número, deveríamos mencionar a base numérica a qual estamos nos referindo. Obviamente, isso não se faz necessário na prática, pois estamos sempre representando os números na base decimal, portanto sabemos exatamente o seu significado. Por exemplo, quando escrevemos o número 1532, o que realmente queremos dizer? Estamos dizendo que esse número representa uma quantidade equivalente a $1 \times 1000 + 5 \times 100 + 3 \times 10 + 2$, ou, escrevendo a base de outra forma, $1 \times 10^3 + 5 \times 10^2 + 3 \times 10^1 + 2 \times 10^0$. Essa é a chamada **representação posicional** de números.

Na base binária, o mecanismo é o mesmo, porém, ao invés de potências de 10, utilizamos potências de 2. Portanto, um número binário como 1011 (lembre-se, do ginásio, que na base binária só existem os algarismos 0 e 1) significa $1 \times 2^3 + 0 \times 2^2 + 1 \times 2^1 + 1 \times 2^0$ que na base 10 é $8+2+1=11$.

Um **número inteiro** apresenta a chamada representação de **ponto fixo**, onde a posição do ponto “decimal” está fixa e todos os dígitos são usados para representar o número em si, com exceção do primeiro dígito usado para representar o sinal desse número. A figura abaixo ilustra essa representação.

Sinal	Dígitos
-------	---------

Para um **número real** qualquer (inteiro ou não inteiro) é utilizada a representação de **ponto flutuante normalizado** (ou simplesmente, representação de **ponto flutuante**), que é dada pela expressão: $\pm (0.d_1d_2d_3...d_t) \times \beta^e$ onde:

$0.d_1d_2d_3...d_t$ é uma fração na base β , também chamada de **mantissa**, com $0 \leq d_i \leq \beta - 1$, para todo $i = 1, 2, 3, ..., t$ e $d_1 \neq 0$ sendo t o número máximo de dígitos da mantissa (também chamado de número de algarismos significativos) que é determinado pelo comprimento da palavra do computador; e é um expoente que varia em um intervalo dado pelos limites da máquina utilizada, assim $-m \leq e \leq M$, onde m é o limite inferior e M é o limite superior da máquina.

Esse tipo de representação é chamado de **ponto flutuante**, pois o ponto da fração “flutua” conforme o número a ser representado e sua posição é expressa pelo expoente e .

Alguns exemplos da representação de ponto flutuante podem ser vistos na tabela a seguir:

Número na base decimal	Base	Representação em ponto flutuante	Mantissa	Expoente
1532	10	0.1532×10^4	0.1532	4
15.32	10	0.1532×10^2	0.1532	2
0.00255	10	0.255×10^{-2}	0.255	-2
10	10	0.10×10^2	0.10	1
10	2	0.1010×2^4	0.1010	4

Arredondamento em Ponto Flutuante

Introdução

Chamamos atenção para o fato de que o conjunto dos números representáveis em qualquer máquina é finito, e portanto discreto, ou seja não é possível representar em uma máquina todos os números de um dado intervalo $[a,b]$. A implicação imediata desse fato é que o resultado de uma simples operação aritmética ou o cálculo de uma função, realizadas com esses números, podem conter erros. A menos que medidas apropriadas sejam tomadas, essas imprecisões causadas, por exemplo, por simplificação no modelo matemático (algumas vezes necessárias para se obter um modelo matemático solúvel); erro de truncamento (troca de uma série infinita por uma finita); erro de arredondamento (devido a própria estrutura da máquina); erro nos dados (dados imprecisos obtidos de experimentos, ou arredondados na entrada); etc, podem diminuir e algumas vezes destruir, a precisão dos resultados.

Assim, nosso objetivo aqui será o de alertar o aluno para os problemas que possam surgir durante a resolução de um problema, bem como dar subsídios para evitá-los e para uma melhor interpretação dos resultados obtidos.

Sistema de Números Discreto no Computador

Inicialmente, descreveremos como os números são representados num computador.

Representação de um Número Inteiro

Em princípio, a representação de um **número inteiro** no computador não apresenta dificuldade.

OBSERVAÇÃO:

O número *zero* pertence a qualquer sistema e é representado com mantissa igual a zero e $e = -m$.

Representação de um Número Real

Fundamentado no que você aprendeu até aqui, resolva o exercício abaixo

EXERCÍCIO:

Escrever os números: $x_1 = 0.35$; $x_2 = -5.172$; $x_3 = 0.0123$; $x_4 = 5391.3$ e $x_5 = 0.0003$, onde todos estão na base $\beta = 10$, em ponto flutuante na forma normalizada.

Para representarmos um sistema de números em ponto flutuante normalizado, na base β , com t dígitos significativos e com limites do expoente m e M , usaremos a notação: $F(\beta, t, m, M)$.

Assim um número não nulo em $F(\beta, t, m, M)$ será representado por:

$$\pm 0.d_1d_2d_3\dots d_t \times \beta^e, \text{ onde } 0 \leq d_i \leq \beta - 1, d_1 \neq 0 \text{ e } -m \leq e \leq M$$

EXERCÍCIO: Considere o sistema $F(10, 3, 2, 2)$. Represente nesse sistema os números do exercício anterior $x_1 = 0.35$; $x_2 = -5.172$; $x_3 = 0.0123$; $x_4 = 5391.3$ e $x_5 = 0.0003$.

Veja abaixo um exercício resolvido bastante interessante.

Seja $f(x)$ uma **função contínua real** definida no intervalo $[a, b]$, $a < b$ e sejam $f(a) < 0$ e $f(b) > 0$. Então de acordo com o teorema do valor intermediário, existe x , $a < x < b$, tal que $f(x) = 0$.

Seja $f(x) = x^3 - 3$. Determinar x tal que $f(x) = 0$.

SOLUÇÃO: Para a função dada, consideremos $t = 10$ e $\beta = 10$. Obtemos então (não se preocupe como conseguimos o resultado abaixo, pois você só aprenderá mais tarde):

$$f(0.1442249570 \times 10^1) = -0.2 \times 10^{-8};$$

$$f(0.1442249571 \times 10^1) = 0.4 \times 10^{-8}.$$

Observe que entre 0.1442249570×10^1 e 0.1442249571×10^1 não existe nenhum número que possa ser representado no sistema dado e que a função f muda de sinal nos extremos desse intervalo. Assim, esta máquina não contém o número x tal que $f(x) = 0$ e portanto a equação dada não possui solução nessa máquina em que $t = 10$ e $\beta = 10$.

Representação de Números no Sistema $F(\beta, t, m, M)$

Sabemos que os números reais podem ser representados por uma reta contínua (isto é, numa reta). Entretanto, na representação em ponto flutuante podemos representar apenas pontos discretos na reta real. Para ilustrar este fato consideremos o seguinte exemplo.

EXEMPLO: Quantos e quais números podem ser representados no sistema $F(2, 3, 1, 2)$?

SOLUÇÃO: Temos que $\beta = 2$ então os dígitos podem ser 0 ou 1; $m = 1$ e $M = 2$ então $-1 \leq e \leq 2$ e $t = 3$. Assim, os números são da forma:

$$\pm 0.d_1d_2d_3 \times \beta^e.$$

Logo temos: duas possibilidades para o sinal, uma possibilidade para d_1 , duas para d_2 , duas para d_3 e quatro para as formas de β^e . Fazendo o produto $2 \times 1 \times 2 \times 2 \times 4$ obtemos 32. Assim neste sistema podemos representar 33 números visto que o zero faz parte de qualquer sistema.

Para responder quais são os números, notemos que as formas da mantissa são: 0.100, 0.101, 0.110 e 0.111 e as formas de β^e são: 2^{-1} , 2^0 , 2^1 , 2^2 . Assim, obtemos os seguintes números:

$$0.100 \times \begin{cases} 2^{-1} = (0.25)_{10} \\ 2^0 = (0.5)_{10} \\ 2^1 = (1.0)_{10} \\ 2^2 = (2.0)_{10} \end{cases},$$

desde que $(0.100)_2 = (0.5)_{10}$;

$$0.101 \times \begin{cases} 2^{-1} = (0.3125)_{10} \\ 2^0 = (0.625)_{10} \\ 2^1 = (1.25)_{10} \\ 2^2 = (2.5)_{10} \end{cases},$$

desde que $(0.101)_2 = (0.625)_{10}$;

$$0.110 \times \begin{cases} 2^{-1} = (0.375)_{10} \\ 2^0 = (0.75)_{10} \\ 2^1 = (1.5)_{10} \\ 2^2 = (3.0)_{10} \end{cases},$$

desde que $(0.110)_2 = (0.75)_{10}$;

$$0.111 \times \begin{cases} 2^{-1} = (0.4375)_{10} \\ 2^0 = (0.875)_{10} \\ 2^1 = (1.75)_{10} \\ 2^2 = (3.5)_{10} \end{cases},$$

desde que $(0.111)_2 = (0.875)_{10}$.

EXERCÍCIO:

Considerando o mesmo sistema do exemplo acima, represente os números: $x_1 = 0.38$, $x_2 = 5.3$ e $x_3 = 0.15$ dados na base 10.

As operações num computador são arredondadas. Para ilustrar este fato, consideremos o seguinte exemplo.

EXEMPLO: Calcular o quociente entre 15 e 7.

SOLUÇÃO:

Usando a calculadora deste computador no qual estou trabalhando neste instante, temos: 2,1428571428571428571428571428571.

Suponha agora que só dispomos do total de 4 dígitos para representar o quociente 15/7.

Daí, $15/7 = 2,142$

Mas não seria melhor aproximarmos 15/7 por 2,143? A resposta é sim e isso significa que o número foi arredondado. Isto sugere a seguinte...

DEFINIÇÃO: Arredondar um número x , por outro com um número menor de dígitos significativos, consiste em encontrar um número $\bar{x}$, pertencente ao sistema de numeração, tal que $|\bar{x} - x|$ seja o menor possível.

IMPORTANTE!

Assim, em linhas gerais, para arredondar um número, na base 10, devemos apenas observar o primeiro dígito a ser descartado. Se este dígito é menor que 5 deixamos os dígitos inalterados e se é maior ou igual a 5 acrescentamos uma unidade ao último algarismo remanescente.

Arredondamento

Trataremos o arredondamento em ponto flutuante com o exemplo abaixo, use os conhecimentos adquiridos até aqui e o seu bom senso e resolva-o.

EXEMPLO: Considere uma máquina que utiliza o sistema $F(10, 3, 5, 5)$, isto é os números tem a seguinte forma $0, d_1 d_2 d_3 \cdot 10^e$ com $-5 \leq e \leq 5$. Represente neste sistema os números:

$x_1 = 1234,56$; $x_2 = -0,00054962$; $x_3 = 0,9995$; $x_4 = 123456,7$ e $x_5 = -0,0000001$.

RESPOSTA:

$x_1 = 0,123 \cdot 10^4$; $x_2 = -0,550 \cdot 10^{-3}$; $x_3 = 0,1 \cdot 10^1$;

x_4 é impossível, pois o expoente é 6, acima da capacidade da máquina (overflow);

x_5 é impossível, pois o expoente é -6 , abaixo da capacidade da máquina (underflow).

OPERAÇÕES ARITMÉTICAS EM PONTO FLUTUANTE

Considere uma máquina qualquer e uma série de operações aritméticas. Pelo fato do arredondamento ser feito após cada operação temos, ao contrário do que é válido para números reais, que as operações aritméticas (adição, subtração, multiplicação e divisão) não são nem associativas e nem distributivas.

Ilustraremos esse fato através de exercícios, os quais você deve tentar fazer.

Nos exercícios abaixo considere o sistema com base $\beta = 10$, e com 3 dígitos significativos.

Efetue as operações indicadas:

i) $(11,4 + 3,18) + 5,05$ e $11,4 + (3,18 + 5,05)$

RESPOSTAS: 19,7 e 19,6

ii) $\frac{3,18 \times 11,4}{5,05}$ e $\left(\frac{3,18}{5,05}\right) \times 11,4$

RESPOSTAS: 7,19 e 7,18

iii) $3,18 \times (5,05 + 11,4)$ e $3,18 \times 5,05 + 3,18 \times 11,4$.

RESPOSTAS: 52,5 e 52,4

BIBLIOGRAFIA:

HUMES, A. F. P. C.; MELO, I. S. H.; YOSHIDA, L. K.; MARTINS, W. T: **Noções de Cálculo numérico**. São Paulo: McGraw-Hill, 1984.

RUGGIERO, M. A. G.; LOPES, V. L. R. **Cálculo numérico**: aspectos teóricos e computacionais. São Paulo: Makron Books, 1996.

MARCELO GAMEIRO E ANTONIO CÉSAR GERMANO. Apostila de Complementos de Cálculo Numérico.

ARTIGOS E LIVROS ENCONTRADOS NA INTERNET

ZEROS REAIS DE FUNÇÕES REAIS

Dizemos que α é um zero (ou raiz) de uma função real de variável real $y = f(x)$ se e somente se $f(\alpha) = 0$.

Usaremos muitas vezes o seguinte teorema (que **localiza** os possíveis zeros de uma função):

TEOREMA DE BOLZANO: Seja f uma função contínua em $[a, b]$ com $f(a) \cdot f(b) < 0$ então a função f tem **pelo menos** um zero em $[a, b]$.

EXEMPLO: Seja a função $f(x) = x \cdot \ln x - 3,2$. Podemos calcular o valor de $f(x)$ para valores arbitrários de x , como mostrado na tabela abaixo (usando apenas duas casas decimais):

x	1	2	3	4
$f(x)$	-3,20	-1,81	0,10	2,36

Pelo **teorema de Bolzano**, concluímos que existe pelo menos uma raiz real no intervalo $[2, 3]$.

Seja f uma função contínua em $[a, b]$ com $f(a) \cdot f(b) < 0$ e suponha ainda que α seja seu **único** zero em $[a, b]$ (mais tarde veremos alguns processos para não só **localizar** zeros de funções mas também para **isolar** esses zeros).

MÉTODO ITERATIVO: Um método é dito **iterativo** (e não iterativo!!!) basicamente quando para calcularmos uma nova aproximação, usamos uma aproximação anterior, obtida por algum método.

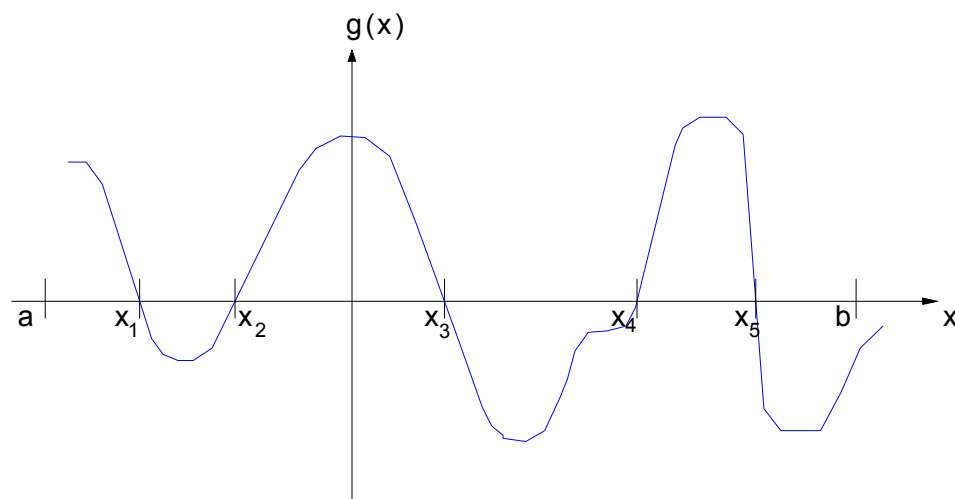
Estudaremos três métodos para determinar alguns zeros reais de funções reais, são eles:

- 1) Método da bissecção ou dicotomia;
- 2) Método de Newton-Raphson;
- 3) Método das aproximações sucessivas ou iteração linear.

RESUMO E ALGUMAS OBSERVAÇÕES

Graficamente, os zeros de uma função $f(x)$ correspondem aos valores de x em que a função intercepta o eixo horizontal do gráfico, como mostrado na figura abaixo.

A função $g(x)$ da figura abaixo tem 5 raízes no intervalo $[a, b]$: x_1, x_2, x_3, x_4, x_5 .



Às vezes as raízes de uma função podem ser encontradas analiticamente, ou seja, resolvendo a equação $f(x)=0$ de maneira exata, como mostrado nos exemplos a seguir:

1-) $f(x) = x - 3$

$x = 3$ é raiz de $f(x)$ pois :

$$f(3) = 3 - 3 = 0$$

2-) $g(x) = \frac{8}{3}x - 4$

$$\frac{8}{3}x - 4 = 0 \Rightarrow \frac{8}{3}x = 4 \Rightarrow x = \frac{12}{8} = \frac{3}{2}$$

$x = \frac{3}{2}$ é a raiz de $g(x)$ pois :

$$g\left(\frac{3}{2}\right) = \frac{8}{3} \cdot \frac{3}{2} - 4 = 0$$

3-) $h(x) = x^2 - 5x + 6$

$$x^2 - 5x + 6 = 0$$

$$\Delta = 25 - 24 = 1$$

$$x = \frac{5 \pm \sqrt{1}}{2}$$

$$x_1 = 3$$

$$x_2 = 2$$

tanto $x = 2$ quanto $x = 3$ são soluções de $h(x)$ pois :

$$h(3) = 3^2 - 5 \cdot 3 + 6$$

$$h(3) = 15 - 15 = 0$$

$$h(2) = 2^2 - 5 \cdot 2 + 6$$

$$h(2) = 10 - 10 = 0$$

Porém, nem sempre é possível se encontrar analiticamente a raiz de uma função, como nos casos a seguir:

1-) $f(x) = x^7 + 2x^2 - x + 1$

2-) $g(x) = \sin x + e^x$

3-) $h(x) = x + \ln x$

Nestes casos precisamos de um método numérico para encontrar **uma estimativa** para a raiz da função estudada, ou seja, um valor tão aproximado quanto se deseje.

Tais métodos devem envolver as seguintes etapas:

- (a) Determinação de um intervalo em x que contenha pelo menos uma raiz da função $f(x)$, ou seja, isolamento das raízes (quando possível);
- (b) Cálculo da raiz aproximada através de um processo iterativo (que será explicado logo a seguir) até a precisão desejada.

Processos Iterativos

Existe um grande número de métodos numéricos que são processos iterativos. Como o próprio nome já diz (consulte um dicionário para verificar o significado de *iterativo*), esses processos se caracterizam pela **repetição** de uma determinada operação. A idéia nesse tipo de processo é repetir um determinado cálculo várias vezes obtendo-se a cada repetição ou **iteração** um resultado mais preciso que aquele obtido na iteração anterior. E, a cada iteração utiliza-se o resultado da iteração anterior como parâmetro de entrada para o cálculo seguinte.

Alguns aspectos comuns a qualquer processo iterativo, são:

- ✓ **Estimativa inicial:** como um processo iterativo se caracteriza pela utilização do resultado da iteração anterior para o cálculo seguinte, a fim de se iniciar um processo iterativo, é preciso que se tenha uma estimativa inicial do resultado do problema. Essa estimativa pode ser conseguida de diferentes formas, conforme o problema que se deseja resolver;
- ✓ **Convergência:** a fim de se obter um resultado próximo do resultado real, é preciso que a cada passo ou iteração, o resultado esteja mais próximo daquele esperado, isto é, é preciso que o método **convirja** para o resultado real. Essa convergência nem sempre é garantida em um processo numérico. Portanto, é muito importante se estar atento a isso e realizar a verificação da convergência do método para um determinado problema antes de tentar resolvê-lo;

- ✓ **Critério de Parada:** obviamente não podemos repetir um processo numérico infinitamente. É preciso pará-lo em um determinado instante. Para isso, devemos utilizar um certo critério, que vai depender do problema a ser resolvido e da precisão que precisamos obter na solução. O critério adotado para **parar** as iterações de um processo numérico é chamado de *critério de parada*.

Para encontrarmos as raízes ou zeros de uma função iremos utilizar métodos numéricos iterativos. Como já mencionado, o primeiro passo para se resolver um processo iterativo corresponde a obtenção de uma estimativa inicial para o resultado do problema. No caso de zeros de funções, usamos a operação chamada de **isolamento de raízes**:

Isolamento de Raízes

Para determinarmos o número e a localização aproximada de raízes de uma função, a fim de obtermos uma estimativa inicial a ser usada nos processo iterativos, podemos examinar o comportamento dessa função através de um esboço gráfico.

Por exemplo, seja uma função $f(x)$ tal que:

$$f(x) = g(x) - h(x)$$

As raízes de $f(x)$, são os valores de x tais que: $g(x) - h(x) = 0$, ou ainda $g(x) = h(x)$

Logo, os valores de x em que **o gráfico de $g(x)$ intercepta o gráfico de $h(x)$** é a raiz de $f(x)$.

Passemos agora ao estudo dos métodos mencionados na página 13, isto é:

- 1) Método da bissecção ou dicotomia;
- 2) Método de Newton-Raphson;
- 3) Método das aproximações sucessivas ou iteração linear.

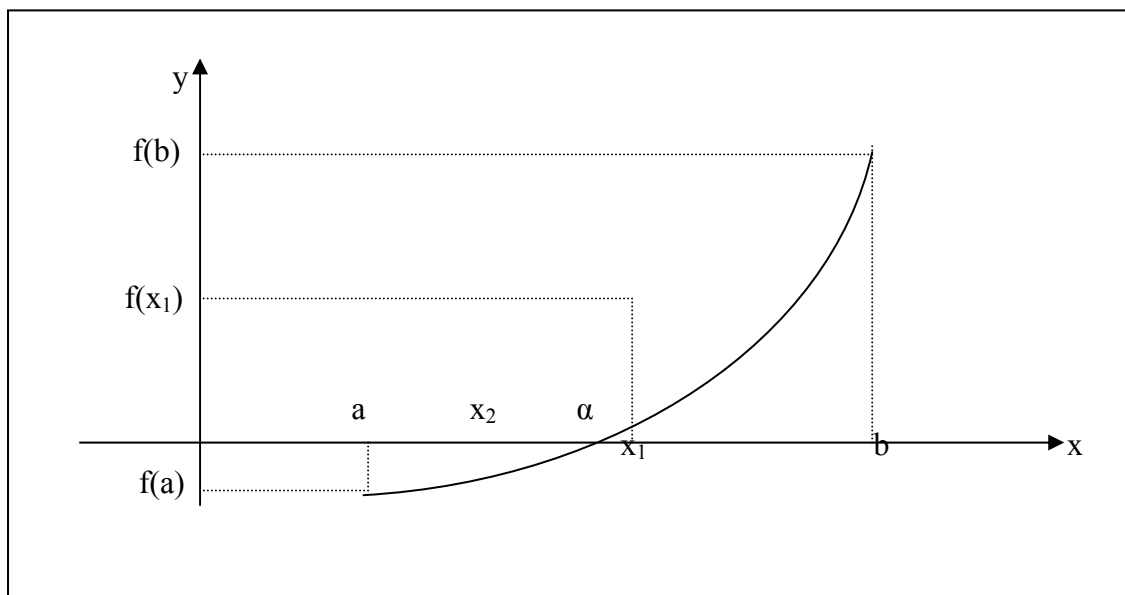
MÉTODO DA BISSECÇÃO OU DICOTOMIA

O método da bissecção consiste em dividir o intervalo $[a, b]$ ao meio, obtendo assim os intervalos $[a, x_1]$ e $[x_1, b]$, onde $x_1 = \frac{a+b}{2}$.

Vamos analisar o que pode ocorrer:
$$\begin{cases} f(x_1) = 0, \text{ assim } x_1 = \alpha \text{ é a raiz procurada (fim!)} \\ f(x_1) > 0 \text{ ou } f(x_1) < 0 \text{ nestes casos, comparamos com } f(a) \\ \text{e } f(b) \text{ para podermos aplicar o teorema de Bolzano...} \end{cases}$$

Se, por exemplo, $f(a) < 0$, $f(b) > 0$ e $f(x_1) > 0$, a raiz procurada α estiver entre a e x_1 , repetimos o processo acima para o intervalo $[a, x_1]$, encontrando o ponto x_2 e assim sucessivamente ... “espremendo” a raiz cada vez mais.

Veja o gráfico abaixo e o exemplo que segue.



Note que o primeiro erro que se comete é menor que $\frac{b-a}{2}$, isto é $|x_1 - \alpha| < \frac{b-a}{2}$, o segundo é

$|x_2 - \alpha| < \frac{b-a}{2^2} = \frac{b-a}{4}$, em geral, na “quebra” do intervalo (ou seja, na iteração) de ordem n , o erro é

menor que $\frac{b-a}{2^n}$, assim $|x_n - \alpha| < \frac{b-a}{2^n}$.

EXEMPLO: Determine o zero positivo de $f(x) = x^2 - 3$, com seis iterações (critério de parada).

RESOLUÇÃO: é claro que estamos procurando o valor de $\sqrt{3}$!

Inicialmente vamos isolar a raiz em um intervalo.

Os métodos clássicos são

i) O método gráfico;

No caso em questão, o gráfico é bastante conhecido, uma parábola, que intercepta o eixo Ox em dois pontos simétricas em relação à origem. Investigando um pouco concluímos que existe um zero da função no intervalo $[1, 2]$.

ii) Dando valores a x e estudando certas características da função (uma delas, por exemplo, a derivada).

Atribuindo valores a x , teremos:

$$\begin{cases} f(0) = 0^2 - 3 = -3 < 0 \\ f(1) = 1^2 - 3 = -2 < 0 \\ f(2) = 2^2 - 3 = 1 > 0 \end{cases} \text{ . Desde que } f(1) \text{ é positivo e } f(2) \text{ é negativo, pelo teorema de Bolzano, há, no}$$

mínimo, um zero de f no intervalo $[1, 2]$. Considerando que $f(x) = x^2 - 3 \Rightarrow f'(x) = 2x$ e restringindo o estudo da derivada ao intervalo $[1, 2]$, vemos que a derivada é positiva, isto é, $f' > 0$ então, com certeza há apenas uma raiz nesse intervalo. Portanto o intervalo inicial que usaremos para aplicar o método da bissecção é $[1, 2]$. Veja a tabela abaixo:

n	a_n	b_n	x_n	$f(x_n)$	<i>Erro (menor que)</i>
1	1,00000	2,00000	1,50000	-0,75000	0,50000
2	1,50000	2,00000	1,75000	0,06250	0,25000
3	1,50000	1,75000	1,62500	-0,35938	0,12500
4	1,62500	1,75000	1,68750	-0,15234	0,06250
5	1,68750	1,75000	1,71875	-0,04590	0,03125
6	1,71875	1,75000	1,73438	0,00806	0,01563

Concluimos que a melhor estimativa para a raiz é 1,73438.

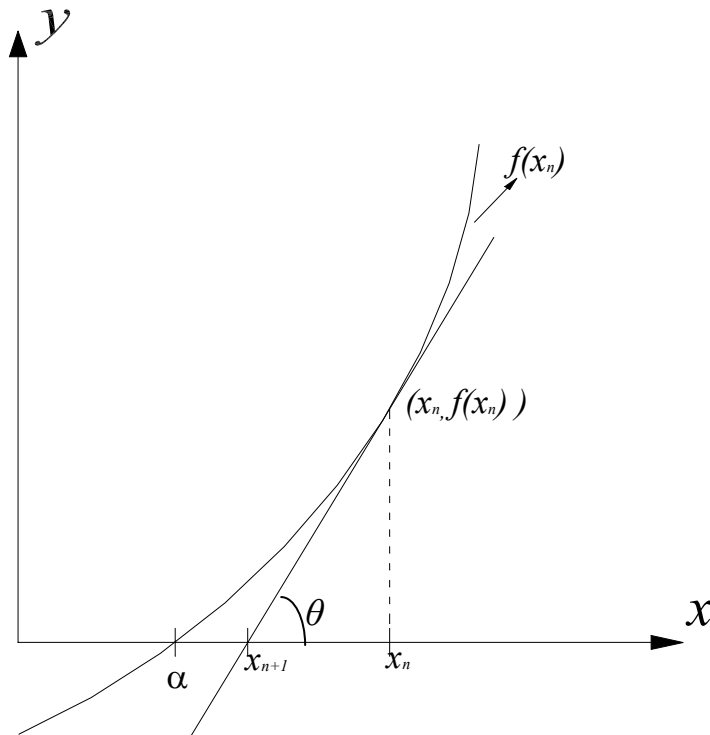
OBSERVAÇÃO: Alguns critérios de parada são:

- 1) Número de iterações;
- 2) $|f(x_n)| < \varepsilon$, onde ε é a tolerância exigida, a depender de cada problema;

- 1) Ache o zero da função $f(x) = x^3 - 9x + 3$ no intervalo $[0,1]$, pelo método da bissecção, utilizando cinco casas decimais e $|f(x_n)| < 0,0001$ como critério de parada.
- 2) Encontre o zero da função $f(x) = 2x^3 - \ln x + 3$, pelo método da bissecção, utilizando cinco casas decimais e como critério de parada *Erro menor que 0,03125*.
- 3) Determine o zero da função $f(x) = x - e^{x-2}$, usando $x_0 = 0,4$, pelo método de Newton-Raphson, utilizando cinco casas decimais e três iterações como critério de parada

MÉTODO DE NEWTON-RAPHSON

Suponha que ocorra a seguinte situação:

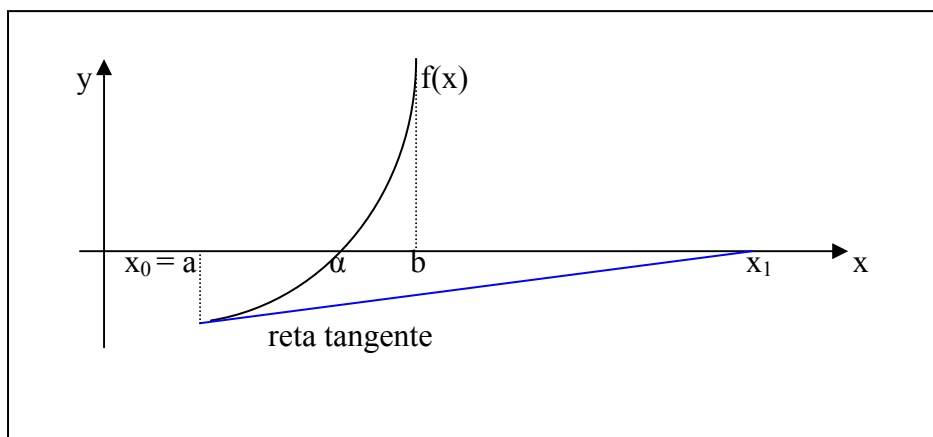


$$\operatorname{tg} \theta = f'(x_n) = \frac{f(x_n)}{x_n - x_{n+1}} \Rightarrow x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$$

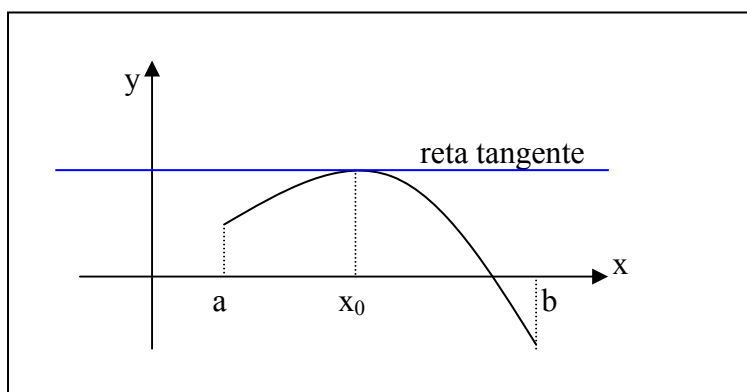
De um modo geral:

Seja f uma função contínua em $[a, b]$ e α seu único zero em $[a, b]$, suponha que $f \in C^2$ (isto é f' e f'' sejam contínuas) e ainda, considere que $f' \neq 0$ em $[a, b]$. Com o objetivo de determinar uma aproximação para a raiz α , usamos o processo iterativo: $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$. (Que, $\odot$, nem sempre converge!). Veja algumas outras possíveis situações abaixo:

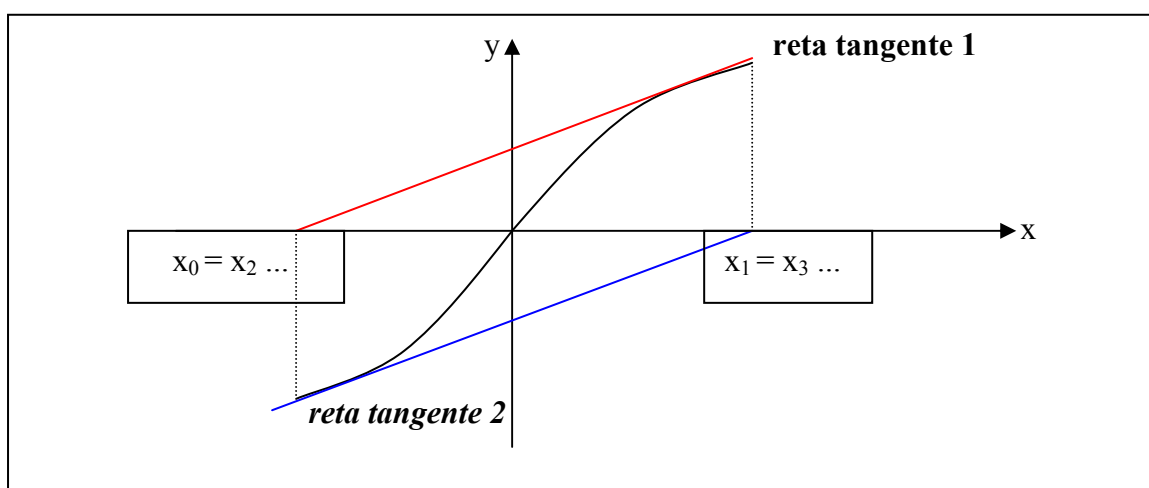
a) Caso em que x_1 “cai” fora do intervalo $[a, b]$ (a escolha do x_0 não foi legal! Possivelmente $x_0 = b$ era uma escolha melhor...)



b) Caso em que $f'(x_0) = 0$ (péssima escolha!)



c) Caso do “loop” infinito



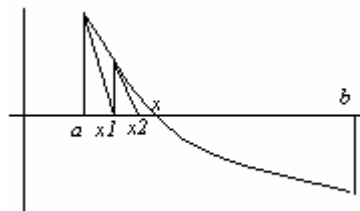
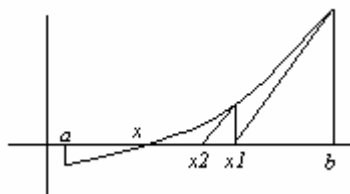
Portanto, de um modo geral, a depender de x_0 , o processo pode convergir ou não!

Critério enunciado por **Barroso**, página 125:

Se f' e f'' são não nulas e preservarem o sinal em $[a, b]$ (onde há uma raiz isolada α de f) e x_0 (valor inicial) seja tal que $f(x_0) \cdot f''(x_0) > 0$ então o método de Newton converge.

Além disso, temos o **Critério de Fourier** para o método de Newton-Raphson:

- Se $f(a) \cdot f''(a) > 0 \Rightarrow x_0 = a$ (veja os gráficos abaixo!)
- Se $f(b) \cdot f''(b) > 0 \Rightarrow x_0 = b$



Outro resultado é o teorema encontrado em Ruggiero e Lopes (página 69):

Sejam $f(x)$, $f'(x)$ e $f''(x)$ contínuas em um intervalo I que contém a raiz isolada $x = \alpha$ de $f(x) = 0$.

Suponhamos que $f'(\alpha) \neq 0$. Então existe um intervalo $J \subset I$ tal que $\forall x_0 \in J$, a sequência $\{x_k\}$ gerada

pela fórmula recursiva $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$ converge para α .

Prezados alunos, como vimos anteriormente o método gráfico pode ser extremamente útil para decidirmos sobre a localização e isolamento da raiz. Também podemos recorrer a ele para a escolha do valor inicial da iteração x_0 .

EXEMPLO 1:

Determine $\sqrt{3}$ pelo método de Newton-Raphson. (Discutiremos o critério de parada e a estimativa de erro no próprio exercício). Usaremos cinco decimais com o arredondamento tradicional.

RESOLUÇÃO: Note que achar $\sqrt{3}$ é o mesmo que determinar o zero positivo de $f(x) = x^2 - 3$,

pois: $x = \sqrt{3} \Rightarrow x^2 = 3 \Rightarrow x^2 - 3 = 0$. Lembre que, no método da bissecção nós fizemos este mesmo problema, usando 6 iterações. O que queremos provar é que usando o método de Newton-Raphson a convergência é muito mais rápida. Podemos usar o método gráfico e também dar valores a x ; de qualquer sorte, como visto anteriormente, nós sabemos que existe uma única raiz no intervalo $[1, 2]$. Testando as hipóteses do critério do Barroso:

$f(x) = x^2 - 3$; $f'(x) = 2x$; $f''(x) = 2$. Considerando o intervalo $[1, 2]$, as derivadas não se anulam e também preservam o sinal (na verdade ambas são sempre positivas). Devemos agora escolher um x_0 conveniente. Pelo critério de Fourier, enunciado acima, tomaremos $x_0 = 2$ (pois $f(2) \cdot f''(2) > 0$).

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$$

n	x_n	$f(x_n) = x_n^2 - 3$	$f'(x_n) = 2x_n$	$\frac{f(x_n)}{f'(x_n)}$	x_{n+1}
0	$x_0 = 2$	1	4	0,25	$x_1 = 1,75$
1	$x_1 = 1,75$	0,0625	3,5	0,01786	$x_2 = 1,73214$
2	$x_2 = 1,73214$	0,00031	3,46428	0,00009	$x_3 = 1,73205$

Com duas iterações tínhamos $x_2 = 1,73214$ (melhor que o resultado obtido pelo método da bissecção, no qual fizemos seis iterações...). Finalmente com três iterações, temos $x_3 = 1,73205$, assim a raiz pedida é aproximadamente 1,73205, isto é, $\alpha \approx 1,73205$.

EXEMPLO 2:

Calcule a raiz de $f(x) = x^2 + x - 6$, usando o método de Newton-Raphson, $x_0 = 3$ como estimativa inicial e como critério de parada $|f(x_n)| \leq 0,020$.

Para encontrar a raiz de $f(x)$ usando o método de Newton-Raphson, devemos ter:

$$x_{n+1} = \varphi(x_n)$$

onde,

$$\varphi(x) = x - \frac{f(x)}{f'(x)}$$

$$\varphi(x) = x - \frac{x^2 + x - 6}{2 \cdot x + 1} = \frac{2 \cdot x^2 + x - x^2 - x + 6}{2 \cdot x + 1} = \frac{x^2 + 6}{2 \cdot x + 1}$$

Portanto, temos que:

x_n	$f(x_n)$	$\varphi(x_n)$
3	6	2,1429
2,1429	0,7349	2,0039
2,0039	0,0195	

A estimativa da raiz de $f(x)$ é: $\bar{x} = 2,0039$

OBSERVAÇÃO: Alguns critérios de parada para o método de Newton-Raphson são:

- 1) Número de iterações;
- 2) $|f(x_n)| < \varepsilon$, onde ε é a tolerância exigida, a depender de cada problema;
- 3) $\left| \frac{f(x_n)}{f'(x_n)} \right| < \varepsilon$;

EXERCÍCIOS

- 1) Ache o zero da função $f(x) = x^3 - 9x + 3$ no intervalo $[0,1]$, usando $x_0 = 0,5$, pelo método de Newton-Raphson, utilizando cinco casas decimais e $|f(x_n)| < 0,0001$ como critério de parada.
- 2) Encontre o zero da função $f(x) = 2x^3 - \ln x + 3$, usando $x_0 = 2$, pelo método de Newton-Raphson, utilizando cinco casas decimais e $|f(x_n)| < 0,0001$ como critério de parada.
- 3) Determine o zero da função $f(x) = x - e^{x-2}$, usando $x_0 = 0,4$, pelo método de Newton-Raphson, utilizando cinco casas decimais e quatro iterações como critério de parada

MÉTODO DAS APROXIMAÇÕES SUCESSIVAS (OU ITERAÇÃO LINEAR)

Seja f uma função contínua em $[a, b]$ e α seu único zero em $[a, b]$.

Por um artifício **sempre** podemos transformar $f(x) = 0$ em $x = \varphi(x)$ (basta considerar, por exemplo, $x + f(x) = \varphi(x)$ ou ainda, $x - f(x) = \varphi(x)$. Note, portanto, que $\varphi(x)$ não é única...).

EXEMPLO: Encontre algumas funções de iteração a partir de $f(x) = x^2 + \ln x - x + 1$.

$f(x) = x^2 + \ln x - x + 1$ fazendo $f(x) = 0$, temos:

$$a) x^2 + \ln x - x + 1 = 0 \therefore x = x^2 + \ln x + 1$$

$$\therefore \varphi_1(x) = x^2 + \ln x + 1$$

$$b) x^2 + \ln x - x + 1 = 0$$

$$\ln x = x - x^2 - 1 \therefore x = e^{(x-x^2-1)}$$

$$\therefore \varphi_2(x) = e^{(x-x^2-1)}$$

$$c) x^2 + \ln x - x + 1 = 0$$

$$x \cdot x = x - \ln x - 1 \therefore x = \frac{x - \ln x - 1}{x}$$

$$\therefore \varphi_3(x) = \frac{x - \ln x - 1}{x}$$

ou ainda

$$d) x^2 + \ln x - x + 1 = 0 \text{ (somando e subtraindo } \cos x \text{)}$$

$$x^2 + \ln x - x + 1 + \cos x - \cos x = 0 \therefore \cos x = \cos x - x^2 - \ln(x) + x - 1$$

$$\therefore x = \arccos(\cos x - x^2 - \ln(x) + x - 1)$$

$$\therefore \varphi_4(x) = \arccos(\cos x - x^2 - \ln(x) + x - 1)$$

Considerando x_0 uma primeira aproximação de α , calcula-se $\varphi(x_0)$. Faz-se $x_1 = \varphi(x_0)$; $x_2 = \varphi(x_1)$;

$x_3 = \varphi(x_2)$, e assim sucessivamente, isto é, $x_{n+1} = \varphi(x_n)$, $n = 0, 1, 2, \dots$. Se a sequência convergir temos

então $\lim_{n \rightarrow +\infty} x_n = \alpha$.

OBSERVAÇÃO: $\varphi(x)$ é dita **uma** função de iteração de $f(x) = 0$.

Vejamos um importante teorema (condição suficiente):

TEOREMA DA CONVERGÊNCIA (Ruggiero e Lopes, página 58): Seja $\alpha \in I$ uma raiz isolada de $f(x) = 0$, onde I é um intervalo centrado em α . Considere $\varphi(x)$ uma função de iteração de $f(x) = 0$ (isto é, $\varphi(\alpha) = \alpha$, onde φ também é definida em I), com $\varphi(x)$ derivável. Se $|\varphi'(x)| \leq M < 1$ para todos os pontos em I e $x_0 \in I$, então os valores em $x_{n+1} = \varphi(x_n)$, $n = 0, 1, 2, \dots$ convergem para α .

DEMONSTRAÇÃO:

PRIMEIRA PARTE: Provaremos que, se $x_{k-1} \in I$ então $x_k \in I$, $\forall k \in \mathbb{N}^*$ (isto garante a aplicação do processo iterativo $x_{n+1} = \varphi(x_n)$, $n = 0, 1, 2, \dots$).

$f(\alpha) = 0 \Leftrightarrow \alpha = \varphi(\alpha)$ e seja $x_k = \varphi(x_{k-1})$ (o caso $x_{k-1} = \alpha$ é trivial!), subtraindo membro a membro:
 $x_k - \alpha = \varphi(x_{k-1}) - \varphi(\alpha)$, desde que, por hipótese φ é derivável, podemos aplicar o teorema do valor médio no intervalo $[x_{k-1}, \alpha]$ (ou possivelmente em $[\alpha, x_{k-1}]$, vamos supor, para fixar idéias $[x_{k-1}, \alpha]$).

Deste modo, $\exists c_k \in [x_{k-1}, \alpha]$ tal que $\varphi(x_{k-1}) - \varphi(\alpha) = \varphi'(c_k) \cdot (x_{k-1} - \alpha)$, ou seja:

$x_k - \alpha = \varphi'(c_k) \cdot (x_{k-1} - \alpha)$, como $|\varphi'(x)| \leq M < 1, \forall x \in I$, podemos escrever:

$|x_k - \alpha| \leq M \cdot |x_{k-1} - \alpha| < |x_{k-1} - \alpha|$, logo x_k está mais próximo de α que x_{k-1} , isto é, $x_k \in I$.

Concluimos que, dado um ponto $x_0 \in I$, o próximo $x_1 = \varphi(x_0) \in I$ e assim sucessivamente...

SEGUNDA PARTE: Provaremos que a seqüência $x_{n+1} = \varphi(x_n)$, $n = 0, 1, 2, \dots$ converge para α , ou seja,

$$\lim_{k \rightarrow +\infty} x_k = \alpha.$$

$$|x_1 - \alpha| \leq M \cdot |x_0 - \alpha|$$

$$|x_2 - \alpha| \leq M \cdot |x_1 - \alpha| \leq M^2 \cdot |x_0 - \alpha|$$

$$|x_3 - \alpha| \leq M \cdot |x_2 - \alpha| \leq M^3 \cdot |x_0 - \alpha|$$

$\vdots$

$$|x_k - \alpha| \leq M \cdot |x_{k-1} - \alpha| \leq M^k \cdot |x_0 - \alpha|$$

Passando ao limite:

$$0 \leq \lim_{k \rightarrow +\infty} |x_k - \alpha| \leq \lim_{k \rightarrow +\infty} \left[M^k \cdot \underbrace{|x_0 - \alpha|}_{\text{constante}} \right], \text{ desde que } \lim_{k \rightarrow +\infty} M^k \text{ pois } 0 \leq M < 1, \text{ temos:}$$

$$\lim_{k \rightarrow +\infty} |x_k - \alpha| = 0 \Rightarrow \left| \lim_{k \rightarrow +\infty} (x_k - \alpha) \right| = 0 \Rightarrow \lim_{k \rightarrow +\infty} (x_k - \alpha) = 0 \Rightarrow \lim_{k \rightarrow +\infty} x_k = \alpha.$$

Como vimos acima, a função de iteração φ não é única. O nosso foco será encontrar uma função de iteração que satisfaça às hipóteses do teorema acima. Vamos praticar...

Seja calcular a **raiz positiva** de $f(x) = x^2 - x - 2$ (é claro que já sabemos que $x' = 2$ e $x'' = -1$).

Primeiramente, vamos listar algumas possibilidades para $\varphi(x)$ (função de iteração):

a) $x^2 - x - 2 = 0 \therefore x = \underbrace{x^2 - 2}_{\varphi_1(x)}$, isto é, $\varphi_1(x) = x^2 - 2$;

b) $x^2 - x - 2 = 0 \therefore x^2 = x + 2 \therefore x \cdot x = x + 2 \therefore x = \underbrace{\frac{x+2}{x}}_{\varphi_2(x)} \therefore x = 1 + \underbrace{\frac{2}{x}}_{\varphi_2(x)}$, ou seja, $\varphi_2(x) = 1 + \frac{2}{x}$, para $x \neq 0$;

c) $x^2 - x - 2 = 0 \therefore x^2 = x + 2 \therefore x = \pm \sqrt{x+2} \therefore \begin{cases} \varphi_3(x) = \sqrt{x+2} \\ \varphi_4(x) = -\sqrt{x+2} \end{cases}$ para $x \geq -2$

Sem testar as hipóteses do teorema acima, vamos trabalhar com $\varphi_1(x) = x^2 - 2$ e $\varphi_3(x) = \sqrt{x+2}$ com a intenção de observar a convergência na vizinhança de $x = 2$ (que nós já sabemos ser a raiz positiva de $f(x) = x^2 - x - 2$). Escolheremos como vizinhança o intervalo $[1,3]$ e como aproximação inicial $x_0 = 2,5$.

(i) Para $\varphi_1(x) = x^2 - 2$, temos o processo iterativo $x_{n+1} = \varphi_1(x_n) = x_n^2 - 2$:

n	x_n	$\varphi_1(x_n) = x_n^2 - 2$	x_{n+1}
0	$x_0 = 2,5$	$\varphi_1(x_0) = x_0^2 - 2 = 2,5^2 - 2 = 4,25$	$x_1 = 4,25$
1	$x_1 = 4,25$	$\varphi_1(x_1) = x_1^2 - 2 = 4,25^2 - 2 = 16,0625$	$x_2 = 16,0625$
2	$x_2 = 16,0625$	$\varphi_1(x_2) = x_2^2 - 2 = 16,0625^2 - 2 = 256,00391$	$x_3 = 256,00391$

Observando a quarta coluna podemos observar que a seqüência $x_0, x_1, x_2, \dots$ não é convergente. Pode pintar uma dúvida em sua mente... será que se o x_0 fosse tomado mais perto de $x = 2$, a função de iteração convergiria? Bom, como primeiro passo, aconselhamos que você experimente começar com $x_0 = 2,01$ e, após algumas iterações, você constatará que não haverá convergência. Tente de novo, agora para $x_0 = 1,99$ e verifique a divergência. Agora concluamos juntos que o “defeito”, no caso em estudo,

não está em x_0 mas possivelmente na escolha da função de iteração $\varphi_1(x)$. Vamos então, considerar a próxima função de iteração $\varphi_3(x) = \sqrt{x+2}$.

(ii) Para $\varphi_3(x) = \sqrt{x+2}$, temos o processo iterativo $x_{n+1} = \varphi_3(x_n) = \sqrt{x_n+2}$:

n	x_n	$x_{n+1} = \varphi_3(x_n) = \sqrt{x_n+2}$	x_{n+1}
0	$x_0 = 2,5$	$\varphi_3(x_0) = \sqrt{x_0+2} = \sqrt{2,5+2} = \sqrt{4,5} = 2,12132$	$x_1 = 2,12132$
1	$x_1 = 2,12132$	$\varphi_3(x_1) = \sqrt{2,12132+2} = \sqrt{4,12132} = 2,03010$	$x_2 = 2,03010$
2	$x_2 = 2,03010$	$\varphi_3(x_1) = \sqrt{2,03010+2} = \sqrt{4,0310} = 2,00751$	$x_3 = 2,00751$

Observando a quarta coluna podemos observar que a seqüência $x_0, x_1, x_2, \dots$ é convergente. Portanto, considerando três iterações, temos $\alpha \approx 2,00751$.

Ufa! Conseguimos! Agora, gostaria de chamar sua atenção para o fato de que a função $\varphi_3(x) = \sqrt{x+2}$ satisfaz às hipóteses do teorema da convergência da página 25 (isto é, a convergência era esperada...)

pois $\varphi'_3(x) = \frac{1}{2\sqrt{x+2}} \leq \frac{1}{2\sqrt{3}} < 1$ em $[1,3]$.

OBSERVAÇÃO: Se desejássemos calcular a **raiz negativa** de $f(x) = x^2 - x - 2$ a função de iteração seria $\varphi_4(x) = -\sqrt{x+2}$ (Porquê ?).

OBSERVAÇÃO: O Método de Newton-Raphson é um caso particular do método de iteração linear.

O método de iteração linear consiste em estimar a raiz de uma função $f(x)$ usando o processo iterativo:

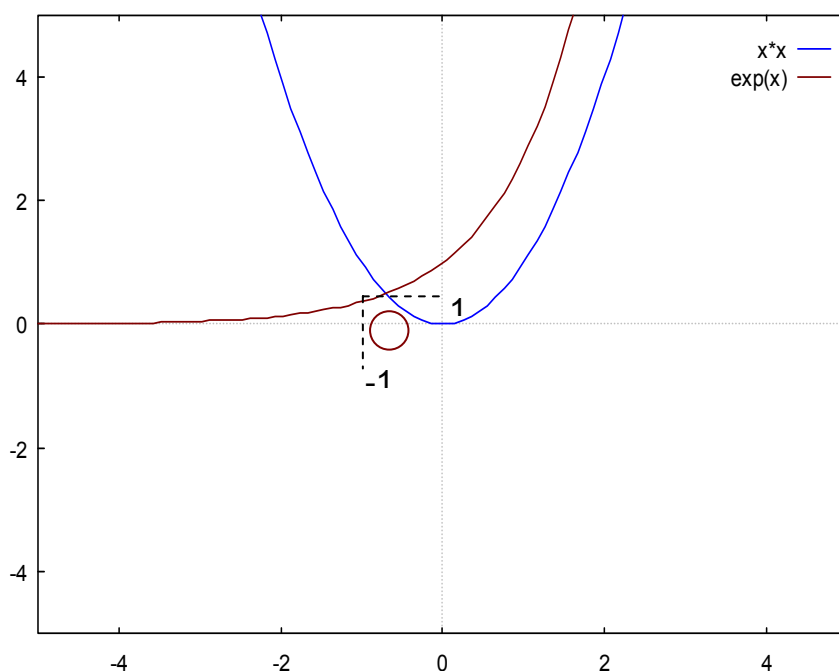
$x_{n+1} = \varphi(x_n)$. Podemos escrever uma forma geral para a função de iteração:

$\varphi(x) = x + A(x).f(x)$ pois para x igual à raiz de $f(x)$, tem-se $f(x)=0$, ou seja $x=\varphi(x)$ para qualquer $A(x) \neq 0$.

EXERCÍCIO RESOLVIDO: Encontre uma estimativa para a **raiz negativa** de $f(x) = x^2 - e^x$ usando o método da iteração linear.

Vamos iniciar a solução encontrando uma boa estimativa inicial para o valor da raiz de $f(x)$. Para isso, vamos usar o método gráfico para o isolamento de raízes. Escrevendo:

$$f(x) = g(x) - h(x) \Rightarrow g(x) = x^2 \text{ e } h(x) = e^x \text{ temos:}$$



A partir do esboço gráfico acima, conclui-se que a **raiz negativa** encontra-se no intervalo $[-1, 0]$.

Devemos agora escolher uma função de iteração $\varphi(x)$. Para isso, escrevemos:

$$f(x) = 0 \therefore x^2 - e^x = 0 \therefore x = \pm \sqrt{e^x}$$

Ou seja, podemos ter como função iteração, os dois casos abaixo:

$$\varphi(x) = \sqrt{e^x}$$

$$\varphi(x) = -\sqrt{e^x}$$

Usando $\varphi(x) = -\sqrt{e^x}$ (pois desejamos o zero negativo!) e $x_0 = -1$, temos:

$$x_0 = -1 \rightarrow \varphi(x_0) = \varphi(-1) = -\sqrt{e^{-1}} = -0,606$$

$$x_1 = -0,606 \rightarrow \varphi(x_1) = \varphi(-0,606) = -\sqrt{e^{-0,606}} = -0,738$$

$$x_2 = -0,738 \rightarrow \varphi(x_2) = \varphi(-0,738) = -\sqrt{e^{-0,738}} = -0,691$$

$$x_3 = -0,691 \rightarrow \varphi(x_3) = \varphi(-0,691) = -\sqrt{e^{-0,691}} = -0,707$$

$$x_4 = -0,707$$

Verifique se as hipóteses do teorema da convergência, da página 26 são satisfeitas, isto é:

TEOREMA DA CONVERGÊNCIA (Ruggiero e Lopes, página 58): Seja $\alpha \in I$ uma raiz isolada de $f(x) = 0$, onde I é um intervalo centrado em α . Considere $\varphi(x)$ uma função de iteração de $f(x) = 0$ (isto é, $\varphi(\alpha) = \alpha$, onde φ também é definida em I), com $\varphi(x)$ derivável. Se $|\varphi'(x)| \leq M < 1$ para todos os pontos em I e $x_0 \in I$, então os valores em $x_{n+1} = \varphi(x_n)$, $n = 0, 1, 2, \dots$ convergem para α .

OBSERVAÇÃO: Substituindo os valores de x_k em $f(x)$ para cada iteração k , vemos que a cada etapa nos aproximamos mais da raiz de $f(x)$, pois o valor dessa função fica mais próximo de zero a cada iteração, como mostrado na tabela abaixo:

x	$f(x) = x^2 - e^x$
-1	0,632
-0,606	-0,178
-0,738	0,067
-0,691	-0,024
-0,707	0,007

PRATIQUE

- 1) Determine a raiz de $f(x) = x^2 - \sin x$ no intervalo $[0,5;1]$, até a quarta iteração ($x_0 = 0,9$).
- 2) Seja $f(x) = x^3 - x - 1$
 - a) Determine um intervalo de f contendo um zero positivo.
 - b) Ache algumas funções de iteração.
 - c) Use uma delas (que convirja, é claro!) e, com três iterações, determine uma aproximação para a raiz que se situa no intervalo obtido no item a).

LISTA DE EXERCÍCIOS

1. Quais são as 3 causas mais importantes de erros numéricos em operações realizadas em computadores e calculadoras?
2. Cite as características básicas de todo processo iterativo.
3. O que é um zero ou raiz de função?
4. Como você poderia usar o método da bissecção para estimar o valor de $\sqrt{7}$? Estime esse valor com uma precisão de (ou erro menor que) 0,1.
5. Dada a função $f(x) = \sin x - x^2 + 4$:
 - (a) Determine o intervalo em x que contém pelo menos uma raiz de $f(x)$ (graficamente ou aritmeticamente usando o Teorema de Bolzano);
 - (b) Partindo-se desse intervalo, utilize o método da bissecção para determinar o valor dessa raiz após 4 iterações.
 - (c) Qual é o erro no seu resultado final?
6. Dada a função $f(x) = e^{-x} + x^2 - 2$:
 - (a) Determine graficamente o intervalo em x que contém pelo menos uma raiz de $f(x)$;
 - (b) Faça a mesma estimativa, mas desta vez aritmeticamente usando o Teorema de Bolzano;
 - (c) Partindo-se desse intervalo, utilize o método da bissecção para determinar o valor dessa raiz com uma precisão de 0,05.
7. O que significa a convergência de um método iterativo? Que condições garantem a convergência no método da iteração linear? O que fazer caso seja constatado que o método da iteração linear não irá convergir para um dado problema?
8. Dada a função $f(x) = \ln x - x^2 + 4$, mostre 3 formas para a função $\phi(x)$ que poderiam ser usadas para se estimar a raiz de $f(x)$.
9. Mostre que as seguintes funções de iteração satisfazem as condições (i) e (ii) do teorema de convergência:

$$(a) \varphi(x) = \frac{x^3}{9} + \frac{1}{3}$$

$$(b) \varphi(x) = \frac{\cos x}{2}$$

$$(c) \varphi(x) = \frac{\exp(x/2)}{2} = \frac{e^{x/2}}{2}$$

$$(d) \varphi(x) = (x+1)^{1/3}$$

Estime as raízes positivas das seguintes funções pelo método de iteração linear, usando o critério de parada como sendo de quatro iterações (Use as funções de iteração do exercício anterior).

$$(a) f(x) = x^3 - 9x + 3$$

$$(b) f(x) = 2x - \cos x$$

$$(c) f(x) = e^x - 4x^2$$

$$(d) f(x) = x^3 - x - 1$$

10. Seja a seguinte função: $f(x) = \frac{1}{x^3} - x^2 + 1$

Use o método de Newton-Raphson para encontrar uma estimativa da raiz de $f(x)$ tal que $|f(x)| < 10^{-4}$. Parta de $x_0 = 1$.

11. Seja a função $f(x) = e^x - 4x^2$.

(a) Encontre o intervalo que deva possuir pelo menos uma raiz de $f(x)$.

(b) Usando $\varphi(x) = \frac{1}{2}e^{x/2}$, estime a raiz de $f(x)$ com $|x_n - x_{n-1}| < 0,001$.

(c) Faça a mesma estimativa usando o método de Newton-Raphson. Qual dos dois métodos converge mais rapidamente?

(d) Um outro critério de parada que poderia ser usado corresponde à verificação se o valor de $f(x)$ está próximo de zero. Qual resultado para a raiz de $f(x)$ se obteria caso se usasse como critério de parada a condição $|f(x)| < 0,001$?

13. Seja: $f(x) = x^3 - 3x + 1$

(a) Mostre que $f(x)$ possui uma raiz em $[0,1]$.

(b) Mostre que $\varphi(x) = \frac{x^3 + 1}{3}$ é uma possível função de iteração obtida a partir de $f(x)$.

(c) Verifique se $\varphi(x)$ satisfaz as condições (i) e (ii) do Teorema de Convergência.

(d) Encontre uma estimativa para a raiz de $f(x)$ através do método da iteração linear e usando a função $\varphi(x)$ do item (c), tal que $|f(x)| < 0,0070$.

(e) Faça a mesma estimativa, mas desta vez ao invés de utilizar a função $\varphi(x)$ do item (c), utilize o método de Newton-Raphson.

SISTEMAS LINEARES

INTRODUÇÃO

As equações não existem por si, ou seja, não são invenções abstratas da Matemática. Muito pelo contrário, decorrem de situações concretas de nosso cotidiano. Veja os seguintes exemplos e suas respectivas representações na linguagem matemática:

a) A diferença entre as idades de Magno (x) e Marcos (y) é de 4 anos: $x - y = 4$

b) Numa fábrica trabalham 532 pessoas entre homens (x) e mulheres (y). O número de homens é o triplo do número de mulheres: $x + y = 532$ e $x = 3y$

Os exemplos citados representam **equações lineares** e, ao conjunto destas, chamamos de **Sistemas Lineares**.

A resolução de sistemas lineares é um problema que surge em diversas áreas do conhecimento e ocorre, na prática, com muita frequência. Por exemplo: cálculo de estruturas na Construção Civil, cálculo do ponto de equilíbrio de mercado na Economia e dimensionamento de redes elétricas.

EQUAÇÃO LINEAR

Entende-se por equação linear toda expressão da forma $a_1x_1 + a_2x_2 + a_3x_3 + \dots + a_nx_n = b$ onde:

$x_1, x_2, x_3, \dots, x_n$ são incógnitas ou termos desconhecidos e $a_1, a_2, a_3, \dots, a_n$ são números reais chamados coeficientes e b é um número real chamado termo independente ou seja, **em cada termo** da equação linear aparece uma única incógnita e seu expoente é sempre igual a 1

EXEMPLOS:

a) $2x_1 + x_2 = 12$ ou $2x + y = 12$
b) $x_1 + 2x_2 - 3x_3 = 15$ ou $x + 2y - 3z = 15$
c) $3x_1 - 4x_2 + x_3 - 5x_4 = 10$ ou $3x - 4y + z - 5w = 10$

SOLUÇÃO DE UMA EQUAÇÃO LINEAR

Uma solução de uma equação linear é uma seqüência de números reais $(k_1, k_2, k_3, \dots, k_n)$, tal que, substituindo-se respectivamente as incógnitas da equação pelos números reais $k_1, \dots, k_n$, na ordem em que se apresentam, verifica-se a igualdade, ou seja,

$$a_1k_1 + a_2k_2 + a_3k_3 + \dots + a_nk_n = b$$

EXEMPLOS:

a) Uma solução da equação $2x + 4y = 22$ é o par $(5, 3)$
b) Uma solução da equação $3x + 2y - 5z = 32$ é a terna $(2, 3, -4)$
c) Uma solução da equação $x + 2y - 4z + w = 3$ é a quadra $(3, 2, 1, 0)$

OBSERVAÇÕES IMPORTANTES

1. Se $a_1 = a_2 = a_3 = \dots = a_n = b = 0$ então qualquer ênupla $(k_1, k_2, k_3, \dots, k_n)$ é solução.

2. Se $a_1 = a_2 = a_3 = \dots = a_n = 0$ e $b \neq 0$ então a equação não admite solução.

SISTEMAS LINEARES

Chama-se **SISTEMA LINEAR** o conjunto de duas ou mais equações lineares.

EXEMPLOS:

$$\begin{array}{ll} a) \begin{cases} x - y = 4 \\ 3x + y = 2 \end{cases} & b) \begin{cases} x + 2y - z = 0 \\ 2x - y + z = 1 \end{cases} \\ c) \begin{cases} x + y - 2z = 0 \\ 2x - y + z = 3 \\ 7x - 2y - z = 9 \end{cases} & d) \begin{cases} x - y = 2 \\ 3x + 2y = 1 \\ 8x + 2y = 6 \end{cases} \end{array}$$

Genericamente, um sistema linear S de **m** equações e **n** incógnitas é escrito:

$$S = \begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 + \dots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 + \dots + a_{2n}x_n = b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 + \dots + a_{3n}x_n = b_3 \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + a_{m3}x_3 + \dots + a_{mn}x_n = b_m \end{cases}$$

Abreviadamente, o sistema linear é representado por: $S = \sum_{j=1}^n a_{ij} x_j = b_i, i = 1, 2, 3, \dots, m$

SOLUÇÃO DE UM SISTEMA LINEAR

Dizemos que a ênupla $(k_1, k_2, k_3, \dots, k_n)$ é solução de um sistema linear se verificar, **simultaneamente**, todas as equações do sistema

EXEMPLOS:

$$\begin{array}{ll} a) \text{ o par } (5, 1) \text{ é solução do sistema } \begin{cases} 2x + 3y = 13 \\ 3x - 5y = 10 \end{cases} \\ b) \text{ o terno } (1, 3, -2) \text{ é solução do sistema } \begin{cases} x + 2y + 3z = 1 \\ 4x - y - z = 3 \\ x + y - z = 6 \end{cases} \end{array}$$

MATRIZES ASSOCIADAS A UM SISTEMA LINEAR

De um modo geral, qualquer sistema linear pode ser escrito na forma matricial: $A \cdot X = B$.

onde $A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & & a_{2n} \\ a_{31} & a_{32} & a_{33} & & a_{3n} \\ \dots & & & & \\ a_{m1} & a_{m2} & a_{m3} & & a_{mn} \end{bmatrix}$, $X = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \\ x_n \end{bmatrix}$ e $B = \begin{bmatrix} b_1 \\ b_2 \\ b_3 \\ \vdots \\ b_m \end{bmatrix}$ são, respectivamente, matriz dos coeficientes (incompleta), matriz das incógnitas (solução) e matriz dos termos independentes.

Pode-se associar, também, a um sistema linear uma matriz denominada **matriz completa** (ou matriz ampliada) que é:

$$M = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & a_{23} & & a_{2n} & b_2 \\ a_{31} & a_{32} & a_{33} & & a_{3n} & b_3 \\ \dots & & & & & \\ a_{m1} & a_{m2} & a_{m3} & & a_{mn} & b_m \end{bmatrix}$$

EXEMPLO: Representar matricialmente o sistema

$$A \cdot X = B \Rightarrow \begin{bmatrix} 2 & 1 & 1 \\ 3 & -1 & 0 \\ 4 & 3 & 7 \end{bmatrix} \cdot \begin{bmatrix} x \\ y \\ z \end{bmatrix} = \begin{bmatrix} 8 \\ 7 \\ 2 \end{bmatrix}$$

OBSERVAÇÃO: Há vários métodos para a resolução de sistemas lineares, alguns deles exigem uma maior quantidade de operações que outros, isto é, há um esforço computacional maior. Pesquisem sobre esse assunto (esforço computacional) no livro do professor Barroso.

MÉTODOS DIRETOS (NÃO ITERATIVOS)

1) MÉTODO DA ELIMINAÇÃO DE GAUSS

Através das operações elementares que vocês aprenderam em Álgebra Linear o nosso objetivo será transformar a matriz dos coeficientes do sistema (*) abaixo, $n \times n$, com solução única e tal que $a_{ii} \neq 0, \forall i = 1, \dots, n$ numa matriz triangular superior (isto é, onde todos os elementos abaixo da diagonal principal são iguais a zero), veja o esquema abaixo:

$$(*) \begin{cases} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n = b_2 \\ \vdots \\ a_{n1}x_1 + a_{n2}x_2 + \cdots + a_{nn}x_n = b_n \end{cases} \Rightarrow \begin{cases} k_{11}x_1 + k_{12}x_2 + \cdots + k_{1n}x_n = w_1 \\ k_{22}x_2 + \cdots + k_{2n}x_n = w_2 \\ \vdots \\ k_{nn}x_n = w_n \end{cases}$$

MÉTODOS DIRETO (NÃO ITERATIVO)

MÉTODO DA ELIMINAÇÃO DE GAUSS

Através das operações elementares que vocês aprenderam em Álgebra Linear o nosso objetivo será transformar a matriz dos coeficientes do sistema (*) abaixo, $A = (a_{ij})_{n \times n}$, com solução única e tal que $a_{ii} \neq 0, \forall i = 1, \dots, n$ numa matriz triangular superior (isto é, onde todos os elementos abaixo da diagonal principal da matriz dos coeficientes são zero) matematicamente $K = (k_{ij})_{n \times n}$, onde $k_{ij} = 0$ se $i > j$, veja o esquema abaixo:

OPERAÇÕES ELEMENTARES SOBRE LINHAS DE UMA MATRIZ

EXEMPLO 1: $A = \begin{bmatrix} 2 & 1 \\ 4 & 5 \\ 0 & 2 \end{bmatrix} \xrightarrow[L'_3=L_2]{L'_2=L_3} B = \begin{bmatrix} 2 & 1 \\ 0 & 2 \\ 4 & 5 \end{bmatrix}$ ou, simplesmente $A = \begin{bmatrix} 2 & 1 \\ 4 & 5 \\ 0 & 2 \end{bmatrix} \xrightarrow{L_2 \leftrightarrow L_3} B = \begin{bmatrix} 2 & 1 \\ 0 & 2 \\ 4 & 5 \end{bmatrix}$.

EXEMPLO 2: $C = \begin{bmatrix} 0 & 1 \\ 3 & 4 \\ 5 & -21 \end{bmatrix} \xrightarrow{L'_3 = \frac{-2}{3}L_3} D = \begin{bmatrix} 0 & 1 \\ 3 & 4 \\ -10/3 & 14 \end{bmatrix}.$

EXEMPLO 3: $E = \begin{bmatrix} 2 & 1 & 0 & 4 \\ 7 & -3 & \pi & 1 \\ 2 & 0 & 1 & 3 \end{bmatrix} \xrightarrow{L_3=L_3+(-2)L_1} F = \begin{bmatrix} 2 & 1 & 0 & 4 \\ 7 & -3 & \pi & 1 \\ -2 & -2 & 1 & -5 \end{bmatrix}.$

Seja resolver o sistema:
$$\begin{cases} 2x_1 - 3x_2 + x_3 = -2 \\ 5x_1 - 4x_2 - 2x_3 = 12 \\ -x_1 + 4x_2 + 4x_3 = -10 \end{cases}$$

$A = \begin{bmatrix} 2 & -3 & 1 & -2 \\ 5 & -4 & -2 & 12 \\ -1 & 4 & 4 & -10 \end{bmatrix}$. Nosso objetivo inicial é tornar nulos os elementos 5 e -1 que estão situados

$L'_i = L_i - \frac{a_{ij}}{a_{jj}}.L_j$, onde os elementos $\frac{a_{ij}}{a_{jj}}$ são chamados **multiplicadores** m_{ij} e os elementos a_{jj} da diagonal principal, considerando a matriz dos coeficientes, são chamados de **pivôs** observe:

$$A = \begin{bmatrix} 2 & -3 & 1 & -2 \\ 5 & -4 & -2 & 12 \\ -1 & 4 & 4 & -10 \end{bmatrix} \xrightarrow{\substack{L'_2 = L_2 - \frac{5}{2}L_1 \\ L'_3 = L_3 - \left(-\frac{1}{2}\right)L_1}} \begin{bmatrix} 2 & -3 & 1 & -2 \\ 0 & 3,5 & -4,5 & 17 \\ 0 & 2,5 & 4,5 & -11 \end{bmatrix}$$

OBSERVAÇÃO: O número 2 é o pivô. O número $\frac{5}{2}$ é o multiplicador m_{21} e o número $-\frac{1}{2}$ é o multiplicador m_{31} .

Segundo passo: Continuando o processo, nosso objetivo agora é tornar nulo o elemento 2,5 que está situado abaixo do segundo elemento da diagonal principal (o número 3,5, que é o novo pivô).

Note que o novo multiplicador é $m_{32} = \frac{2,5}{3,5}$. Assim:

$$\begin{bmatrix} 2 & -3 & 1 & -2 \\ 0 & 3,5 & -4,5 & 17 \\ 0 & 2,5 & 4,5 & -11 \end{bmatrix} \xrightarrow{L'_3 = L_3 - \frac{2,5}{3,5}L_2} \begin{bmatrix} 2 & -3 & 1 & -2 \\ 0 & 3,5 & -4,5 & 17 \\ 0 & 0 & 7,7144 & -23,1431 \end{bmatrix}$$

Fazendo a substituição retroativa ou retro-substituição (do fim para o começo):

$$7,7144 \cdot x_3 = -23,1431 \Rightarrow x_3 = -3;$$

$$3,5 \cdot x_2 - 4,5 \cdot (-3) = 17 \Rightarrow x_2 = 1;$$

$$2 \cdot x_1 - 3 \cdot 1 + (-3) = -2 \Rightarrow x_1 = 2. \text{ Portanto a solução é: } S = \{(2, 1, -3)\}$$

RESOLVA OS SISTEMAS ABAIXO:

$$\text{a) } \begin{cases} 2x_1 + 3x_2 + 3x_3 = 9 \\ x_1 + x_2 + 2x_3 = 2 \\ 4x_1 + 3x_2 - 2x_3 = 3 \end{cases} \quad \text{RESPOSTA: } S = \{(-3, 5, 0)\}$$

$$\text{b) } \begin{cases} 10x_1 + 2x_2 + x_3 = 7 \\ x_1 + 5x_2 + x_3 = -8 \\ 2x_1 + 3x_2 + 10x_3 = 6 \end{cases} \quad \text{RESPOSTA: } S = \{(1, -2, 1)\}$$

$$\text{c) } \begin{cases} x_1 + 2x_2 + x_3 = 9 \\ 2x_1 + x_2 - x_3 = 3 \\ 3x_1 - x_2 - 2x_3 = -4 \end{cases} \quad \text{RESPOSTA: } S = \{(1, 3, 2)\}$$

$$\text{d) } \begin{cases} -x_1 + 2x_2 - x_3 = 3 \\ 2x_1 + x_2 + 3x_3 = 5 \\ x_1 + 3x_2 + 5x_3 = 13 \end{cases} \quad \text{RESPOSTA: } S = \left\{ \left(-\frac{14}{15}, \frac{28}{15}, \frac{5}{3} \right) \right\}$$

$$e) \begin{cases} 2x + y + z = 7 \\ x - y + z = 2 \\ 5x - 3y + z = 2 \end{cases}$$

$$\text{RESPOSTA: } S = \{(1, 2, 3)\}$$

$$f) \begin{cases} \frac{1}{x} + \frac{2}{y} + \frac{4}{z} = 1 \\ \frac{2}{x} + \frac{3}{y} + \frac{8}{z} = 0 \\ \frac{-1}{x} + \frac{9}{y} + \frac{10}{z} = 5 \end{cases}$$

$$\text{RESPOSTA: } S = \left\{ \left(\frac{7}{11}, \frac{1}{2}, \frac{-7}{8} \right) \right\}$$

TEOREMA: $A_{n \times n}$ é inversível se, e somente se, $A \sim I_n$.

(Ver demonstração em qualquer livro da bibliografia)

Uma importante aplicação do teorema acima é o cálculo da matriz inversa de uma matriz A (se, é claro, A for inversível). Veja o exemplo abaixo.

Suponha que queremos achar a inversa da matriz $A = \begin{pmatrix} 2 & 1 & 0 & 0 \\ 1 & 0 & -1 & 1 \\ 0 & 1 & 1 & 1 \\ -1 & 0 & 0 & 3 \end{pmatrix}$.

PRIMEIRO PASSO: Coloque a matriz identidade de ordem 4 ao lado da matriz A acima:

$$\left(\begin{array}{cccc|cccc} 2 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & -1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 & 1 & 0 \\ -1 & 0 & 0 & 3 & 0 & 0 & 0 & 1 \end{array} \right) \quad \left(\begin{array}{cccc|cccc} 2 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & -1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 & 1 & 0 \\ -1 & 0 & 0 & 3 & 0 & 0 & 0 & 1 \end{array} \right)$$

SEGUNDO PASSO: Aplique as operações elementares sobre linhas para reduzir a matriz dada (a matriz à esquerda) à forma escalonada; ao mesmo tempo efetue as mesmas operações à matriz identidade (a matriz à direita). Se ao final das operações elementares você conseguir transformar a matriz A na matriz identidade, neste momento observe a matriz à direita, ela é a matriz inversa procurada.

Vamos mostrar as operações elementares e achar a matriz inversa de A :

i) Trocando de lugar a primeira linha e a segunda linha, teremos:

$$\left(\begin{array}{cccc|cccc} 1 & 0 & -1 & 1 & 0 & 1 & 0 & 0 \\ 2 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 & 1 & 0 \\ -1 & 0 & 0 & 3 & 0 & 0 & 0 & 1 \end{array} \right)$$

ii) Somando à quarta a primeira e à segunda, a primeira linha multiplicada por -2 , obteremos:

$$\left(\begin{array}{cccc|cccc} 1 & 0 & -1 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 2 & -2 & 1 & -2 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 & 1 & 0 \\ 0 & 0 & -1 & 4 & 0 & 1 & 0 & 1 \end{array} \right)$$

iii) Subtraindo a segunda linha da terceira:

$$\left(\begin{array}{cccc|cccc} 1 & 0 & -1 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 2 & -2 & 1 & -2 & 0 & 0 \\ 0 & 0 & -1 & 3 & -1 & 2 & 1 & 0 \\ 0 & 0 & -1 & 4 & 0 & 1 & 0 & 1 \end{array} \right)$$

iv) Trocando o sinal da terceira linha e, depois, anulando os outros elementos da terceira coluna da matriz à esquerda:

$$\left(\begin{array}{cccc|cccc} 1 & 0 & 0 & -2 & 1 & -1 & -1 & 0 \\ 0 & 1 & 0 & 4 & -1 & 2 & -2 & 0 \\ 0 & 0 & 1 & -3 & 1 & -2 & -1 & 0 \\ 0 & 0 & 0 & 1 & 1 & -1 & -1 & 1 \end{array} \right)$$

v) Finalmente, anulando os elementos da quarta coluna (exceto o último) da matriz à esquerda, obteremos a matriz identidade à esquerda e a matriz inversa de A denotada por A^{-1} à direita:

$$\left(\begin{array}{cccc|cccc} 1 & 0 & 0 & 0 & 3 & -3 & -3 & 2 \\ 0 & 1 & 0 & 0 & -5 & 6 & 2 & -4 \\ 0 & 0 & 1 & 0 & 4 & -5 & -4 & 3 \\ 0 & 0 & 0 & 1 & 1 & -1 & -1 & 1 \end{array} \right)$$

assim:

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

1) Determine, se possível, a inversa das matrizes abaixo:

a) $A = \begin{pmatrix} 1 & 3 \\ 4 & 2 \end{pmatrix}$

RESPOSTA: $A^{-1} = \begin{pmatrix} -1/5 & 3/10 \\ 2/5 & -1/10 \end{pmatrix}$

b) $A = \begin{pmatrix} -1 & 2 \\ 4 & -8 \end{pmatrix}$

RESPOSTA: A matriz A não é inversível

$$\text{c) } A = \begin{pmatrix} 3 & -3 & 5 \\ 2 & 4 & -6 \\ 4 & -1 & 0 \end{pmatrix} \quad \text{RESPOSTA: } A^{-1} = \begin{pmatrix} 1/6 & 5/36 & 1/18 \\ 2/3 & 5/9 & -7/9 \\ 1/2 & 1/4 & -1/2 \end{pmatrix}$$

2) Ache o conjunto solução do sistema:

$$\begin{cases} 3x - 3y + 5z = 1 \\ 2x + 4y - 6z = 0 \\ 4x - y = 2 \end{cases} \quad \text{RESPOSTA: } S = \left\{ \left(\frac{5}{18}, \frac{-8}{9}, \frac{-1}{2} \right) \right\}$$

3) **Prove que** se A , B e C são matrizes inversíveis de ordem n , então $(ABC)^{-1} = C^{-1}B^{-1}A^{-1}$.

4) Sejam A , B e C matrizes inversíveis de mesma ordem. Resolva as equações em X , sabendo-se que:

$$\text{a) } A.X.B = C \quad \text{b) } A.(B+X) = A \quad \text{c) } A.C.X.B = C$$

$$\text{RESPOSTAS: a) } X = A^{-1}.C.B^{-1} \quad \text{b) } X = I - B \quad \text{c) } X = C^{-1}.A^{-1}.C.B^{-1}$$

MÉTODO ITERATIVO

Considere o sistema (*) abaixo, de ordem n , isto é $n \times n$, com solução única e onde os elementos da

[illegible]

MÉTODO DE GAUSS SEIDEL

Para implementarmos o método de Gauss-Seidel, devemos partir de uma aproximação inicial (qualquer) $x^{(0)} = (x_1^{(0)}, x_2^{(0)}, \dots, x_n^{(0)})$ (usualmente a aproximação inicial é a nula $x^{(0)} = (0, 0, \dots, 0)$) e, utilizamos a cada nova iteração o resultado da iteração imediatamente anterior. Matematicamente:

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

VEJA OS EXEMPLOS QUE SEGUEM ABAIXO.

EXEMPLO 1:

$$\begin{cases} 10x_1 + 2x_2 + x_3 = 7 \\ x_1 + 5x_2 + x_3 = -8 \\ 2x_1 + 3x_2 + 10x_3 = 6 \end{cases} \Rightarrow \text{Inicialmente "tiramos" o valor de } x_1, x_2 \text{ e } x_3, \text{ do seguinte modo:}$$

$$\begin{cases} x_1 = \frac{7 - 2x_2 - x_3}{10} \\ x_2 = \frac{-8 - x_1 - x_3}{5} \\ x_3 = \frac{6 - 2x_1 - 3x_2}{10} \end{cases}, \text{ que dá origem ao seguinte processo iterativo:}$$

$$\begin{cases} x_1^{(k+1)} = \frac{7 - 2x_2^{(k)} - x_3^{(k)}}{10} \\ x_2^{(k+1)} = \frac{-8 - x_1^{(k+1)} - x_3^{(k)}}{5} \\ x_3^{(k+1)} = \frac{6 - 2x_1^{(k+1)} - 3x_2^{(k+1)}}{10} \end{cases}, \text{ substituindo, a princípio, } k=0, \text{ temos}$$

$$\begin{cases} x_1^{(1)} = \frac{7 - 2x_2^{(0)} - x_3^{(0)}}{10} \\ x_2^{(1)} = \frac{-8 - x_1^{(1)} - x_3^{(0)}}{5} \\ x_3^{(1)} = \frac{6 - 2x_1^{(1)} - 3x_2^{(1)}}{10} \end{cases}, \text{ utilizando como aproximação inicial } x^{(0)} = (x_1^{(0)}, x_2^{(0)}, x_3^{(0)}) = (0, 0, 0), \text{ temos:}$$

$$x_1^{(1)} = 0,7; x_2^{(1)} = -1,74; x_3^{(1)} = 0,982$$

$$\text{Fazendo agora } k=1, \text{ obteremos: } x_1^{(2)} = 0,9498; x_2^{(2)} = -1,9864; x_3^{(2)} = 1,0059$$

Podemos continuar o processo indefinidamente, uma tabela feita no Excel (tente construí-la!), nos dá:

Resolução por Gauss-Seidel

0,0000	0,0000	0,0000
0,7000	-1,7400	0,9820
0,9498	-1,9864	1,0059
0,9967	-2,0005	1,0008
1,0000	-2,0002	1,0000
1,0000	-2,0000	1,0000

Desse modo, com cinco iterações (pois a solução nula não conta, foi um “chute inicial” e não uma iteração), temos a solução $(1, -2, 1)$ (substitua no sistema e verifique que, realmente esta é a solução procurada!)

EXEMPLO 2:

$$\begin{cases} 20x_1 + x_2 + x_3 + 2x_4 = 33 \\ x_1 + 10x_2 + 2x_3 + 4x_4 = 38,4 \\ x_1 + 2x_2 + 10x_3 + x_4 = 43,5 \\ 2x_1 + 4x_2 + x_3 + 20x_4 = 45,6 \end{cases}$$

Começando com a solução nula $(0, 0, 0, 0)$ e usando o Excel:

Resolução por Gauss-Seidel

0,0000	0,0000	0,0000	0,0000
1,6500	3,6750	3,4500	1,2075
1,1730	2,5497	3,6020	1,4727
1,1951	2,4110	3,6010	1,4982
1,1996	2,4005	3,6001	1,4999
1,2000	2,4000	3,6000	1,5000
1,2000	2,4000	3,6000	1,5000

Portanto a solução, (como tudo indica!) é a quádrupla $(1, 2; 2, 4; 3, 6; 1, 5)$.

Será que não fica, para vocês, meus curiosos alunos, uma pergunta no ar:

Professor, este método de Gauss Seidel sempre converge, isto é, cada vez mais os valores obtidos na tabela se aproximam da solução do sistema?

A resposta é... **NÃO!**

Veremos abaixo alguns critérios que se satisfeitos, garantem a convergência do processo de Gauss-Seidel. Esses critérios são chamados de condições suficientes, para o Matemático isso significa que se esses critérios **não** forem satisfeitos pode haver ou não convergência, ou seja, nada podemos afirmar!

CRITÉRIOS DE CONVERGÊNCIA PARA O MÉTODO ITERATIVO GAUSS SEIDEL (GS)

CRITÉRIO DAS LINHAS:

Considerando o sistema (*) nas condições estabelecidas acima temos o

TEOREMA: Se $|a_{jj}| > \sum_{\substack{k \neq j \\ k=1}}^n |a_{jk}|, \forall j = 1, \dots, n$ então o método de GS gera uma sequência convergente

$\{x^{(k)}\}, k = 0, 1, 2, \dots$, para a solução do sistema, independentemente da escolha do valor inicial $x^{(0)}$.

EXEMPLOS

a) O sistema
$$\begin{cases} 10x_1 + 2x_2 + x_3 = 7 \\ x_1 + 5x_2 + x_3 = -8 \\ 2x_1 + 3x_2 + 10x_3 = 6 \end{cases}$$
 tem convergência garantida pois:
$$\begin{cases} |10| > |2| + |1| \\ |5| > |1| + |1| \\ |10| > |2| + |3| \end{cases}$$

b) Já o sistema $\begin{cases} x_1 + x_2 = 3 \\ x_1 - 3x_2 = -3 \end{cases}$ não tem convergência garantida pelo critério das linhas pois não

o satisfaz, confira: $\begin{cases} |1| > |1| \text{ (Falso)} \\ |-3| > |1| \end{cases}$. Pelo fato do critério das linhas não ser satisfeito **nada podemos**

afirmar em relação à convergência por GS. Se você quiser arriscar, verá que após algumas iterações, usando o valor inicial $(0,0)$, observará que haverá convergência para a solução $(1,5;1,5)$.

c) O sistema $\begin{cases} x_1 + 3x_2 + x_3 = -2 \\ 5x_1 + 2x_2 + 2x_3 = 3 \\ 6x_2 + 8x_3 = -6 \end{cases}$ não satisfaz o critério das linhas, mas se permutarmos a linha 2 com a

linha 1, teremos $\begin{cases} 5x_1 + 2x_2 + 2x_3 = 3 \\ x_1 + 3x_2 + x_3 = -2 \\ 6x_2 + 8x_3 = -6 \end{cases}$ que, agora satisfaz o critério das linhas pois: $\begin{cases} |5| > |2| + |2| \\ |3| > |1| + |1| \\ |8| > |0| + |6| \end{cases}$ e,

portanto, tem convergência garantida.

Há outros critérios de convergência. Um deles que é bem simples é o critério das colunas (análogo ao critério das linhas!). Veja abaixo:

CRITÉRIO DAS COLUNAS:

Considerando o sistema (*) nas condições estabelecidas anteriormente vamos enunciar o

TEOREMA: Se $|a_{jj}| > \sum_{\substack{k \neq j \\ k=1}}^n |a_{kj}|, \forall j = 1, \dots, n$ então o método de GS gera uma sequência convergente

$\{x^{(k)}\}, k = 0, 1, 2, \dots$, para a solução do sistema, independentemente da escolha do valor inicial $x^{(0)}$.

EXEMPLOS

a) Como vimos, o sistema $\begin{cases} 10x_1 + 2x_2 + x_3 = 7 \\ x_1 + 5x_2 + x_3 = -8 \\ 2x_1 + 3x_2 + 10x_3 = 6 \end{cases}$ satisfaz o critério das linhas e tem convergência garantida

apesar de que não satisfaz o critério das colunas pois: $\begin{cases} |10| > |2| + |1| \\ |5| > |2| + |3| \text{ (Falso)} \\ |10| > |1| + |1| \end{cases}$.

Note que se multiplicarmos a segunda equação por 2, teremos um sistema que convergirá por qualquer dos dois critérios.

b) O sistema $\begin{cases} x_1 + x_2 = 3 \\ x_1 - 3x_2 = -3 \end{cases}$ não tem convergência garantida pois: $\begin{cases} |1| > |1| \text{ (Falso)} \\ |-3| > |1| \end{cases}$. Nada podemos

afirmar usando o critério das colunas.

APLICAÇÕES

Resolva os sistemas lineares que seguem, pelo método de Gauss-Seidel (GS). Antes de resolver verifique se algum critério de convergência é satisfeito, caso não seja, tente alguma mudança de linhas e/ou colunas de modo que a convergência esteja assegurada. Após resolver usando uma calculadora científica “normal”, procure resolver todos os sistemas lineares na planilha Excel. Em todos os exercícios inicie com a “solução nula”.

a) $\begin{cases} 3x_1 - x_2 = 1 \\ -2x_1 - 4x_2 = 3 \end{cases}$. Com 7 iterações para GS (utilize quatro decimais).

b) $\begin{cases} 2x_1 - x_2 = 1 \\ x_1 + 2x_2 = 3 \end{cases}$. Com 10 iterações para GS (utilize quatro decimais).

c) $\begin{cases} x_1 + x_2 = 3 \\ x_1 - 3x_2 = -3 \end{cases}$. Com 13 iterações para GS (utilize quatro decimais).

d) $\begin{cases} 3x_1 + 0,15x_2 - 0,09x_3 = 6 \\ 0,08x_1 + 4x_2 - 0,16x_3 = 12 \\ 0,05x_1 + 0,3x_2 + 5x_3 = 20 \end{cases}$. Com 6 iterações para GS (cinco decimais).

e) $\begin{cases} 5x_1 + x_2 + x_3 = 5 \\ 3x_1 + 4x_2 + x_3 = 6 \\ 3x_1 + 3x_2 + 6x_3 = 0 \end{cases}$. Com 8 iterações para GS (cinco decimais).

f) $\begin{cases} 3x_1 + 2x_2 + x_3 = 0 \\ x_1 + 4x_2 + 4x_3 = -1 \\ x_1 - x_2 + 6x_3 = 3 \end{cases}$. Com 7 iterações para GS (cinco decimais).

g) $\begin{cases} 20x_1 + x_2 + x_3 + 2x_4 = 33 \\ x_1 + 10x_2 + 2x_3 + 4x_4 = 38,4 \\ x_1 + 2x_2 + 10x_3 + x_4 = 43,5 \\ 2x_1 + 4x_2 + x_3 + 20x_4 = 45,6 \end{cases}$. Com 7 iterações para GS (quatro decimais).

h) $\begin{cases} x_1 + 3x_3 = 3 \\ -x_2 + x_3 = 1 \\ 2x_1 + x_2 + 3x_3 = 9 \end{cases}$. Com 7 iterações para GS.

FAÇA UMA PESQUISA SOBRE OS SEGUINTE TEMAS:

1) MÉTODO DA ELIMINAÇÃO DE GAUSS COM CONDENSAÇÃO PIVOTAL

Neste caso, a cada passo, escolhe-se, a cada passo, o elemento de maior valor absoluto como elemento da diagonal da matriz dos coeficientes. Teoricamente isso diminui os erros de arredondamento. Dê exemplos numéricos e explique o mais simples possível esse método.

2) REFINAMENTO DA SOLUÇÃO.

O que significa e em que casos deve ser aplicado. Dê exemplos numéricos.