



UNIVERSIDADE DA BEIRA INTERIOR
Covilhã | Portugal

Introdução ao Método dos Elementos Finitos

Mecânica Estrutural (10391/10411)

2022

Pedro V. Gamboa

Departamento de Ciências Aeroespaciais

1. Integração Numérica de Equações Diferenciais

- Neste capítulo abordam-se técnicas para integração numérica de equações diferenciais ordinárias (EDO) que permitem a simulação temporal de sistemas físicos em computadores.
- Os métodos aqui apresentados são gerais, na medida em que se aplicam a modelos compostos por qualquer número de equações de movimento, lineares ou não.

1.1. Equações Diferenciais

As equações diferenciais aparecem com grande frequência em diversos problemas de modelação de fenómenos físicos. Exemplos são equações que descrevem escoamento de fluidos, transferência de calor e massa, química, dinâmica e vibrações em sistemas mecânicos, etc..

Uma equação diferencial é definida como uma equação que envolve derivadas de funções. A ordem de uma equação diferencial é descrita em função da maior ordem p da derivada envolvida.

Dois tipos básicos podem aparecer, o primeiro envolve equações diferenciais ditas ordinárias. Neste caso existe apenas uma variável independente, $y(x)$:

$$\frac{dy}{dx} = x + y$$

1.1. Equações Diferenciais

As equações diferenciais ordinárias contêm parâmetros físicos concentrados.

O segundo tipo acontece quando existe mais de uma variável independente, por exemplo $u(x,y)$ representando o deslocamento numa placa em função de x e y :

$$\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} = \nabla^2 u = 0$$

onde ∇^2 é o Laplaciano. Esta equação é um exemplo de equação diferencial parcial. Este tipo de equação envolve parâmetros distribuídos.

Neste caso iremos focar apenas a solução numérica de equações diferenciais ordinárias (EDO).

1.1. Equações Diferenciais

Um facto interessante é constatar que as EDOs não possuem apenas uma solução e sim uma família ou conjunto de soluções possíveis. Para particularizar a solução de uma EDO é essencial definir valores de condições suplementares.

Caso essas condições sejam especificadas no mesmo ponto, tem-se uma condição inicial e, neste contexto, o problema é classificado como de valor inicial (**PVI**). Por outro lado, se forem especificadas em mais de um ponto, tem-se um problema de valor de contorno (**PVC**).

As equações diferenciais podem ser lineares ou não-lineares, dependendo se é válido ou não o princípio da sobreposição. Um exemplo de equação diferencial ordinária não-linear é:

$$u''(x) + u'^2(x) = 1$$

1.1. Equações Diferenciais

A grande preocupação dos matemáticos é garantir a existência e unicidade da solução de PVI e PVC. Um problema de VC normalmente é mais complexo, pois em inúmeros exemplos não se garante unicidade da solução.

Em problemas de dinâmica de sistemas mecânicos a aplicação da 2ª lei de Newton gera sistemas de EDOs que são essencialmente não-lineares.

Exceto para casos bem particulares, em geral linearizados e com aplicação de hipóteses simplificadoras, a solução analítica destas equações é inviável.

Assim, justifica-se a aplicação e implementação de métodos numéricos.

1.1. Equações Diferenciais

A ideia básica de grande parte destes métodos numéricos é ser capaz de construir uma solução para uma equação do tipo $x_0(t)=f(x,t)$ dada uma condição $x(t_0)=x_0$.

O que se procura é definir uma sequência de valores $t_1, t_2, \dots, t_n$, não necessariamente espaçados uniformemente e calcular aproximações numéricas para $x_i(t_i)$ baseado em informações passadas.

Se apenas uma informação anterior é empregue o método é conhecido como sendo da classe **passo simples**.

Por outro lado, se usarmos vários valores anteriores, o método é de **passo múltiplo**. Alguns métodos clássicos usados envolvem a aproximação numérica da série de Taylor, como será apresentado.

1.2. Solução de EDs por Métodos Numéricos

As classes fundamentais de problemas envolvendo equações diferenciais ordinárias são mostradas abaixo (Tao, 1988):

- **Problemas de valor inicial:** são sistemas espacialmente homogêneos cujas propriedades são assumidas como uniformes, tratados como sistemas de parâmetros concentrados e que variam no tempo; e
- **Problemas de valores de contorno:** envolvem sistemas em estado estacionário que têm gradientes internos (variações espaciais).

Para resolver problemas de valor inicial há necessidade de conhecer o valor das variáveis no instante inicial, ao passo que para resolver problemas de valor de contorno são necessários os valores das variáveis nas fronteiras do sistema.

1.2.1. Tipos de métodos numéricos para resolver EDOs

Os métodos seguintes são utilizados na solução de equações diferenciais ordinárias (Tao, 1988):

- métodos de passo simples
- métodos de passo múltiplo

Ambos os métodos estimam a solução como sendo uma série de valores a intervalos específicos que geram uma função que satisfaz a equação diferencial.

1.2.1. Tipos de métodos numéricos para resolver EDOs

CARACTERÍSTICAS DOS MÉTODOS DE PASSO SIMPLES

- São métodos que necessitam apenas de um valor para estimar a solução em cada intervalo de tempo.
- Essa estimativa é então usada para encontrar o valor da função no próximo intervalo de tempo.
- Exemplo: a fórmula para encontrar y_n depende explicitamente apenas de y_{n-1} .

1.2.1. Tipos de métodos numéricos para resolver EDOs

CARACTERÍSTICAS DOS MÉTODOS DE PASSO MÚLTIPLO

- São métodos em que o cálculo de y_n depende explicitamente de dois ou mais valores anteriores.
- Por exemplo, no método de passo duplo, o cálculo de y_n depende dos seguintes valores: $y_n = f(y_{n-1}, y_{n-2})$
- Esses métodos são também conhecidos como preditor-corretor.

1.2.1. Tipos de métodos numéricos para resolver EDOs

COMPARAÇÃO ENTRE MÉTODOS DE PASSO SIMPLES/MÚLTIPLO

- Ambos são precisos, embora os métodos de passo múltiplo normalmente apresentem melhores resultados quando a função é descontínua em certos pontos ou as suas derivadas variam muito rapidamente (sistemas *stiff*).
- A maioria dos problemas de valor inicial não têm múltiplas condições iniciais, sendo necessário gerá-las por um método de passo simples para inicializar o método de passo múltiplo.
- Os requisitos computacionais (capacidade de memória) dos métodos de passo múltiplo normalmente excedem os de métodos de passo simples.

1.2.2. Métodos Numéricos de Integração de Passo Simples para Resolver Problemas de Valor Inicial

Discute-se aqui a solução para sistemas de 1ª ordem.
Forma geral de sistemas de 1ª ordem:

$$\dot{y} = f(t, y_a) \quad \text{com} \quad y(0) = y_0$$

Objetivo:

- Estimar $y(t)$; ou
- Estimar t correspondente a um dado valor de y .

O método básico aplicado neste tipo de algoritmo é descrito a seguir.

Iniciando nas condições iniciais (t_0), é feita uma estimativa de $y(t)$ no próximo valor de t (t_1) correspondente a $t_1 = t_0 + h$, onde h = passo de integração.

1.2.2. Métodos Numéricos de Integração de Passo Simples para Resolver Problemas de Valor Inicial

O novo valor é usado para estimar o próximo ponto e assim por diante.

A base matemática por trás deste método consiste em expandir a função $y(t)$ numa *série de Taylor*. Estima-se o valor da função a uma pequena distância de um valor conhecido, usando as derivadas da função calculadas no valor conhecido da mesma.

Assim, dada a função $g(t)$ com derivadas contínuas na região em torno de $t=a$, o valor de $g(a+h)$, onde h é pequeno, é dado por:

$$g(a+h) = g(a) + hg'(a) + \frac{1}{2!}h^2g''(a) + \frac{1}{3!}h^3g'''(a) + \dots$$

Repetindo este cálculo sequencialmente para valores de $t_n=a+nh$, onde $n=1,2,3\dots$, resulta numa série de estimativas de $g(t)$.

1.2.2. Métodos Numéricos de Integração de Passo Simples para Resolver Problemas de Valor Inicial

Se em cada cálculo é incluído um número infinito de termos, obtém-se a solução exata. No entanto, na prática o número de termos na série deve ser finito. Assim, aceita-se algum erro.

A ordem de magnitude desse erro de truncagem é definida pela potência do último termo incluído. Assim, se o último termo é de ordem h^5 , o erro resultante é no máximo de ordem h^5 e a ordem de magnitude desse erro é abreviada por $O(h^5)$. A inclusão de mais termos reduz o erro de truncagem.

São descritos a seguir dois métodos numéricos de integração que se enquadram nessa filosofia: Euler e Runge-Kutta (Carnahan *et al.*, 1969).

1.2.3. Método de Euler

Consiste em utilizar a 1ª derivada já conhecida e usar um passo pequeno de integração h , truncando termos de ordem ≥ 2 :

$$y(t+h) = y(t) + h \cdot \dot{y}(t) = y(t) + h \cdot f[t, y(t)]$$

onde

$$\dot{y}(t) = f(t, y)$$

ou

$$u_1 = h \cdot f(t_n, y_n)$$

$$y_{n+1} = y_n + u_1$$

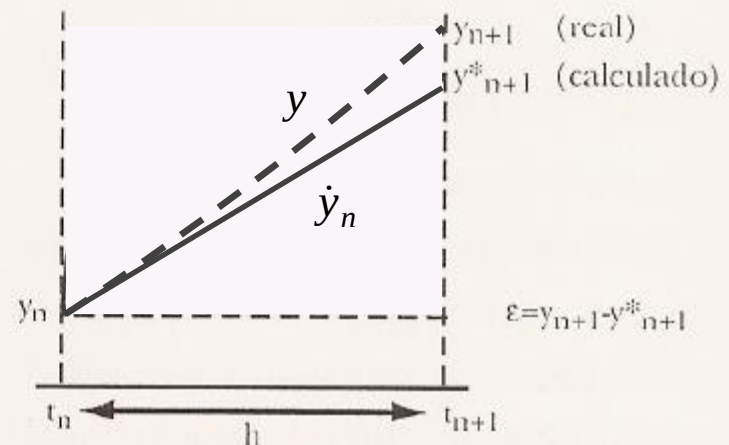
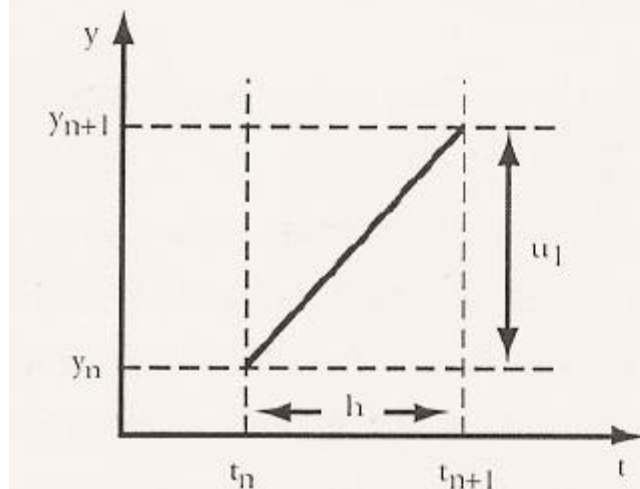
$$t_{n+1} = t_n + h$$

$$y = y_n \quad \text{quando} \quad t = t_n; \quad \text{procura-se} \quad y_{n+1} \quad \text{quando} \quad t = t_n + h$$

1.2.3. Método de Euler

Avalia-se a equação com as condições iniciais conhecidas e realiza-se a extrapolação segundo a tangente à curva nesse ponto.

Na figura abaixo está a interpretação gráfica do método de Euler.



1.2.3. Método de Euler

Algoritmo para implementar o Método de Euler:

Dados:

- tamanho do passo de integração h
- $y(0) = y_0$
- número de iterações N
- $\dot{y}(t) = f(t, y)$

Para $n=0$ a $n=N-1$ faz-se:

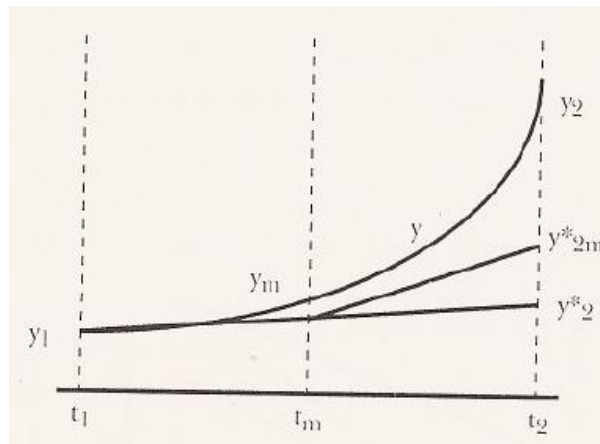
- $y_{n+1} = y_n + h \cdot f(t_n, y_n)$
- $t_{n+1} = t_n + h$
- saída: y_{n+1}

O erro de truncagem é $O(h)$, sendo proporcional ao tamanho h do passo.

1.2.3. Método de Euler

Como o número de cálculos é inversamente proporcional ao tamanho do passo, uma solução precisa requer um grande número de computações.

A relação entre a magnitude do erro e o tamanho do passo no método de Euler é ilustrada pela figura abaixo (Franks, 1972), que mostra a melhoria que ocorre quando um passo único de t_1 a t_2 é comparado com dois meio-passos: t_1 a t_m e t_m a t_2 .



1.2.3. Método de Euler

Para o caso com dois meio-passos, começando com a derivada $\dot{y}_1$ em t_1 , avança-se meio passo até t_m , onde a derivada $\dot{y}_m$ é reavaliada.

Verifica-se que o resultado final y_{2m}^* está mais próximo do valor correto y_2 que o valor y_2^* obtido com um passo único.

Para um método de integração de 1ª ordem, os erros numéricos são diretamente proporcionais ao tamanho do passo ($\varepsilon = K.h$), de forma que a tolerância desejada, isto é, o máximo erro aceitável, determina o tamanho do passo a ser usado.

1.2.4. Método de Runge-Kutta de 2ª Ordem

Os métodos de Runge-Kutta (RK) imitam os termos da série de Taylor sem, no entanto, derivar a equação original.

O método de Runge-Kutta de 1ª ordem equivale ao método de Euler.

No método de Runge-Kutta de 2ª ordem a função é avaliada nos pontos extremos do intervalo h :

$$y(0) = y_0$$

$$u_1 = h \cdot f(t_n, y_n)$$

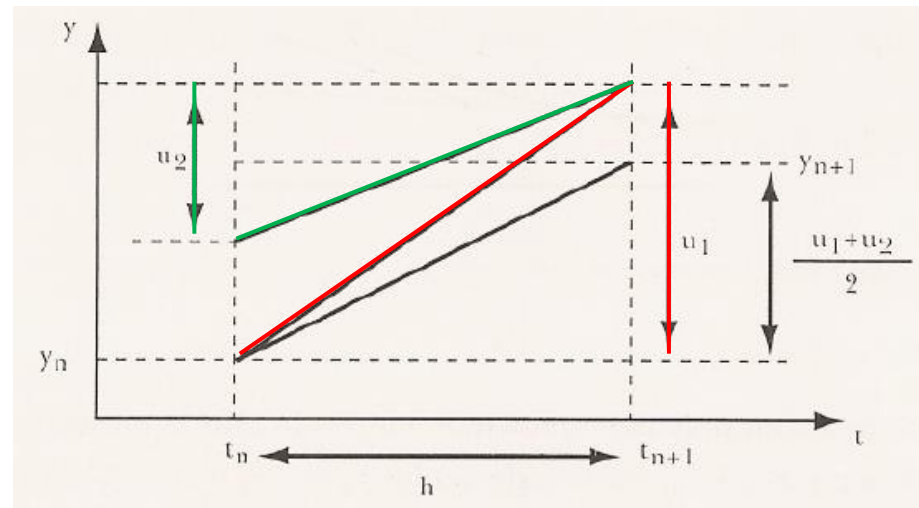
$$u_2 = h \cdot f(t_n + h, y_n + u_1)$$

u_2 é avaliado no ponto definido pelo método de 1ª ordem (Euler).

1.2.4. Método de Runge-Kutta de 2ª Ordem

$$y_{n+1} = y_n + \frac{u_1 + u_2}{2}$$

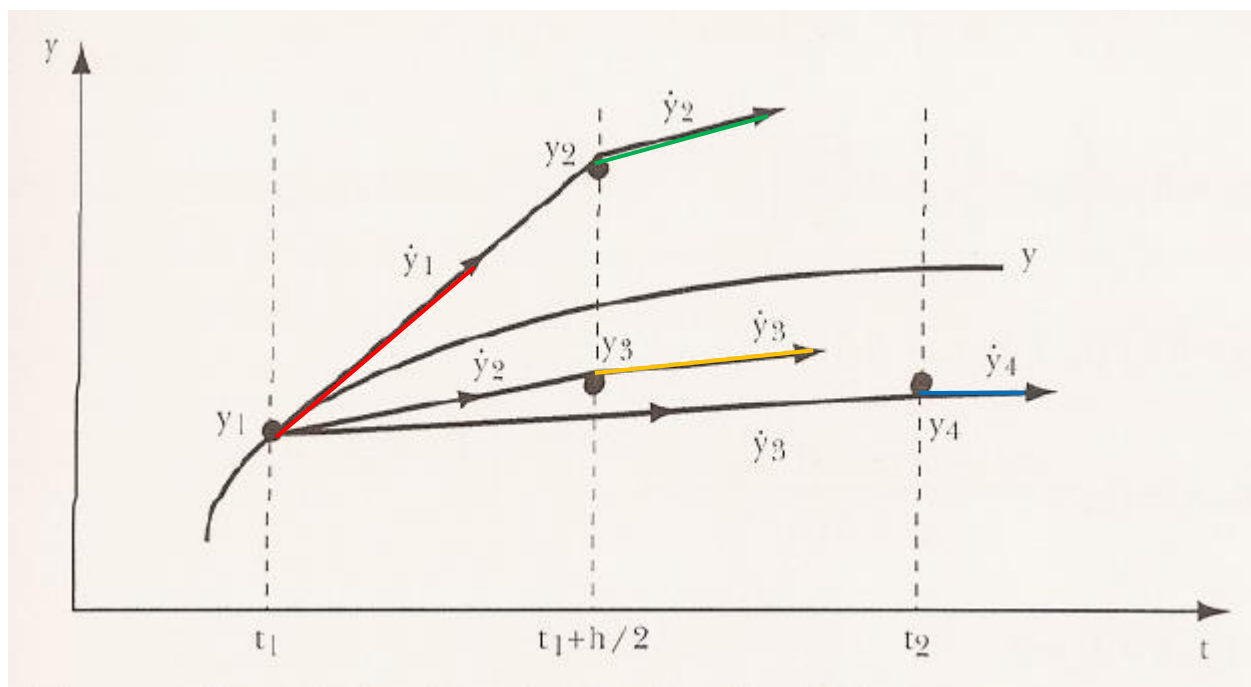
$$t_{n+1} = t_n + h$$



Visto que o erro de truncagem é $O(h^2)$, um método de Runge-Kutta de ordem superior é necessário se for desejado um resultado mais preciso.

1.2.5. Método de Runge-Kutta de 4ª Ordem

Este é o método RK mais utilizado. Consiste em avaliar as derivadas no início, meio e fim do intervalo de integração. O último passo é fazer uma soma ponderada dessas derivadas.



1.2.5. Método de Runge-Kutta de 4ª Ordem

Passos:

- a derivada $\dot{y}_1$ é avaliada em t_1 e, usando o método de Euler, o valor da função é calculado em $(t_1+h/2)$, resultando y_2 ;
- a derivada é avaliada em $(t_1+h/2)$ resultando $\dot{y}_2$;
- iniciando em (y_1, t_1) , a função em $t_1+h/2$ é recalculada usando a derivada $\dot{y}_2$ para gerar y_3 ;
- a derivada $\dot{y}_3$ é avaliada em $t_1+h/2$;
- iniciando em (y_1, t_1) , a função é calculada em $t_2=t_1+h$ usando a derivada $\dot{y}_3$ para gerar y_4 ;
- a derivada $\dot{y}_4$ é avaliada em t_2 ; e

1.2.5. Método de Runge-Kutta de 4ª Ordem

- usando as derivadas $\dot{y}_1$ a $\dot{y}_4$, o valor da função $y(t_2)$ é calculado por:

$$y(t_2) = y(t_1) + \frac{h}{6}(\dot{y}_1 + 2\dot{y}_2 + 2\dot{y}_3 + \dot{y}_4)$$

Algoritmo para implementar o Método RK de 4ª ordem:

Dados:

- tamanho do passo de integração h
- $y(0) = y_0$
- número de iterações N
- $\dot{y}(t) = f(t, y)$

1.2.5. Método de Runge-Kutta de 4ª Ordem

Para $n=0$ a $n=N-1$ faz-se:

$$u_1 = h \cdot f(t_n, y_n)$$

$$u_2 = h \cdot f\left(t_n + \frac{h}{2}, y_n + \frac{u_1}{2}\right)$$

$$u_3 = h \cdot f\left(t_n + \frac{h}{2}, y_n + \frac{u_2}{2}\right)$$

$$u_4 = h \cdot f(t_n + h, y_n + u_3)$$

$$y_{n+1} = y_n + \frac{u_1 + 2u_2 + 2u_3 + u_4}{6}$$

$$t_{n+1} = t_n + h$$

Saída

$$y_{n+1}$$

1.2.5. Método de Runge-Kutta de 4ª Ordem

Um esquema alternativo para terminar o algoritmo é especificar t e não N . Visto que o erro de truncagem é $O(h^4)$, o erro é muito reduzido quando comparado ao método de Euler. No entanto, são necessários quatro cálculos da função por passo, resultando que uma redução no tamanho do passo aumenta muito o número de cálculos.

Outra forma para reduzir o número de cálculos é ajustar o tamanho do passo para manter o erro de truncagem especificado. Em regiões de baixas taxas de variação de $\dot{y}(t)$, o tamanho do passo pode crescer, reduzindo o número de computações, e o contrário ocorre para altas taxas de variação de $\dot{y}(t)$.

Existe um esquema muito utilizado para verificação do erro de truncagem, atribuído a Fehlberg.

1.2.6. Método de Newmark

O sistema de equações diferenciais de segunda ordem em dinâmica estrutural pode ser resolvido por qualquer método considerando a existência de alguma excitação F externa sendo aplicada no sistema ou mesmo uma condição inicial de deslocamento e velocidade nalgum nó.

Entre estes métodos, o de Newmark é um dos mais versáteis e popular para solução de grandes sistemas de equações diferenciais de segunda ordem. Aqui não será dada nenhuma prova. Apenas é apresentado sucintamente o método e mostrado um algoritmo efetivo para a solução do sistema de EDOs.

1.2.6. Método de Newmark

Considere-se a equação do movimento do sistema descrita pelas matrizes de massa e rigidez e com o amortecimento sendo do tipo proporcional à massa e/ou rigidez:

$$M\ddot{x} + C\dot{x} + Kx = F$$

sendo $\ddot{x}$, $\dot{x}$ e x os vetores aceleração, velocidade e deslocamento, respetivamente.

A equação acima pode ser integrada usando um método numérico. Em essência, a integração numérica direta é baseada em duas ideias.

Na primeira, ao invés de tentar satisfazer a equação acima em todos os instantes t , procura-se satisfazê-la apenas em intervalos discretos de tempo Δt .

1.2.6. Método de Newmark

A segunda ideia consiste em variar os deslocamentos, velocidades e acelerações dentro do intervalo de tempo Δt assumido.

Em seguida, considera-se que os vetores deslocamento, velocidade e aceleração no instante inicial t_0 , denotados por $x(0)$, $\dot{x}(0)$ e $\ddot{x}(0)$, respetivamente, são conhecidos e implementa-se a solução das equações de equilíbrio para um tempo de t_0 até t_N .

Na solução, o tempo total considerado é dividido em N intervalos iguais Δt ($\Delta t = t_N/N$) e o esquema de integração empregue estabelece uma solução aproximada para os instantes Δt , $2\Delta t$, $3\Delta t$, ..., t , $t+\Delta t$, ..., T_N .

1.2.6. Método de Newmark

O esquema geral no método de Newmark assume que:

$$\begin{aligned}\dot{x}(t + \Delta t) &= \dot{x}(t) + [(1 - \gamma)\ddot{x}(t) + \gamma\ddot{x}(t + \Delta t)]\Delta t \\ x(t + \Delta t) &= x(t) + \dot{x}(t)\Delta t + \left[\left(\frac{1}{2} - \beta \right) \ddot{x}(t) + \beta\ddot{x}(t + \Delta t) \right] \Delta t^2\end{aligned}$$

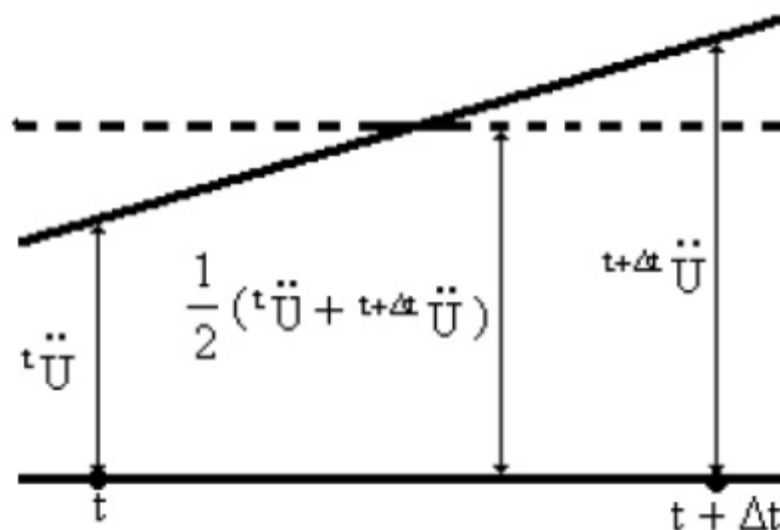
As constantes γ e β são conhecidas como parâmetros de Newmark e são determinados visando obter exatidão e estabilidade numérica.

Na literatura existem muitas variações deste algoritmo. Newmark originalmente propôs o esquema conhecido como aceleração média constante, conhecida como regra trapezoidal, em que neste caso $\gamma=1/2$ e $\beta=1/6$.

1.2.6. Método de Newmark

A figura mostra o esquema de integração.

Porém outros esquemas podem ser usados, como por exemplo $\gamma = 1/2$ e $\beta = 1/4$.



1.2.6. Método de Newmark

A ideia é fazer com que a equação do movimento seja válida nos intervalos de tempo de 0 até t_N :

$$M\ddot{x}(0) + C\dot{x}(0) + Kx(0) = F(0)$$

$$\vdots$$

$$M\ddot{x}(t) + C\dot{x}(t) + Kx(t) = F(t)$$

$$M\ddot{x}(t + \Delta t) + C\dot{x}(t + \Delta t) + Kx(t + \Delta t) = F(t + \Delta t)$$

$$\vdots$$

$$M\ddot{x}(t_N) + C\dot{x}(t_N) + Kx(t_N) = F(t_N)$$

1.2.6. Método de Newmark

Com base nesta ideia e no esquema de integração de Newmark pode escrever-se um algoritmo computacional para integração de equações diferenciais de segunda ordem de sistemas lineares descrito por quatro passos básicos:

- Inicialização
- Predição
- Equação de equilíbrio
- Correção

1.2.6. Método de Newmark

Escrevendo explicitamente cada passo temos:

1. Dados do problema: M, C, K
2. Inicialização

$$\ddot{x}(0) = M^{-1} \{F(0) - C\dot{x}(0) - Kx(0)\}$$

3. Incremento temporal

$$t_{k+1} = t_k + \Delta t$$

4. Predição (usando informação em t)

$$\dot{x}_{t_{k+1}} = \dot{x}_{t_k} + (1 - \gamma) \Delta t \ddot{x}_{t_k}$$

$$x_{t_{k+1}} = x_{t_k} + \Delta t \dot{x}_{t_k} + \left(\frac{1}{2} - \beta \right) \Delta t^2 \ddot{x}_{t_k}$$

1.2.6. Método de Newmark

5. Equação de equilíbrio

$$S = M + \gamma \Delta t C + \beta \Delta t^2 K$$

$$\ddot{x}_{t_{k+1}} = S^{-1} (F_{t_k} - C \dot{x}_{t_k} - K x_{t_k})$$

6. Correção (usando informação em $t+h$)

$$\dot{x}_{t_{k+1}} = \dot{x}_{t_k} + \Delta t \gamma \ddot{x}_{t_{k+1}}$$

$$x_{t_{k+1}} = x_{t_k} + \Delta t^2 \beta \ddot{x}_{t_{k+1}}$$

7. Critério de conclusão: atingir t_N .

1.2. Solução de EDs por Métodos Numéricos

Exemplo 1.01: Considere um pêndulo constituído por uma massa rígida $m=0,2\text{kg}$ e por um fio rígido de pequena massa e comprimento $l=2\text{m}$. O pêndulo está suspenso numa articulação sem atrito e oscila livremente. Considere também que qualquer resistência do ar é desprezável. Determine o movimento do pêndulo em função do tempo usando como grau de liberdade o ângulo que o fio faz com a vertical θ . As condições iniciais são $\theta_0=10^\circ$ e $\dot{\theta}_0=0$. Obtenha este movimento usando três métodos:

- a) Solução analítica linearizando a equação diferencial;
- b) Solução numérica com o método de Euler;
- c) Solução numérica com o método de Runge-Kutta de 4ª ordem;
- d) Compare as soluções.

2. Método dos Elementos Finitos

- Vários tipos de problemas físicos encontrados nas ciências e nas engenharias são descritos matematicamente na forma de equações diferenciais ordinárias e parciais.
- A solução exata é, usualmente, fruto de um método de solução analítica encontrado através de métodos algébricos e diferenciais aplicados a geometrias e condições de contorno particulares.
- A aplicação generalizada dos métodos analíticos para diferentes geometrias e condições de contorno (fronteira) torna impraticável ou até mesmo impossível a obtenção de soluções analíticas exatas.

2. Método dos Elementos Finitos

- O chamado *Método dos Elementos Finitos (MEF)* consiste em diferentes métodos numéricos que aproximam a solução de problemas de valor de fronteira descritos tanto por equações diferenciais ordinárias quanto por equações diferenciais parciais através da subdivisão da geometria do problema em elementos menores, chamados *elementos finitos*, nos quais a aproximação da solução exata pode ser obtida por interpolação de uma solução aproximada.
- Atualmente o MEF encontra aplicação em praticamente todas as áreas de engenharia, como na análise de tensões e deformações, transferência de calor, mecânica dos fluídos e reologia, eletromagnetismo, etc.

2.1. Resumo histórico

- O MEF foi originalmente concebido pelo matemático Courant durante a 2ª guerra mundial através da publicação de um artigo em 1943.
- Como nessa época ainda não haviam sido desenvolvidos computadores capazes de realizar uma grande quantidade de cálculos matemáticos, o método matemático foi ignorado pela academia durante vários anos.
- Na década de 1950 engenheiros e investigadores envolvidos no desenvolvimento de aviões a jato na Boeing iniciaram os primeiros trabalhos práticos no estabelecimento do MEF aplicados à indústria aeronáutica.

2.1. Resumo histórico

- M. J. Turner, R. W. Clough, H. C. Martin e L. J. Topp publicaram em 1956 um dos primeiros artigos que delinearam as principais ideias do MEF, entre elas a formulação matemática dos elementos e a montagem da matriz de elementos.
- Mas no artigo ainda não se fazia referência ao nome *elementos finitos* para designar os elementos de discretização da geometria do problema físico.
- O segundo co-autor do artigo, Ray Clough era na época professor em Berkeley e, durante o período de férias escolares, trabalhou na Boeing e descreveu o método com o nome de método dos elementos finitos num artigo publicado posteriormente.

2.1. Resumo histórico

- Os seus trabalhos deram início à investigação intensa em Berkeley por outros professores, entre eles E. Wilson e R. L. Taylor, juntamente com os estudantes de pós-graduação T. J. R. Hughes, C. Felippa e K. J. Bathe.
- Durante muitos anos, Berkeley foi o principal centro de investigação do MEF.
- Essa investigação coincidiu com a rápida disseminação de computadores eletrónicos nas universidades e institutos de investigação, que levaram o método a tornar-se amplamente utilizado em áreas estratégicas à segurança americana durante o período da Guerra Fria, tais como pesquisa nuclear, defesa, indústria automóvel e aeroespacial.
- E. Wilson desenvolveu um dos primeiros programas de computador de cálculo pelo MEF.

2.1. Resumo histórico

- A sua popularidade foi possível pela disponibilização gratuita do *software*, facto bastante comum nos anos 1960, pois o valor comercial de programas de computadores ainda não era reconhecido nessa época.
- Em 1965, a agência espacial norte-americana NASA financiou um projeto liderado por Dick McNeal para desenvolver um programa de cálculo pelo MEF de uso geral.
- Este programa, batizado de NASTRAN, incluía uma grande capacidade de manipulação de dados e permitia análise de tensão e deformação, cálculo de vigas, de problemas de cascas e placas, análise de estruturas complexas como asas de aviões e análise de vibrações em duas e três dimensões.
- O programa inicial foi colocado em domínio público, porém continha muitos *bugs* de programação.

2.1. Resumo histórico

- Logo após o término do projeto, Dick MacNeal e Bruce McCormick criaram uma empresa de *software* que corrigiu a maioria dos bugs e comercializaram essa versão corrigida com o nome MS-NASTRAN.
- Na mesma época, John Swanson estava a desenvolver um programa de MEF na Westinghouse para a análise de reatores nucleares.
- Em 1969, Swanson deixou a Westinghouse para comercializar o programa ANSYS.
- O programa tinha capacidade de análise de problemas lineares e não-lineares e essas características tornariam o *software* ANSYS um dos programas de elementos finitos comerciais mais utilizados atualmente.

2.1. Resumo histórico

- Outros programas comerciais desenvolvidos desde então foram o LS-DYNA usado para análises não-lineares tais como teste de colisão, conformação de metais e simulação de protótipos; ALGOR, ABAQUS e COSMOS como programas de MEF de uso geral; sendo que todos os programas possuem versões para microcomputadores e alguns versões mais potentes para sistemas computacionais paralelos e “clusters”.

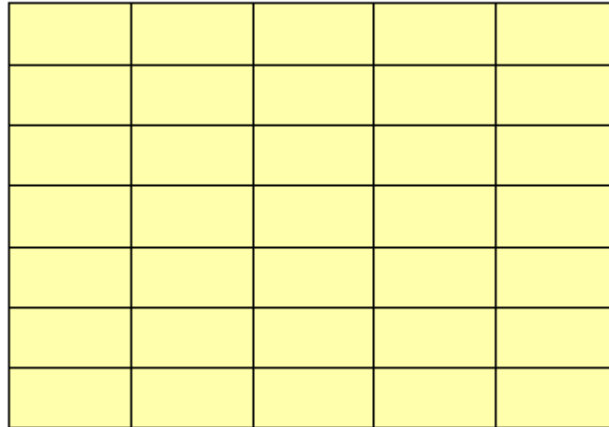
2.2. Diferenças entre o MDF e o MEF

- As diferenças entre o Método das Diferenças Finitas (MDF) e o MEF residem no facto de no MDF serem aplicadas aproximações nas derivadas das equações diferenciais, reduzindo-as a um problema de sistemas de equações lineares que fornecem a solução em pontos (nós) discretos no interior do domínio do problema.
- No MEF, a solução das equações diferenciais que governam o problema físico pode ser obtida por funções de aproximação que satisfazem condições descritas por equações integrais no domínio do problema.
- Essas funções de aproximação podem ser funções polinomiais com grau razoável de ajuste em elementos discretizados a partir da geometria do problema satisfazendo as equações integrais em cada elemento discreto ou elemento finito.

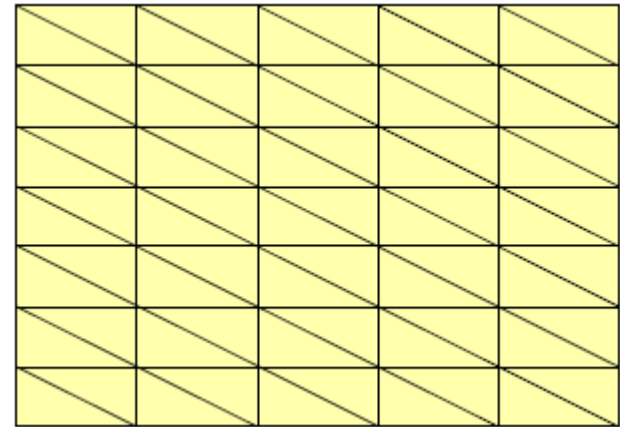
2.2. Diferenças entre o MDF e o MEF

- Assim, tal como no MDF, no MEF ocorre um processo de discretização do domínio, mas diferente daquele, pois o MEF resulta em soluções descritas por polinómios conhecidos em todo o domínio e não apenas em nós da malha de diferenças finitas.
- Outra diferença marcante entre o MDF e o MEF está na topologia de discretização do domínio.
- No MDF 2D empregam-se geralmente malhas de topologia triangular ou retangular estruturada.
- Na malha estruturada os intervalos entre nós adjacentes nas direções x e y são constantes, como pode ser observado na figura.

2.2. Diferenças entre o MDF e o MEF



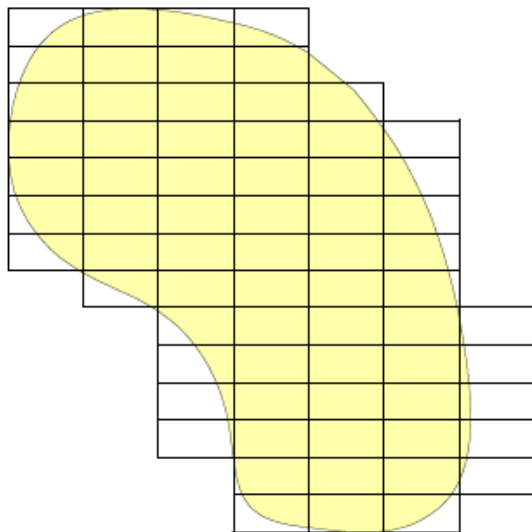
malha estruturada
retangular aplicada a um
polígono regular



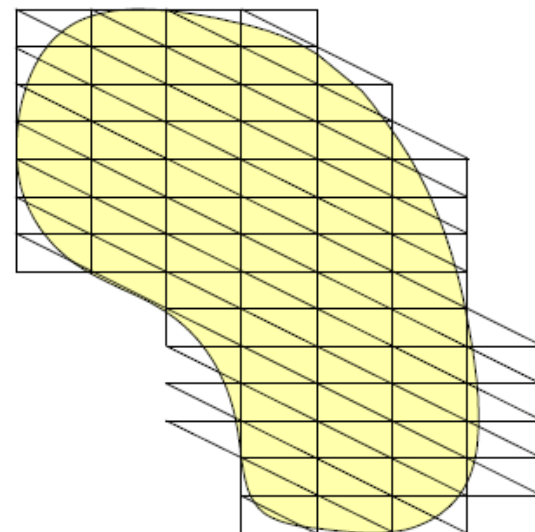
malha estruturada
triangular aplicada a um
polígono regular

2.2. Diferenças entre o MDF e o MEF

- O emprego de malhas estruturadas dificulta a descrição de geometrias irregulares e por essa razão a aplicação do MDF em problemas com geometria irregular resulta em problemas numéricos de aproximação da fronteira.



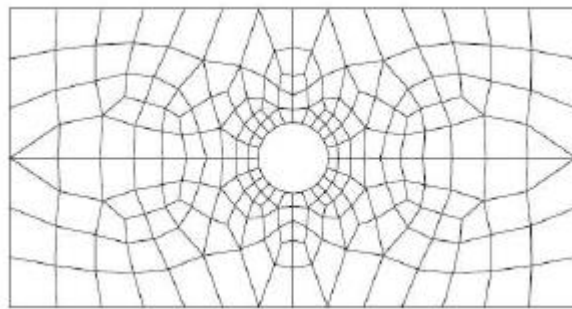
malha estruturada retangular
aplicada a uma figura arbitrária



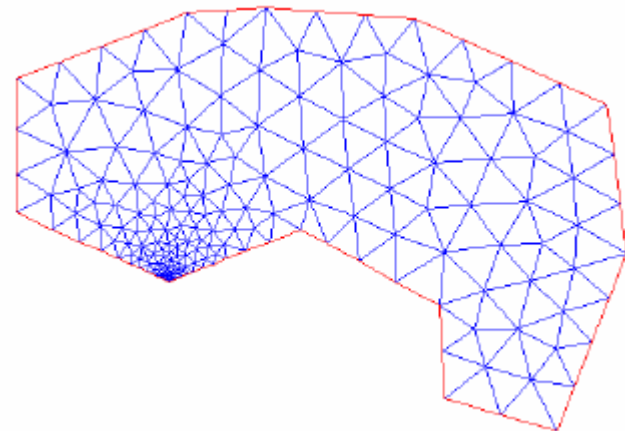
malha estruturada triangular
aplicada a uma figura arbitrária

2.2. Diferenças entre o MDF e o MEF

- O MEF, por sua vez, não requer uma topologia de malha estruturada e, como usualmente emprega uma aproximação polinomial aos valores interiores aos elementos discretizados, pode utilizar para descrever problemas com geometria 2D elementos triangulares ou trapezoidais não estruturados.



malha não estruturada trapezoidal



malha não estruturada triangular

2.3. Forma forte do MEF

- No MEF desenvolveram-se duas formas de resolução de problemas descritos por EDOs e por EDPs.
 - A chamada “forma forte” consiste na resolução direta das equações que governam o problema físico e as suas condições de contorno.
 - A “forma fraca” evoluiu de métodos numéricos aproximados que são representações integrais das equações diferenciais que governam o problema físico.
- A forma forte, em contraste com a forma fraca, requer continuidade nas soluções das variáveis dependentes do potencial.
- Independentemente das funções que definem essas variáveis, elas devem ser diferenciáveis pelo menos até à ordem da equação diferencial que define o problema.

2.3. Forma forte do MEF

- A obtenção da solução exata pela forma forte é, em geral, difícil e limitada a casos especiais.
- O MDF pode ser aplicado na obtenção da solução aproximada de problemas pela forma forte; entretanto, o MDF funciona bem apenas para problemas com geometrias e condições de contorno regulares.
- A forma fraca permite a aplicação de um método único para resolver diferentes tipos de problemas físicos, na medida em que os métodos para transformação das equações diferenciais para a forma integral são genéricos e podem ser usados em diversos tipos de equações diferenciais.
- Os principais métodos usados na resolução pela forma fraca são o método variacional e os métodos dos resíduos ponderados.

2.3.1. Resolução pela forma forte da equação de difusão de calor 1D

A transferência de calor em regime permanente numa barra de comprimento L submetida ao aquecimento q é um problema de valor de fronteira 1D, descrito pela equação de difusão de calor

$$\frac{d}{dx} \left(kA \frac{dT}{dx} \right) - q = 0, \quad 0 < x < L$$

Considerando as condições de contorno homogêneas

$$T(0) = T(L) = 0$$

A solução da equação na forma forte pode ser obtida pela sua integração no intervalo $0 < x < L$ e, atendendo às condições de contorno, resulta em

$$T(x) = \frac{q}{2kA} x(x - L)$$

2.3.1. Resolução pela forma forte da equação de difusão de calor 1D

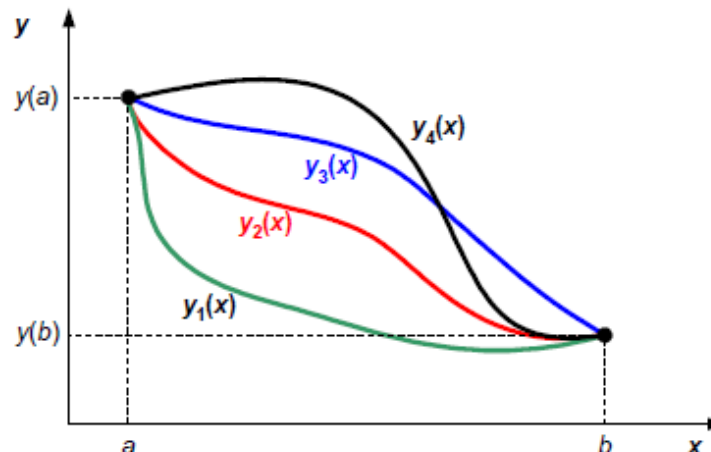
Esta solução analítica representa a solução pela forma forte da equação de difusão de calor. A forma gráfica da equação é uma curva parabólica com um máximo em $x=L/2$.

2.4. Forma fraca do MEF

- Os diversos métodos matemáticos de resolução de problemas de valor de fronteira podem ser classificados em dois métodos principais:
 - Método variacional ou de Rayleigh-Ritz;
 - Método dos resíduos ponderados.

2.4.1. Método Variacional ou Método de Rayleigh-Ritz

O método variacional, desenvolvido independentemente por W. Ritz (1908) e por Lord Rayleigh, é um método analítico no qual se minimiza um funcional que descreve a distância de um caminho limitado nas extremidades $[a,b]$ por uma função $y(x)$. A figura mostra diferentes funções que representam caminhos entre os limites $[a,b]$. O caminho mínimo será determinado pela minimização do funcional $I[y]$.



2.4.1. Método Variacional ou Método de Rayleigh-Ritz

Considere-se a equação diferencial ordinária linear de 2ª ordem:

$$y'' - Q(x)y = F(x)$$

com as condições de contorno: $y(a)=y_a$, $y(b)=y_b$.

O funcional que descreve esta equação diferencial é

$$I[u] = \int_a^b \left[\left(\frac{du}{dx} \right)^2 - Qu^2 + 2Fu \right] dx$$

A relação entre o funcional e a equação diferencial é estabelecida pela condição de Euler-Lagrange:

$$\frac{\partial}{\partial x} \left[\frac{\partial}{\partial y'} F(x, y, y') \right] = \frac{\partial}{\partial y} F(x, y, y')$$

2.4.1. Método Variacional ou Método de Rayleigh-Ritz

na qual a equação diferencial de 2ª ordem é expressa na forma da função $F(x, y, y')$.

A minimização do funcional

$$I[u] = \int_a^b F(x, y, y') dx$$

corresponde à condição que minimiza a função (ou caminho) entre os valores de fronteira $[a, b]$ descrito pela solução da equação diferencial.

2.4.1. Método Variacional ou Método de Rayleigh-Ritz

- Exemplo:

Verificar que o funcional

$$I[u] = \frac{1}{2} \int \left[a \left(\frac{du}{dx} \right)^2 + ku^2 \right] dx$$

é equivalente à equação diferencial

$$a \frac{d^2 u}{dx^2} - ku = 0$$

através do critério de Euler-Lagrange.

2.4.1. Método Variacional ou Método de Rayleigh-Ritz

- Solução:

O integrando do funcional pode ser escrito como

$$F(x, u, u') = au'^2 + ku^2$$

As derivadas de $F(x, u, u')$ são

$$\frac{\partial F}{\partial u'} = 2au' \quad ; \quad \frac{\partial F}{\partial u} = 2ku \quad ; \quad \frac{\partial}{\partial x} \frac{\partial F}{\partial u'} = 2au''$$

Substituindo na condição de Euler-Lagrange obtém-se

$$2au'' = 2ku \quad \Rightarrow \quad au'' - ku = 0$$

que corresponde à equação diferencial inicial.

2.4.2. Distribuição de temperatura numa barra 1D

Vamos aplicar o método variacional para encontrar a distribuição de temperatura em regime permanente numa barra unidimensional de comprimento L submetida ao aquecimento q descrito pela equação de difusão de calor

$$\frac{d}{dx} \left(kA \frac{dT}{dx} \right) - q = 0, \quad 0 < x < L$$

Considerando as condições de contorno homogêneas

$$T(0) = T(L) = 0$$

o funcional da equação diferencial é descrito por

$$I[T] = \int_0^L \left[\left(\frac{dT}{dx} \right)^2 + \frac{2q}{kA} T \right] dx$$

2.4.2. Distribuição de temperatura numa barra 1D

Tendo em conta que a equação diferencial é de 2ª ordem, vamos considerar que a solução tentativa seja descrita pela equação algébrica de 2º grau:

$$T(x) = ax + bx^2$$

Calculando a derivada $dT/dx = a + 2bx$ e substituindo na equação de $I[T]$, vem que

$$I[T] = \int_0^L \left[(a + 2bx)^2 + \frac{2q}{kA} (ax + bx^2) \right] dx$$

Integrando esta equação tem-se

$$I[T] = \left(a^2x + 2abx^2 + \frac{4}{3}b^2x^3 \right) + \frac{2q}{kA} \left(\frac{ax^2}{2} + \frac{bx^3}{3} \right) \Bigg|_0^L$$

2.4.2. Distribuição de temperatura numa barra 1D

ou

$$I[T] = a^2 L + 2a \left(b + \frac{q}{2kA} \right) L^2 + \frac{b}{3} \left(4b + \frac{2q}{kA} \right) L^3$$

2.4.2. Distribuição de temperatura numa barra 1D

Os coeficientes a e b serão determinados pela minimização do funcional $I[T]$ em relação aos coeficientes, isto é, fazendo $\partial I/\partial a=0$ e $\partial I/\partial b=0$.

Aplicando as derivadas parciais de $I[T]$ em função de a e b , vem que

$$\frac{\partial I}{\partial a} = 2aL + 2bL^2 + \frac{qL^2}{kA} = 0$$

$$\frac{\partial I}{\partial b} = 2aL^2 + \frac{8}{3}bL^3 + \frac{2qL^3}{3kA} = 0$$

Resolvendo o sistema de equações acima, obtém-se os coeficientes a e b da solução tentativa.

2.4.2. Distribuição de temperatura numa barra 1D

A solução final fica

$$T(x) = \frac{q}{2kA} x(x - L)$$

A solução descrita por esta equação é idêntica à solução analítica da EDO e das condições de contorno.

Desta forma, mostramos neste exemplo particular que a solução pela forma fraca obtida através do método variacional possui o mesmo resultado da solução analítica na forma forte da EDO.

2.4.3. Método dos Resíduos Ponderados

O funcional também satisfaz as condições de contorno naturais, $du/dx=0$ numa extremidade na qual as condições de contorno essenciais, $u=u_0$, não são aplicadas.

O método dos resíduos ponderados inicia-se com uma equação diferencial genérica na forma

$$Lu = f$$

na qual L é um operador diferencial qualquer.

Este método evita a procura de uma expressão variacional equivalente. Admite-se uma solução aproximada u^* e substitui-se esta solução na equação diferencial.

2.4.3. Método dos Resíduos Ponderados

Como esta é uma solução aproximada, a operação resulta num erro residual na equação diferencial:

$$Lu^* - f = r$$

Não se pode forçar o resíduo r a desaparecer diretamente da equação, mas pode forçar-se, para um integral ponderado sobre o domínio Ω da solução, que o resíduo desapareça.

Isto quer dizer que a solução em Ω da solução do produto do termo residual por uma função peso w tem que ser igual a zero:

$$I = \int_{\Omega} r w d\Omega = 0$$

2.4.3. Método dos Resíduos Ponderados

Substituindo funções de interpolação pela solução aproximada u^* e pela função peso w , resulta num conjunto de equações algébricas que podem ser resolvidas para n coeficientes indeterminados da função de interpolação.

Uma das formas utilizadas para tornar o resíduo $r=Lu^*-f$ pequeno é anular o integral, isto é, anular o resíduo pela média.

Considere que a função peso w é uma função que testa o resíduo, de modo que ela também é conhecida como função teste. A classe de funções teste é tal que o integral possa ser escrito na forma

$$\int_{\Omega} (Lu^*)w d\Omega = \int_{\Omega} f w d\Omega$$

2.4.3. Método dos Resíduos Ponderados

Geralmente, a formulação matemática original baseada na equação diferencial

$$Lu = f$$

denomina-se forma clássica ou forte e a formulação baseada no método dos resíduos ponderados por forma fraca.

Pode demonstrar-se que para funções teste r pertencentes ao subespaço das funções aproximadas u^* , as formulações clássica e fraca são equivalentes e que, portanto, conduzem às mesmas soluções.

2.4.4. Funções de Aproximação

Podem ser obtidas diversas formas de aproximação da função u . Entretanto, as condições estabelecidas para que as formulações forte e fraca sejam equivalentes restringem a forma e o número de aproximações que podem ser utilizadas para as funções u .

O problema consiste em obter-se uma aproximação de uma função real no intervalo $[a,b]$, na forma:

$$u_n(x) = c_1\phi_1(x) + c_2\phi_2(x) + \dots c_n\phi_n(x) = \sum_{j=1}^n c_j\phi_j(x)$$

As funções $\phi_j(x)$ são conhecidas e supostas linearmente independentes e os coeficientes c_j são parâmetros a determinar.

2.4.4. Funções de Aproximação

A equação diferencial pode ser escrita numa outra forma geral como:

$$D[x, y] = 0, \quad a < x < b$$

sujeita às condições de contorno homogéneas

$$T(0) = T(L) = 0$$

No método dos resíduos ponderados escreve-se uma solução na forma

$$u(x) = \sum_{i=1}^n c_i N_i(x)$$

2.4.4. Funções de Aproximação

na qual $u(x)$ é a solução aproximada e exprime como o produto de coeficientes constantes c_i a serem determinados e $N_i(x)$ são funções tentativas (*trial functions*).

Os requisitos das funções tentativas são que sejam contínuas no domínio do problema e que satisfaçam as condições de contorno exatamente.

A escolha das funções tentativas é definida pelo tipo de problema físico descrito pelo problema de valor de fronteira (PVF).

O resíduo $r(x)$ é calculado pela equação

$$r(x) = D[x, u(x)] \neq 0$$

2.4.4. Funções de Aproximação

O método dos resíduos ponderados requer que os coeficientes c_i sejam avaliados de forma que

$$\int_a^b w_i(x)r(x)dx = 0 \quad i = 1, 2, \dots, n$$

onde $w_i(x)$ representam n funções peso que minimizam o integral.

A escolha da função peso $w_i(x)$ define o tipo do método de resíduo ponderado a ser utilizado, de acordo com os seguintes critérios:

- Método de Galerkin: critério $w_i(x) = N_i(x)$
- Método dos Mínimos Quadrados: critério $w_i(x) = \partial u / \partial c_i$
- Método da Colocação: critério $w_i(x) = \delta(x - x_i)$ (*função delta de Dirac*)

2.4.4. Funções de Aproximação

O uso da integração por partes com o método de Galerkin normalmente reduz os requisitos de continuidade das funções de aproximação.

Se o funcional variacional existir, o método de Galerkin fornecerá a mesma aproximação algébrica. Assim, ela oferece sempre uma estimativa de erro ótima para a solução por elementos finitos.

2.4.5. Método de Galerkin

No método dos resíduos ponderados pelo critério de Galerkin também conhecido como *Método de Galerkin*, a função tentativa $N_i(x)$ é igualada à função peso $w_i(x)$, de modo que o sistema de equações lineares é determinado pelo integral

$$\int_a^b w_i(x)r(x)dx = \int_a^b N_i(x)r(x)dx = 0 \quad i = 1, 2, \dots, n$$

Veremos no exemplo seguinte a aplicação do método de Galerkin.

2.4.5. Método de Galerkin

- Exemplo:

Resolver o problema de valor de fronteira descrito pela equação diferencial ordinária

$$\frac{d^2 y}{dx^2} - 10x^2 = 5$$

Sujeita às condições de fronteira homogêneas

$$y(0) = y(1) = 0$$

2.4.5. Método de Galerkin

- Solução:

A presença do termo quadrático na EDO sugere que funções tentativas polinomiais possam ser usadas.

Para as condições de contorno homogêneas em $x=a$ e $x=b$, será usada a seguinte função tentativa

$$N(x) = (x-a)^p (x-b)^q$$

onde as constantes p e q são valores estritamente positivos e inteiros.

Essa função tentativa satisfaz as condições de contorno e é contínua no intervalo $a \leq x \leq b$.

2.4.5. Método de Galerkin

A função tentativa mais simples obtém-se quando se coloca $p=q=1$

$$N_1(x) = x(x-1)$$

Usando esta função tentativa na solução aproximada da EDO (acetato 71)

$$u(x) = c_1 x(x-1)$$

de onde se obtém a primeira e a segunda derivadas

$$\frac{du}{dx} = c_1(2x-1) \quad ; \quad \frac{d^2u}{dx^2} = 2c_1$$

2.4.5. Método de Galerkin

Observamos neste ponto que a solução escolhida não corresponde à solução “física” do PVF, pois a segunda derivada acima é constante, enquanto que na EDO que descreve o problema, a segunda derivada é função da variável x^2 .

Entretanto, continuaremos com o cálculo do problema para ilustrar o método de Galerkin.

Substituindo a segunda derivada de $u(x)$ na equação para o cálculo do resíduo (acetato 67), resulta em

$$r(x) = 2c_1 - 10x^2 - 5$$

que, claramente, é não-nulo.

2.4.5. Método de Galerkin

Substituindo no integral (acetato 75)

$$\int_0^1 x(x-1)(2c_1 - 10x^2 - 5)dx = 0$$

Integrando a equação acima, vem que $c_1=4$, de modo que a solução aproximada resulta em

$$u(x) = 4x(x-1)$$

Para este exemplo simples, podemos encontrar a solução analítica através da integração sucessiva da EDO

$$\frac{dy}{dx} = \int \frac{d^2y}{dx^2} dx = \int (10x^2 + 5) dx = \frac{10}{3} x^3 + 5x + C_1$$

onde C_1 é uma constante de integração.

2.4.5. Método de Galerkin

Integrando novamente

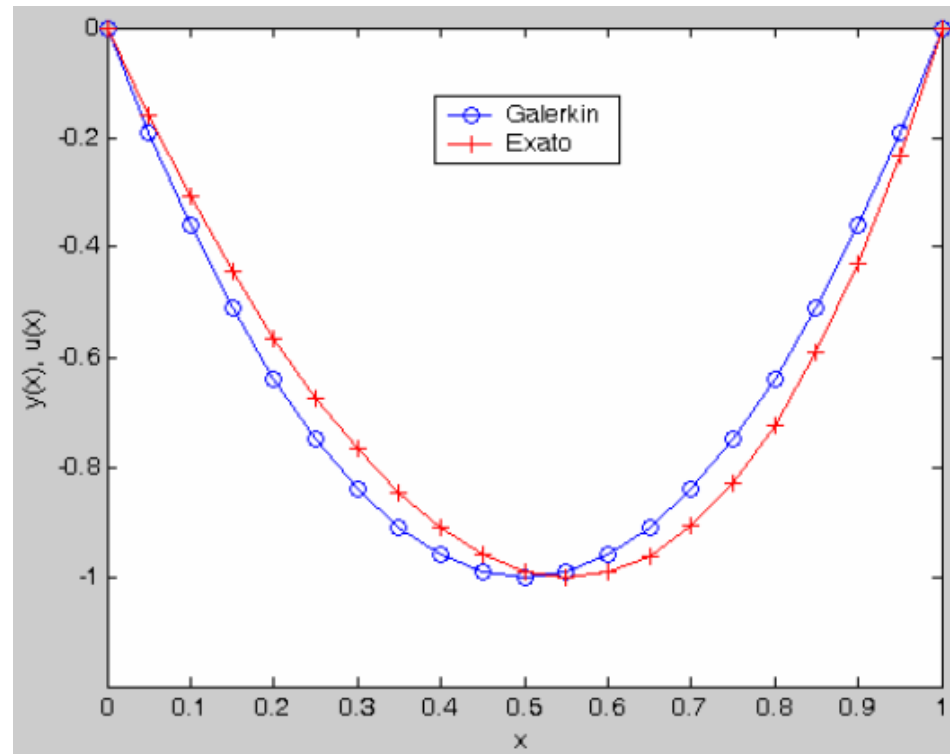
$$y = \int \frac{dy}{dx} dx = \int \left(\frac{10}{3} x^3 + 5x + C_1 \right) dx = \frac{5}{6} x^4 + \frac{5}{2} x^2 + C_1 x + C_2$$

Aplicando a condição de contorno $y(0)=0$, obtém-se $C_2=0$, ao passo que a condição de contorno $y(1)=0$ faz com que $C_1=-10/3$, de maneira que a solução exata seja

$$y = \frac{5}{6} x^4 + \frac{5}{2} x^2 - \frac{10}{3} x$$

2.4.5. Método de Galerkin

A figura abaixo mostra as curvas da solução aproximada pelo método de Galerkin e da solução analítica exata da EDO.



2.4.6. Os diferentes métodos de análise

O diagrama mostrado na figura apresenta os principais métodos analíticos e numéricos para a solução de PVF de equações diferenciais.

Embora o método dos elementos finitos seja uma técnica essencialmente numérica, podem utilizar-se métodos analíticos na sua *forma fraca*.

