

ONDA ELÁSTICA - GALERKIN COM DIREÇÕES ALTERNADAS

Este exemplar corresponde a redação final da tese devidamente corrigida e defendida pelo Sr. JOSÉ AUGUSTO NUNES FERNANDES ^{7/11} e aprovada pela Comissão Julgadora.

Campinas, 08 de Dezembro de 1988.



Prof.ª Dra. MARIA CRISTINA CUNHA BEZERRA ^{7/11}
(ORIENTADORA)

Dissertação apresentada ao Instituto de Matemática, Estatística e Ciência da Computação, UNICAMP, como requisito parcial para obtenção do Título de Mestre em Matemática Aplicada.

À Tinha, Xexa e Guto pelo carinho
e compreensão que sempre me dedicaram.

AGRADECIMENTOS

- À Prof^a. Dr^a. Maria Cristina Cunha Bezerra, minha orientadora, pela dedicação, flexibilidade e segurança que tornaram possível a realização desse trabalho.
- Aos meus pais, Sr. Zé e D^a. Dalva, e aos irmãos: Zé Fernandes, Zé Airton, Zé Luiz, Decíola, Edna, Izolino e Zé Wilson pelo total apoio demonstrado ao longo de minha vida.
- Ao amigo e procurador, professor Raimundo José de Siqueira Mendes, que muito me auxiliou na realização do mestrado.
- Aos amigo Geólogo Luís Augusto e aos professores Mauro Sobrinho, Manoel Alves e Otávio Valle pelos constantes incentivos.
- Aos professores do Núcleo Pedagógico Integrado e do Departamento de Matemática da U.F.P^a que assumiram minhas atividades durante o curso de Mestrado.

- Aos companheiros da turma de mestrado, ingressantes no ano de 1986 : Jônatas, Chico, Socorro, Vinícius, Jurandir, Siomara e Fernando pelo saudável período de convivência que irá deixar saudades.
- Aos amigos Kako, Beto e Alexandre com quem convivi nos últimos meses da realização desse trabalho, e que o tornaram menos árduo.
- Ao amigo Meco que um dia partiu deixando boas recordações como um grande amigo e incentivador.
- Aos meus tios Jersey de Brito Nunes e Roderick Castelo Branco; professores exemplares.
- Ao Professor Petrônio Pulino, pelo exemplo de competência e senso de pesquisa.
- À CAPES, via PICD - U.F.P^a - PROPESP pelo indispensável apoio.

ÍNDICE

NOTAÇÃO GERAL.

I - INTRODUÇÃO.	1
II - ONDA ELÁSTICA.	
2.1 - INTRODUÇÃO.	6
2.2 - EQUAÇÃO DA ONDA ELÁSTICA.	7
2.3 - DECOMPOSIÇÃO DA ONDA ELÁSTICA.	9
2.4 - ONDA ELÁSTICA EM DUAS DIMENSÕES ESPACIAIS.	12
2.5 - CONDIÇÃO DE CONTORNO.	14
2.6 - CONDIÇÃO INICIAL.	14
2.7 - ARGUMENTAÇÃO MATEMÁTICA.	15
2.8 - FONTES PARA O PROBLEMA.	18
2.9 - COND. DE CONTORNO E COND. INICIAIS PARA OS POTENCIAIS.	19
2.10 - EQUAÇÕES A RESOLVER.	20
III - GALERKIN COM DIREÇÕES ALTERNADAS.	
3.1 - INTRODUÇÃO.	22
3.2 - GALERKIN - DIREÇÕES ALTERNADAS.	23
3.3 - PROCEDIMENTO GERAL.	33
3.4 - NUMERAÇÃO DA MALHA.	34
3.5 - BASE PARA O ELEMENTO FINITO.	34
3.6 - RESOLUÇÃO DOS SISTEMAS LINEARES.	39

3.7 - ARMAZENAMENTO.	44
3.8 - PROBLEMA COM C. DE CONTORNO TIPO NEUMANN HOMOGÊNEO NA FRONTEIRA.	46
IV - FERRAMENTAS COMPUTACIONAIS.	
4.1 - INTRODUÇÃO.	49
4.2 - ALGORITMO.	50
4.3 - FLUXOGRAMA PARA ONDAS ACÚSTICAS.	51
4.4 - O PROGRAMA DE COMPUTADOR.	53
CONCLUSÃO.	96
SUGESTÕES PARA TRABALHOS FUTUROS.	98
BIBLIOGRAFIA.	99
APÊNDICE - MÉTODO DA COLOCAÇÃO	101

NOTAÇÃO GERAL

Neste nosso trabalho, adotaremos a seguinte notação :

Capítulos - A numeração será em algarismo romano, precedida de CAP.

Tópicos e Subtópicos - A numeração será em letras do tipo orador , constando o número do capítulo, seguido primeiramente do número do tópico dentro daquele capítulo e posteriormente do número do subtópico. Exemplo : 4.3.1 - refere-se ao primeiro subtópico do terceiro tópico do quarto capítulo.

Figuras e Tabelas - Serão identificadas por números pequenos que identificarão a figura, ou tabela, e o capítulo da qual ela faz parte, a numeração das figuras serão precedidas de fig. enquanto as tabelas terão TAB. a frente da numeração.

Usualmente teremos neste trabalho :

$$\Omega = [a , b] \times [c , d]$$

$$\bar{\Omega} = \text{Fecho de } \Omega.$$

$$\partial \Omega = \text{Fronteira de } \Omega.$$

$$h_x = x_j - x_{j-1}$$

$$h_y = y_j - y_{j-1}$$

$$\Delta t = t^{n+1} - t^n$$

F , G , H ... - Funções escalares F , G , H ,... de n variáveis.

$\hat{U}$, $\hat{F}$,... - Vetores pertencentes ao $\mathbb{R}^n$.

F_i - A i -ésima componente do vetor F .

∇F - Vetor gradiente de F , tal que $F_i = \left(\frac{\partial F}{\partial x_i} \right)_{i=1}^n$.

$\nabla \cdot \hat{U}$ - Divergente do vetor $\hat{U} = \sum_{i=1}^n \left(\frac{\partial U_i}{\partial x_i} \right)$

$\nabla \times \hat{U}$ - Rotacional do vetor U . Se $U \in \mathbb{R}^3$, $\nabla \times U = \begin{pmatrix} \frac{\partial U_3}{\partial x_2} - \frac{\partial U_2}{\partial x_3} \\ \frac{\partial U_1}{\partial x_3} - \frac{\partial U_3}{\partial x_1} \\ \frac{\partial U_2}{\partial x_1} - \frac{\partial U_1}{\partial x_2} \end{pmatrix}$

ΔF - Operador Laplaciano = $\left(\frac{\partial^2}{\partial x_1^2}, \frac{\partial^2}{\partial x_2^2}, \dots, \frac{\partial^2}{\partial x_n^2} \right)$

$L^2(\Pi)$ - Espaço das funções de quadrados integráveis sobre Π .

H_0^1 - Subespaço de Sobolev, das funções que tem gradiente em $L^2(\pi)$, e anulam-se na fronteira.

H^1 - Subespaço de Sobolev, das funções que tem gradiente em $L^2(\pi)$.

$A \otimes B$ - Produto tensorial entre A e B .

$\langle f, g \rangle = \iint_{\Omega} f \cdot g \, dx \, dy$

1.1 - INTRODUÇÃO.

No tempo em que vivemos, ainda é crescente o interesse na descoberta de jazidas de petróleo. Fontes alternativas de energia, no que pesem suas eficiências, podem acarretar danos que não compensem a sua utilização. A título de exemplo podemos citar duas alternativas e suas deficiências: no caso da energia nuclear está demonstrado que o risco de catástrofes e o lixo atômico podem comprometer seriamente sua utilização. Por outro lado a geração de energia hidrelétrica está associada a construções de hiper-estruturas com conseqüentes desmatamentos e inundações que podem modificar o sistema ecológico de vastas regiões e provocar danos imprevisíveis.

Geofísicos com interesse na prospecção de jazidas de petróleo, desenvolvem métodos sofisticados para detectar e dimensionar estas jazidas. Muitos desses podem ser classificados numa sub classe que é conhecida como " Métodos Sísmicos ".

O estudo da propagação de ondas em meios acústicos e elásticos tem despertado muito interesse por parte dos geofísicos. Os Métodos Sísmicos são gerados pela modelagem matemática deste fenômeno físico; basicamente trata-se de um sistema de equações diferenciais parciais com condições iniciais e condições de contorno adequadas.

O princípio básico dos métodos sísmicos pode ser visualizado pela figura simplificada 1.1 .

PROPAGACAO DE ONDA EM MEIO ELASTICO

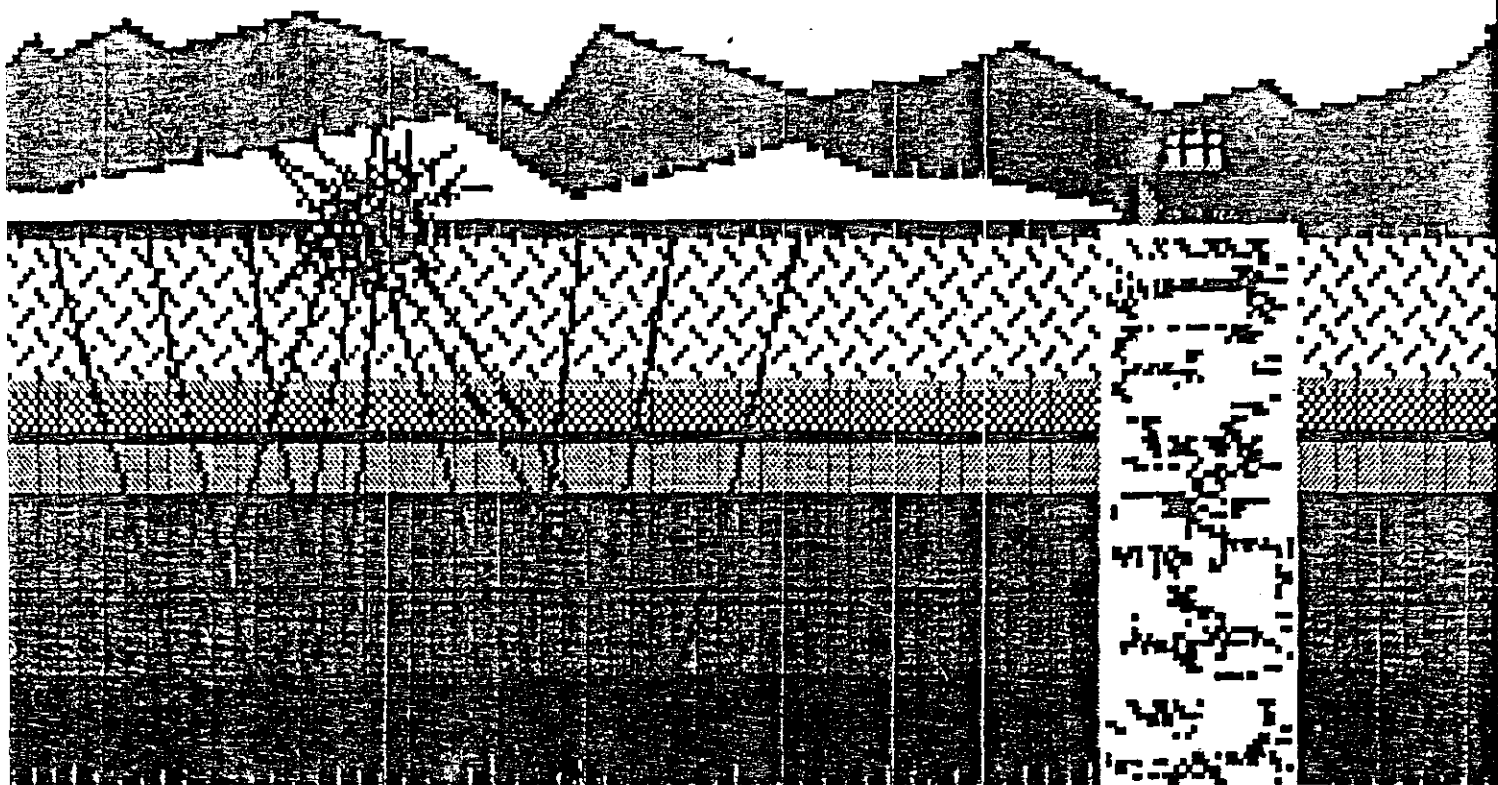


FIG . 1.1

Numa simples descrição, um abalo é produzido por uma ou mais fontes de propagação sísmica (dinamite por

exemplo), posicionadas numa determinada localização da terra (que pode estar na superfície ou no subsolo) . Ondas propagam-se a partir da localização da fonte e são parcialmente refletidas para a superfície produzindo movimento superficial. A componente vertical do movimento é detectada por sismômetros localizados também na superfície ou no interior da terra; esses aparelhos gravam os sinais recebidos para posterior análise.

Matematicamente podemos caracterizar dois grandes problemas associados à interpretação das respostas obtidas para o modelo e/ou dos sinais captados pelos sismômetros. O primeiro usualmente é denominado de problema inverso em que se tenta, através de métodos matemáticos, descobrir os parâmetros característicos dos meios (densidade , módulo de elasticidade, etc...). O segundo conhecido como problema direto, no qual se parte do pressuposto de que os parâmetros do meio são conhecidos e tenta-se através da resolução de equações diferenciais parciais obter o deslocamento no caso elástico, ou a pressão no caso acústico, decorrente do abalo sísmico. Neste segundo caso as soluções das equações diferenciais que modelam o movimento podem ser obtidas tanto por meios analíticos quanto por meios numéricos (diferenças finitas, elementos finitos, colocação, etc...), ou por uma combinação de métodos analíticos e numéricos.

O empenho nestes tipos de problemas é de fundamental importância e podemos dizer que eles interagem no sentido

de que ao resolvermos um problema inverso devemos constatar a exatidão dos resultados pelo método direto, e vice-versa.

Podemos considerar este nosso estudo como introdutório no campo da pesquisa de métodos numéricos para problemas diretos, e aqui obteremos solução discreta pelo uso do método de Galerkin.

Nosso ponto de partida foi o estudo da teoria da elasticidade até a dedução da equação do movimento para um meio elástico homogêneo isotrópico, equação esta que servirá de modelo para a obtenção do deslocamento. Estudamos ainda as soluções analíticas da onda livre, ou seja, aquela em que o termo fonte é desprezado.

Na busca de soluções discretas para a equação do movimento na presença do termo fonte (situação esta de maior sentido prático), tentamos inicialmente atacar diretamente esta equação pelo método de Galerkin e chegamos a um sistema de equações diferenciais que seria de complexa resolução. Resolvemos então adotar a sugestão proposta em [1],[2],[5],[6], onde o deslocamento e a fonte são decompostos de maneira adequada. Chega-se então a duas equações da onda acústica em termos de potenciais que, desde que conhecidos, nos permitem calcular o deslocamento através de simples derivações.

Com relação aos métodos numéricos de resolução de ondas acústicas, inicialmente tentamos o método de Galerkin e tivemos problemas quanto a espaço de memória, quer seja a nível de microcomputadores pessoais tipo IBM-PC, ou mesmo no computador VAX/VMS

Versão 4.5 da UNICAMP. Optamos então pelo método da colocação natural com B-Splines e pela colocação ortogonal nos nós gaussianos da partição [13], neste caminho nos deparamos com problemas de instabilidade numérica.

Fizemos uso então do método de Galerkin com direções alternadas [3]; o problema de instabilidade é contornado pela introdução de um novo termo envolvendo um parâmetro independente da partição no espaço ou no tempo. Este método mostrou-se eficiente para nossos propósitos.

Desenvolvemos também um software correspondente que pode ser usado nas duas equações acústicas surgidas em decorrência da decomposição em ondas longitudinal e transversal. Tomando por base as proposições de [3] e [4], apresentaremos a solução para essas equações em função dos potenciais, obtendo daí o deslocamento correspondente a uma determinada fonte, para uma única camada, com condição de contorno explicitada e condições iniciais características do nosso problema.

Encerrando o trabalho apresentaremos as conclusões e sugestões que acreditamos serão de validade para trabalhos posteriores.

2.1 - INTRODUÇÃO.

Muitos fenômenos físicos são observados a partir de uma perturbação inicial, que chamaremos condições iniciais ou condições em $t = 0$. Nesta ocasião, são produzidos movimentos ondulatórios com origem no local da perturbação. Estes fenômenos são conhecidos como radiação e dentre as várias situações físicas desta natureza podemos citar os abalos provocados por terremotos ou explosões nucleares subterrâneas. Um problema típico de radiação é o resultante de um abalo produzido em um meio ilimitado por uma ou mais fontes distribuídas sobre uma região finita do meio. O movimento decorrente das perturbações provocadas por estas fontes atuando num meio elástico homogêneo e isotrópico, é matematicamente modelado de uma maneira conveniente pela equação do movimento; descreveremos este modelo nesse capítulo.

Com base no modelo, admitindo que o deslocamento (que será a incógnita da equação do movimento) e a força possam ser decompostos convenientemente em somas de campos vetoriais, é possível estabelecer uma nova formulação que simplifica o cálculo do deslocamento. Para tornar o problema matematicamente bem colocado, condições de contorno e condições iniciais serão juntadas às equações diferenciais. Trataremos também da validade desse tipo de decomposição, assim como das implicações dela decorrentes quando da resolução numérica.

2.2 - EQUAÇÃO DA ONDA ELÁSTICA.

O sistema de equações que rege o movimento de um dado corpo ocupando uma região regular V de um meio elástico homogêneo e isotrópico, de fecho $\bar{V}$ e contorno S é obtido a partir das seguintes relações [6]:

- i) Equação do movimento de tensões.
- ii) Lei de Hooke.
- iii) Relação deformação-deslocamento.

Se a relação deformação-deslocamento for substituída na lei de Hooke, obtém-se um sistema completo de equações diferenciais parciais para as tensões e deformações, se substituirmos a equação resultante em i) teremos o sistema de equações para determinação do deslocamento, que é matematicamente descrito por:

$$\rho \frac{\partial^2 \hat{U}}{\partial t^2} = \mu \Delta \hat{U} + (\lambda + \mu) \nabla \nabla \cdot \hat{U} + \hat{F} \quad (2.1)$$

onde:

$\hat{U}$ = vetor deslocamento.

ρ = densidade do meio.

λ e μ = constantes de Lamé (características do meio).

$\hat{F}$ = forças do corpo (por unidade de volume).

Δ = operador Laplaciano.

De um modo geral, devemos ter:

$$\hat{U}(x_1, x_2, x_3, t) \in C^2(V \times T) \cap C^1(\bar{V} \times T)$$

e

$$\hat{F}(x_1, x_2, x_3, t) \in C(\bar{V} \times T).$$

Devemos ainda lhe prescrever condições na fronteira S . As mais comuns utilizadas são :

i) Condição sobre o Deslocamento.

Quando as componentes do vetor deslocamento devem ser estipuladas a priori na região de fronteira.

ii) Condição de Tração.

Neste caso, são impostas condições de tração, o que corresponde a pré determinar as componentes do tensor tensão.

iii) Condição Mista

Corresponde a estabelecer condições do tipo i) numa região S_1 de S , e do tipo ii) em outra parte $S - S_1$.

Há necessidade também de se estabelecer condições iniciais. Devemos então ter informações no tempo $t = 0$, sobre o deslocamento e sobre sua derivada com relação a variável tempo.

A equação (2.1) pode ser reescrita utilizando a identidade vetorial :

$$\Delta \hat{U} \equiv \nabla \nabla \cdot \hat{U} - \nabla \times \nabla \times \hat{U}$$

$$\frac{\partial^2 \hat{U}}{\partial t^2} = \frac{\lambda + 2\mu}{\rho} \nabla \nabla \cdot \hat{U} - \frac{\mu}{\rho} \nabla \times \nabla \times \hat{U} + \frac{1}{\rho} \hat{F} \quad (2.2)$$

2.3 - DECOMPOSIÇÃO DA ONDA ELÁSTICA.

Se tentarmos atacar a equação (2.2), com as suas condições de contorno e condições iniciais, veremos que será necessário resolver um sistema de equações diferenciais não muito simples. Para contornar este problema, optamos por decompor o vetor deslocamento, assim como o vetor da força, de uma maneira adequada em termos de potenciais. Esses, quando desacoplados, satisfazem as equações da onda, daí afirmarmos que a equação (2.1) descreve um movimento ondulatório.

Dessa forma, o vetor deslocamento é decomposto numa soma vetorial:

$$\hat{U} = \hat{U}_L + \hat{U}_T \quad (2.3)$$

onde:

$$\hat{U}_L = \nabla \varphi \quad \text{por nós denominada de Onda Longitudinal.}$$

e

$$\hat{U}_T = \nabla \times \hat{\psi} \quad \text{por nós denominada de Onda Transversal.}$$

Para φ e $\hat{\psi}$, funções do tempo e variáveis no espaço, usualmente também chamadas de potenciais escalar e vetorial respectivamente.

De maneira idêntica (à da decomposição proposta para o vetor deslocamento) procederemos em relação ao vetor $\hat{F}$, sendo:

$$\hat{F} = \nabla \sigma + \nabla \times \hat{\chi} \quad (2.4)$$

Substituindo (2.3) e (2.4) em (2.2), teremos:

$$\begin{aligned} \frac{\partial^2}{\partial t^2} [\nabla \varphi + \nabla \times \hat{\psi}] = c_L^2 \nabla \nabla \cdot [\nabla \varphi + \nabla \times \hat{\psi}] - c_T^2 \nabla \times \nabla \times [\nabla \varphi \\ + \nabla \times \hat{\psi}] + \frac{1}{\rho} [\nabla \sigma + \nabla \times \hat{\chi}] \end{aligned} \quad (2.5)$$

onde:

$$c_L^2 = \frac{\lambda + 2\mu}{\rho}$$

e

$$(2.6)$$

$$c_T^2 = \frac{\mu}{\rho}$$

e assim:

$$\nabla [\frac{\partial^2 \varphi}{\partial t^2} - c_L^2 \Delta \varphi - \frac{1}{\rho} \sigma] + \nabla \times [\frac{\partial^2 \hat{\psi}}{\partial t^2} - c_T^2 \Delta \hat{\psi} - \frac{1}{\rho} \hat{\chi}] = 0 \quad (2.7)$$

que satisfaz a equação do movimento (2.1) se simultaneamente tivermos:

$$\left\{ \begin{array}{l} \frac{\partial^2 \varphi}{\partial t^2} - c_L^2 \Delta \varphi = \frac{1}{\rho} \sigma \quad (2.8) \\ \frac{\partial^2 \hat{\psi}}{\partial t^2} - c_T^2 \Delta \hat{\psi} = \frac{1}{\rho} \hat{\chi} \quad (2.9) \end{array} \right.$$

Assim os potenciais φ e $\hat{\psi}$ são soluções das equações das ondas (2.8) e (2.9) respectivamente, onde:

$$c_L = \sqrt{(\lambda + 2\mu)/\rho} \quad \text{é a velocidade da onda longitudinal.}$$

e

$$c_T = \sqrt{\mu/\rho} \quad \text{é a velocidade da onda transversal.}$$

Observe que das expressões de c_L e c_T tiramos:

$$c_L / c_T = \sqrt{(\lambda + 2\mu)/\mu}$$

portanto, como $\mu > 0$ e $\lambda \geq 0$ para materiais reais, temos que:

$$c_L \geq \sqrt{2} c_T \quad (2.10)$$

ou seja, a onda longitudinal propaga-se com maior velocidade do que a onda transversal.

A seguir mostramos uma tabela com alguns valores aproximados para a densidade e para as velocidades das ondas longitudinal (c_L) e transversal (c_T) de alguns meios:

Meio	ρ (Kg/m ³)	C_L (m/s)	C_T (m/s)
Ar	1.2	340	-
Água	1000	1480	-
Aço	7800	5900	3200
Cobre	8900	4600	2300
Alumínio	2700	6300	3100
Vidro	2500	5800	3400

TAB. 2.1

Fisicamente, verifica-se que a decomposição proposta em (2.3) e (2.4) é procedente, e que a equação do movimento (2.1) contém implicitamente dois tipos de onda: longitudinal e transversal.

A onda longitudinal tem movimento paralelo à direção de propagação, e a onda transversal tem seu movimento perpendicular à direção de propagação.

As ondas, longitudinal e transversal, também recebem outras denominações.

Pela decomposição proposta em (2.3), podemos observar que $\nabla \times \hat{U}_L = 0$, o que torna natural a sua denominação de onda irrotacional, outros nomes ainda lhe são atribuídos: dilatacional, compressional, primária ou onda P.

Verificamos também que se em (2.3) $\hat{U}_T = \nabla \times \hat{\psi}$ teremos que $\nabla \cdot \hat{U}_T = 0$, daí a natural nomenclatura de onda solenoidal, mas também a encontramos com os nomes de : rotacional, cisalhante, secundária ou onda S [4].

2.4 - ONDA ELÁSTICA EM DUAS DIMENSÕES ESPACIAIS.

Em várias situações físicas o modelo bidimensional pode ser utilizado na análise do movimento; nesse caso, a representação de $\hat{U}$ é mais simplificada. De (2.3), podemos obter as expressões das duas únicas componentes de $\hat{U}$:

$$U_1 = \frac{\partial \varphi}{\partial x_1} + \frac{\partial \hat{\psi}}{\partial x_2}$$

e

(2.11)

$$U_2 = \frac{\partial \varphi}{\partial x_2} - \frac{\partial \hat{\psi}}{\partial x_1}$$

O mesmo ocorrerá com as componentes da força, que devido a (2.4) serão :

$$F_1 = \frac{\partial \sigma}{\partial x_1} + \frac{\partial \hat{\chi}}{\partial x_2}$$

e

(2.12)

$$F_2 = \frac{\partial \sigma}{\partial x_2} - \frac{\partial \hat{\chi}}{\partial x_1}$$

Como observamos anteriormente, os potenciais, escalar e vetorial, são funções variáveis no espaço e no tempo, e aqui, o potencial vetorial é do tipo $\hat{\psi}(0,0,\psi(x_1,x_2,t))$, o mesmo ocorrendo com a componente vetorial da força, assim sendo, daqui por diante, quando estivermos tratando de problemas bidimensionais no espaço, faremos referência apenas às funções ψ e χ , respectivamente.

2.5 - CONDIÇÕES DE CONTORNO.

Se o meio é homogêneo, ou seja, a densidade, assim como as constantes de Lamé não variam, as ondas longitudinal e transversal não interagem enquanto se propagam. Por outro lado, se as constantes variam no meio, o que se costuma fazer é estratificá-lo como se possuissem diversas camadas homogêneas. Nesse caso, a onda longitudinal ao entrar em contacto com uma outra camada irá gerar uma onda transversal e vice-versa; aí então, essas interações devem estar contidas implicitamente nas condições de contorno. Em [1] encontramos os diversos tipos de condições de contorno correspondentes a situações físicas reais. No nosso estudo, nos limitaremos a um único meio homogêneo em contacto com uma superfície S , infinitamente rígida, , ou seja de um altíssimo módulo de elasticidade; nesse caso :

$$U_j|_S = 0 \quad j = 1,2 \quad (2.13)$$

Matematicamente essa é uma condição de contorno do tipo Dirichlet, ou seja, $\hat{U} \equiv 0$ na fronteira.

2.6 - CONDIÇÕES INICIAIS.

Como citamos anteriormente, para que um problema de evolução, matematicamente caracterizado como uma equação hiperbólica de segunda ordem, esteja bem posto há de se estabelecer duas condições iniciais: o valor da solução e de sua derivada com

relação ao tempo no instante inicial $t = 0$. Considerando que no tempo zero estaremos em repouso completo, admitiremos condições iniciais homogêneas :

$$U(x_1, x_2, 0) = 0 \quad (2.14)$$

e

$$U_t(x_1, x_2, 0) = 0 \quad (2.15)$$

2.7 - ARGUMENTAÇÃO MATEMÁTICA.

Antes de prosseguirmos na efetiva resolução do modelo proposto, analisaremos alguns aspectos que nos garantem a existência de uma decomposição vetorial do tipo (2.3) e (2.4), bem como a existência de solução para o problema matemático (2.1) com suas respectivas condições de contorno e condições iniciais.

A questão da decomposição foi demonstrada primeiramente há mais de um século por Clebsch em 1863 (de fato num trabalho datado de 1861), embora seus argumentos tenham sido considerados insatisfatórios do ponto de vista matemático. Duhem provou de modo aceitável em 1898 [5]. Em [8] encontramos a demonstração do teorema a seguir que nos garante a decomposição proposta.

Teo. 2.1 - Para cada função $\hat{U}$ do $[L^2(\Omega)]^n$, existem uma função $\varphi \in H^1(\Omega)$ e uma função $\psi \in H^1(\Omega)$ se $n = 2$ (respectivamente $\hat{\psi} \in (H^1(\Omega))^3$ se $n = 3$), tais que :

$$\left\{ \begin{array}{l} \hat{U} = \overrightarrow{\text{grad}} \varphi + \begin{cases} \overrightarrow{\text{rot}} \psi & \text{se } n = 2 \\ \overrightarrow{\text{rot}} \hat{\psi} & \text{se } n = 3 \end{cases} \\ (\hat{U} - \overrightarrow{\text{grad}} \varphi) \cdot \vec{\nu} = 0 \end{array} \right.$$

onde :

$$\overrightarrow{\text{rot}} \psi = \left(-\frac{\partial \psi}{\partial x_2}, -\frac{\partial \psi}{\partial x_1} \right) \quad (2.16)$$

e

$$\overrightarrow{\text{rot}} \hat{\psi} = \left(\frac{\partial \psi_3}{\partial x_2} - \frac{\partial \psi_2}{\partial x_3}, \frac{\partial \psi_1}{\partial x_3} - \frac{\partial \psi_3}{\partial x_1}, \frac{\partial \psi_2}{\partial x_1} - \frac{\partial \psi_1}{\partial x_2} \right)$$

e $\vec{\nu}$ é o vetor unitário normal exterior a Ω .

Como nosso atual estudo está sendo feito para duas coordenadas espaciais, estaremos tratando do caso em que $n = 2$.

Dessa forma teremos uma única função escalar φ (a menos de uma constante aditiva), que é determinada por:

$$\left\{ \begin{array}{ll} \Delta \varphi = \text{div} \cdot \hat{U} & \text{em } \Omega \\ \frac{\partial \varphi}{\partial \nu} = \hat{U} \cdot \vec{\nu} & \text{em } \partial \Omega \end{array} \right. \quad (2.17)$$

e uma única função ψ , tal que $\psi = 0$ em $\partial \Omega$, que é determinada por

$$\left\{ \begin{array}{ll} -\Delta \psi = \text{rot} \hat{U} & \text{em } \Omega \\ \psi = 0 & \text{em } \partial \Omega \end{array} \right. \quad (2.18)$$

A equação (2.17) é prontamente obtida quando tomamos as equações de (2.11) e fazemos:

$$\frac{\partial U_1}{\partial x_1} + \frac{\partial U_2}{\partial x_2} \quad (2.19)$$

enquanto que as condições de contorno surgem como consequência da condição de existência da decomposição vetorial que será imposta no próximo teorema 2.2.

Para obtermos a equação (2.18) devemos então tomar as equações de (2.11) e fazer:

$$\frac{\partial U_1}{\partial x_2} - \frac{\partial U_2}{\partial x_1} \quad (2.20)$$

Em [8] encontramos ainda, como ficam as condições de contorno para (2.17) e (2.18) quando tratamos de diversas camadas.

Quanto à questão da existência de solução para (2.1), baseada na aplicação do teorema 2.1 e outros argumentos matemáticos, encontramos em [6 pag. 85] a demonstração do seguinte teorema:

Teo. 2.2 - Sejam $\hat{U}(\hat{x}, t)$ e $\hat{F}(\hat{x}, t)$, satisfazendo as condições: $\hat{U} \in C^2(V \times T)$ e $\hat{F} \in C(V \times T)$ e a equação :

$$\rho \frac{\partial^2 \hat{U}}{\partial t^2} = \mu \Delta \hat{U} + (\lambda + \mu) \nabla \nabla \cdot \hat{U} + \hat{F}$$

Onde $\hat{F} = \nabla \sigma + \nabla \chi$. Então existem uma função escalar $\varphi(\hat{x}, t)$ e uma função vetorial $\hat{\psi}(\hat{x}, t)$, tais que $\hat{U}(\hat{x}, t)$, é representado por $\hat{U} = \nabla \varphi + \nabla \chi \hat{\psi}$ onde $\nabla \cdot \varphi = 0$, e onde $\varphi(\hat{x}, t)$ e $\hat{\psi}(\hat{x}, t)$ satisfazem as equações das ondas não homogêneas :

$$\left\{ \begin{array}{ll} \frac{\partial^2 \varphi}{\partial t^2} - c_L^2 \Delta \varphi = \frac{1}{\rho} \sigma & c_L^2 = \frac{\lambda + 2\mu}{\rho} \\ \frac{\partial^2 \hat{\psi}}{\partial t^2} - c_T^2 \Delta \hat{\psi} = \frac{1}{\rho} \hat{\chi} & c_T^2 = \frac{\mu}{\rho} \end{array} \right.$$

2.8 - FONTES PARA O PROBLEMA

A decomposição da fonte de maneira análoga à do deslocamento, requer o conhecimento de funções σ e χ que satisfaçam as condições estabelecidas em (2.17) e (2.18) para a decomposição de $\hat{F}$ em $\nabla \sigma + \nabla \chi$. Tomaremos então funções σ e χ que com as seguintes condições de contorno:

$$\chi = 0 \quad \text{e} \quad \frac{\partial \sigma}{\partial \nu} = \hat{F} \cdot \vec{\nu} \quad \text{em} \quad \partial \Omega \quad (2.21)$$

de posse dessas funções σ e χ as componentes da fonte serão :

$$F_1 = \frac{\partial \sigma}{\partial x_1} + \frac{\partial \chi}{\partial x_2}$$

e

(2.22)

$$F_2 = \frac{\partial \sigma}{\partial x_2} - \frac{\partial \chi}{\partial x_1}$$

respectivamente.

2.9 - CONDIÇÕES DE CONTORNO E CONDIÇÕES INICIAIS PARA OS POTENCIAIS.

Em 2.5 estabelecemos as condições de contorno para o deslocamento. Há no entanto, necessidade de conhecermos as condições de contorno dos potenciais φ e ψ com as quais resolveremos as equações das ondas acústicas não homogêneas estabelecidas em (2.8) e (2.9).

O potencial vetorial ψ terá condição do tipo Dirichlet, identicamente nulo na fronteira, em razão de sua unicidade quando da decomposição, ou seja:

$$\psi \equiv 0 \quad \text{em } \partial \Omega. \quad (2.23)$$

Para existência da decomposição, o potencial escalar tem que obedecer a seguinte condição de fronteira:

$$\frac{\partial \varphi}{\partial \nu} = \hat{U} \cdot \vec{\nu} \quad \text{em } \partial \Omega \quad (2.24)$$

e como $\hat{U} \equiv 0$ em $\partial \Omega$, teremos então :

$$\frac{\partial \varphi}{\partial \nu} = 0 \quad \text{em } \partial \Omega. \quad (2.25)$$

Portanto o potencial escalar φ terá

condição de fronteira do tipo Neumann homogêneo.

As condições iniciais para os potenciais serão idênticas às do vetor deslocamento uma vez que os mesmos surgiram como desdobramentos no espaço e não no tempo; portanto, os potenciais, assim como as suas derivadas primeiras em relação à variável tempo serão nulas no instante inicial.

2.10 - EQUAÇÕES A RESOLVER.

Para obtermos o vetor deslocamento, solução da equação do movimento para um meio Elástico homogêneo e isotrópico, modelado matematicamente em (2.1) devemos então resolver as seguintes equações das ondas acústicas para os potenciais φ e ψ :

$$\left\{ \begin{array}{l} \frac{\partial^2 \varphi}{\partial t^2} - c_L^2 \Delta \varphi = \frac{1}{\rho} \sigma \quad \text{em } \Omega \in \mathbb{R}^2 \quad t \in [0, T] \\ \varphi(x_1, x_2, 0) = \varphi_t(x_1, x_2, 0) = 0 \quad (x_1, x_2) \in \Omega \\ \frac{\partial \varphi}{\partial \nu} = 0 \quad \text{em } \partial \Omega \end{array} \right. \quad (2.26)$$

e

$$\left\{ \begin{array}{l} \frac{\partial^2 \psi}{\partial t^2} - c_T^2 \Delta \psi = \frac{1}{\rho} \chi \quad \text{em } \Omega \in \mathbb{R}^2 \quad t \in [0, T] \\ \psi(x_1, x_2, 0) = \psi_t(x_1, x_2, 0) = 0 \quad (x_1, x_2) \in \Omega \\ \psi = 0 \quad \text{em } \partial \Omega \end{array} \right. \quad (2.27)$$

De posse dos potenciais ϕ e ψ , calculamos o deslocamento correspondente à fonte por nós imposta em 2.9, por meio das equações de (2.11).

3.1 - INTRODUÇÃO.

No segundo capítulo verificamos que para determinarmos o vetor deslocamento em meios elásticos homogêneos e isotrópicos, modelado matematicamente pela equação (2.1), de um modo que nos facilitasse a implementação de métodos numéricos, fizemos uma decomposição dos vetores fonte e deslocamento; em razão disso obtivemos duas equações de ondas acústicas para os potenciais escalar e vetorial. De posse dos potenciais, através das equações (2.11) podemos obter o deslocamento correspondente.

Nesse capítulo nos deteremos no tratamento numérico usado na obtenção da solução discreta das referidas equações acústicas, isto é, as equações hiperbólicas do tipo (2.25) e (2.26). Trataremos aqui de um caso mais genérico, ou seja, resolveremos não somente para os casos de condições iniciais homogêneas (como nos nossos problemas), mas para as situações em que as condições iniciais sejam funções quaisquer.

De um modo geral, nos basearemos nas proposições das referências [3] e [4]. Na primeira Douglas elabora um estudo completo de como o método de direções alternadas pode ser aplicado na formulação de Galerkin, para problemas elípticos, parabólicos e hiperbólicos quando o domínio de definição das variáveis espaciais é retangular. Por outro lado, na segunda referência Fairweather generaliza as idéias de Douglas, para regiões que possam ser decompostas em retângulos (regiões em forma de L ou U por exemplo). Aqui pouco nos deteremos nas demonstrações das

proposições contidas nestas duas referências.

O método de Galerkin com direções alternadas apresentado por Douglas mostrou-se de uma complexidade computacional não muito elevada, e de uma boa eficiência, principalmente no tocante a espaço de memória, o que permite a implementação a nível de micro computadores. Quando já se ouve falar em multi-processadores como algo de fácil acesso dentro em breve, a nível nacional, acreditamos também que esse método deva adaptar-se com vantagens às máquinas com essas características, uma vez que diversas tarefas podem ser realizadas simultaneamente.

Abordaremos inicialmente a equação (2.25), mostrando depois como atacar (2.26), de maneira que ambas possam ser resolvidas por um único programa de computador.

3.2 - GALERKIN COM DIREÇÕES ALTERNADAS.

Consideremos o problema hiperbólico de segunda ordem:

Encontrar $U = U(x,y,t)$ tal que :

$$\left\{ \begin{array}{l} \frac{\partial^2 U}{\partial t^2} - a(x,y) \Delta U = f(x,y,t) \quad (x,y) \in \Omega \subset \mathbb{R}^2 \quad 0 < t \leq T \\ U = 0 \quad (x,y) \in \partial \Omega \\ U(x,y,0) = U_0(x,y) \\ U_t(x,y,0) = U_1(x,y) \end{array} \right\} \quad (x,y) \in \Omega \quad (3.1)$$

onde $a(x,y)$ é uma função conhecida, tal que :

$$0 < a_{\min} \leq a(x,y) \leq a_{\max} < \infty$$

e $f(x,y,t)$, $U_0(x,y)$ e $U_1(x,y)$ também são funções conhecidas.

A seguir introduziremos alguns espaços de dimensão finita que usaremos nas discretizações posteriores.

O método das direções alternadas consiste basicamente em reduzir um problema multi-dimensional para um conjunto de problemas uni-dimensionais. Para que isso seja possível, devemos fazer uma escolha apropriada de um subespaço $\mathcal{M}$, que será representado por um produto tensorial de subespaços de dimensão finita, o qual utilizaremos no método de Galerkin. Assim sendo, tomemos $\mathcal{M}$ um subespaço de dimensão finita de $H_0^1(\Omega)$ tal que $\partial^2 v / \partial x \partial y \in L^2(\Omega)$ para todo $v \in \mathcal{M}$.

No que segue, estaremos considerando $\Omega = [0,1] \times [0,1]$ e, $\mathcal{M} = \mathcal{M}_x \otimes \mathcal{M}_y$, onde $\mathcal{M}_x$ e $\mathcal{M}_y$ são também subespaços de dimensão finita de $H_0^1([0,1])$, cujas bases são $\{\alpha_p\}_{p=1}^{N_x}$ e $\{\beta_q\}_{q=1}^{N_y}$ respectivamente, ou seja:

$$\mathcal{M}_x = \text{span} (\alpha_1, \dots, \alpha_{N_x}),$$

e

$$\mathcal{M}_y = \text{span} (\beta_1, \dots, \beta_{N_y})$$

portanto,

$$\mathcal{M} = \mathcal{M}_x \otimes \mathcal{M}_y = (\alpha_1 \beta_1, \alpha_1 \beta_2, \dots, \alpha_{N_x} \beta_{N_y}).$$

Observe que $\mathcal{M}$ é obtido através de um produto tensorial em que o primeiro fator estará na direção x enquanto o segundo estará na direção y .

A formulação fraca para esse problema pode ser obtida da maneira usual, isto é, partindo da formulação clássica (3.1), aplicando a identidade de Green e as condições de contorno. Na variável tempo usando diferenças finitas centrais, tem-se :

FORMULAÇÃO FRACA

$$\left\langle \frac{U^{n+1} - 2U^n + U^{n-1}}{(\Delta t)^2}, V \right\rangle + \langle a \nabla U^n, \nabla V \rangle = \langle f^n, V \rangle \quad (3.2)$$

$$\forall V \in \mathcal{M}$$

$$\text{onde: } \langle g, h \rangle = \iint_{\Omega} g \cdot h \, dx \, dy,$$

$$U^n = U(x, y, t_n),$$

$$f^n = f(x, y, t_n) \quad e$$

$$\Delta t = t_{n+1} - t_n \quad \text{para } n = 1, 2, \dots$$

Seguindo as idéias de Douglas introduzidas em [3], ao método das direções alternadas está associada uma perturbação da equação (3.2) através da introdução de dois termos :

$$\begin{aligned}
& \left\langle \frac{z}{(\Delta t)^2}, V \right\rangle + \lambda \langle \nabla z, \nabla V \rangle + \lambda^2 (\Delta t)^2 \left\langle \frac{\partial^2 z}{\partial x \partial y}, \frac{\partial^2 V}{\partial x \partial y} \right\rangle + \\
& \langle a \nabla U^n, \nabla V \rangle = \langle f^n, V \rangle
\end{aligned} \tag{3.3}$$

para $n \geq 1$ e $\forall V \in \mathcal{M}$

onde $z = U^{n+1} - 2U^n + U^{n-1}$

A modificação da equação nos permite fatorar (3.3) num produto tensorial. Como veremos a seguir, o sistema linear associado à discretização poderá então ser decomposto em dois sistemas, um em cada direção.

O escalar λ surgido em (3.3) é um parâmetro requisitado para estabilidade do método, que é obtida quando :

$$\lambda > \frac{1}{4} a_{\max} \tag{3.4}$$

Escrevendo a aproximação da solução a cada nível do tempo, $U^n(x,y)$, em termos da base proposta para $\mathcal{M}$, como :

$$U^n(x,y) = \sum_{ij} U_{ij} \alpha_i(x) \otimes \beta_j(y) \tag{3.5}$$

tomando a equação (3.3) para $V = \alpha_p(x) \beta_q(y)$ e agrupando adequadamente os termos obtidos, podemos reescrever aquela equação agora numa forma decomposta:

$$(C^x + \lambda (\Delta t)^2 A^x) \otimes (C^y + \lambda (\Delta t)^2 A^y) z = (\Delta t)^2 \tilde{f}^n \quad (3.6)$$

$$\text{onde: } (\tilde{f}^n)_{pq} = - \langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle + \langle f^n, \alpha_p \beta_q \rangle \quad (3.7)$$

$$e \quad C_{ij}^x = \int_0^1 \alpha_i(x) \alpha_j(x) dx \quad C_{ij}^y = \int_0^1 \beta_i(y) \beta_j(y) dy$$

(3.8)

$$A_{ij}^x = \int_0^1 \alpha_i'(x) \alpha_j'(x) dx \quad A_{ij}^y = \int_0^1 \beta_i'(y) \beta_j'(y) dy$$

Como podemos observar, a matriz do sistema resultante (3.6) é um produto possuindo um fator em cada direção.

Logo, desde que tenhamos U^0 e $U^1 \in \mathcal{M}$ para iniciar o processo, obtemos $U^2, U^3, \dots, U^n$, através de (3.6), onde (como vimos anteriormente) $z = U^{n+1} - 2U^n + U^{n-1}$, e portanto, teremos a aproximação da solução pelo uso de (3.5).

Se chamarmos $e^n = U(x, y, t_n) - U(x, y)$, escolhermos funções teste convenientemente; manipularmos algebricamente a equação em cada nível de tempo; somarmos as equações correspondentes aos diversos níveis e aplicarmos a desigualdade de Gronwall podemos mostrar (procedimento adotado por Douglas em [3]) que:

$$\begin{aligned} & \| \Delta_t e \|_{L^2 \times L^\infty}^2 + \| e \|_{H_0^1 \times L^\infty}^2 + (\Delta t)^4 \| \frac{\partial^2}{\partial x \partial y} \Delta_t e \|_{L^2 \times L^\infty}^2 \leq \\ & C \left[(\Delta t)^4 + \| e^0 \|_{H_0^1}^2 + \| e^1 \|_{H_0^1}^2 + \| \Delta_t e^0 \|_{L^2}^2 + (\Delta t)^4 \| \frac{\partial^2}{\partial x \partial y} \Delta_t e^0 \|_{L^2}^2 + \right. \\ & \left. \| \partial \delta_t u \|_{H_0^1 \times L^2}^2 + \| \partial \delta_t u \|_{L^2 \times L^\infty}^2 + \| \Delta_t^2 \delta u \|_{L^2 \times L^2}^2 + \right. \\ & \left. (\Delta t)^4 \left\{ \| \frac{\partial^2}{\partial x \partial y} \partial \delta_t u \|_{L^2 \times L^\infty}^2 + \| \frac{\partial^2}{\partial x \partial y} \Delta_t^2 \delta u \|_{L^2 \times L^2}^2 \right\} \right] \end{aligned} \quad (3.9)$$

$$\begin{aligned} & \| \partial \delta_t u \|_{H_0^1 \times L^2}^2 + \| \partial \delta_t u \|_{L^2 \times L^\infty}^2 + \| \Delta_t^2 \delta u \|_{L^2 \times L^2}^2 + \\ & (\Delta t)^4 \left\{ \| \frac{\partial^2}{\partial x \partial y} \partial \delta_t u \|_{L^2 \times L^\infty}^2 + \| \frac{\partial^2}{\partial x \partial y} \Delta_t^2 \delta u \|_{L^2 \times L^2}^2 \right\} \end{aligned}$$

Observamos ainda que a aplicação do lema de Gronwall só é possível com a hipótese fundamental $\lambda > \frac{1}{4} a_{\max}$.

Veremos a seguir que é possível escolher as aproximações iniciais U^0 e U^1 , de tal modo que :

$$\| e^0 \|_{H_0^1}^2 + \| e^1 \|_{H_0^1}^2 + \| \Delta_t e^0 \|_{L^2}^2 + (\Delta t)^4 \| \frac{\partial^2}{\partial x \partial y} \Delta_t e^0 \|_{L^2}^2$$

$$\text{seja da ordem de } (\Delta t)^4 [0 ((\Delta t)^4)]. \quad (3.10)$$

Procederemos de modo a explorar a mesma estrutura de (3.6), o que facilita a implementação computacional no sentido de que poucos acréscimos ao programa serão necessários.

Desse modo, adotamos os procedimentos propostos por Douglas em [3].

No tempo zero, fazemos:

$$(C^x + A^x) \otimes (C^y + A^y) U^0 = \Phi^0 \quad (3.11)$$

onde :

$$(\Phi^0)_{pq} = \langle U_0, \alpha_p \beta_q \rangle + \langle \nabla U_0, \nabla \alpha_p \beta_q \rangle + \langle \frac{\partial^2 U_0}{\partial x \partial y}, \frac{\partial^2}{\partial x \partial y} \alpha_p \beta_q \rangle$$

Com erro na norma L^2 sendo $O((\Delta t)^2)$. Logo, não há incompatibilidade com o erro do procedimento previsto em (3.6).

No tempo 1, fazemos:

$$[C^x \otimes C^y + \Delta t (A^x \otimes C^y + C^x \otimes A^y) + (\Delta t)^2 A^x \otimes A^y] (U^1 - U^0) = \Phi^1 \quad (3.12)$$

onde:

$$(\Phi^1)_{pq}^1 = \langle S, \alpha_p \beta_q \rangle + \langle \nabla S, \nabla \alpha_p \beta_q \rangle + \langle \frac{\partial^2 S}{\partial x \partial y}, \frac{\partial^2}{\partial x \partial y} \alpha_p \beta_q \rangle$$

e

$$S = \Delta t U_1 + \frac{(\Delta t)^2}{2} [\nabla \cdot (a \nabla U_0) + f^0]$$

Observa-se que a expressão (3.12) não pode ser fatorada. Assim, sugere-se um processo iterativo (convergência comprovada em [3]), tal que, para $(U^1 - U^0)$ arbitrário, teremos:

$$\begin{aligned}
 (C^x + \Delta t A^x) \otimes (C^y + \Delta t A^y) (U^1 - U^0)^{q+1} = \\
 (\Delta t)^2 (1 - (\Delta t)^2) A^x \otimes A^y (U^1 - U^0)^q + \Phi^1 \quad (3.13)
 \end{aligned}$$

onde Φ^1 é calculado como proposto em (3.12).

Uma razoável aproximação inicial para a solução do sistema (3.13) é tomá-la igual a zero, o que significa dizer que U^1 é igual a U^0 , portanto, desde que não tenhamos grandes variações, a convergência será atingida em poucas iterações.

Voltando a examinar a estimativa de erro das aproximações através do método de Galerkin-Direções Alternadas obtida em (3.9) verificamos que o espaço de aproximação $\mathcal{M}$ pode então ser escolhido apropriadamente em termos de $h = 1/(\text{Nro. de elementos numa dada direção})$. De fato, a análise da precisão do método de Galerkin está intimamente ligada ao estudo das propriedades de aproximação dos espaços escolhidos $\mathcal{M}$. Muitos matemáticos têm dedicado especial atenção a este problema, e os espaços conhecidos por Splines ou polinômios de Hermite por partes têm sido usados com frequência. As propriedades de aproximações e a possibilidade de construção de bases locais destes espaços os deixam em posição de destaque.

A título de exemplo, consideraremos os espaços dos polinômios de Hermite associados a uma partição Π do intervalo I . Estes espaços podem ser caracterizados por :

$$H^{(m)}(\Pi) = \left\{ v \in C^{m-1}(I) \text{ tal que } v|_{I_j} \text{ é polinômio de grau } 2m-1 \right\} \quad (3.14)$$

Nesta notação, $C^{m-1}(I)$ representa o conjunto das funções com derivadas até a ordem $m-1$ contínuas em I e os I_j são os diversos subintervalos associados à partição Π . Observe que $H^{(1)}(\Pi)$ nada mais é do que o conjunto dos splines lineares gerados pelas funções "chapéu" das quais nos utilizaremos nos nossos experimentos computacionais. Por outro lado, $m = 2$ nos dará as funções de Hermite também amplamente conhecidas e para as quais é possível estabelecer bases locais com ótimas propriedades [11] e [12].

Os espaços $H^{(m)}(\Pi)$ possuem a propriedade de que para qualquer função $u \in H^j(I)$ (espaço de Sobolev de ordem j) $\tilde{u} \in H^{(m)}(\Pi)$ e uma constante C , independente de h e de u tal que :

$$\| D^l(u - \tilde{u}) \|_2 \leq C h^{j-l} \| u \|_j \quad (3.15)$$

para todo $0 \leq l \leq m$, $l \leq j \leq 2m$.

Assim, como exemplo, os Hermite cúbicos ($m = 2$) teremos:

$$l = 0 \quad \| u - \tilde{u} \| \leq C h^j \| u \|_j \quad (3.16)$$

isto é, as aproximações são da $O(h^4)$ para funções $u \in H^4(I)$,

$$l = 1 \quad || D(u - \tilde{u}) || \leq C h^{j-1} || u ||_j \quad (3.17)$$

ou seja, a aproximação da derivada é da $O(h^3)$ para funções de $H^4(I)$,

$$l = 2 \quad || D^2(u - \tilde{u}) || \leq C h^{j-2} || u ||_j \quad (3.18)$$

nesse caso, as aproximações para a derivada segunda são da $O(h^2)$ para funções de $H^4(I)$.

Em resumo, obtivemos um resultado que nos fornece estimativas de erro para as aproximações calculadas por (3.3). De fato, as estimativas mostram que se :

- i) u é solução de (3.2) e U é solução de (3.3)
- ii) u, u_t, u_{tt} são suficientemente regulares
- iii) tomarmos $M = H^{(m)}(\Pi)$

então existe uma constante C independente de h e Δt tal que :

$$|| \Delta_t e ||_{L^2 \times L^\infty}^2 + || e ||_{H_0^1 \times L^\infty}^2 + (\Delta t)^4 || \frac{\partial^2}{\partial x \partial y} \Delta_t e ||_{L^2 \times L^\infty}^2 \leq C \left[(\Delta t)^2 + h^{2m-2} \right] \quad (3.19)$$

Em particular, se $m = 2$ (Hermite cúbico)

$$\|e\|_{H_0^1 \times L^\infty}^2 \leq C \left[(\Delta t)^2 + h^2 \right] \quad (3.20).$$

3.3 - PROCEDIMENTO GERAL

Em resumo, estabelecemos um algoritmo que permitirá o cálculo de soluções aproximadas para a equação (3.1). Neste procedimento, a cada nível de tempo, precisamos resolver os seguintes sistemas lineares :

$$(C^x + \mu A^x) \otimes (C^y + \mu A^y) z = b \quad (3.21)$$

onde:

No tempo zero teremos:

$$\mu = 1$$

$$z = U^0$$

b como previsto em (3.11).

No tempo 1 teremos:

$$\mu = \Delta t$$

$$z = U^1 - U^0$$

b como previsto em (3.13).

Nos demais tempos:

$$\mu = \lambda (\Delta t)^2$$

$$z = U^{n+1} - 2 U^n + U^{n-1}$$

b como previsto em (3.7).

3.4 - NUMERAÇÃO DA MALHA.

Como o problema que estamos tratando (3.1) tem fronteira nula, consideramos apenas os nós internos; a numeração surge de um modo natural pela maneira como resolveremos os sistemas lineares, isto é, procederemos a enumeração de baixo para cima e da esquerda para a direita :

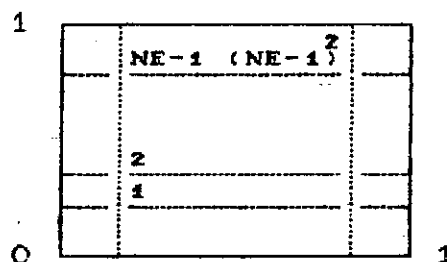


FIG. 3.1

Para efeito de simplificação, como estamos trabalhando em domínios quadrangulares, estaremos considerando uma malha possuindo a mesma quantidade de elementos em ambas as direções. Denominaremos de NE o número de elementos em cada direção.

3.5 - BASE PARA O ELEMENTO FINITO.

Passando agora à fase de cálculos efetivos, precisamos exibir as matrizes A e C de (3.8). Nesta etapa é necessário escolher uma base conveniente para M . Sabemos que no método do elemento finito (M.E.F) o espaço M é gerado por funções que são polinômios por partes com um certo grau de regularidade no domínio. De acordo com o que vimos na seção 3.2, aumentando o grau dos polinômios que compõe a base é possível obter melhores aproximações; por outro lado, as matrizes A e C terão estrutura com grau também crescente de complexidade, assim como o próprio cálculo dos elementos a_{ij} e c_{ij} deve ser mais criterioso. Portanto, no método dos elementos finitos há um forte compromisso custos x benefícios. Do lado dos custos estão a complexidade do programa, a capacidade de memória para executá-lo e o tempo de processamento. Nos benefícios estão principalmente a regularidade e o grau de aproximação das soluções obtidas.

Os resultados teóricos, de estimativas para erros associados ao M.E.F, são desenvolvidos principalmente para equações diferenciais que admitam soluções suficientemente regulares e para subespaços de dimensão finita que permitam um grau de aproximação suficientemente bom.

No nosso trabalho optamos pela simplicidade ao escolhermos a base para M . Tomamos então elementos lineares em cada direção, ou seja, $M_x = M_y = H^{(1)}(\pi)$. A seguir descreveremos rapidamente estes espaços.

Seja $\hat{\Omega}$ o elemento padrão:

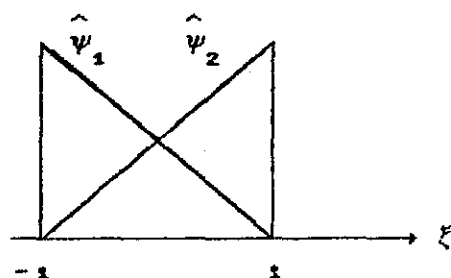


FIG. 3.2

ou seja :

$$\hat{\psi}_1(\xi) = \frac{1}{2} (1 - \xi)$$

$$\hat{\psi}_2(\xi) = \frac{1}{2} (1 + \xi)$$

(3.22)

No elemento genérico Ω^k teremos associados a $\hat{\psi}_1$ e $\hat{\psi}_2$ as funções ψ_1 e ψ_2 .

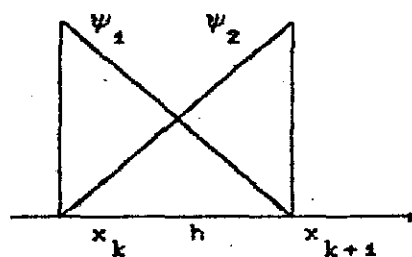


FIG. 3.3

Para escrever as funções ψ_1 e ψ_2 usamos a transformação que leva do elemento padrão no elemento genérico :

$$F(\xi) = \frac{h}{2} \xi + \frac{x_{k+1} + x_k}{2} ; \quad (3.23)$$

assim,

$$F^{-1}(x) = \frac{2(x - x_k)}{h} - 1 \quad (3.24)$$

desse modo, teremos no elemento genérico,

$$\psi_1(x) = \frac{x_{k+1} - x}{h} \quad (3.25)$$

e

$$\psi_2(x) = \frac{x - x_k}{h}$$

Definida a base global; tomando as expressões que definem as matrizes A e C, e calculando as integrais nos elementos genéricos teremos :

$$\begin{aligned}
 C_{ij}^x &= C_{ij}^y = \int_{x_k}^{x_{k+1}} \psi_i(x) \psi_j(x) dx = \int_{y_k}^{y_{k+1}} \psi_i(y) \psi_j(y) dy = \begin{cases} \frac{h}{3} & \text{se } i = j \\ \frac{h}{6} & \text{se } |i-j|=1 \\ 0 & \text{se } |i-j|>1 \end{cases} \\
 A_{ij}^x &= A_{ij}^y = \int_{x_k}^{x_{k+1}} \psi_i'(x) \psi_j'(x) dx = \int_{y_k}^{y_{k+1}} \psi_i'(y) \psi_j'(y) dy = \begin{cases} \frac{1}{h} & \text{se } i = j \\ -\frac{1}{h} & \text{se } |i-j|=1 \\ 0 & \text{se } |i-j|>1 \end{cases}
 \end{aligned}
 \tag{3.26}$$

Ainda na notação usual do método dos elementos finitos, a matriz de rigidez por elemento será :

$$K^e = \left[C^x + \mu A^x \right]^e = \left[C^y + \mu A^y \right]^e = \begin{bmatrix} \frac{h}{3} + \frac{\mu}{h} & \frac{h}{6} - \frac{\mu}{h} \\ \frac{h}{6} - \frac{\mu}{h} & \frac{h}{3} + \frac{\mu}{h} \end{bmatrix} \tag{3.27}$$

2 x 2

onde μ é o parâmetro definido em (3.21).

Procedendo o agrupamento dos elementos (denominado Assembly na literatura especializada), a matriz de rigidez global terá uma estrutura muito simples :

$$K = C^x + \mu A^x = C^y + \mu A^y = \begin{bmatrix} 2*\left(\frac{h}{3} + \frac{\mu}{h}\right) & \frac{h}{6} - \frac{\mu}{h} & & & & & \\ & \frac{h}{6} - \frac{\mu}{h} & . & & & & \\ & & . & . & & & \\ & & & . & . & . & \\ & & & & . & . & \\ & & & & & . & \\ & & & & & & \frac{h}{6} - \frac{\mu}{h} \\ & & & & & & & 2*\left(\frac{h}{3} + \frac{\mu}{h}\right) \end{bmatrix}$$

(3.28)

Em outras palavras, K é uma matriz de dimensão $(NE - 1)$, simétrica e tridiagonal, sendo que todos os elementos de uma mesma diagonal são iguais.; o fator dois surgido na diagonal é devido às contribuições de elementos vizinhos.

3.6 - RESOLUÇÃO DOS SISTEMAS LINEARES.

Como vimos na seção 3.2, procedida a discretização, teremos de resolver um sistema linear que possui a seguinte estrutura:

$$(C^x + \mu A^x) \otimes (C^y + \mu A^y) z = b \quad (3.29)$$

Por definição, o produto tensorial entre duas matrizes $A_{m \times n}$ e $B_{r \times s}$ é:

$$A \otimes B = \begin{bmatrix} a_{11}^B & a_{12}^B & \dots & a_{1n}^B \\ \vdots & \vdots & & \vdots \\ a_{m1}^B & a_{m2}^B & \dots & a_{mn}^B \end{bmatrix} \quad (3.30)$$

observa-se que a matriz resultante tem $m \times r$ linhas e $n \times s$ colunas. Da definição acima algumas propriedades do produto tensorial podem ser estabelecidas [16] :

$$(i) \quad A \otimes B \otimes C = (A \otimes B) \otimes C = A \otimes (B \otimes C)$$

$$(ii) \quad (A + B) \otimes (C + D) = A \otimes C + A \otimes D + B \otimes C + B \otimes D$$

$$(iii) \quad (A \otimes B) (C \otimes D) = (AC) \otimes (BD)$$

Observe que tanto em (ii) quanto em (iii) existem restrições nas dimensões das matrizes envolvidas de modo a permitir as operações adição e multiplicação.

Usando em (3.29) a propriedade (iii) com $B = I$, $C = I$, $A = C^x + \mu A^x$ e $D = C^y + \mu A^y$ verificamos que aquela expressão é equivalente a :

$$[(C^x + \mu A^x) \otimes I] [I \otimes (C^y + \mu A^y)] z = b \quad (3.31)$$

Esta nova maneira de escrever o sistema se torna útil por permitir a decomposição em dois sistemas mais simples :

$$[(C^x + \mu A^x) \otimes I] \gamma = b \quad (3.32)$$

e depois:

$$[I \otimes (C^y + \mu A^y)] z = \gamma \quad (3.33)$$

Chamando $[(C^x + \mu A^x) \otimes I]$ de $\bar{A}$ e $[I \otimes (C^y + \mu A^y)]$ de $\bar{B}$, teremos que resolver

$$\bar{A} \gamma = b \quad (3.34)$$

e

$$\bar{B} z = \gamma \quad (3.35)$$

Olhando para a expressão que define $\bar{A}$ e para (3.28), vemos que :

$$\bar{A} = \begin{bmatrix} A_1 & A_2 & & & & & & & & \\ A_2 & A_1 & A_2 & & & & & & & \\ & A_2 & A_1 & A_2 & & & & & & \\ & & & & \ddots & & & & & \\ & & & & & \ddots & & & & \\ & & & & & & \ddots & & & \\ & & & & & & & \ddots & & \\ & & & & & & & & \ddots & \\ & & & & & & & & & A_2 \\ & & & & & & & & & A_2 & A_1 \end{bmatrix} \quad (3.36)$$

Assim, essa matriz é simétrica e tridiagonal por blocos, cujos blocos (A_1 e A_2) são matrizes diagonais. Nos blocos das diagonais secundárias (A_2), os elementos serão iguais a

$h/6 - \mu/h$, enquanto nos blocos da diagonal principal (A_1), os elementos serão iguais a $2*(h/3 + \mu/h)$.

Procedendo analogamente em relação a $\bar{B}$, verificamos que ela terá a seguinte estrutura :

$$\bar{B} = \begin{bmatrix} B & & & & & \\ & B & & & & \\ & & B & & & \\ & & & \ddots & & \\ & & & & 0 & \\ & & & & & \ddots \\ & & & & & & B & \\ & & & & & & & B \end{bmatrix} \quad (3.37)$$

Portanto, a matriz do sistema linear (3.35), $\bar{B}$, é diagonal por blocos; cada um desses blocos (B) é uma matriz simétrica e tridiagonal, cujos elementos das diagonais secundárias são iguais a $h/6 - \mu/h$ e cujos elementos da diagonal principal são iguais a $2*(h/3 + \mu/h)$.

Como vimos anteriormente, em razão de termos tomados o mesmo número de elementos por direção, $\bar{A}$ e $\bar{B}$ são matrizes de ordem $(NE - 1)^2$, enquanto os blocos A_1 , A_2 e B são de ordem $(NE - 1)$.

Em [4], é apresentada a sugestão de efetuarmos $P\bar{A}P^t$, onde P é uma matriz de permutação de ordem $(NE-1)^2$ de modo a transformar o sistema (3.34), num sistema diagonal por blocos, ou seja, com mesma estrutura de (3.35). Optamos por outro caminho ;

julgamos melhor efetuar $P \bar{A}$ transformando o sistema (3.34) num sistema triangular superior por blocos com a seguinte estrutura:

$$P \bar{A} = \bar{\bar{A}} = \begin{bmatrix} I & A_2 A_1^{-1} & & & 0 \\ & I & A_3^{-1} & & \\ & & I & A_4^{-1} & \\ & 0 & & \ddots & \vdots \\ & & & & A_n^{-1} \\ & & & & & I \end{bmatrix} \quad (3.38)$$

onde n é igual a $NE - 1$

Como citamos anteriormente, as matrizes A_1 e A_2 são diagonais e possuem todos os elementos iguais, portanto a inversa de A_1 (A_1^{-1}) será facilmente calculada; é também uma matriz diagonal com todos os elementos iguais ao inverso multiplicativo dos elementos de A_1 .

Tomando $Q = A_2 A_1^{-1}$, e $M = A_1 A_2^{-1}$ temos então :

$$\begin{aligned} A_3 &= M - Q \\ A_4 &= M - A_3^{-1} \\ A_5 &= M - A_4^{-1} \\ &\vdots \\ A_n &= M - (A_{n-1})^{-1} \end{aligned} \quad (3.39)$$

A matriz de permutação que tomamos, responsável pela triangularização de $\bar{A}$ e que devemos multiplicar pelo vetor b ao resolver o sistema (3.34), é triangular inferior, com a

seguinte estrutura:

$$P = \begin{bmatrix} A_1^{-1} & & & & & & \\ -[A_3 A_1]^{-1} & [A_3 A_2]^{-1} & & & & & \\ [A_4 A_3 A_1]^{-1} & -[A_4 A_3 A_2]^{-1} & [A_4 A_2]^{-1} & & & & \\ \vdots & \vdots & \vdots & \ddots & \vdots & & \\ \vdots & \vdots & \vdots & \vdots & \vdots & \ddots & \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \ddots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ -[A_{n+1} \dots A_3 A_1]^{-1} & \dots & \dots & \dots & \dots & -[A_{n+1} A_n A_2]^{-1} & [A_{n+1} A_2]^{-1} \end{bmatrix} \quad (3.40)$$

Os blocos que ficam nas subdiagonais ímpares têm sinal negativo.

Dessa maneira, o sistema (3.34) ficou transformado em:

$$\bar{A} \gamma = \bar{b}$$

$$\text{Onde : } \begin{cases} \bar{A} = P A \\ e \\ \bar{b} = P b \end{cases} \quad (3.41)$$

Se tomarmos o vetor b com a seguinte estrutura de blocos :

$$b = \begin{bmatrix} B_1 \\ B_2 \\ \vdots \\ \vdots \\ \vdots \\ \vdots \\ B_N \end{bmatrix} \quad (3.42)$$

então $\bar{b}$ será :

$$\bar{b} = P b = \begin{bmatrix} A_1^{-1} B_1 \\ -[A_3 A_1]^{-1} B_1 + [A_3 A_2]^{-1} B_2 \\ \vdots \\ -[A_{n+1} \dots A_3 A_1]^{-1} B_1 + \dots - [A_{n+1} \dots A_3 A_2]^{-1} B_{n-1} + [A_{n+1} A_2]^{-1} B_n \end{bmatrix} \quad (3.43)$$

3.7 - ARMAZENAMENTO.

Do exposto anteriormente, estamos em condições de analisar aquilo que consideramos um ganho na opção pelo método elementos finitos com direções alternadas em contraste com o método elementos finitos na sua formulação tradicional. Como sabemos, na formulação tradicional do método elementos finitos a discretização nos conduz a um sistema de $(NE - 1)^2$ equações, até aqui temos sempre nos referido às matrizes dos sistemas lineares que iremos resolver com a ordem de $(NE-1)^2$, portanto não havendo vantagem quanto a espaço de memória requerido. Na verdade, nosso procedimento será feito de um modo diferente.

De fato, apesar das matrizes dos sistemas (3.34) e (3.35) serem de ordem $(NE-1)^2$ resolveremos sistemas de ordem $NE-1$, constituídos pelos blocos que descrevemos anteriormente, e a

solução será armazenada no próprio vetor do lado direito do sistema.

Assim, para resolver (3.34) precisaremos de espaço para armazenar o vetor b ocupando $(NE-1)^2$ posições de memória, (na verdade $\bar{b}$) e a matriz de permutação que, por ser constituída de blocos que são diagonais com todos os elementos iguais, será armazenada num vetor de $NE-2$ posições. Como citamos anteriormente, a solução será armazenada no próprio vetor $\bar{b}$. Resolvido o sistema proveniente de (3.34), cujos bloco, matrizes dos coeficientes, são da ordem $NE-1$, deveremos resolver os sistemas de (3.35).

Em (3.35) temos todos os blocos iguais (veja 3.37); além disso, observamos que eles são tridiagonais, simétricos e positivos definidos, utilizamos então a decomposição de Cholesky uma única vez a cada intervalo de tempo (de fato somente nos três primeiros tempos, pois nos demais tempos não haverá necessidade de refazer a decomposição uma vez que a matriz dos coeficientes não se alterará). Resolvemos então dois sistemas triangulares obtendo a solução para aquele determinado tempo.

Podemos então verificar a validade da afirmação que fizemos anteriormente de que a numeração dos nós surge de uma maneira natural quando da resolução dos sistemas. De fato, toda vez que resolvermos um bloco de (3.34) estaremos determinando a solução na direção de y , ou seja, os resultados nos nós verticais; como resolveremos de trás para a frente, resolver um bloco de (3.35) significa retroceder um nó na direção de x .

Em termos de armazenamento para resolver (3.35), seguimos a sugestão de [11], em que, dado um sistema linear cuja matriz dos coeficientes é de ordem N , simétrica e de banda m , devemos armazená-la numa matriz retangular de ordem $[(m+1) \times N]$, de modo que a diagonal fique na primeira coluna e as subdiagonais (superiores ou inferiores) nas colunas posteriores. Esse tipo de procedimento armazena $(m+m^2)/2$ zeros, o que no nosso caso não será muito. Como nossos blocos são tridiagonais ($m=1$) armazenamos apenas a diagonal e uma subdiagonal, ou seja, guardamos numa matriz cuja dimensão é $[(NE-1) \times 2]$; armazenamos portanto um único zero. Usamos este procedimento tanto na hora da decomposição quanto na hora de resolver os sistemas triangulares.

3.8 - PROBLEMA COM CONDIÇÃO DE NEUMANN NA FRONTEIRA

No capítulo anterior vimos que temos de resolver dois problemas da equação da onda : um com condição de fronteira do tipo Dirichlet nulo e outro com condição de Neumann homogêneo na fronteira.

Até aqui, nos detivemos no problema com condições de Dirichlet, que por ter a solução valor nulo na fronteira, não numeramos os nós do bordo da malha.

Ao resolvermos o problema com condição de Neumann homogênea na fronteira, devemos estar atentos ao fato de que a

solução não será nula na fronteira, mas sim a sua derivada normal exterior. Portanto teremos de acrescentar, de alguma forma, os nós da fronteira.

A formulação fraca para esse problema é a mesma de (3.2); no entanto algumas alterações surgem nos passos seguintes. Em primeiro lugar observamos que o espaço de aproximação, nesse problema não mais poderá ser H_0^1 , pois a solução não será nula na fronteira; tomaremos o espaço de aproximação H^1 .

Nosso raciocínio foi no sentido de não acrescentarmos muitas alterações no programa de computador desenvolvido para resolver o problema anterior. Assim, sugerimos o procedimento a seguir.

Simulamos uma ampliação do domínio, de modo que nesse novo domínio a solução assuma valores nulos na fronteira. Observamos que esta ampliação é equivalente a usar H^1 no domínio não ampliado, quando, como é o nosso caso, a base para o H^1 for composta pelos splines lineares. Em termos práticos este procedimento corresponde a aumentar em dois o número de nós em cada direção. Nesta nova malha podemos utilizar o mesmo programa desde que em algumas partes do mesmo fique especificada a identificação de qual dos dois problemas estará tratando naquele instante.

Dessa forma, nosso domínio foi ampliado para :

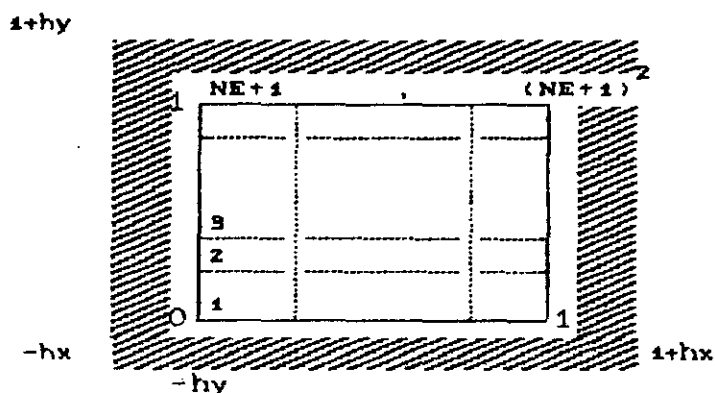


FIG. 3.4

Deveremos estar atentos ao efetuar os cálculos do lado direito dos sistemas lineares. No cálculo das integrais, deve-se identificar quando estaremos na região da expansão da fronteira (área hachureada); nesse caso, a contribuição será nula.

4.1 - INTRODUÇÃO.

No capítulo anterior fornecemos informações que nos possibilitam elaborar um programa que implementado em um computador operacionalize os procedimentos ali prescritos, ou seja, que nos dê condições de efetivamente resolver as equações (2.26) e (2.27), descritas no capítulo II. Neste capítulo descreveremos a maneira como realizamos esta etapa de programação do esquema de cálculo proposto.

Nosso programa será desenvolvida na linguagem FORTRAN, e foi implementado inicialmente em microcomputadores pessoais do tipo IBM-PC. Com a finalidade de obtermos soluções no menor tempo possível, assim como pela possibilidade de avaliação de malhas com um maior número de elementos por direção, fizemos posterior implementação no computador VAX/VMS Versão 4.5 da UNICAMP, e os resultados aqui apresentados foram obtidos nesse tipo de equipamento.

Inicialmente descreveremos o algoritmo utilizado para determinar o vetor deslocamento em (2.1) tomando por base as decomposições propostas em (2.3) e (2.47). Apresentaremos o fluxograma com os procedimentos adotados ao resolver (2.26) ou (2.27).

Quando resolvemos numericamente uma equação cuja solução é conhecida, podemos então analisar o comportamento da aproximação pela medida do erro em relação à solução exata. Neste capítulo apresentaremos alguns exemplos de equações em que a solução é conhecida; forneceremos medidas para os erros cometidos para estes

exemplos.

Por fim apresentamos resultados da equação das ondas acústicas (quando não conhecemos a solução exata). Exibiremos ainda gráficos para as ondas longitudinal e transversal .

4.2 - ALGORITMO.

Como vimos no capítulo dois, deveremos resolver duas equações uma para a onda longitudinal e outra para a onda transversal; a primeira tem condição de fronteira do tipo Neumann homogêneo e a segunda tem condição do tipo Dirichlet nulo no bordo. De posse dos potenciais obtidos das resoluções das equações supra-mencionadas, podemos obter o vetor deslocamento pela diferenciação citada em (2.11).

Apresentamos a seguir o algoritmo utilizado nesse nosso trabalho.

1. Entrada dos dados. Deve-se fornecer os seguintes dados :

1.1 - Número de elementos por direção ?

1.2 - Número de intervalos de tempo ?

1.3 - Valor de λ da perturbação do método de Galerkin com direções alternadas ?

2. Resolução da equação da onda acústica (2.26).

- 2.1 - Identifica o problema.
- 2.2 - Questiona se tem solução exata.
- 2.3 - Solução em toda a malha ?
- 2.4 - Aproximação para o processo iterativo surgido no tempo 1 ?

3. Resolução da equação da onda acústica (2.27).

- 3.1 - Identifica o problema.
- 3.2 - Questiona se tem solução exata.
- 3.3 - Solução em toda a malha ?
- 3.4 - Aproximação para o processo iterativo surgido no tempo 1 ?

4.3 - FLUXOGRAMA PARA ONDAS ACÚSTICAS.

Apresentamos a seguir o fluxograma com os procedimentos utilizados quando da resolução das equações das ondas acústicas (2.26) e (2.27).

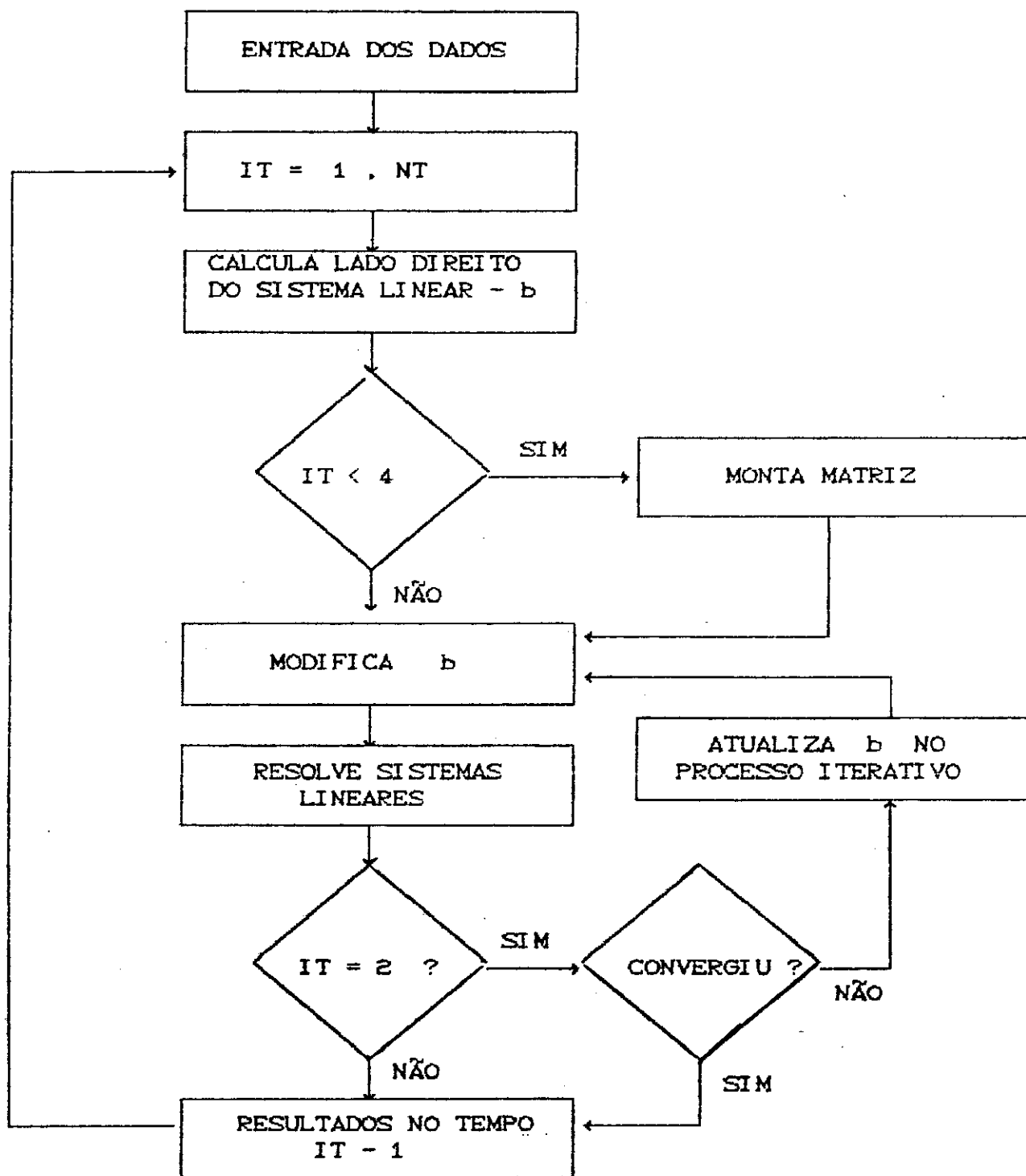


FIG. 4.1

IT = Indicador do nível de tempo.

NT = Número de intervalos de Tempo.

4.4 - O PROGRAMA DE COMPUTADOR.

Passamos agora a descrição do programa computacional que efetiva os procedimentos previstos na figura 4.1, e que tem a seguinte estrutura:

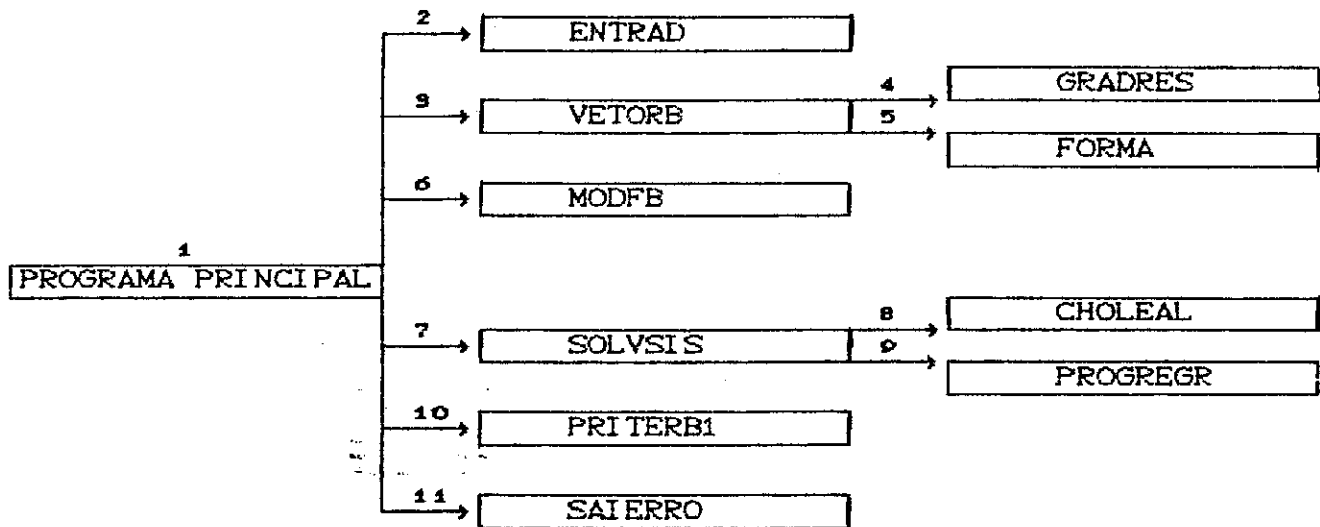


FIG. 4.2

4.4.1 - VARIÁVEIS COMUNS.

Nesses nossos programas temos os seguintes blocos comuns :

BL1/NE, NN, NI, ND, NT

e

BL2/H, DT, XLAMB, AP, P1, P2

onde :

NE → Número de elementos por direção

$NN \rightarrow NE - 1$ para condição de contorno do tipo Dirichlet nulo no bordo; se condição de contorno for do tipo Neumann homogêneo $NN = NE + 1$.

$NI \rightarrow NE - 2$ para condição de contorno do tipo Dirichlet nulo no bordo, se condição de contorno for do tipo Neumann homogêneo $NI = NE$.

$ND \rightarrow (NN)^2$

$NT \rightarrow$ Número de intervalos de tempo

$H \rightarrow$ Comprimento do sub-intervalo no espaço

$DT \rightarrow$ Comprimento do sub-intervalo no tempo

$XLAMB \rightarrow$ Parâmetro requisitado para estabilidade λ

$AP \rightarrow$ Velocidade do meio

$P1$ e $P2 \rightarrow$ Nós gaussianos num elemento genérico

4.4.2 - PROGRAMA PRINCIPAL.

Nesse programa temos a destacar :

AA e BC que servirão para compor as matrizes. Observe que isso somente é feito até o terceiro nível de tempo, uma vez que nos demais ela não se alterará.

AK e $EMAX$ erro relativo e erro relativo máximo ocorridos no processo iterativo do tempo 1.

EPS é o critério de parada para o processo iterativo.

SUB-ROTINA PRITERB1 calculará a primeira parcela do processo iterativo ocorrido no tempo 1 (V1). É de se notar que se ITER (contador das iterações neste processo) for igual a zero não chamaremos esta sub-rotina, isto ocorre porque (como citamos no capítulo III) a aproximação inicial é nula.

Vetor B guarda o lado direito do sistema linear e também a solução.

Vetor RES1 guarda a solução de duas iterações anteriores.

Vetor RES2 guarda a solução da iteração imediatamente anterior.

```

C*****
C  RESOLVE EQ. DA ONDA BIDIMENSIONAL PELO METODO DAS DIRECOES *
C  ALTERNADAS COM CONDICAO DE FRONTEIRA TIPO DIRICHLET OU      *
C  NEUMANN HOMOGENEOS.                                         *
C*****
C  PARAMETER (MZ=100)
C  PARAMETER (MM=MZ*MZ)
C  REAL*8 BC(MM),H,XX(MZ),YY(MZ),DT,XT(MZ),XLAMB,AA,BC
C  REAL*8 EPS,V1(MM),PC(MZ),UMI,P1,P2
C  REAL*8 RES1(MM),RES2(MM),AP,EMAX,B1(MM),B2(MM),AC(MZ,2)
C  COMMON/BL1/NE,NN,NI,ND,NT
C  COMMON/BL2/H,DT,XLAMB,AP,P1,P2
C
C  **** ENTRADA DOS DADOS DO PROBLEMA ****
C
C  CALL ENTRADCXX,YY,XT,IND,ISE,ISSD
C
C  **** LACO NO TEMPO ****
C
C  DO 40 IT=1,NT
C
C  **** COMPOE O LADO DIREITO ****
C

```

```
CALL VETORBC(XX,YY,XT,IT,RES2,B,IND)
```

```
NO TEMPO 1(IT=2) B1(I) E'2A. PARCELA CALCULADA PELA INTEGRAL  
B2(I) SOL. NO PASSO ANTERIOR E A APROXIMACAO INICIAL E'NULA
```

```
DO 5 I=1,ND  
  IF(IT.EQ.2)THEN  
    B2(I)=0.DO  
    B1(I)=B(I)  
  ELSEIF(IT.GT.2)THEN  
    B(I)=B(I)*DT*DT  
  ENDIF  
CONTINUE
```

```
**** MONTAGEM DAS MATRIZES DO SISTEMA ****
```

```
IF(IT.LT.4)THEN  
  IF(IT.EQ.1) THEN  
    UMI=1.DO  
  ELSEIF(IT.EQ.2)THEN  
    ITER=0  
    WRITE(*,*)'EPSLON P/ PARADA'  
    READ(*,*)EPS  
    UMI=DT  
  ELSE  
    UMI=XLAMB*DT*DT  
  ENDIF  
  AA=(H/3.DO+UMI/H)*2.DO  
  BC=H/6.DO-UMI/H  
ENDIF
```

```
**** MODIFICA O LADO DIREITO ****
```

```
CALL MODFBC(AA,BC,P,B)
```

```
**** RESOLVE OS SISTEMAS ****
```

```
CALL SOLVSISC(IT,A,AA,BC,P,B)
```

```
IF(IT.EQ.2)THEN
```

```
**** PROCESSO ITERATIVO NO TEMPO 1 ****
```

```
IF(ITER.EQ.0)GOTO 21
```

```
EMAX=0.DO
```

```
IF(ITER.GT.0)THEN
```

```
DO 12 J=1,ND
```

```
  AK=ABS((B2(J)-B(J))/B(J))
```

```
  IF(AK.GT.EMAX)EMAX=AK
```

```
ENDIF
```

```

21      IF(EMAX.GT.EPS.OR.ITER.EQ.0) THEN
C
C      ****      CALCULA O PRIMEIRO TERMO DO B1 (V1)      ****
C
C      CALL PRITERB1(B,V1)
C
C      ****      ATUALIZA PROCESSO ITERATIVO      ****
C
C      DO 14      I=1,ND
C      B2(I)=B(I)
14      B(I)=V1(I)+B1(I)
C      ITER=ITER+1
C      ENDIF
C      IF(EMAX.GT.EPS.OR.ITER.EQ.1)GOTO 11
C      WRITE(*,*)'NA ITERACAO ',ITER,' ERRO MAX. REL. ',EMAX
C      ENDIF
C
C      ****      SAIDA DOS RESULTADOS E DO ERRO      ****
C
C      CALL SAIERRO(B,XX,YY,XT,IT,RES1,RES2,ISS,ISE)
C
C      ****      ATUALIZA RESULTADOS      ****
C
C      DO 15      IP=1,ND
C      RES1(IP)=RES2(IP)
15      RES2(IP)=B(IP)
40      CONTINUE
C      END

```

4.4.3 - ENTRADA DOS DADOS DO PROBLEMA.

Consideramos essa sub-rotina auto explicativa; ela nada mais contém do que os dados de entrada, os cálculos das partições assim como das coordenadas no espaço e no tempo e os nós gaussianos. Aqui também é feita a identificação de que problemas estaremos tratando, se com condição de fronteira do tipo Neumann ou Dirichlet homogêneos. Acreditamos que os próprios comentários sejam suficientes para esclarecer qualquer dúvida.

C ENTRADA DOS DADOS *

C*****

SUBROUTINE ENTRADXX,YY,XT,IND,ISE,ISS)

PARAMETER (MZ=100)

REAL*8 XI,XF,H,F,X,Y,XX(MZ),YY(MZ),XT(MZ),DT,XLAMB,AP,PR,P1,P2

COMMON/BL1/NE,NN,NI,ND,NT

COMMON/BL2/H,DT,XLAMB,AP,P1,P2

C

XI=0.DO

XF=1.DO

WRITE(*,*)' NRO. DE ELEMENTOS EM X E Y ,(DEVEM SER IGUAIS)'

READ(*,*)NE

WRITE(65,*)' Nro. DE ELEMENTOS. ',NE

WRITE(*,*)' NRO DE TEMPOS'

READ(*,*)NT

WRITE(65,*)' Nro. DE TEMPOS',NT

DT=1.DO/DBLE((NT-1))

WRITE(*,*)' VALOR DA VELOCIDADE <A>'

READ(*,*)AP

WRITE(65,*)' VALOR DA VELOC. <A>',AP

WRITE(*,*)' LAMBDA'

READ(*,*)XLAMB

WRITE(65,*)' LAMBDA',XLAMB

H=(XF-XI)/DBLE(NE)

WRITE(*,*)' C. DE CONTOURNO TIPO NEUMANN ? (1) OU DIRICHLET'

READ(*,*)IND

WRITE(*,*)' POSSUI SOL. EXATA SIM=1'

READ(*,*)ISE

IF(ISE.EQ.1)THEN

WRITE(*,*)' RESULTADO EM TODA MALHA (S=1)'

READ(*,*)ISS

ELSE

ISS=1

ENDIF

C

C .CONDICAO DE NEUMAN NN=NE+1 E NI=NE,SE DIRICHLET NN=NE-1 E NI=NE-2

C NEUMANN O DOMINIO E' AMPLIADO E LACO E' MAIOR

C

IF(IND.EQ.1)THEN

NN=NE+1

NI=NE

XX(1)=XI-H

YY(1)=XI-H

NLACO=NE+3

ELSE

NN=NE-1

NI=NE-2

XX(1)=XI

YY(1)=XI

```

      NLACO=NE+1
      ENDIF
      ND=NN*NN
C
C      DIRICHLET O LACO EH ATEH NE+1 DE NEUMAN EH ATEH NE+3
C
      DO 10 I=2,NLACO
        J=I-1
        XX(I)=XX(J)+H
10     YY(I)=XX(I)
        XTC(I)=0.DO
        DO 20 J=2,NT+1
          I=J-1
20     XTC(J)=XTC(I)+DT
C      NOS GAUSSIANS
      PR=1.DO/SQRT(3.DO)
      P1=.5DO*H*(1.DO-PR)
      P2=.5DO*H*(1.DO+PR)
      RETURN
      END

```

4.4.4 - CÁLCULO DAS INTEGRAIS.

Ao calcularmos a matriz do sistema linear resultante da discretização, não tivemos muito trabalho e os cálculos envolvidos foram razoavelmente simples e feitos precisamente em função de h , λ e Δt . Agora precisamos calcular as integrais para obtenção do lado direito do sistema linear.

Para o cálculo, fizemos o percurso no sentido de baixo para cima e da esquerda para a direita, como é feita a própria enumeração das nossas malhas.

No tempo zero e no tempo um, nossos cálculos foram feitos exclusivamente usando Quadratura Gaussiana com dois nós em cada direção por elemento.

Para exemplo, suponhamos que estejamos no nó pq , onde $p = 1, \dots, NN$ e $q = 1, \dots, NN$, ou seja, que queiramos calcular

o elemento $(p-1) \times (q-1)$ do vetor do lado direito do sistema linear. Tomemos então a figura abaixo para melhor visualização

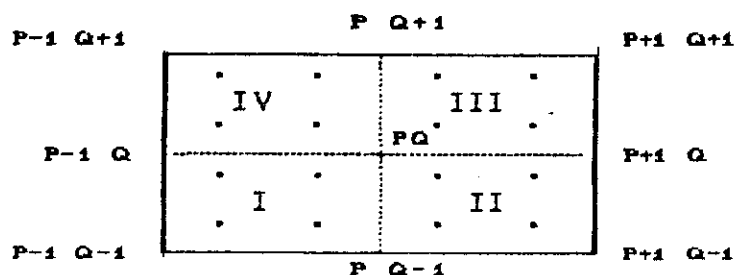


FIG. 4.9

Estando no nó pq, começaremos identificando os 4 nós gaussianos na direção de x e na direção de y, partindo do nó p-1, q-1. Nosso caminho será no sentido anti-horário percorrendo os quadrantes I, II, III, e IV, daí o laço que é dado em I de 1 até 4.

Como está comentado, Kx e Ky servirão para identificar quais nós gaussianos, XG e YG, devem ser tomados num determinado quadrante.

Assim procedendo teremos a integral num determinado quadrante I igual a :

$$\langle H, V \rangle = \iint_I H \cdot V \, dx \, dy = \frac{h^2}{4} \sum_{i=1}^2 \sum_{j=1}^2 H(XG(i), YG(j)) \cdot V(XG(i), YG(j)) \quad (4.1)$$

Tomando V para as funções de forma, efetuamos os cálculos de :

No tempo zero

$$\langle U_0, \alpha_p \beta_q \rangle, \langle \nabla U_0, \nabla \alpha_p \beta_q \rangle \text{ e } \langle \frac{\partial^2 U_0}{\partial x \partial y}, \frac{\partial^2}{\partial x \partial y} \alpha_p \beta_q \rangle \quad (4.2)$$

onde U_0 é a condição inicial da solução.

No tempo 1

$$\langle S, \alpha_p \beta_q \rangle, \langle \nabla S, \nabla \alpha_p \beta_q \rangle \text{ e } \langle \frac{\partial^2 S}{\partial x \partial y}, \frac{\partial^2}{\partial x \partial y} \alpha_p \beta_q \rangle \quad (4.3)$$

$$\text{onde } S = \Delta t U_1 + \frac{(\Delta t)^2}{2} [\nabla \cdot (a \nabla U_0) + f^0]$$

e U_1 é a condição inicial da derivada com relação ao tempo.

Nos demais tempos, teremos que calcular as seguintes integrais :

$$- \langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle + \langle f^n, \alpha_p \beta_q \rangle ; \quad (4.4)$$

observamos que o cálculo da segunda integral nestes tempos é feita de modo similar às que tínhamos nos casos anteriores, o mesmo não se pode dizer em relação a primeira delas, uma vez que U^n é a solução no tempo imediatamente anterior portanto um resultado numérico e não uma expressão algébrica.

Uma saída para a situação criada acima seria fazer uma interpolação bidimensional de modo a obter uma aproximação para o gradiente da solução no tempo anterior nos nós gaussianos, e

assim utilizar quadratura. No entanto, com receio de comprometer a precisão optamos por desenvolver o produto interno analiticamente em termos das funções de forma e assim obter expressões para a primeira integral de (4.4); portanto, não usamos quadratura gaussiana no cálculo dessa integral e o resultado será fornecido pela sub-rotina GRADRES.

ICON é um contador que será utilizado quando o problema que estivermos tratando possuir condição de fronteira do tipo Neumann homogênea. No final do capítulo anterior chamamos a atenção de que a região de ampliação fictícia da fronteira não deve contribuir no cálculo das integrais. Quando estivermos nos nós da fronteira devemos saber em que elementos devemos efetuar os cálculos para compor o lado direito do sistema linear, isso ocorrerá, segundo esta estrutura de dados que utilizamos, quando ICON for igual a zero; nos elementos dos cantos ICON será igual a 2 e nos outros elementos marginais ICON será igual a 1, nessas duas situações não efetuamos os cálculos.

```

C*****
C   CALCULA AS INTEGRAIS   *
C*****
      SUBROUTINE VETORBC(XX,YY,XT,IT,RES2,B,IND)
      PARAMETER (MZ=100)
      PARAMETER (MM=MZ*MZ)
      REAL*8 S1,S2,X1,X2,Y1,Y2,H,XG(4),YG(4),PP,PQ,F,FI(4)
      REAL*8 XX(MZ),YY(MZ),S3,S4,G(2),DXF,DYF,Font1,XLAMB,DT
      REAL*8 DFI(2,4),D2FI(4),RES2(MM),XT(MZ),BC(MM),RR,ARR(4),AP
      REAL*8 DDF,D3,D4,DX11,DY11,DD11,F11,P1,P2
      COMMON/BL1/NE,NN,NI,ND,NT
      COMMON/BL2/H,DT,XLAMB,AP,P1,P2
C
      DO 50 IM=1,NN
        X1=XX(IM)
        X2=XX(IM+1)
        XG(1)=X1+P1

```

```

XG(2)=X1+P2
XG(3)=X2+P1
XG(4)=X2+P2
DO 50 IN=1,NN
  JJ=NN*(IM-1)+IN

```

C
C
C

K E' UTILIZADO NA SUBROTINA GRADRES

```

K=JJ-NN-1
Y1=YY(IN)
Y2=YY(IN+1)
YG(1)=Y1+P1
YG(2)=Y1+P2
YG(3)=Y2+P1
YG(4)=Y2+P2
S1=0. DO
DO 20 I=1,4

```

C
C
C

ICON DIFERENTE DE ZERO IDENT. QUADRANTE FORA DO DOM. REAL

ICON=0

C
C
C
C
C

KX E KY SERVEM P/ SABER QUAIS NOS GAUSSIANOS DEVEM SER
TOMADOS E P/ INDICAR SE OS MESMOS ESTAO FORA DO DOMINIO
REAL QUANDO USAMOS CONDICAO DE NEUMANN

```

IFC(1.EQ.1.OR.1.EQ.4)THEN
  KX=0
ELSE
  KX=2
ENDIF
IFC(1.LT.3)THEN
  KY=0
ELSE
  KY=2
ENDIF

```

C
C
C

COM C. DE NEUMANN VERIFICA SE ESTA' FORA DO DOMINIO REAL

```

IFC(IND.EQ.1)THEN
IFC(IN.EQ.1.OR.IM.EQ.1.OR.IN.EQ.NN.OR.IM.EQ.NN)THEN
  KW=KX+1
  KZ=KY+1
  IF(XG(KW).LT.0.OR.XG(KW).GT.1)ICON=ICON+1
  IF(YG(KZ).LT.0.OR.YG(KZ).GT.1)ICON=ICON+1
ENDIF
IF(ICON.GT.0)GOTO 20
ENDIF

```

C
C
C

S2 PRIMEIRA INTEGRAL S3 A SEGUNDA E S4 A TERCEIRA

```

S2=0. D0
S3=0. D0
S4=0. D0

```

```

C
C
C

```

```

CHAMA ROTINA QUE IDENTIFICA SOL. NO TEMPO ANTERIOR

```

```

IFCIT.LT.3)GOTO 5

```

```

C
C
S

```

```

CALL GRADRES(RES2,I,JJ,K,IN,RR,ARR)

```

```

DO 10 JX=1,2
DO 10 JY=1,2
PP=XG(KX+JX)
PQ=YG(KY+JY)

```

```

C
C
C

```

```

CHAMA AS FUNCOES DE FORMA

```

```

CALL FORMA1(I,IT,IN,IM,XX,YY,PP,PQ,FI,DFI,D2FI)
IFCIT.EQ.1)THEN
GC1)=DXFCPP,PQ)
GC2)=DYFCPP,PQ)
D3=FCPP,PQ)
D4=DDFCPP,PQ)
ELSE IFCIT.EQ.2)THEN
GC1)=DX11(CPP,PQ,AP,DT)
GC2)=DY11(CPP,PQ,AP,DT)
D3=F11(CPP,PQ,AP,DT)
D4=DD11(CPP,PQ,AP,DT)

```

```

ELSE
D3=FONT1(CPP,PQ,XT(IT-1))
ENDIF
IFCIT.LT.3)THEN
S3=S3+GC1)*DFI(1,I)+GC2)*DFI(2,I)
S4=S4+D4*D2FI(I)
ENDIF

```

```

10

```

```

S2=S2+D3*FI(I)
IFCIT.EQ.1)THEN
S1=S1+.25D0*S2+.25D0*S3+.25D0*S4
ELSE IFCIT.EQ.2)THEN
S1=S1+.25D0*S2+.25D0*S3*DT+.25D0*S4*DT*DT*DT*DT
ELSE
S1=S1+.25D0*S2-AP*RR
ENDIF

```

```

20

```

```

CONTINUE

```

```

50

```

```

B(JJ)=S1
CONTINUE
RETURN
END

```

4.4.5 - FUNÇÕES DE FORMA

Como comentamos no capítulo III, adotamos $H^{(1)}(\Pi)$ para nossos espaços de aproximação $\mathcal{M}_x$ e $\mathcal{M}_y$. Assim, no elemento genérico, teremos para ψ_1 e ψ_2 .

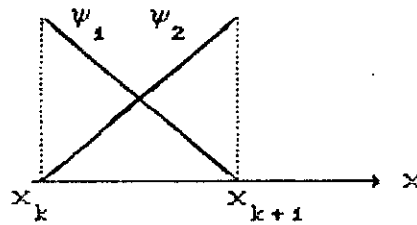


FIG. 4.4

$$\psi_1(x) = \frac{x_{k+1} - x}{h}$$

e

$$\psi_2(x) = \frac{x - x_k}{h}$$

Olhando para a figura 4.3 quando da colocação das funções de forma, teremos a seguinte perspectiva:

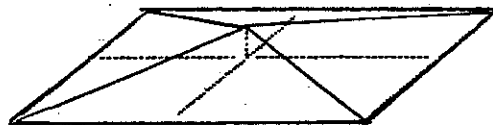


FIG. 4.5

Assim teremos :

No Primeiro Quadrante :

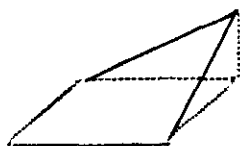


FIG. 4.6

$$FI(1) = \psi_2(x) \psi_2(y) = \frac{1}{h^2} (x - x_k)(y - y_k) \quad (4.5)$$

No Segundo quadrante :

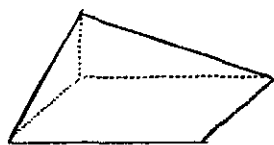


FIG. 4.7

$$FI(2) = \psi_1(x) \psi_2(y) = \frac{1}{h^2} (x_{k+1} - x)(y - y_k) \quad (4.6)$$

No Terceiro quadrante :

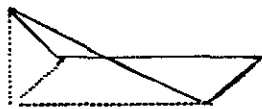


FIG. 4.8

$$FI(3) = \psi_1(x) \psi_1(y) = \frac{1}{h^2} (x_{k+1} - x)(y_{k+1} - y) \quad (4.7)$$

No Quarto Quadrante :

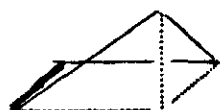


FIG. 4.9

$$FI(4) = \psi_2(x) \psi_1(y) = \frac{1}{h^2} (x - x_k)(y_{k+1} - y) \quad (4.8)$$

A, B, C e D servem exatamente para identificar os nós anteriores e posteriores ao elemento pq da figura 4.3.

I informa em que quadrante estamos, o desvio será para 10, 20, 30, ou 40 , para I igual a 1, 2, 3 ou 4 respectivamente, aí os cálculos são feitos somente naquele quadrante e retorna imediatamente para a sub-rotina VETORB.

Somente nas 2 primeiras iterações é que se calcula DFI E D2FI.

```

C*****
C  FUNCOES DE FORMA  *
C*****
      SUBROUTINE FORMA1(I,IT,IN,IM,XX,YY,S,T,FI,DFI,D2FI)
      PARAMETER (MZ=100)
      REAL*8 S,T,A,B,C,D,FI(4),DFI(2,4),XX(MZ),YY(MZ),D2FI(4)

C
      A=XX(IM)
      B=XX(IM+2)
      C=YY(IN)
      D=YY(IN+2)

C
C  FI=FUNCOES DE FORMA  DFI=DERIVADA  D2FI=DERIV. SEGUNDA
C
      GOTO(10,20,30,40)I
10    FI(1)=(S-A)*(T-C)
      IF(IT.LT.3)THEN
          DFI(1,1)=(T-C)
          DFI(2,1)=(S-A)
          D2FI(1)=1.D0
      ENDIF
      RETURN
20    FI(2)=(-S+B)*(T-C)
      IF(IT.LT.3)THEN
          DFI(1,2)=-(T-C)
          DFI(2,2)=(-S+B)
          D2FI(2)=-1.D0
      ENDIF
      RETURN
30    FI(3)=(-S+B)*(-T+D)
      IF(IT.LT.3)THEN
          DFI(1,3)=-(-T+D)
          DFI(2,3)=(-S+B)
          D2FI(3)=1.D0
      ENDIF
      RETURN
40    FI(4)=(S-A)*(-T+D)
      IF(IT.LT.3)THEN

```

```

      DFI(1,4)=(-T+D)
      DFI(2,4)=-(-S-A)
      D2FI(4)=-1.D0
    ENDIF
    RETURN
C
C   IF PORQUE SO NOS DOIS PRIMEIROS TEMPOS CALCULA-SE DFI E D2FI.
C
    END

```

4.4.6 - INTEGRAL ENVOLVENDO GRADIENTE DA SOLUÇÃO NA ITERAÇÃO ANTERIOR.

Como citamos na seção anterior, nas iterações maiores que dois, há uma integração que, em troca da precisão, não utilizávamos quadratura gaussiana; isso ocorre quando temos que calcular : $\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle$.

O que fizemos foi desenvolver esse produto interno em termos de :

$$U^n = \sum_{i,j} U_{ij} \alpha_i(x) \beta_j(y) \quad (4.9)$$

Dessa forma, teremos :

$$\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle = a \left\{ \sum_j U_{ij} \left[\alpha'_i \alpha'_p \beta_j \beta_p + \alpha_i \alpha_p \beta'_j \beta'_p \right] \right\} \quad (4.10)$$

onde os índices mais abaixo (x e y) evidenciam em relação a que variável é feita a diferenciação.

Ainda tomando como referência a figura 4.3 teremos :

No Primeiro Quadrante :

$$U^n = U_{p-1 \ q-1} \alpha_1 \beta_1 + U_{p \ q-1} \alpha_2 \beta_1 + U_{pq} \alpha_2 \beta_2 + U_{p-1 \ q} \alpha_1 \beta_2 \quad (4.11)$$

efetuando os devidos cálculos encontramos :

$$\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle = a \left[-\frac{1}{3} U_{p-1 \ q-1} - \frac{1}{6} U_{p \ q-1} + \frac{2}{3} U_{pq} - \frac{1}{6} U_{p-1 \ q} \right] \quad (4.12)$$

No Segundo Quadrante :

$$U^n = U_{p \ q-1} \alpha_1 \beta_1 + U_{p+1 \ q-1} \alpha_2 \beta_1 + U_{p+1 \ q} \alpha_2 \beta_2 + U_{p \ q} \alpha_1 \beta_2 \quad (4.13)$$

e a integral :

$$\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle = a \left[-\frac{1}{6} U_{p \ q-1} - \frac{1}{3} U_{p+1 \ q-1} - \frac{1}{6} U_{p+1 \ q} + \frac{2}{3} U_{pq} \right] \quad (4.14)$$

No Terceiro Quadrante :

$$U^n = U_{p \ q} \alpha_1 \beta_1 + U_{p+1 \ q} \alpha_2 \beta_1 + U_{p+1 \ q+1} \alpha_2 \beta_2 + U_{p \ q+1} \alpha_1 \beta_2 \quad (4.15)$$

e a integral :

$$\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle = a \left[\frac{2}{3} U_{pq} - \frac{1}{6} U_{p+1 \ q} - \frac{1}{3} U_{p+1 \ q+1} - \frac{1}{6} U_{p \ q+1} \right] \quad (4.16)$$

No Quarto Quadrante :

$$U^n = U_{p-1 \ q-1} \alpha_1 \beta_1 + U_{p+1 \ q-1} \alpha_2 \beta_1 + U_{pq} \alpha_2 \beta_2 + U_{p-1 \ q} \alpha_1 \beta_2 \quad (4.17)$$

e a integral :

$$\langle a \nabla U^n, \nabla \alpha_p \beta_q \rangle = a \left[-\frac{1}{6} U_{p-1 \ q-1} + \frac{2}{3} U_{p+1 \ q-1} - \frac{1}{6} U_{pq} - \frac{1}{3} U_{p-1 \ q} \right] \quad (4.18)$$

Observe que os valores $-\frac{1}{6}$, $\frac{2}{3}$, $-\frac{1}{6}$ e $\frac{1}{3}$ repetem-se alternadamente; o vetor ARR guarda esses valores e utilizamos uma estrutura de dados que toma as seqüências prescritas acima, sendo que o vetor IB será o apontador das mudanças que devem ocorrer de um quadrante para o outro.

PS é um vetor que estando num quadrante localiza os quatro valores da solução no tempo anterior.

AQ é um vetor auxiliar que guarda temporariamente o valor de ARR.

RR é o valor da integral no quadrante.

```
C*****
C   PROGRAMA PARA CALCULAR A INTEGRAL DO PRODUTO      *
C   ENTRE O GRAD. DA SOLUCAO NA ITERACAO ANTERIOR     *
C   E O GRADIENTE DAS FUNCOES DE FORMA              *
C*****
C   SUBROUTINE GRADRES(CRES2,I,JJ,K,IN,RR,ARR)
C   PARAMETER (MM=10000)
C   REAL*8 RES2(MM,PS(4),RR,ARR(4),AQ(4))
C   INTEGER IB(4)
C   COMMON/BL1/NE,NN,NI,ND,NT
C
C   IB(1)=0
C   IB(2)=NN
C   IB(3)=1
```

```

      IBC(4)=-NN
      K=K+IBC(1)
      KK=K
      DO 30 J=1,4
        KK=KK+IBC(J)
        I1=IN-1+I/3+J/3
        I2=IN+I/3+J/3-NN-2
      IF((KK.GT.0).AND.(KK.LE.ND).AND.(I1.GT.0).AND.(I2.NE.0))THEN
        PSC(J)=RES2(KK)
      ELSE
        PSC(J)=0. DO
      ENDIF
30    CONTINUE
      RR=0. DO
      IF(I.GT.1)GOTO 32
      ARR(1)=-1. DO/3. DO
      ARR(2)=-1. DO/6. DO
      ARR(3)=2. DO/3. DO
      ARR(4)=ARR(2)
32    DO 34 IX=1,4
34    AQC(IX)=ARR(IX)
      DO 40 JJ=1,4
        IF(I.EQ.1)GOTO 35
        KP=JJ-I+5
        IF(KP.GT.4)THEN
          KP=MOD(KP,4)
        ENDIF
        ARR(JJ)=AQC(KP)
35    RR=RR+ARR(JJ)*PSC(JJ)
40    CONTINUE
      RETURN
      END

```

4.4.7 - MODIFICA O VETOR B.

No capítulo anterior, verificamos que ao resolver o sistema linear (3.34) fazíamos uma pré multiplicação por uma matriz P , e dessa forma nossa matriz $\bar{A}$ ficava transformada em $\bar{\bar{A}}$ e o vetor b do lado direito do sistema passava a ser $\bar{b}$. Aqui fazemos essa transformação.

$C1$ corresponde a M .

$P(1)$ será correspondente a A_2 e $P(I) = A_{I+2}$.

Convém observar que M não será tratado como uma matriz, o mesmo ocorrendo com A , isso porque tratam-se de matrizes diagonais com todos os elementos iguais; por essa razão armazenamos esses valores e os deles decorrentes num vetor.

```

C*****
C      MODIFICA O VETOR b      *
C*****
      SUBROUTINE MODFBC(AA,BC,P,B)
      PARAMETER (MZ=100)
      PARAMETER (MM=MZ*MZ)
      REAL*8 B(MM),P(MZ),C1,AA,BC
      COMMON/BL1/NE,NN,NI,ND,NT
C
      C1=AA/BC
      P(1)=C1-BC/AA
      DO 20 I=2,NI
20    P(I)=C1-1.00/P(I-1)
      DO 30 I=1,NN
30    B(I)=B(I)/AA
      DO 40 KC=1,NI
          K1=NN*(KC-1)
          K2=NN*KC
          DO 40 J=1,NN
              K=K1+J
              JK=K2+J
40    B(JK)=(-B(K)+B(JK)/BC)/P(KC)
      RETURN
      END

```

4.4.8 - RESOLUÇÃO DOS SISTEMAS LINEARES.

Aqui temos a destacar a montagem de um bloco

da matriz $\bar{B}$ do sistema (3.35) (são todos iguais), observe que essa montagem só é efetivada nas três primeiras iterações, uma vez que nas demais ele não se modifica.

Como $\bar{A}$ decorrente de (3.34) foi construída de maneira a possuir os blocos na diagonal iguais à identidade e os outros blocos da diagonal superior são todos iguais e obtidos em função do vetor P , ao resolvermos o primeiro sistema não montamos a matriz e o resultado (como comentado no capítulo anterior) ficará armazenado no próprio vetor b (que depois da sub-rotina anterior passou a ser $\bar{b}$).

Observamos que a matriz $\bar{B}$ é armazenada na forma retangular com duas colunas e NN (igual a $NE-1$ para condição de contorno de Dirichlet homogêneo e igual a $NE+1$ para condição de contorno do tipo Neumann homogêneo) linhas. Verificamos também que a decomposição de Cholesky só é feita uma vez a cada nível de tempo e somente até a terceira iteração, depois, como já citamos, a matriz não se modifica.

```

C*****
C   RESOLVE OS SISTEMAS LINEARES          *
C*****
      SUBROUTINE SOLVSI SCIT,A,AA,BC,P,B)
      PARAMETER (MZ=100)
      PARAMETER (MM=MZ*MZ)
      REAL*8 ACMZ,2),BCMMD,PCMZO,XCMZO,BBCMZO,C1,RJ,AA,BC
      COMMON/BL1/NE,NN,NI,ND,NT
C
C   COMPOE A MATRIZ A NAS 3 PRIMEIRAS ITERACOES
C
      IF (IT.LT.4) THEN

```

```

DO 10 I=1,NN
  AC(I,1)=AA
10 AC(I,2)=BC
  AC(NN,2)=0
ENDIF
C
C RESOLVE O PRIMEIRO SISTEMA
C
DO 80 IK=1,NN
  IF( IK.EQ.1) THEN
    RJ=0. DO
  ELSE IF( IK.EQ.NN) THEN
    RJ=BC/AA
  ELSE
    RJ=1. DO/PC(NN-1K)
  ENDIF
  DO 50 K=1,NN
    NP=NN-K+1
    NQ=ND-1K*NN+NP
50 BB(NP)=BC(NQ)-RJ*BB(NP)
C
C RESOLVE O SEGUNDO SISTEMA
C DECOMPOE UMA UNICA VEZ ATE A 3A. ITERACAO, DEPOIS USA A DECOMP.
C
  IF( IK.GT.1.OR.IT.GT.3) GOTO 55
C
C DECOMPOE CHOLESKI NA PRIMEIRA VEZ
C
  CALL CHOLEALCA)
C
C RESOLVE SISTEMAS TRIANGULARES
C
55 CALL PROGREGRC(A,BB,X)
C
C ARMAZENA EM B O RESULTADO
C
  DO 60 JJ=1,NN
    NR=NN-JJ+1
    NS=ND-1K*NN+NR
60 BC(NS)=X(NR)
80 CONTINUE
  RETURN
  END

```

4.4.9 - DECOMPOSIÇÃO DE CHOLESKY.

Essa sub-rotina foi feita para efetuar a

decomposição de Cholesky em matrizes simétricas, positivas definidas, de banda M, e que estejam armazenadas numa forma retangular contendo a diagonal principal na primeira coluna e as demais subdiagonais (somente as inferiores ou somente as superiores), nas colunas restantes. Em virtude de termos utilizado funções "chapéu" para a base do elemento finito, nossos blocos B são tridiagonais, portanto a banda é 1, se usarmos outro tipo de base haverá variação no valor de M, e o programa continua válido desde que se ajuste esse parâmetro.

```

C*****
C  FAZ A DECOMPOSICAO DE CHOLESKY NA MATRIZ RETANGULAR *
C*****
      SUBROUTINE CHOLEAL(A)
      PARAMETER (MZ=100)
      REAL*8 A(MZ,2)
      COMMON/BL1/NE,NN,NI,ND,NT

C
C  M E' A BANDA, ELEM. LIN. M=1
C
      M=1
      DO 50 MP=1,NN
        A(MP,1)=SQRT(A(MP,1))
        DO 10 J=2,M+1
          A(MP,J)=A(MP,J)/A(MP,1)
          IF(NN.LT.MP+M) THEN
            MI=NN
          ELSE
            MI=MP+M-1
          ENDIF
10      CONTINUE
        DO 20 I=MP+1,MI+1
          IF(I.GT.NN) GOTO 20
          MK=M+MP-I+1
          MJ=I-MP
          DO 15 J=1,MK
15      A(I,J)=A(I,J)-A(MP,MJ+1)*A(MP,MJ+J)
20      CONTINUE
50      CONTINUE
      RETURN
      END

```

4.4.10 - RESOLUÇÃO DOS SISTEMAS TRIANGULARES.

Essa sub-rotina resolve os dois sistemas triangulares surgidos na decomposição de Cholesky; aqui também consideramos as matrizes dos sistemas triangulares dispostos na forma retangular que foi descrita na seção anterior, e o mesmo comentários ali contidos são válidos para o parâmetro M.

```

C*****
C   RESOLVE OS SISTEMAS TRIANGULARES   *
C*****
      SUBROUTINE PROGREGR(A,BB,X)
      PARAMETER (MZ=100)
      REAL*8 A(MZ,2),BB(MZ),X(MZ),S
      COMMON/BL1/NE,NN,NI,ND,NT
C
C   M E' A BANDA, ELEM. LINEARES M=1
C
      M=1
C
C   SUBST. PROGRESSIVA
C
      X(1)=BB(1)/A(1,1)
      DO 20 J=2,NN
        S=0. DO
          DO 10 I=1,J-1
            IF(J-I.GT.M)GOTO 10
            S=S+A(I,J-I+1)*X(I)
10          CONTINUE
20        X(J)=(BB(J)-S)/A(J,1)
C
C   SUBST. REGRESSIVA
C
      X(NN)=X(NN)/A(NN,1)
      DO 40 I=NI,1,-1
        S=0. DO
          DO 30 J=I+1,NN
            IF(J-I.GT.M)GOTO 30
            S=S+A(I,J-I+1)*X(J)
30          CONTINUE
40        X(I)=(X(I)-S)/A(I,1)
      RETURN

```

END

4.4.11 - CÁLCULO DO PRIMEIRO TERMO NO PROCESSO ITERATIVO.

No tempo 1 temos um processo iterativo e o lado direito do sistema é determinado por :

$$(\Delta t)^2 (1 - (\Delta t)^2) A^x \otimes A^y (U^1 - U^0)^{q+1} + \mathbb{F}^1 \quad (4.19)$$

onde $\mathbb{F}^1$ é calculado pela soma das integrais de (4.3).

Dando para $(U^1 - U^0)^{q+1}$ (vetor de $(NN)^2$ posições) a estrutura de blocos :

$$U = \begin{bmatrix} U_1 \\ U_2 \\ U_3 \\ \vdots \\ U_{N-1} \\ U_N \end{bmatrix} \quad (4.20)$$

e tomando $k = (\Delta t)^2 (1 - (\Delta t)^2)$; teremos o primeiro termo do processo iterativo igual a :

$$\frac{k}{h} \cdot \begin{bmatrix} 2 A^y U_1 - A^y U_2 \\ - A^y U_1 + 2 A^y U_2 - A^y U_3 \\ \vdots \\ - A^y U_{I-1} + 2 A^y U_I - A^y U_{I+1} \\ \vdots \\ - A^y U_{N-1} + 2 A^y U_N \end{bmatrix} \quad (4.21)$$

Lembrando sempre que (devido a base por nós adotada) :

$$A_{ij}^x = A_{ij}^y = \begin{cases} \frac{2}{h} & \text{se } i = j \\ -\frac{1}{h} & \text{se } |i-j| = 1 \\ 0 & \text{se } |i-j| > 1 \end{cases} \quad (4.22)$$

PM no programa é o valor que acima chamamos de k.

V é o vetor do lado direito já atualizado e modificado, e corresponderá a um bloco que descrevemos acima.

V1 é o vetor de saída com os valores que serão adicionados a $\tilde{F}^1$.

```

C*****
C   CALCULA O PRIMEIRO TERMO DE B NO TEMPO 1      *
C*****
SUBROUTINE PRITERB1(V,V1)
PARAMETER (MZ=100)
PARAMETER (MM=MZ*MZ)
REAL*8 A1,B1,C1,A2,B2,C2,VCMM,BCMM,PM,H,V1CMM
REAL*8 AP,XLAMB,DT,P1,P2
COMMON/BL1/NE,NN,NI,ND,NT
COMMON/BL2/H,DT,XLAMB,AP,P1,P2
C
PM=DT*DT*(1.0D0-(DT*DT))/(H*H)

```

```

DO 10 I=1,NN
DO 10 J=1,NN
IK=(I-1)*NN+J
MK=IK-NN
NK=IK+NN
IF(I.EQ.1) THEN
  A1=0. DO
  C1=V(NK)
ELSE IF(I.EQ.NN) THEN
  A1=V(MK)
  C1=0. DO
ELSE
  A1=V(MK)
  C1=V(NK)
ENDIF
B1=V(IK)
10 B(IK)=-A1+2. DO*B1-C1
DO 20 I=1,NN
DO 20 J=1,NN
JK=(I-1)*NN+J
IF(J.EQ.1) THEN
  A2=0. DO
  C2=B(JK+1)
ELSE IF(J.EQ.NN) THEN
  A2=B(JK-1)
  C2=0. DO
ELSE
  A2=B(JK-1)
  C2=B(JK+1)
ENDIF
B2=B(JK)
V(JK)=(-A2+2. DO*B2-C2)*PM
20 CONTINUE
RETURN
END

```

4.4.12 - SAÍDA DOS RESULTADOS E ANÁLISE DO ERRO.

Essa sub-rotina processa os resultado a cada iteração, pois como vimos na seção 3.3, as soluções dos sistemas lineares (Z) poderão não ser a solução do nosso problema, mas algo obtido em sua função; assim teremos:

No Tempo zero

$$U^0 = Z \quad (4.23)$$

No Tempo 1

$$U^1 = Z + U^0 \quad (4.24)$$

Nos demais Tempos

$$U^{n+1} = Z - U^{n-1} + 2 U^n \quad (4.25)$$

A efetua os cálculos da solução em cada nó.

B é um vetor que armazena a solução a cada nível do tempo.

Nessa sub-rotina colocamos também a opção de avaliar a medida do erro no caso de se conhecer a solução exata. A medida foi feita na norma do L^2 , e a integração realizada pelo método de Simpson [17]; utilizamos este método e não quadratura para ganharmos no fator tempo.

ERR calcula o erro entre A e UEX (solução exata) em cada nó.

SSS é o acumulador do erro.

```

C*****
C  SAIDA DOS RESULTADOS E DO ERRO RELATIVO  *
C*****
SUBROUTINE SAIERROCB,XX,YY,XT,IT,RES1,RES2,ISS,ISE
PARAMETER (MZ=100)
PARAMETER (MM=MZ*MZ)
REAL*8 BCMMD,ER,SSS,UEX,XX(MZ),YY(MZ),XT(MZ),AB
REAL*8 RES1(MMD),RES2(MMD),A,P1,P2,H,DT,XLAMB,AP
COMMON/BL2/H,DT,XLAMB,AP,P1,P2
COMMON/BL1/NE,NN,NI,ND,NT

C
SSS=0. DO
IF(ISS.EQ.1.AND.ISE.EQ.1) THEN
WRITE(*,40)
ELSEIF(ISS.EQ.1.AND.ISE.NE.1) THEN
WRITE(*,45)
ENDIF

```

```

DO 10 I=1,NN
DO 10 J=1,NN
  JN=NN*(I-1)+J
  IP=MOD(I,2)
  IJ=MOD(JN,2)
  IF(I*J.EQ.1) THEN
    A=BC(JN)
  ELSE IF(I*J.EQ.2) THEN
    A=BC(JN)+RES2(JN)
  ELSE
    A=BC(JN)+2.DO*RES2(JN)-RES1(JN)
  ENDIF

```

C
C
C

```

SE POSSUIR SOL. EXATA ANALISA O ERRO

```

```

  IF(I*SE.EQ.1) THEN
    AB=UEX(XX(I+1),YY(J+1),XT(IT))
    ER=A-AB
    IF(I*P.EQ.1) THEN
      IF(I*J.EQ.1) THEN
        ER=ER*ER*16.OOO
      ELSE
        ER=ER*ER*8.OOO
      ENDIF
    ELSE IF(I*P.EQ.0) THEN
      IF(I*J.EQ.0) THEN
        ER=ER*ER*8.OOO
      ELSE
        ER=ER*ER*4.OOO
      ENDIF
    ENDIF
    IF(I*SS.EQ.1) THEN
      WRITE(*,50)XX(I+1),YY(J+1),A,AB
    ENDIF
    SSS=SSS+ER
  ELSE
    IF(I*SS.EQ.1) THEN
      WRITE(*,60)XX(I+1),YY(J+1),A
    ENDIF
  ENDIF

```

```

  WRITE(98,*)A
  BC(JN)=A

```

10

```

CONTINUE
IF(I*SE.EQ.1) THEN
  SSS=SQRT(SSS*H*H/9.OOO)
  WRITE(*,*)' ERRO NA NORMA L2 NO TEMPO',IT-1,' E ',SSS
  WRITE(65,*)' ERRO NA NORMA L2 NO TEMPO',IT-1,' E ',SSS
ENDIF
40 FORMAT(//,' X',7X,' Y',8X,' SOL. APROX',4X,' SOL. EXATA')
45 FORMAT(//,' X',7X,' Y',8X,' SOL. APROX')
50 FORMAT(1X,F6.4,3X,F6.4,3X,F11.8,3X,F11.8)

```

```

60  FORMAT(1X,F6.4,3X,F6.4,3X,F11.8)
    RETURN
    END

```

4.4.13 - FUNÇÕES.

Há necessidade de se acrescentar ao programa as condições iniciais, assim como outras funções que serão utilizadas nos cálculos das integrais; colocamos abaixo essas funções sendo que aqui os valores apresentados servem para o seguinte exemplo :

$$\left\{ \begin{array}{l}
 \frac{\partial^2 U}{\partial t^2} - \Delta U = 0 \quad \Omega = [0,1] \times [0,1] \quad t \in [0,T] \\
 U(X,Y,0) = \text{SEN}(\text{PI} * X) \text{SEN}(\text{PI} * Y) \\
 U_t(X,Y,0) = 0 \\
 U(X,Y,t) = 0 \quad \forall (X,Y) \in \partial \Omega
 \end{array} \right.$$

A solução exata para esse problema é :

$$U = \text{SEN}(\text{PI} * X) \text{SEN}(\text{PI} * Y) \text{COS}(\sqrt{2} * \text{PI} * t)$$

```

C*****TEMPO ZERO IT=1*****
C      VALOR INICIAL DA SOLUCAO  Uo
      REAL*8 FUNCTION F(X,Y)
      REAL*8 X,Y
      F=DSIN(3.141592653D0*X)*DSIN(3.141592653D0*Y)
      RETURN
      END
C      DERIVADA EM REL. A X DE Uo
      REAL*8 FUNCTION DXF(X,Y)
      REAL*8 X,Y
      DXF=3.141592653D0*DCOS(3.141592653D0*X)*DSIN(3.141592653D0*Y)
      RETURN
      END
C      DERIVADA EM REL. A Y DE Uo
      REAL*8 FUNCTION DYF(X,Y)

```

```

      REAL*8 X,Y
      DYF=3.141592653D0*DSIN(3.141592653D0*X)*DCOS(3.141592653D0*Y)
      RETURN
      END
C      DERIVADA 2a. EM REL. A X E Y DE Uo
      REAL*8 FUNCTION DDF(X,Y)
      REAL*8 X,Y
      DDF=9.8696044D0*DCOS(3.141592653D0*X)*DCOS(3.141592653D0*Y)
      RETURN
      END
C*****TEMPO UM IT=2*****
C      COND. INICIAL DA DERIV.
      REAL*8 FUNCTION G(X,Y)
      REAL*8 X,Y
      G=0.D0
      RETURN
      END
C      LAPLACIANO DA COND. INICIAL DA FUNCAO
      REAL*8 FUNCTION LAPL(X,Y)
      REAL*8 X,Y
      LAPL=-19.7392088D0*DSIN(3.141592653D0*X)*DSIN(3.141592653D0*Y)
      RETURN
      END
C      FONTE NO TEMPO ZERO
      REAL*8 FUNCTION FONTE(X,Y)
      REAL*8 X,Y
      FONTE=0.D0
      RETURN
      END
C      DERIV. DA COND. INICIAL DA DER. G EM REL. A X
      REAL*8 FUNCTION DGX(X,Y)
      REAL*8 X,Y
      DGX=0.D0
      RETURN
      END
C      DERIV. DO LAPLAC. EM REL A X
      REAL*8 FUNCTION LAPX(X,Y)
      REAL*8 X,Y
      LAPX=-62.0125534D0*DCOS(3.141592653D0*X)*DSIN(3.141592653D0*Y)
      RETURN
      END
C      DERIV. DA FONTE NO Tzero EM REL. A X
      REAL*8 FUNCTION FOTX(X,Y)
      REAL*8 X,Y
      FOTX=0.D0
      RETURN
      END
C      DERIV. DA COND. INICIAL DA DER. G EM REL. A Y
      REAL*8 FUNCTION DGY(X,Y)
      REAL*8 X,Y
      DGY=0.D0

```

```

RETURN
END
C      DERIV. DO LAPLAC. EM REL. A Y
REAL*8 FUNCTION LAPY(X,Y)
REAL*8 X,Y
LAPY=-62.0125534D0*DSINC(3.141592653D0*X)*DCOS(3.141592653D0*Y)
RETURN
END
C      DERIV. DA FONTE NO Tzero EM REL. A Y
REAL*8 FUNCTION FOTY(X,Y)
REAL*8 X,Y
FOTY=0.D0
RETURN
END
C      DERIV. 2a. DA COND. INICIAL DA DERIV. EM X e Y
REAL*8 FUNCTION DGXY(X,Y)
REAL*8 X,Y
DGXY=0.D0
RETURN
END
C      DERIV. 2a. DO LAPL. EM X e Y
REAL*8 FUNCTION LAPXY(X,Y)
REAL*8 X,Y
LAPXY=-194.818182D0*DCOS(3.141592653D0*X)*DCOS(3.141592653D0*Y)
RETURN
END
C      DERIV. 2a. DA FONTE NO Tzero EM REL. A X e Y
REAL*8 FUNCTION FOTXY(X,Y)
REAL*8 X,Y
FOTXY=0.D0
RETURN
END
C      PRIMEIRO TERMO (S) NO TEMPO 1
REAL*8 FUNCTION F11(X,Y,A,DT)
REAL*8 X,Y,DT,A,G,LAPL,FONTE
F11=DT*G(X,Y)+.5D0*DT*DT*(A*LAPL(X,Y)+FONTE(X,Y))
RETURN
END
C      DERIVADA DE S EM REL A X
REAL*8 FUNCTION DX11(X,Y,A,DT)
REAL*8 X,Y,DT,A,DGX,LAPX,FOTX
DX11=DT*DGX(X,Y)+.5D0*DT*DT*(A*LAPX(X,Y)+FOTX(X,Y))
RETURN
END
C      DERIVADA DE S EM REL. A Y
REAL*8 FUNCTION DY11(X,Y,A,DT)
REAL*8 X,Y,DT,A,DGY,LAPY,FOTY
DY11=DT*DGY(X,Y)+.5D0*DT*DT*(A*LAPY(X,Y)+FOTY(X,Y))
RETURN
END
C      DERIVADA 2a. DE S EM REL. A X E Y

```

```

      REAL*8 FUNCTION DD11(X,Y,A,DT)
      REAL*8 X,Y,DT,A,DGXY,LAPXY,FOTXY
      DD11=DT*DGXY(X,Y)+.5D0*DT*DT*(A*LAPXY(X,Y)+FOTXY(X,Y))
      RETURN
      END
C*****TEMPO N IT>=3*****
C      FONTE DO PROBLEMA EM TODOS OS TEMPOS
      REAL*8 FUNCTION FONT1(X,Y,T)
      REAL*8 X,Y,T
      FONT1=0.D0
      RETURN
      END
C      SOL. EXATA DO PROBLEMA
      REAL*8 FUNCTION UEX(X,Y,T)
      REAL*8 X,Y,T
      UEX=DSIN(3.141592653D0*X)*DSIN(3.141592653D0*Y)
      $*DCOS(4.442882938D0*T)

```

Para esse exemplo com as seguintes entradas :

Nro. de Elementos por direção = 40
 Tempo Final = 1 s
 Nro. de Tempos = 101
 Coeficiente (a) = 1
 Parâmetro de Estabilidade (λ) = 10

obtivemos os seguintes resultados:

ERRO NA NORMA L2	NO TEMPO	0.00 s	E	4.727E-05
ERRO NA NORMA L2	NO TEMPO	0.10 s	E	1.017E-03
ERRO NA NORMA L2	NO TEMPO	0.20 s	E	3.378E-03
ERRO NA NORMA L2	NO TEMPO	0.30 s	E	6.165E-03
ERRO NA NORMA L2	NO TEMPO	0.40 s	E	8.122E-03
ERRO NA NORMA L2	NO TEMPO	0.50 s	E	8.086E-03
ERRO NA NORMA L2	NO TEMPO	0.60 s	E	5.358E-03

ERRO NA NORMA L2	NO TEMPO	0.70 s	E	2.064E-05
ERRO NA NORMA L2	NO TEMPO	0.80 s	E	7.191E-03
ERRO NA NORMA L2	NO TEMPO	0.90 s	E	1.460E-02
ERRO NA NORMA L2	NO TEMPO	1.00 s	E	1.539E-02

Testamos o programa com codição do tipo Dirichlet homogêneo para o problema cuja solução exata é $U = x^2 y^2 (x-1)(y-1)t$, no domínio $[0,1] \times [0,1]$, com as seguintes entradas:

Nro. de Elementos por direção = 80
Tempo Final = 1 s
Nro. de Tempos = 101
Coeficiente (a) = 1
Parâmetro de Estabilidade (λ) = 10

obtivemos os seguintes resultados:

ERRO NA NORMA L2	NO TEMPO	0.00 s	E	0.000E+00
ERRO NA NORMA L2	NO TEMPO	0.05 s	E	1.887E-07
ERRO NA NORMA L2	NO TEMPO	0.10 s	E	3.344E-07
ERRO NA NORMA L2	NO TEMPO	0.15 s	E	4.629E-07
ERRO NA NORMA L2	NO TEMPO	0.20 s	E	5.708E-07
ERRO NA NORMA L2	NO TEMPO	0.25 s	E	6.571E-07
ERRO NA NORMA L2	NO TEMPO	0.30 s	E	7.229E-07
ERRO NA NORMA L2	NO TEMPO	0.35 s	E	7.715E-07
ERRO NA NORMA L2	NO TEMPO	0.40 s	E	8.072E-07
ERRO NA NORMA L2	NO TEMPO	0.45 s	E	8.349E-07
ERRO NA NORMA L2	NO TEMPO	0.50 s	E	8.599E-07
ERRO NA NORMA L2	NO TEMPO	0.55 s	E	8.862E-07
ERRO NA NORMA L2	NO TEMPO	0.60 s	E	9.170E-07
ERRO NA NORMA L2	NO TEMPO	0.65 s	E	9.547E-07
ERRO NA NORMA L2	NO TEMPO	0.70 s	E	1.001E-06
ERRO NA NORMA L2	NO TEMPO	0.75 s	E	1.060E-06
ERRO NA NORMA L2	NO TEMPO	0.80 s	E	1.132E-06
ERRO NA NORMA L2	NO TEMPO	0.85 s	E	1.217E-06
ERRO NA NORMA L2	NO TEMPO	0.90 s	E	1.312E-06
ERRO NA NORMA L2	NO TEMPO	0.95 s	E	1.413E-06
ERRO NA NORMA L2	NO TEMPO	1.00 s	E	1.525E-06

Exemplos como esses últimos, são de pouca aplicabilidade física, podemos até dizer que suas finalidades são mais de caráter acadêmico, não somente por conhecermos a solução exata mas

também por não simularem situações reais.

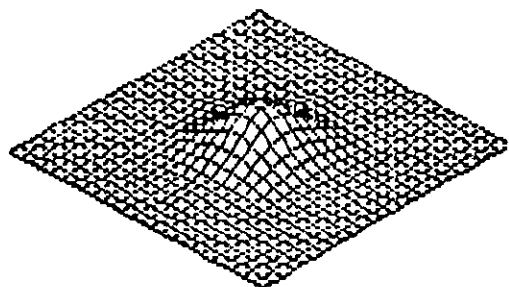
Apresentaremos a seguir um exemplo que melhor aproxima o fenômeno físico, em que se tem uma fonte puntual localizada no centro de um domínio quadrangular, apresentando um decaimento brusco com o tempo. Matematicamente essa fonte pode ser vista como uma função delta de Dirac.

$$\left\{ \begin{array}{l} \frac{\partial^2 U}{\partial t^2} - \Delta U = 100 e^{-50[(x-.5)^2 + (y-.5)^2]} e^{-40t^2} \\ U(X,Y,0) = 0 \\ U_t(X,Y,0) = 0 \\ U(X,Y,t) = 0 \end{array} \right. \quad \begin{array}{l} \Omega = [0,1] \times [0,1] \quad t \in [0,T] \\ \\ \forall (X,Y) \in \partial \Omega \end{array}$$

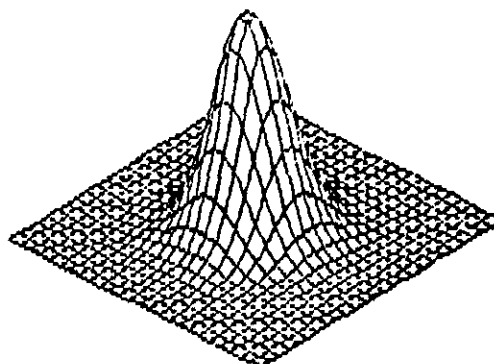
Os gráficos que apresentaremos a seguir foram obtidos com o auxílio do aplicativo Energraph. Nesse primeiro exemplo tivemos as seguintes entradas :

Nro. de Elementos por direção = 25
 Tempo Final = 1.5 s
 Nro. de Tempos = 101
 Coeficiente (a) = .25
 Parâmetro de Estabilidade (λ) = 10

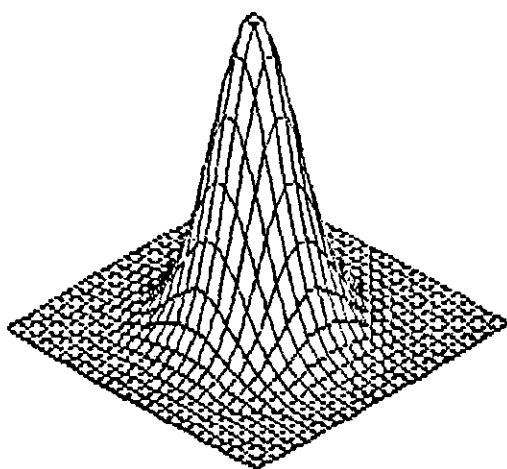
Os valores abaixo de cada gráfico representam o tempo (em segundos) em que o mesmo foi traçado.



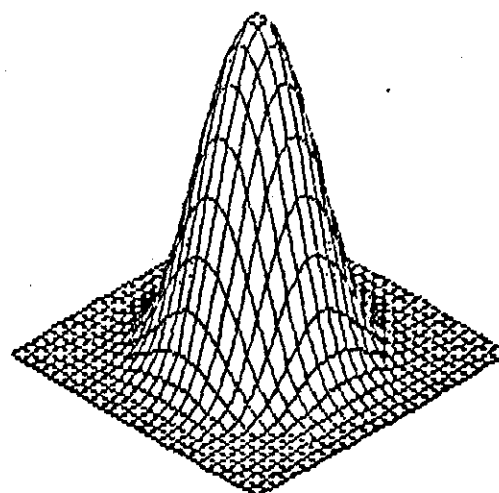
0.10



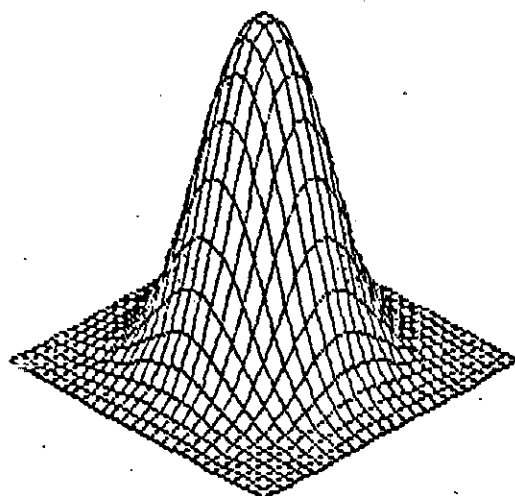
0.20



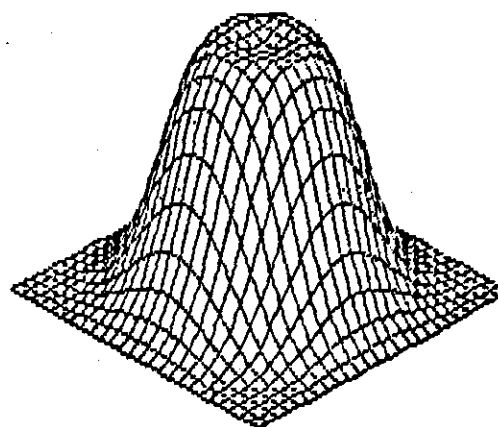
0.30



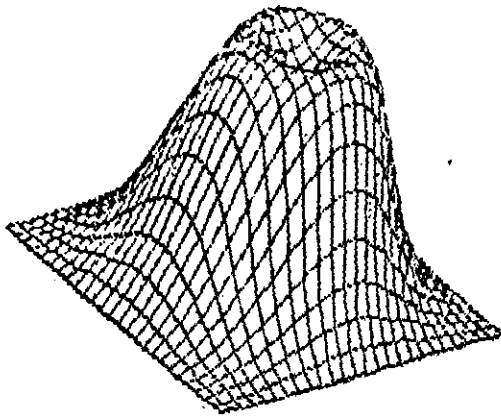
0.40



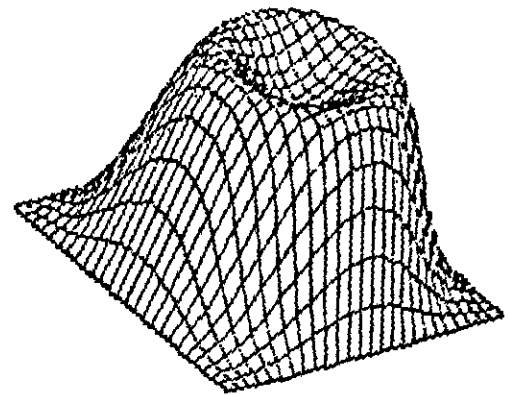
0.50



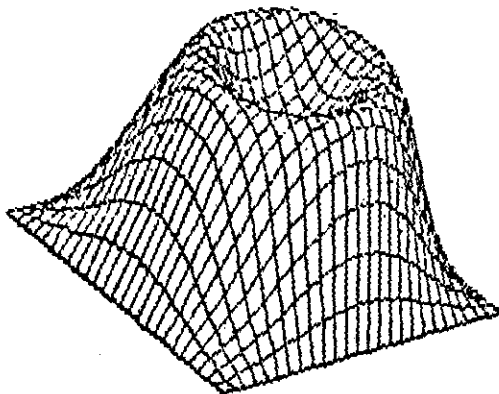
0.60



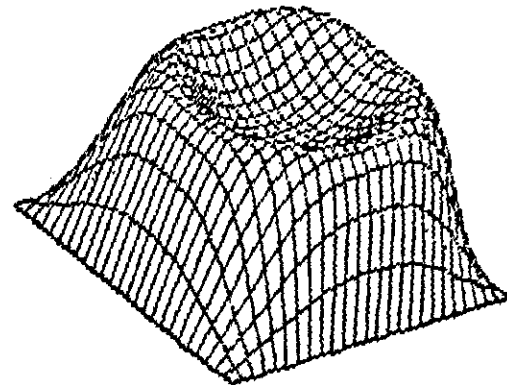
0.70



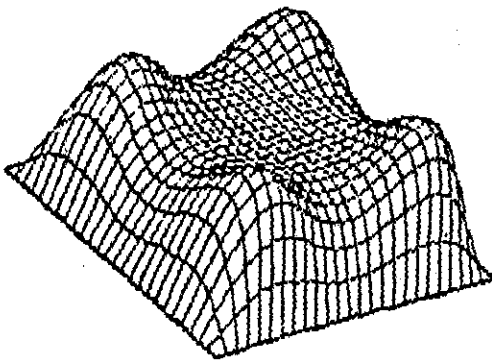
0.80



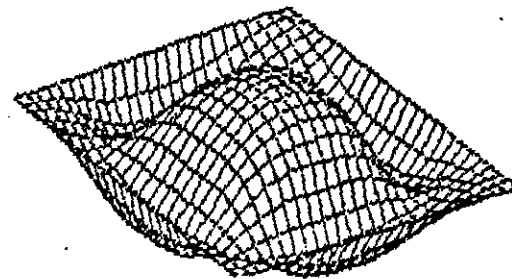
0.90



1.00



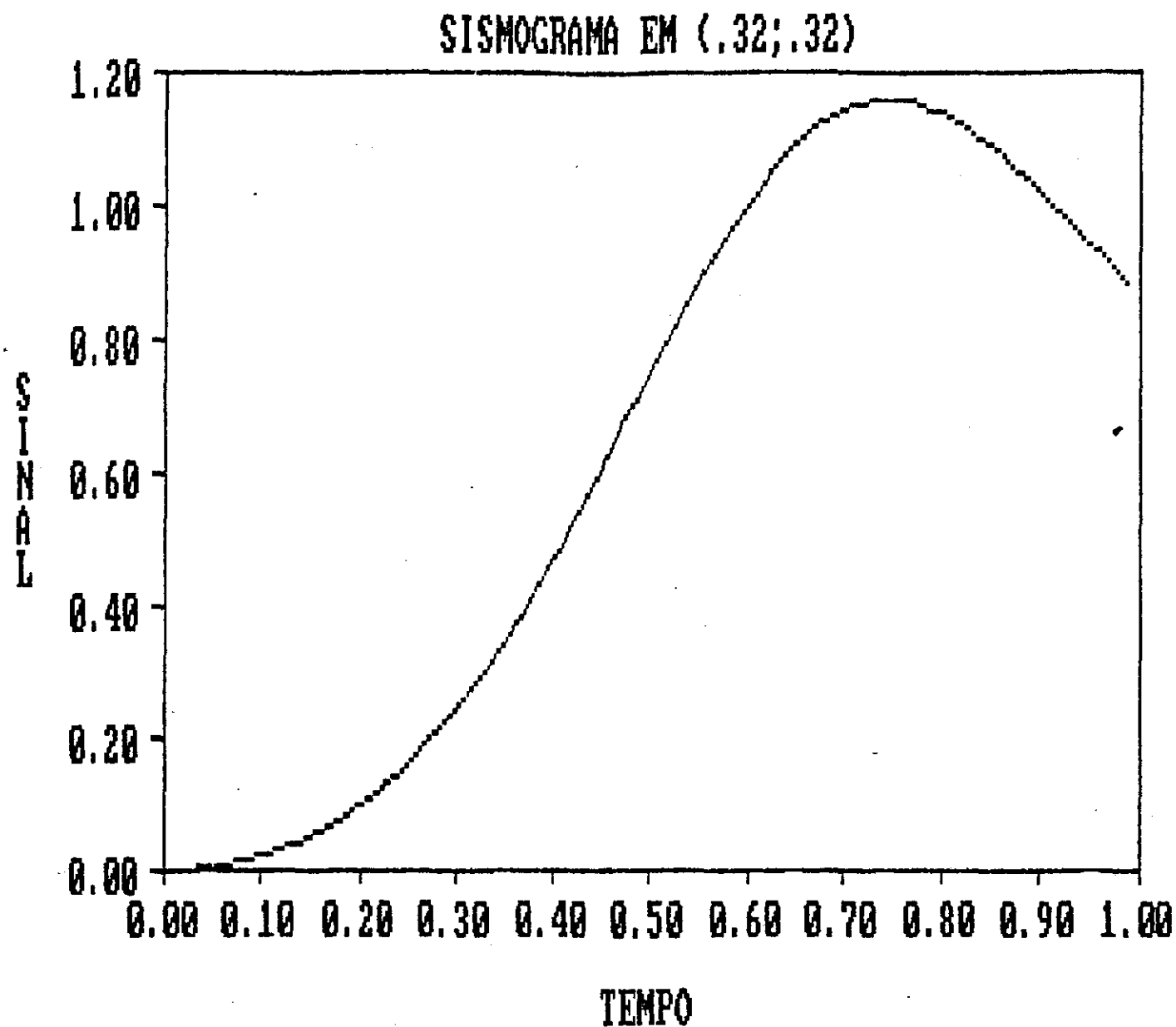
1.20



1.50

Para esse exemplo; fixamos uma posição no espaço, no

caso (.32;.32) e simulamos o seguinte sismograma.



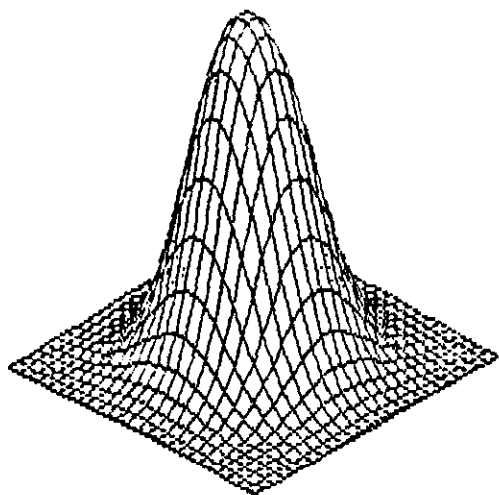
Para a mesma fonte e condições iniciais idênticas ao do problema anterior, executamos nosso programa dessa feita com condição

de fronteira do tipo Neumann homogêneo; nesse caso obtivemos os próximos resultados .

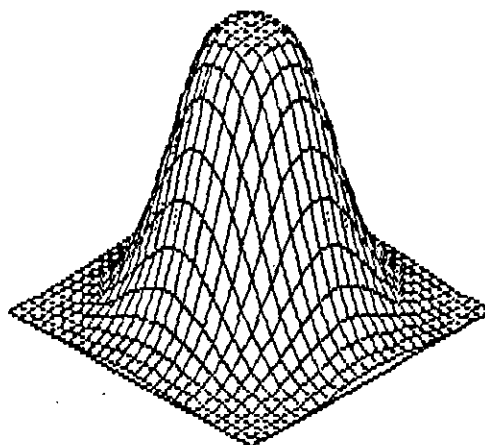
Chamamos a atenção para o significado físico que possui essa situação. Quando impomos condição do tipo Neumann, fisicamente estamos tratando do fluxo que estará saindo do nosso domínio. Se a derivada normal nos limites de nossa região for nula (condição de contorno tipo Neumann homogêneo), estaremos tratando de uma situação em que o fluxo não passa para o exterior. Essa situação é bem visualizada pelos próximos gráficos, obtidos quando resolvemos :

$$\left\{ \begin{array}{l} \frac{\partial^2 U}{\partial t^2} - \Delta U = 100 e^{-50 [(x-.5)^2 + (y-.5)^2]} - 40t^2 \\ U(X,Y,0) = 0 \\ U_t(X,Y,0) = 0 \\ \frac{\partial U(X,Y,t)}{\partial n} = 0 \quad \forall (X,Y) \in \partial \Omega \end{array} \right. \quad \Omega = [0,1] \times [0,1] \quad t \in [0,T]$$

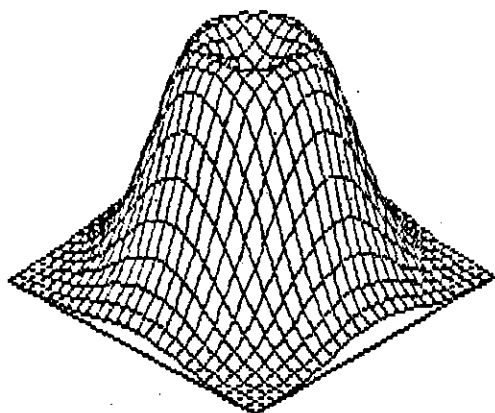
com as mesmas entradas do problema anterior.



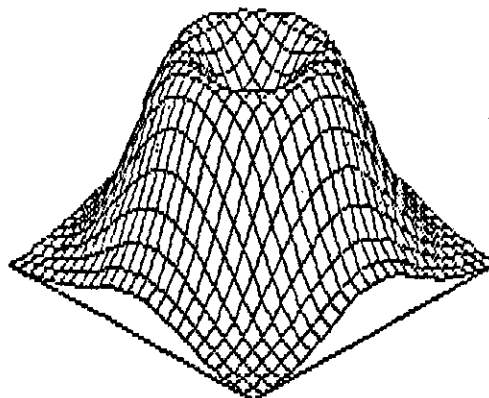
0.50



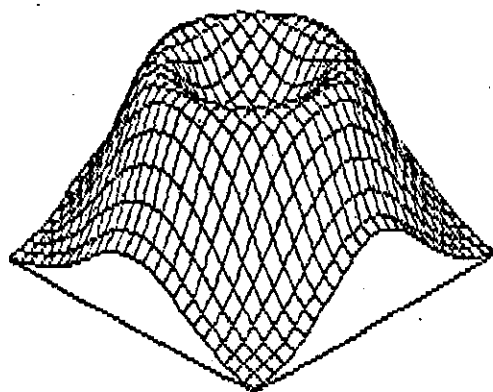
0.60



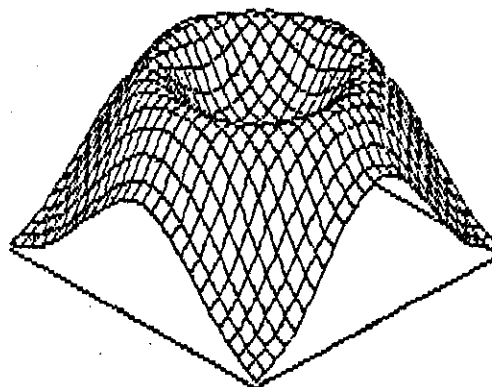
0.70



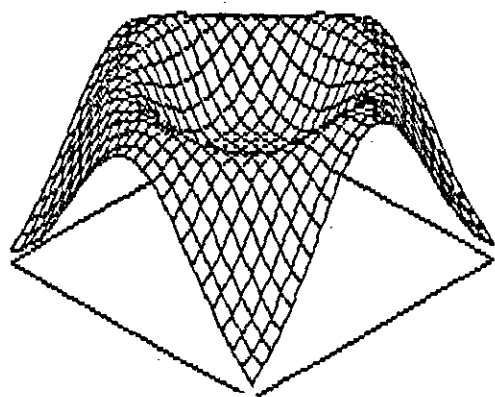
0.80



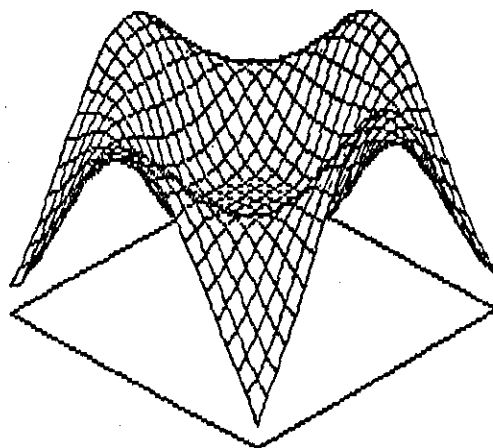
0.90



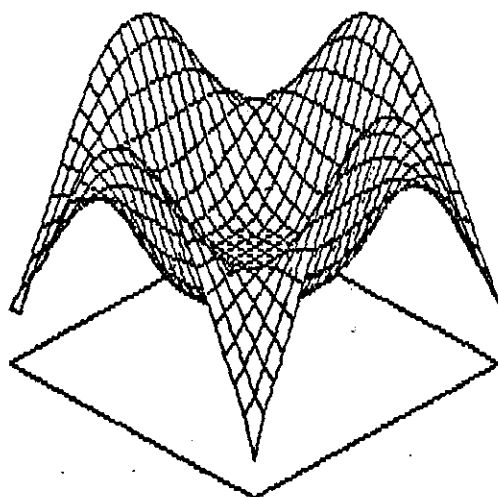
1.00



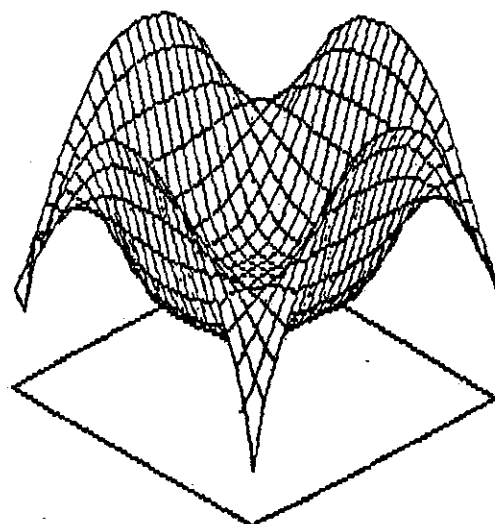
1.10



1.20



1.30



1.40

CONCLUSÃO.

Ao término desse trabalho podemos verificar quão grande é a amplitude dos campos do conhecimento científico, em que os métodos da análise aplicada estão inseridos.

Pelo estudo por nós efetuado verificamos que as respostas provenientes do Método dos Elementos Finitos constituem-se de uma boa aproximação. Nos foi possível observar também que a perspectiva é de um aperfeiçoamento cada vez maior no sentido de que seja possível obter melhores resultados dispendendo menor esforço a cada aprimoramento do método. Outra vantagem do método é, sem dúvida, a possibilidade de obtenção de resultados em regiões de complexa geometria, o que por certo melhor aproxima os resultados das situações reais.

A estratégia das Direções Alternadas por sua vez mostrou-se eficiente, minimizando o problema de espaço de memória computacional requerida quando da aplicação do Método dos Elementos Finitos, o que torna possível sua implementação a nível de micro computadores pessoais, o que não deixa de ser uma grande vantagem.

Acreditamos que o objetivo desse nosso trabalho, (estudo do problema direto dos Métodos Sísmicos com o uso do Método de Galerkin com Direções Alternadas), foi atingido e que trabalhos posteriores virão aprofundar o conhecimento até aqui adquirido.

Portanto concluímos que o problema por nós tratado continua em aberto , deixando espaço para muitos outros trabalhos nessa área de conhecimento, o que não deixa de ser interessante. Esperamos que as colocações feitas aqui sejam úteis quando da elaboração de outras pesquisas.

É de se evidenciar ainda que todo passo adiante, nesse campo, tem validade, principalmente se considerarmos o objetivo final a que ele se propõe.

SUGESTÕES.

Como citamos, o problema por nós abordado nesse trabalho continua em aberto, ou seja, ainda não conseguimos atingir a plenitude dos resultados esperados. Numa visão otimista, acreditamos que isso possa ocorrer num futuro próximo desde que disponhamos fundamentalmente de pesquisadores da área de Matemática Aplicada atuando em conjunto com pesquisadores da Geofísica, além de um razoável suporte de equipamentos.

Pelo realizado nesse trabalho, podemos dar algumas sugestões que poderão ser úteis a pesquisas posteriores a essa. Fundamentalmente sugerimos o seguinte:

- a) Resolução do problema em diversas camadas.
- b) Utilização de outras bases para os Elementos Finitos.
- c) Resolução do problema evitando as reflexões espúrias, surgidas quando se limita o domínio, criando-se com isso fronteiras que fisicamente não existem.
- d) Realização de um estudo no sentido de obtenção de um parâmetro de estabilidade (λ), ótimo.
- e) Aprimoramento do programa computacional para aplicação em equipamentos que disponham de processadores em paralelo.

BIBLIOGRAFIA

- [1] Brekhovskikh, L., and Goncharov, V., Mechanics of Continua and Wave Dynamics -Springer-Verlag - New York, 1982.
- [2] Aki, K., and Richard, P.G., Quantitative Seismology Theory and Methods, Vol. I - Freeman and Company, 1980.
- [3] Douglas, J. J., and Dupont., Alternating - Direction Galerkin Methods on Rectangles, B. Hubbard, vol II, New York, 1971.
- [4] Dendy, J. E. J. and Fairweather, G., Alternating - Direction Galerkin methods for Parabolic and Hyperbolic Problems on Rectangular Polygons, SIAM J. Numer. Anal. vol 12, N. 2, 1975.
- [5] Hudson, J. A. , The Excitation and Propagation of Elastic Waves, Cambridge University Press, 1980.
- [6] Achenbach, J. D. Wave Propagation in Elastic Solids, North-Holland Publishing Company, Amsterdam, 1975.
- [7] Tijonov, A. N. and Samarsky, A. A., Ecuaciones de la Fisica Matemática - 2a. Edição - Editora Mir , Moscou, 1980.
- [8] Girault, V. and Raviart, P.-A, Finite Elements Approximation of the Navier-Stokes Equations, Lecture Notes in Mathematics 749, Springer-Verlag, New York, 1981.
- [9] Carey, G. F. and Oden, J. T., Finite Elements - Vol. I - A

Introduction, Prentice-Hall, New Jersey, 1984.

- [10] Carey, G. F. and Oden, J. T., Finite Elements - Vol.III - Computational Aspects, Prentice-Hall, New Jersey, 1984.
- [11] Schwarz, H. R. and Stiefel, R., Numerical Analysis of Simetric Matrices , Prentice-Hall, U.S.A, 1973.
- [12] Schultz, M. H., Splines Analysis , Prentice-Hall, U.S.A, 1973.
- [13] Prenter, P. M. , Splines and Variational Methods , A Wiley-Interscience Publication, U.S.A, 1975.
- [14] Fairweather, G., Finite Elemente Galerkin Methods for Diferential Equations, Dekken, New York, 1978.
- [15] Kelly, K. R. et aut. , Aplication of Finite Difference Methods to Exploration Seismology, Topics In Numerical Analysis, Miller, Vol III - 1977
- [16] Bellman, R - Introduction to Matrix analysis, McGraw - Hill, New York, 1970.
- [17] Démidovitch, B. et Maron, I. , Éléments de Calcul Numérique, Éditions MIR, Moscou, 1973.

APÊNDICE - MÉTODO DA COLOCAÇÃO

Na busca de soluções para as equações hiperbólicas (2.25) e (2.26), tentamos fazer uso do "Método da Colocação". A razão para essa aplicação foi principalmente a simplicidade da implementação computacional, uma vez que assim procedendo não resolveríamos as integrais que somos forçados a resolver quando usamos o Método de Galerkin. Descreveremos aqui em linhas gerais os procedimentos adotados quando fazemos uso da Colocação.

Vamos supor que estejamos procurando uma solução para um problema envolvendo a seguinte equação :

$$\begin{cases} L(U) = f & \Omega = [0,1] \times [0,1] \\ U = 0 & (x,y) \in \partial \Omega \end{cases} \quad (A.1)$$

Escolhendo espaços de dimensão finita $\varphi(x)$ e $\varphi(y)$ de modo que $\varphi_{ij}(x,y) = \varphi_i(x) \varphi_j(y)$. O Método da Colocação consiste em encontrar uma aproximação : $U_N = \sum \alpha_{ij} \varphi_{ij}$ (A.2)

onde os coeficientes α_{ij} serão obtidos de modo que :

$$L(U_N(x_i, y_j)) = f(x_i, y_j) \quad (A.3)$$

e $0 = x_1 < x_2 < \dots < x_n = 1$ e $0 = y_1 < y_2 < \dots < y_n = 1$ são pontos do domínio Ω ; aqui denominados de nós naturais da partição.

Se U_N existe, então dizemos que ela "coloca" $f(x_i, y_j)$ nos pontos (x_i, y_j) de modo a satisfazer a equação (A.1). Daí a razão do método receber o nome de Colocação.

A obtenção das aproximações depende fundamentalmente da escolha dos espaços de aproximação e dos pontos escolhidos para se fazer a colocação.

Quando a colocação é feita nos nós naturais da partição, uma escolha razoável para as bases $\phi(x)$ e $\phi(y)$ são as funções conhecidas como B-Splines.

A literatura especializada denomina de "Colocação Ortogonal" aquela em que os pontos escolhidos são os nós gaussianos da partição ; nesse caso, a escolha natural para os espaços de aproximação são os polinômios de Hermite cúbico, e comprovadamente a aproximação é melhor do que a obtida no caso em que se utiliza B-Splines com a colocação sendo feita nos nós naturais da partição [13].

Em problemas de evolução os coeficientes α_{ij} são funções da variável tempo.

É conveniente observar que para o sucesso na aplicação desse método em problemas de evolução, há necessidade de se possuir um esquema estável para a discretização da variável tempo, de modo a não comprometer a aproximação.

A falta de resultados teóricos para equações hiperbólicas, e a instabilidade surgida quando discretizamos a variável tempo nos fez procurar outro método que contornasse esse problema surgido. Dessa forma chegamos ao método de Galerkin com direções alternadas.