

2 BASE TEÓRICA

Este capítulo tem por finalidade apresentar a formulação teórica do processo adaptativo adotado neste trabalho. Existem quatro aspectos fundamentais de um processo adaptativo para o problema da análise estrutural: a estratégia adaptativa, a técnica de estimativa de erro, os métodos de suavização utilizados e a técnica utilizada para a geração das malhas de elementos finitos.

2.1. Estratégia Adaptativa

No contexto de estratégias adaptativas, métodos de extensão têm sido preferencialmente escolhidos em relação a outros enfoques, conforme citado por Cavalcante-Neto [12,13]. Estes métodos incluem extensões h , p e r (Figura 3). Na extensão h a malha é refinada quando o indicador de erro excede uma tolerância pré-estabelecida de tal forma que se aumente o número de elementos da malha, diminuindo seus tamanhos, sem aumentar o grau dos seus polinômios. A extensão p é geralmente empregada em uma malha fixa. Neste caso se o erro de um determinado elemento exceder a tolerância pré estabelecida, então a ordem da função do elemento em questão é aumentada para reduzir o erro. A extensão r (de *relocação*) é empregada em um número fixo de nós e tenta mover os nós da malha para áreas que apresentem erros elevados, fazendo com que esta área fique mais refinada mas permanecendo com o mesmo número de elementos da malha original. Qualquer uma destas extensões pode ser também combinada em uma estratégia especial, como por exemplo, h - p , r - h , etc. A estratégia adaptativa utilizada neste trabalho foi a de extensão h .

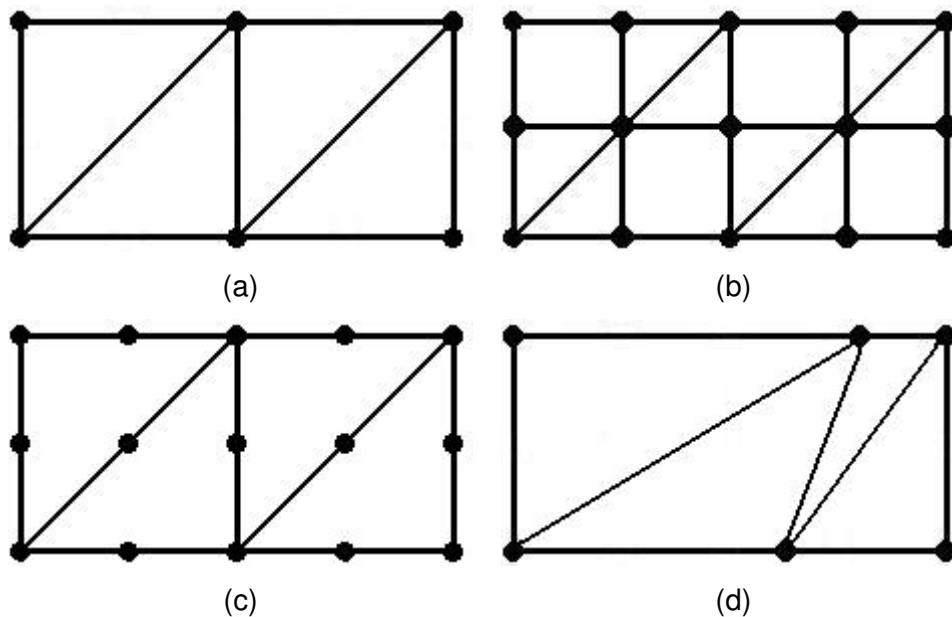


Figura 3 (a) malha original, (b) malha após o uso da estratégia do tipo h , (c) malha após o uso da estratégia do tipo p , (d) malha após o uso da estratégia do tipo r .

2.2. Estimativa de Erro

O conceito de adaptatividade e a própria essência do método dos elementos finitos estão intrinsicamente ligados à definição de erro, pois o método dos elementos finitos é um método numérico, ou seja, resolve problemas de forma aproximada e a idéia do uso da adaptatividade surgiu para melhorar as características das malhas que tivessem gerado resultados ruins na análise de elementos finitos. Além disso, é justamente a análise do erro que controla o critério de convergência do processo adaptativo.

Em geral existem dois tipos de estimativas de erro de discretização: *a-posteriori* e *a-priori* [12,57]. O que se procura, em estratégias *a-priori*, é tentar garantir de todas as formas que a malha gerada inicialmente já seja de boa qualidade, o que é válido e importante mas nem sempre é possível. Já a estratégia *a-posteriori* se preocupa em melhorar os resultados da análise de erro (qualidade da malha) a partir dos resultados de uma dada malha inicial. A estratégia utilizada neste trabalho é a *a-posteriori*, pois é a partir dos resultados da análise de elementos finitos da malha inicial que se obtém os parâmetros necessários para a geração da nova malha, resultando numa diminuição no valor do erro.

As análises de erros de discretizações têm se tornado cada vez mais importantes para aumentar a confiabilidade de modernos métodos numéricos aplicados a problemas de engenharia. Sempre que um método numérico, como o MEF, é usado para solucionar equações diferenciais relacionadas a um funcional definido para um problema discreto, erros são introduzidos pelo processo de discretização que reduz o modelo matemático contínuo para um modelo numérico discreto com um número finito de graus de liberdade.

Esta seção descreve de forma sucinta o procedimento normalmente adotado (Zienkiewicz [66]) para estimativa *a-posteriori* de erro de discretização em um modelo de elementos finitos.

Primeiramente será apresentada a formulação pelo MEF de um problema linear elástico com deslocamentos u definidos em um domínio Ω e com um contorno $\Gamma = \Gamma_u \cup \Gamma_t$, onde Γ_u é a parte do contorno em que os deslocamentos são prescritos e Γ_t é a parte do contorno em que os carregamentos são prescritos.

A equação abaixo apresenta a relação entre deslocamentos e deformações.

$$\varepsilon = Su \quad (1)$$

onde ε são as deformações, u são os deslocamentos e S é o operador diferencial que relaciona as deformações em função dos deslocamentos.

A equação abaixo apresenta a relação constitutiva.

$$\sigma = D\varepsilon \quad (2)$$

onde σ são as tensões e D é a matriz que define as tensões em função das deformações.

De posse das equações (1) e (2) pode-se colocar o problema segundo uma equação diferencial de equilíbrio a um carregamento q conforme a equação abaixo.

$$Lu - q \equiv S^T D S u - q = 0 \quad (3)$$

onde L é o operador diferencial auto-adjunto.

Fazendo-se uma transformação de coordenadas, pode-se obter uma relação entre as tensões σ em Ω com o carregamento t_p em Γ_t através da seguinte equação.

$$GDSu = t_p \quad (4)$$

onde G é um operador linear que relaciona o campo de tensões e as forças de superfície no contorno do domínio.

Numa aproximação por elementos finitos, tem-se

$$u \approx \hat{u} = N\tilde{u} \quad (5)$$

onde $\tilde{u}$ é a solução numérica nodal para o campo de deslocamentos, N são as funções de forma, que relacionam deslocamentos nodais com os deslocamentos em qualquer ponto do interior do domínio de um elemento e $\hat{u}$ é a solução, em termos de deslocamento, no interior do elemento finito. Obtém-se as equações de aproximação (por um processo padrão de Galerkin ou equivalentemente por minimização da energia potencial) para obter

$$K\tilde{u} - f = 0 \quad (6)$$

onde K é a matriz de rigidez e f é o vetor de forças nodais. Estes dois parâmetros são obtidos através das seguintes equações

$$K = \int_{\Omega} (SN)^T D(SN) d\Omega \quad (7)$$

$$f = \int_{\Omega} N^T q d\Omega + \int_{\Gamma_t} N^T t_p d\Gamma \quad (8)$$

Zienkiewicz [66] mostra em seu livro dois tipos distintos de estimadores de erro para problemas de elasticidade: os estimadores de erro baseados no processo de recuperação de *patches* e os fundamentados em resíduos. Mas antes de se elucidar as técnicas de recuperação utilizadas neste trabalho é necessário definir primeiramente o conceito de erro. Neste livro, Zienkiewicz, caracteriza o erro como sendo a diferença entre a solução exata e uma solução aproximada, conforme a equação abaixo.

$$e = u - \hat{u} \quad (9)$$

onde u é o resultado analítico e $\hat{u}$ é o resultado aproximado. Neste contexto u é análogo a uma resposta (por exemplo, deslocamentos) em um procedimento típico de solução numérica.

De uma forma similar pode-se analisar o erro em análise estrutural através de comparações entre deformações ou tensões, tendo-se as seguintes equações como resultado.

$$e_\varepsilon = \varepsilon - \hat{\varepsilon} \quad (10)$$

$$e_\sigma = \sigma - \hat{\sigma} \quad (11)$$

onde ε é o valor exato das deformações, $\hat{\varepsilon}$ é o valor aproximado, e_ε é o valor do erro para deformações, σ é o valor exato das tensões, $\hat{\sigma}$ é o valor aproximado e e_σ é o erro referente às tensões.

Zienkiewicz [66] comenta que estas definições de erro podem muitas vezes levar a resultados ruins, pois para pontos abaixo de cargas pontuais os valores das deformações e tensões tenderiam ao infinito. Situações similares ocorrem em problemas que apresentam cantos, onde, como é bem conhecido, ocorrem singularidades nas tensões em uma análise elástica. Por estas razões é que foram introduzidos vários tipos de ‘normas’ para avaliar o erro da análise de elementos finitos. Estas normas podem ser expressas analiticamente de uma das três formas abaixo.

$$\|e\| = \left(\int_{\Omega} (Se)^T D(Se) d\Omega \right)^{1/2} = \left(\int_{\Omega} (e_\varepsilon)^T D(e_\varepsilon) d\Omega \right)^{1/2} = \left(\int_{\Omega} (e_\sigma)^T D^{-1}(e_\sigma) d\Omega \right)^{1/2} \quad (12)$$

Uma medida mais direta é chamada de norma L_2 , que pode ser associada com erros em qualquer quantificação. Para as tensões, a norma L_2 para o erro é

$$\|e\|_{L_2} = \left(\int_{\Omega} (e_\sigma)^T (e_\sigma) d\Omega \right)^{1/2} \quad (13)$$

Embora estas normas estejam definidas para todo o domínio, nota-se que o quadrado de cada norma pode ser obtido somando as contribuições dos elementos de tal forma que se tenha

$$\|e\|^2 = \sum_{i=1}^m \|e\|_i^2 \quad (14)$$

onde i representa um elemento e m o número total de elementos do modelo. Em uma malha considerada “ótima”, tenta-se fazer com que as contribuições para este quadrado da norma sejam iguais para todos os elementos. Entretanto, o valor absoluto da norma de energia ou da norma L_2 tem pouco significado físico. Desta forma, é preferível adotar um erro relativo

$$\eta = \frac{\|e\|}{\|u\|} \quad (15)$$

que pode ser determinado para todo o domínio Ω ou para subdomínios de elementos, onde $\|u\|$ é o valor positivo da raiz quadrada do dobro da energia de deformação, podendo ser expressa através da seguinte equação

$$\|u\| = \left(\int_{\Omega} (\sigma)^T D^{-1}(\sigma) d\Omega \right)^{1/2} \quad (16)$$

O erro relativo da norma de energia (η) é usado em estratégias adaptativas para avaliar o erro de discretização do problema em questão e para guiar para um refinamento que melhore a qualidade da malha e conseqüentemente os valores dos resultados obtidos. Em uma estratégia adaptativa de refinamento baseada em uma extensão h , o estimador de erro irá definir como o modelo discreto será refinado ou “desrefinado”. Um critério simples para se atingir o erro da solução dentro de um nível aceitável para todo o domínio pode ser estabelecido da seguinte forma

$$\bar{\eta} \leq \eta^* \quad (17)$$

onde η^* é o erro máximo permitido e $\bar{\eta}$ é o erro corrente da análise obtido por

$$\bar{\eta} = \frac{\|\bar{e}\|}{\sqrt{\|\hat{u}\|^2 + \|\bar{e}\|^2}} \quad (18)$$

onde $\|\hat{u}\|$ é a norma de energia obtida da solução por elementos finitos.

Um critério razoável utilizado para uma “malha ótima” é forçar que a norma de energia do erro ($\|\bar{e}\|_i$) seja igualmente distribuída entre os elementos, desta forma para um dado elemento i , tem-se

$$\|\bar{e}\|_i \leq \left(\eta^* \left[\frac{\|\hat{u}\|^2 + \|\bar{e}\|^2}{m} \right]^{1/2} \right) = \bar{e}_m \quad (19)$$

onde m é o número total de elementos.

Pode-se então, desta forma, definir a razão de erro (ξ_i) como sendo

$$\xi_i = \frac{\|\bar{e}\|_i}{\bar{e}_m} \quad (20)$$

É evidente que o refinamento será necessário se

$$\xi_i > 1 \quad (21)$$

Um procedimento mais eficiente consiste em designar completamente uma nova malha (estratégia adaptativa de extensão do tipo h) que satisfaça o seguinte requisito

$$\xi_i \leq 1 \quad (22)$$

No limite do refinamento da malha, assumindo uma certa taxa de convergência [66], o valor da razão de erro ξ_i é então usado para determinar o novo tamanho dos elementos que serão gerados na nova malha conforme a equação abaixo

$$h = \frac{h_i}{\xi_i^{1/p}} \quad (23)$$

onde h_i é o tamanho inicial do elemento, h é o tamanho final e p é a ordem polinomial da função de aproximação.

2.3. Métodos de Suavização

Esta seção descreve os procedimentos de recuperação de *patches* usados para se obter o campo $\bar{\sigma}$, já que o campo de tensões exato σ geralmente é desconhecido. Inicialmente, é descrita a técnica da suavização por médias nodais *Z2-HC* [62] e em seguida é mostrada a técnica *SPR* [63-66] fundamentada em sistemas de mínimos quadrados ponderados.

Vale ressaltar que neste trabalho a parte de implementação referente aos métodos de suavização encontram-se no programa chamado *FEMOOP* (Martha *et al* [43]), que está implementado seguindo uma programação orientada a objetos em C++.

O *FEMOOP* [43] é o módulo de análise utilizado tanto em 2D quanto em 3D. No caso 2D, o modelador geométrico utilizado para criar o modelo geométrico, criar e aplicar os atributos de análise às entidades geométricas, gerar as malhas de elementos finitos e responsável pela visualização dos resultados é o *MTOOL* [47]. Já para o caso 3D, o modelador geométrico é o *MG* [48] e visualizador é o *POS3D* [16].

2.3.1. Zienkiewicz e Zhu – Hinton e Campbell (Z2-HC)

Este método consiste em estimar valores nodais contínuos para nós internos da malha, podendo ser utilizada uma suavização global ou local^{*}. Uma breve descrição do método é apresentada abaixo.

Após se obter os resultados das tensões nos pontos de Gauss, deseja-se estimar uma representação global das tensões, de preferência na forma de valores médios nodais. Portanto os valores de tensões obtidos nos pontos de Gauss devem ser extrapolados para os nós e suavizados, o que pode ser feito de forma global ou local.

Na suavização global de tensões, os valores nodais de tensões são obtidos de forma a minimizar o erro global (de todo o modelo) na avaliação das tensões nos pontos de Gauss.

O procedimento usual, no entanto, é uma suavização local de tensões. Neste caso os valores das tensões nos pontos de Gauss são extrapolados para os nós de cada elemento. Após este passo os valores das tensões para um determinado nó são diferentes em cada um dos elementos que contém este nó. O passo seguinte, portanto, é a suavização dos valores nodais, o que normalmente é feito através de uma média dos valores obtidos de cada elemento adjacente ao nó. No entanto há casos em que deve existir uma descontinuidade na tensão, como por exemplo na interface entre dois materiais diferentes, onde deve haver um valor médio de tensão em um nó de interface para cada grupo de elementos que forem do mesmo material.

A extrapolação dos valores nos pontos de Gauss para os nós pode ser feita de duas formas: ou os valores são interpolados ou eles são obtidos através de um ajuste usando mínimos quadrados.

A Figura 4 mostra de uma forma simplificada a obtenção dos valores nodais através da contribuição de todos os elementos adjacentes ao nó, conforme a idéia do método.

^{*} para mais detalhes ver [26]

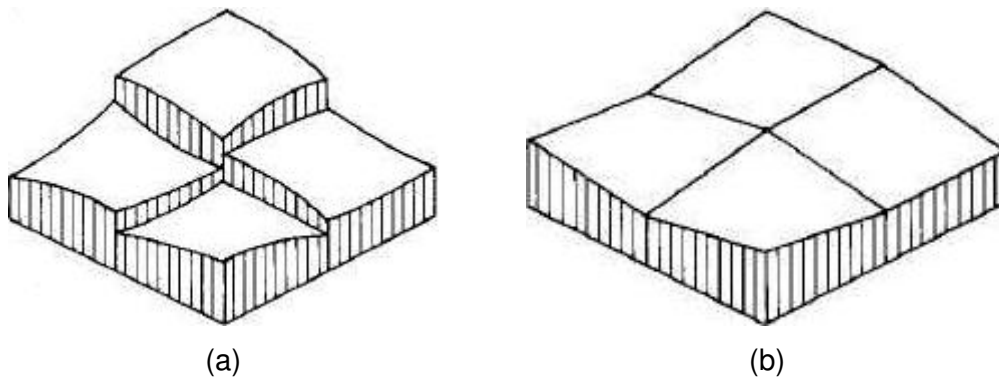


Figura 4 (a) mostra a contribuição de cada elemento, quadrilateral, para o valor de tensão, não suavizada, no nó adjacente aos mesmos, (b) mostra o valor da tensão suavizada no nó adjacente aos elementos quadrilaterais (Hinton & Campbell - [26]).

2.3.2. Superconvergent patch recovery (SPR)

Um campo genérico (por exemplo, tensão) pode ser aproximado pela seguinte expansão polinomial

$$\bar{\sigma} = Pa \quad (24)$$

onde P contém os termos polinomiais apropriados e a é um conjunto de parâmetros desconhecidos. Pode-se observar que esta expansão é usada para cada componente do tensor de tensões. Por exemplo, para problemas bidimensionais usando o elemento finito isoparamétrico quadrático de 8 nós (ver Figura 5.b) a seguinte aproximação é utilizada

$$P = [1, x, y, x^2, xy, y^2] \quad (25)$$

$$a = [a_0, a_1, a_2, a_3, a_4, a_5]^T \quad (26)$$

Os coeficientes desconhecidos a_i podem ser determinados por ajuste através do método dos mínimos quadrados ponderados da expansão polinomial (24) com relação aos valores de σ obtidos da solução de elementos finitos nos pontos de amostragem, isto é, $\hat{\sigma}$. Pequenos grupos (*patches*) de elementos são usados para realizar ajustes locais de mínimos quadrados e um peso (w_i) é considerado aqui para enfatizar a influência dos pontos de amostragem que estão mais próximos do nó que forma o *patch* (ver Figura 5.b). Desta forma, o peso (w_i) pode ser obtido de seguinte maneira

$$\varpi_i = \frac{1}{\rho_i^p}, \quad (27)$$

onde ρ_i é a distância Euclidiana entre o ponto de amostragem i e o nó que forma o *patch* e p é um inteiro. O uso de pesos maiores que zero pode ser efetivo para resolver problemas com gradientes bem elevados e também melhora o problema do mau condicionamento, que acontece em alguns problemas que usam o *SPR* [63-66].

Objetivando ilustrar a recuperação superconvergente discreta [66], considere um *patch* de elementos contendo m pontos de amostragem, como ilustrado na Figura 5, onde existe um ponto de amostragem i qualquer com coordenadas Cartesianas (x_i, y_i) em relação aos eixos globais. Assim sendo, o problema de mínimos quadrados ponderados é reduzido a minimização do seguinte funcional

$$\Lambda = \sum_{i=1}^m \varpi_i^2 [\hat{\sigma}(x_i, y_i) - \bar{\sigma}(x_i, y_i)]^2 \quad (28)$$

Substituindo a equação (24) na equação (28) obtém-se

$$\Lambda = \sum_{i=1}^m \varpi_i^2 [\hat{\sigma}(x_i, y_i) - P(x_i, y_i)a]^2 \quad (29)$$

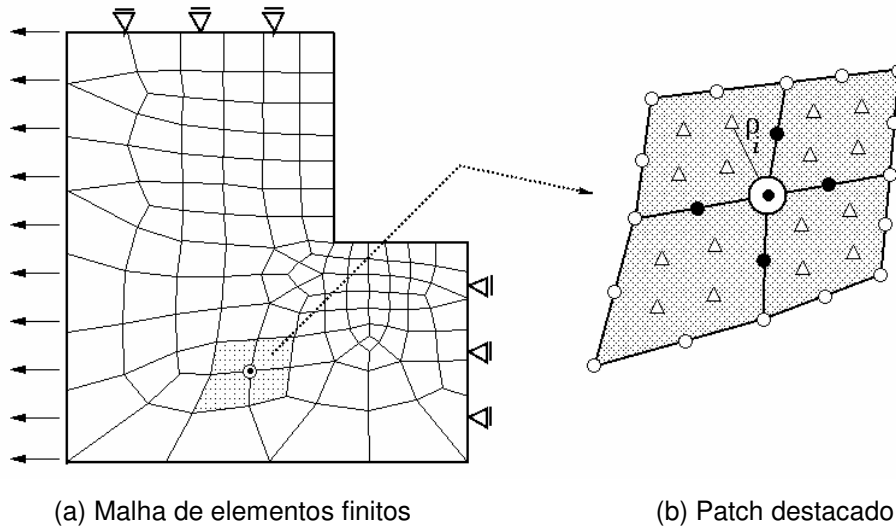


Figura 5 Notação da recuperação por patch: $\odot$ nó que forma o patch; Δ pontos de amostragem; $\bullet$ valores nodais determinados pelo procedimento de recuperação; $\circ$ pontos nodais; ρ_i distância entre o ponto de amostragem i e o nó que forma o patch; Ω denota o domínio do problema e Ω_p indica o domínio do patch (Cavalcante Neto J.B.-1998 [13]).

A Figura 5 mostra uma malha gerada em um exemplo clássico da teoria da elasticidade, destacando a formação do patch.

O problema de minimização é resolvido fazendo-se $\partial\Lambda/\partial a = 0$, levando ao seguinte conjunto de equações lineares

$$Aa = b \quad (30)$$

onde

$$A = \sum_{i=1}^m \varpi_i^2 P^T(x_i, y_i) P(x_i, y_i) \quad (31)$$

$$b = \sum_{i=1}^m \varpi_i^2 P^T(x_i, y_i) \hat{\sigma}(x_i, y_i) \quad (32)$$

Resolvendo-se o sistema de equações em (30) para a , pode-se substituir seu valor na equação (24) e obter o valor das tensões.

2.4. Geração de Malhas

Esta seção apresenta as idéias gerais das principais técnicas de geração de malhas de elementos finitos. A técnica de geração utilizada neste trabalho foi, fundamentalmente, a mesma técnica usada no trabalho de Cavalcante-Neto [13], que combina duas das técnicas apresentadas, tal como descrito no próximo capítulo.

Usualmente, as malhas são classificadas em três grupos principais, os quais são: estruturadas, não-estruturadas e as híbridas. Entretanto, não há um consenso quanto à definição de cada tipo. Uma das sugestões é fazer a diferenciação baseada na topologia da vizinhança dos elementos. Desta forma, as malhas estruturadas são caracterizadas por seus nós internos terem um número constante de elementos adjacentes (ver Figura 6.a). Já as malhas não estruturadas não apresentam número de elementos adjacentes aos nós internos de forma constante (ver Figura 6.b). As malhas híbridas, como o próprio nome já diz, apresentam em algumas regiões do seu domínio características das malhas estruturadas e em outras partes, características das malhas não-estruturadas (ver Figura 6.c).

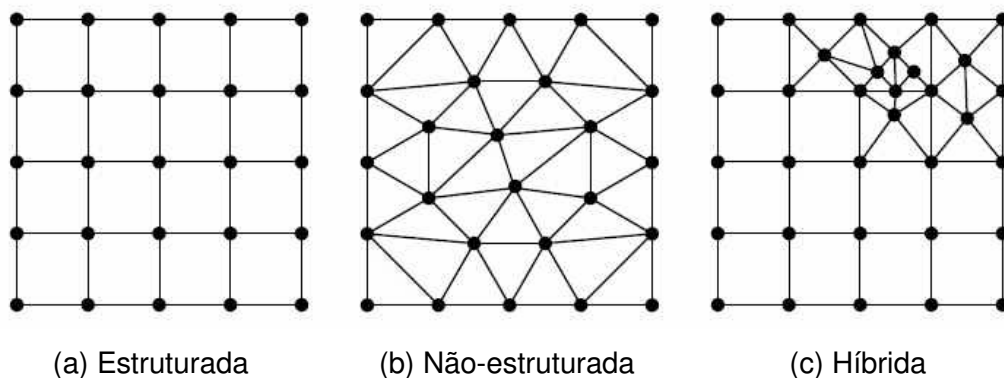


Figura 6 Tipos de malha (Batista V.H.F.-2005 [6]).

Malhas estruturadas discretizam o domínio em elementos possuindo conectividade implícita, ou seja, precisa-se apenas das coordenadas dos nós dos elementos para determinar todas as relações de conectividade existente, facilitando o processo de geração. Já nas malhas não-estruturadas as relações de conectividade entre os elementos têm que ser obtidas explicitamente através de uma tabela que especifique claramente a conectividade de cada elemento. Nas malhas híbridas a conectividade dependerá da combinação final entre as malhas estruturadas e as não-estruturadas.

As malhas não-estruturadas são preferíveis quando se tratar de discretização de domínios que apresentem geometrias arbitrárias (complexas), pois este tipo de malha se ajusta melhor ao contorno, permitem também a aplicação de refinamento local e adaptativo, o que é impossível em malhas estruturadas e muito difícil nas híbridas. As malhas híbridas também são usadas em domínios com geometria arbitrária, mas o processo de geração de malhas é mais complicado que o processo que usa malhas não-estruturadas.

Os elementos mais freqüentes em malhas estruturadas são quadriláteros em 2D e hexaedros em 3D. Nas malhas não-estruturadas são os triângulos e tetraedros, apesar de também ser possível utilizar elementos do tipo quadriláteros, prismas, pirâmides e hexaedros. Já nas malhas híbridas, podem ser usados elementos de formas variadas.

O uso preferencial de triângulos e tetraedros nas malhas não-estruturadas é justificado pelo fato de que estes elementos são capazes de se adaptar a contornos de domínios complexos e permitem suave transição de tamanho entre os elementos. Porém, a quantidade de elementos necessários para discretizar um domínio qualquer usando triângulos e tetraedros é maior quando comparado ao uso de quadriláteros e hexaedros. Apesar disso, a utilização de triângulos

para problemas 2D continuou sendo usada pelo fato de existirem inúmeros trabalhos de pesquisa eficientes tratando o processo de geração de malhas utilizando elementos triangulares [38]; para o caso 3D só em meados da década de 80 e início da, de 90 é que se começou a pesquisar sobre algoritmos de geração de malhas sólidas. Devido a este fato, hexaedros eram usados preferencialmente em detrimento aos tetraedros. Como resultado destas pesquisas já se têm um número razoável de algoritmos de geração de malhas 3D, porém poucos se destacam por possibilitar precisão, robustez, flexibilidade quanto às formas do domínio e elevado nível de automatização, além é óbvio, de resultar em malhas de “boa qualidade”. Carey [11] faz um resumo do avanço da pesquisa no estudo de técnicas de geração de malhas.

É de grande importância citar que esta seção do trabalho, foi baseada no trabalho de Batista [6], que apresentou uma enriquecedora revisão bibliográfica sobre as técnicas de geração de malhas, além de propor uma nova técnica para a geração de malhas não-estruturadas tetraédricas fundamentada em avanço de fronteira. Muitas referências citadas no trabalho de Batista [6] também foram pesquisadas como requisito básico para o entendimento do assunto abordado nesta seção, bem como para a elaboração deste trabalho.

2.4.1. Geração de malhas não-estruturadas

A capacidade de adequar-se a domínios arbitrários e a relativa facilidade de automatização impulsionaram a aplicação das malhas não-estruturadas na mecânica dos sólidos computacional, na dinâmica dos fluidos computacional e na geometria computacional. O grande interesse nesta área levou com que fossem desenvolvidos métodos que tratassem diferentes tipos de domínios (planos, superfícies tridimensionais e sólidos).

Os principais métodos de geração de malhas não-estruturadas são baseados na triangulação de Delaunay, Avanço de Fronteira e Decomposição Espacial Recursiva. Os métodos baseados na triangulação de Delaunay possuem um sólido embasamento matemático, mas apresentam dificuldades quando existem restrições tais como o contorno do domínio. Os métodos baseados no Avanço de Fronteira respeitam o contorno e têm um controle total da inserção dos nós internos. Já os métodos baseados na Decomposição Espacial Recursiva possibilitam acoplamento aos modeladores geométricos, aumentando o grau de automatização dos métodos.

2.4.1.1. Triangulação de Delaunay

A triangulação de Delaunay decompõe poliedros convexos gerados por conjuntos de pontos de modo único, satisfazendo ao critério da “esfera vazia”, que assegura que o círculo circunscrito a qualquer triângulo pertencente a triangulação não contenha pontos em seu interior (ver Figura 7).

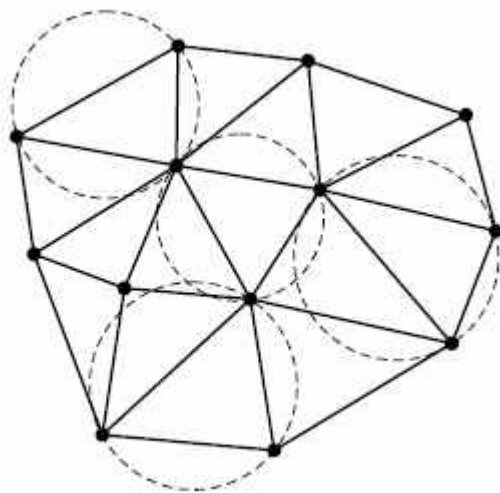


Figura 7 Malha de elementos finitos gerada por triangulação de Delaunay (Batista [6]).

Objetivando forçar a presença de arestas em triangulações de Delaunay, Chew [17] desenvolveu um algoritmo que constrói a chamada triangulação de Delaunay com restrição. A partir deste instante, diferentes algoritmos surgiram, servindo de base para novos algoritmos de geração de malhas.

Cavendish *et al* em 1985 [15] foram os primeiros a aplicarem a triangulação de Delaunay para a geração de malhas tridimensionais.

As principais características de geradores de malhas fundamentados na triangulação de Delaunay são: a produção de mais de um elemento por nó inserido, problemas de arredondamento no teste da “esfera vazia” e o surgimento de *slivers*, que são tetraedros cujos os quatro vértices estão praticamente no mesmo plano.

2.4.1.2. Decomposição Espacial Recursiva

O emprego desta técnica permite a representação de domínios com formas complexas de modo compacto e hierárquico. A idéia básica envolvida é similar a um dos paradigmas fundamentais na construção de algoritmos (*dividir para conquistar*). *Quadrees* e *octrees* são exemplos de estruturas baseadas neste conceito e que permitem a modelagem geométrica de domínios bi e tridimensionais. Yerry & Shephard [60,61] apresentaram um algoritmo para geração de malhas 2D, onde as malhas são geradas a partir de *quadrees*. O processo pode ser resumido da seguinte forma (vide Figura 8). Primeiramente define-se um quadrado englobando todo o domínio, que será a primeira célula da *quadtree*. Em seguida este quadrado é dividido em quatro partes iguais. Cada parte (nova célula) pode ser classificada da seguinte forma: cheia quando o quadrado estiver inteiramente dentro do domínio, vazia quando nenhuma parte do quadrado estiver dentro do domínio ou parcial quando pelo menos uma parte da célula estiver dentro do domínio. Cada célula nova é subdividida em outras quatro células, que também serão classificadas de acordo com o mesmo critério descrito anteriormente. Este procedimento continua até se atingir o nível de refinamento desejado, ou no caso adaptativo, até se atingir o erro requerido. Após se atingir este critério de parada da subdivisão, as células do tipo parcial são seccionadas para satisfazer o contorno e a *quadtree* é modificada para forçar um único nível de diferença entre células adjacentes, o que suaviza a transição entre os tamanhos dos elementos. Finalmente converte-se os quadrados em triângulos e ajusta-se os pontos do contorno. Estes ajustes são responsáveis pela distribuição irregular dos pontos ao longo do contorno, o que é uma das características do método.

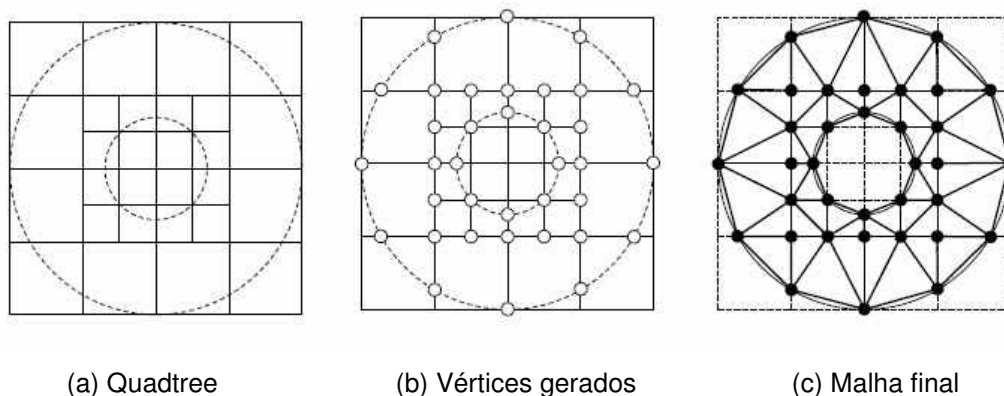


Figura 8 Exemplo de geração de uma malha 2D de elementos finitos usando uma *quadtree* (Batista [6]).

A Figura 9 mostra a *quadtree* em forma de diagrama de árvore referente à *quadtree* mostrada na Figura 8.a, com as células que estão parcialmente dentro do domínio, células que estão totalmente dentro e células que estão fora do domínio.

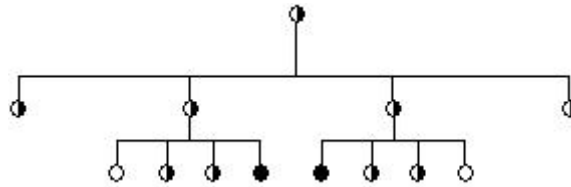


Figura 9 Estrutura *quadtree* em forma de diagrama.

2.4.1.3. Avanço de Fronteira

Esta técnica foi concebida no trabalho de Lo em 1985 [38] surgindo com o intuito de gerar malhas em domínios bidimensionais arbitrários. Inicialmente o contorno do domínio é representado pela união disjunta de arestas delimitando um polígono fechado simples. Os nós pertencentes ao contorno externo são ordenados em sentido anti-horário. Já os nós internos são numerados em sentido contrário. Devido a esta ordenação fica definido que a área a ser triangulada fica sempre à esquerda das arestas do domínio. Antes de se iniciar a discretização são calculadas as coordenadas dos nós internos do domínio obedecendo o espaçamento nodal estabelecido. A etapa de triangulação inicia com a definição de um fronte de geração exatamente igual ao contorno do domínio. Então selecionam-se dois nós do domínio que sejam capazes de gerar um elemento com cada aresta do fronte. O nó que resultar em um elemento com “melhor qualidade”, segundo o critério aplicado, será o nó escolhido. Desta forma o fronte vai avançando até que não haja mais aresta que não pertença a nenhum elemento. A Figura 10 ilustra a descrição do método feita acima.

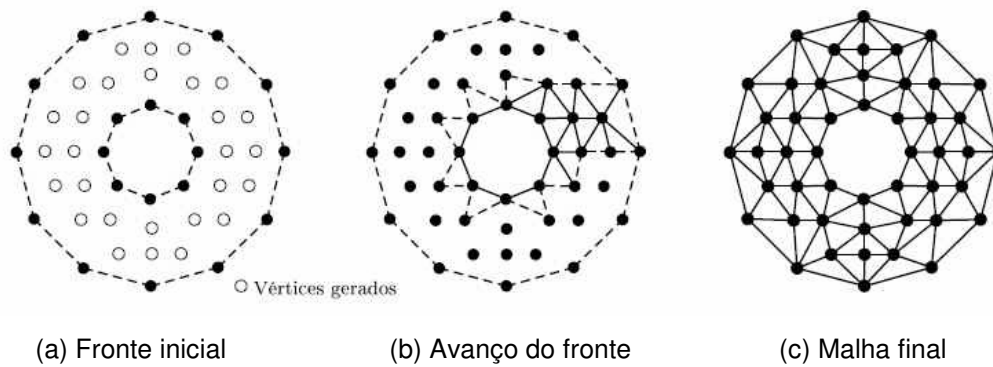


Figura 10 Exemplo de geração de uma malha 2D de elementos finitos usando avanço de fronteira (Batista V.H.F.-2005 [6]).

Apesar da concepção do algoritmo ser atribuída a Lo foi o algoritmo idealizado por Peraire *et al* em 1987 [50] que popularizou o método, pois a inovação proposta por Peraire *et al* permite maior controle na inserção de nós e o acoplamento a processos adaptativos. Nesse trabalho, Peraire *et al* montam uma lista com nós candidatos para formar o elemento, com um ponto ideal e com os nós vizinhos à aresta selecionada. O nó escolhido será aquele que formar um elemento, cujo o interior não contenha outro nó da lista dos candidatos e cuja as arestas não interceptem nenhuma aresta do fronte.