

ALGORITMOS RÁPIDOS DE CLASSIFICAÇÃO
DE IMAGENS MULTIESPECTRAIS

Carlos Roberto de Souza
Celso de Renna e Souza
Flávio Roberto Dias Velasco
Luciano Vieira Dutra
Nelson Delfino d'Ávila Mascarenhas
Nelson I. Tanaka
Renato L. Pedrosa
Ricardo C.M. de Souza

INSTITUTO DE PESQUISAS ESPACIAIS
Conselho Nacional de Desenvolvimento Científico e Tecnológico
Caixa Postal 515 - (12200) - São José dos Campos, SP - Telefone: 22-9977

SUMÁRIO

O presente trabalho divide-se em duas partes. Na primeira delas, é feita uma apresentação de caráter de revisão tutorial das técnicas propostas no INPE para implementação rápida de algoritmos classificatórios de imagens multiespectrais. São incluídos métodos baseados em transformação de matrizes e seleção da ordem de apresentação dos atributos no esquema de classificação de máxima verossimilhança. Abordam-se também as técnicas de utilização de tabelas nesse mesmo esquema. A segunda parte do trabalho apresenta uma comparação do desempenho de algumas dessas técnicas implementadas no sistema de classificação multi-espectral I-100 do INPE.

O presente trabalho divide-se em duas partes. Na primeira delas, é feita uma apresentação de caráter de revisão tutorial das técnicas propostas no INPE para implementação rápida de algoritmos classificatórios de imagens multiespectrais. São incluídos métodos baseados em transformação de matrizes e seleção da ordem de apresentação dos atributos no esquema de classificação de máxima verossimilhança. Abordam-se também as técnicas de utilização de tabelas nesse mesmo esquema.

A segunda parte do trabalho apresenta uma comparação do desempenho de algumas dessas técnicas implementadas no sistema de classificação multiespectral I-100 do INPE.

1. REVISÃO TUTORIAL

1.1 - INTRODUÇÃO

O grande volume de dados contido em imagens multiespectrais de recursos naturais, tipo LANDSAT, tem motivado um contínuo interesse no desenvolvimento de técnicas eficientes para a classificação dessas imagens. Este trabalho visa dois objetivos: 1) fazer uma apresentação dos principais métodos propostos no INPE para o aumento da eficiência dos algoritmos de classificação e 2) apresentar uma comparação entre algumas dessas técnicas quando implementadas no sistema de classificação multiespectral I-100 do Instituto de Pesquisas Espaciais.

1.2 - ESQUEMAS BASEADOS EM TRANSFORMAÇÃO DE MATRIZES

Um dos métodos mais comumente empregados na classificação de recursos naturais por técnicas de reconhecimento de padrões baseia-se na hipótese de distribuição multivariada gaussiana. Utilizando-se o critério de máxima verossimilhança, o problema consiste em, dada um vetor x a ser classificado, se escolher a classe j para a qual a função densidade probabilidade

$$f_j(\underline{x}) = \frac{1}{(2\pi)^{N/2} |\underline{C}_j|^{1/2}} e^{-\frac{1}{2} (\underline{x} - \underline{\eta}_j)^T \underline{C}_j^{-1} (\underline{x} - \underline{\eta}_j)}, \quad (1)$$

onde $\underline{C}_j$ é a matriz de covariância da classe e $\underline{\eta}_j$ sua média, é máxima.

É possível tomar-se o logaritmo de $f_j(\underline{x})$ e então a decisão final assume a forma:

$$\text{se } d_j(\underline{x}) < d_i(\underline{x}) + c_{ij} \quad (2)$$

então para todo $j \neq i$, classifica-se $\underline{x}$ na classe j , onde

$$d_j(\underline{x}) = (\underline{x} - \underline{\eta}_j)^T \underline{C}_j^{-1} (\underline{x} - \underline{\eta}_j) \quad (3)$$

$$\text{e } c_{ij} = \ln(|\underline{C}_i|) - \ln(|\underline{C}_j|) \quad (4)$$

A primeira técnica para aumentar a eficiência do cômputo dessas funções de verossimilhança, sem alterar o resultado de classificação, consiste em se efetuar uma rotação dos eixos do espaço de atributos por meio da chamada transformação de Karhunen - Loeve ou Hotelling, que é dada por

$$\underline{C}'_j = \underline{V}_j^T \underline{A} \underline{V}_j \quad (5)$$

onde $\underline{V}_j$ é uma matriz unitária cujas linhas $\underline{v}_{j\ell}^T$ são os vetores prôprios da matriz de covariância $\underline{C}_j$; e $\underline{A}$ é uma matriz diagonal cujos elementos $\sigma_{j\ell}^2$ são os valores próprios da matriz $\underline{C}_j$.

Feita essa transformação, a expressão de $d_j(\underline{x})$ pode ser computada por

$$d_j(\underline{x}) = \sum_{\ell=1}^N \frac{y_{j\ell}^2}{\sigma_{j\ell}^2} \quad (6)$$

onde $y_{j\ell} = v_{j\ell}^T (\underline{x} - \underline{n}_j)$ (7)

Definindo-se $y_{j\ell}^{2'} = \frac{y_{j\ell}^2}{\sigma_{j\ell}^2}$ (8)

o problema de classificação pode ser reduzido, a menos da constante c_{ij} em (2), a se computar a mínima distância euclidiana no espaço y' .

Da regra de decisão dada por (2), se $d_i(\underline{x}) > d_j(\underline{x}) - c_{ij}$ então a classe i não é selecionada, ou é "eliminada", o que ocorre quando, separando-se a somatória em (6) em duas partes, se tem:

$$\sum_{\ell=1}^n \frac{y_{i\ell}^2}{\sigma_{i\ell}^2} + \sum_{\ell=n+1}^N \frac{y_{i\ell}^2}{\sigma_{i\ell}^2} > d_j(\underline{x}) - c_{ij} \quad (9)$$

onde $1 \leq n < N$.

Tendo em vista o fato de que a segunda somatória é sempre positiva, pode-se eliminar a classe i assim que:

$$\bar{d}_i = \sum_{\ell=1}^n \frac{y_{i\ell}^2}{\sigma_{ji}^2} > d_j(\underline{x}) - c_{ij} \quad (10)$$

Esta formulação do problema de classificação por máxima verossimilhança, proposta por Dye [1], permite cessar o cômputo das parcelas em (9) logo que for atingido um índice n tal que (10) seja verificado. As classes podem ser testadas em ordem decrescente de probabilidades "a priori". Pode-se tomar como candidata inicial a classe do "pixel" adjacente. A eficiência do algoritmo decorre de se usar, somente, tantos canais quanto necessários para realizar a discriminação, sem contudo resultar em perda alguma de informação.

Uma modificação no esquema anterior foi proposta por Eppler [2] na qual $\hat{d}_i$ tende a ser aumentado pelo maior incremento possível $\frac{y_{i\ell}^2}{\sigma_{i\ell}^2}$ para cada valor de $1 \leq \ell \leq N$. Tomando-se o valor esperado do ℓ ésimo incremento; dada a classe j , é possível mostrar que

$$\left\langle \frac{y_{i\ell}^2}{\sigma_{i\ell}^2} \right\rangle_j = \frac{\underline{v}_{i\ell}^T}{\sigma_{i\ell}} \left[\underline{C}_j + (\underline{n}_j - \underline{n}_i) (\underline{n}_j - \underline{n}_i)^T \right] \frac{\underline{v}_{i\ell}}{\sigma_{i\ell}} \quad (11)$$

Para valores especificados de i e j , pode-se calcular a expressão (11) para $1 \leq \ell \leq N$ e ordenar os vetores $\frac{\underline{v}_{i\ell}^T}{\sigma_{i\ell}}$ por valores decrescentes de $\left\langle \frac{y_{i\ell}^2}{\sigma_{i\ell}^2} \right\rangle_j$

Uma segunda técnica, baseada em transformação de matrizes, utiliza a decomposição de Cholesky de uma matriz simétrica no produto de uma matriz triangular inferior e de uma matriz triangular superior, isto é,

$$\underline{C}_j^{-1} = \underline{L}_j^T \underline{L}_j \quad (12)$$

Usando-se essa transformação é possível expressar (3) em termos de

$$d_j(x) = \left[\underline{L}_j(\underline{x} - \underline{n}_j) \right]^T \left[\underline{L}_j(\underline{x} - \underline{n}_j) \right] \quad (13)$$

e definindo-se

$$\hat{y}_{j\ell} = \underline{\hat{v}}_{j\ell}^T (\underline{x} - \underline{n}_j) \quad (14)$$

onde $\underline{\hat{v}}_{j\ell}^T$ é a ℓ ésima linha da matriz $\underline{L}_j$, pode-se calcular $d_j(x)$ por

$$d_j(\underline{x}) = \sum_{\ell=1}^N \hat{y}_{j\ell}^2 \quad (15)$$

Do mesmo modo que em (9), pode-se aqui subdividir a so matória em duas partes e cessar o processo de adição pela mesma técni ca de (10).

Eppler [2] propos também uma ordenação dos atributos a serem encaminhados, de modo a que o valor esperado do incremento $\hat{y}_{i\ell}^2$ dada a classe j , $\langle \hat{y}_{i\ell}^2 \rangle_j$, dado por

$$\langle \hat{y}_{i\ell}^2 \rangle_j = \hat{\underline{v}}_{i\ell}^T \left[\underline{C}_j + (\underline{\eta}_j - \underline{\eta}_i) (\underline{\eta}_j - \underline{\eta}_i)^T \right] \hat{\underline{v}}_{i\ell} \quad (16)$$

seja ordenado em ordem decendente em ℓ , implicando num aumento de $d_j(\underline{x})$ o mais rápido possível, de modo que uma hipótese incorreta possa ser descartada com o menor número de termos.

1.3 - ESQUEMAS BASEADOS EM USO DE TABELAS

O esforço de computar funções de verossimilhança, para cada "pixel", pode ser minimizado adotando-se a técnica de armazenar os resultados de classificações anteriores para cada um dos possíveis valores que podem ser assumidos por esse elemento.

O problema de armazenamento pode se complicar seriamen te ao se tentar preservar o número total de tons de cinza que os "pixels" de rastreadores multiespectrais normalmente apresentam.

Considerando-se, por exemplo, imagens do satélite LANDSAT A e B, com 64 níveis e 4 canais, haveria necessidade de se armazenar 64^4 possíveis vetores, dando aproximadamente 17 milhões

de posições de memória, o que está certamente além da capacidade das memórias de acesso rápido atualmente existentes.

Uma das possíveis tentativas de se minorar esse problema, consiste em se reduzir o número de tons de cinza por uma maior quantização desses tons e, simultaneamente, reduzir o número de canais examinados. Esse método foi usado por Brooner e outros [4] reduzindo de 64 para 10 o número de tons de cinza e considerando três canais de imagens multibanda, contendo verde, vermelho e infravermelho. Isso diminui o número de posições de memória para $10^3 = 1000$, o que é perfeitamente possível de ser armazenado na memória principal de computadores de pequeno porte.

Reconhecendo o efeito negativo da quantização sobre a classificação, Eppler [5] propôs um método com as seguintes características:

- 1) Seleção dos canais mais adequados para cada classe. Esses canais são escolhidos de modo a maximizar a mínima divergência entre a categoria em questão e todas as outras categorias;
- 2) O armazenamento nesses canais apenas das regiões de interesse para cada categoria, isto é, regiões retangulares que contêm a assinatura da classe na fase de treinamento;
- 3) Compressão da região de interesse pela redução dos níveis de quantização.

O processo é implementado pela utilização de duas tabelas. Na primeira delas, são armazenadas as regiões de interesse para cada classe do espaço espectral, utilizando indicadores que variam de 1 a 12 na região e zero fora dela. Saindo da primeira tabela, são utilizados quatro indicadores que apontam para uma posição de memória onde cada "bit" corresponde a uma classe.

Na implementação de Eppler, feita num IBM 360/44, a primeira tabela ocupa 4 (nº de canais) x 9 (nº de categorias) x 255 (nº de níveis de quantização) meias palavras (de 16 "bits"). A segunda tabela ocupa 13^4 meias palavras (a que permitiria a classificação de até 16 categorias). O total da memória necessária é portanto $9216 + 28561 = 37777$ meias palavras de 16 "bits", ou cerca de 20.000 palavras.

É conveniente, neste ponto, mencionar uma outra técnica proposta por Eppler [6], baseada no armazenamento, em tabelas, dos limites de cada região atribuída a uma classe, no espaço de medidas, satisfazendo uma condição de regularidade.

Um conjunto R no espaço E^n , é dito ser regular se e somente se existem constantes L_1 e H_1 e funções $L_2, L_3, \dots, L_N, H_2, H_3, \dots, H_N$, tais que

$$R = \{(x_1, \dots, x_N) \in E^n \mid L_1 \leq x_1 \leq H_1, L_2(x_1) \leq x_2 \leq H_2(x_1), \\ \vdots \\ L_N(x_1, x_2, \dots, x_{N-1}) \leq x_N \leq H_N(x_1, x_2, \dots, x_{N-1})\}$$

A redução do esforço computacional, neste tipo de tabela, é obtida armazenando-se os limites das regiões em tabelas unidimensionais, evitando-se assim o uso de tabelas de dimensões três ou quatro. A desvantagem do método consiste na hipótese de regularidade das regiões consideradas, a qual impede que se considere todas as possibilidades da regra de máxima verossimilhança baseada na hipótese gaussiana.

A Figura 1 mostra um exemplo de espaço de atributos bidimensional (com 2 classes sob hipótese gaussiana e utilizando um limiar de rejeição) em que a hipótese de regularidade não é satisfeita.

O esquema de tabela proposta por Schlien [7] baseia-se no fato de imagens multiespectrais de satélite LANDSAT apresentarem no máximo cerca de 10.000 vetores diferentes, dos 16 milhões possíveis, (para 64 níveis) devido à alta correlação espectral e espacial das amostras. Como resultado, é possível iniciar o processo computando as funções de verossimilhança, armazenando os resultados numa tabela e, então, se um mesmo valor do "pixel" ocorrer novamente, não é necessário repetir a computação mas apenas acessar o resultado já armazenado na tabela. O método escolhido para a estruturação da tabela foi o de "hashing", devido ao acesso relativamente fácil, requerimentos modestos de espaços de armazenagem e possibilidade de manipular eficientemente um número crescente de vetores.

O trabalho de Schlien apresenta também uma previsão do número de diferentes vetores que uma imagem do LANDSAT apresenta, através de um modelo gaussiano. Os contornos da função densidade de probabilidade gaussiana constante são obtidos igualando-se a expressão da distância de Mahalanobis a uma constante

$$(\underline{x} - \underline{n})^T \underline{C}^{-1} (\underline{x} - \underline{n}) = K^2 \quad (17)$$

Essa expressão é uma forma quadrática e define a superfície de um hiperelipsóide, que, em duas dimensões, se reduz a uma elipse. Aumentando-se o valor de K^2 , o volume do hiperelipsóide aumenta, i.e., diminui a probabilidade α de se observar um vetor fora dessa região.

A variável aleatória definida pela expressão (17) tem uma distribuição do tipo qui-quadrado, com N graus de liberdade, onde N é a dimensão do vetor de atributos. Através dessa distribuição é possível calcular, para um valor de K , a probabilidade α de um vetor cair fora do hiperelipsóide. Por outro lado, o volume do hiperelipsóide, arredondado para o inteiro mais próximo, dá aproximadamente o número de vetores desse hiperelipsóide.

Pode-se provar que esse volume é dado por

$$V = \frac{(\sqrt{\pi K})^N}{\gamma(N/2 + 1)} |\underline{C}|^{1/2} \quad (18)$$

onde γ denota a função gamma. Deve-se observar que esse volume é proporcional à raiz quadrada do determinante da matriz de covariância da distribuição. Tal determinante pode ser dado pelo produto das variâncias dos canais individuais, de modo que, quanto maiores essas variâncias, maior será o número de vetores e maior será o tamanho da tabela necessário.

As técnicas de obtenção de algoritmos rápidos, usando tabelas, podem ser divididas em dois tipos, conforme mostra a Figura 2. O primeiro tipo (Figura 2a) consiste em se construir a tabela após o treinamento e posteriormente utilizá-la na classificação. No segundo tipo (Figura 2b), a tabela vai sendo elaborada durante a própria classificação, sendo acessada para cada novo ponto, para se determinar se ele já está classificado, caso contrário classificando-o e guardando-o na tabela. Os métodos de Brooner e outros [4] e Eppler [5 e 6] devem ser considerados como pertencendo ao tipo a) enquanto que o método proposto por Schlien [7] enquadra-se no tipo b).

2. IMPLEMENTAÇÃO DOS ALGORITMOS NO INPE

2.1 - TRANSFORMAÇÕES DE KARHUNEN-LOEVE E DE CHOLSKY

Foi feito um trabalho de simulação comparativa dos algoritmos de classificação baseados nas transformações de Karhunen-Loeve (algoritmo I) e de Cholesky (algoritmo II). Foram utilizadas 8 classes, num espaço quadridimensional, envolvendo 3000 pontos.

O método convencional de máxima verossimilhança exigiria o cômputo de $8 \times 4 \times 3000 = 96000$ eixos. Utilizando-se tanto o algorit

mo I como o algoritmo II, o número mínimo de eixos a serem pesquisados é dado por $8 \times 3000 = 24000$ eixos, mas isso equivale à situação irreal de todos os pontos serem rejeitados. Admitindo-se que todos os pontos pertençam a uma classe, o número de eixos é dado por $7 \times 3000 + 1 \times 4 \times 3000 = 33000$ eixos. A tabela abaixo, retirada da ref. [3], mostra os resultados da execução dos algoritmos de classificação no computador B - 6700 em linguagem ALGOL.

MÉTODO	TEMPO DE PROCESSAMENTO - (SEGUNDOS)	NÚMERO TOTAL DE EIXOS
Algoritmo I	15.1	45 735
Algoritmo II	12.6	62 778

Portanto, o algoritmo I reduz a cerca de metade o número de eixos pesquisados. No método II, embora o número de eixos seja bem maior, o tempo de processamento fica reduzido. Isso se deve basicamente ao fato de que a operação dada por (14) é mais rápida do que aquela dada por (7), tendo em vista a estrutura da matriz triangular inferior.

O mesmo tipo de experimento foi feito utilizando-se o computador PDP - 11/45 do Sistema I - 100 do INPE em linguagem Assembler. Os novos resultados vieram a confirmar os anteriores, pois com o algoritmo I (Karhunen - Loeve) foi utilizada uma média de 2 eixos por "pixel" e por classe, processando 3000 "pixels" em 4,7 seg. enquanto que para o algoritmo II (Cholesky) esses valores foram 2,5 eixos por "pixel" e por classe e 2,5 seg.

2.2 - MÉTODOS DAS TABELAS

Um método semelhante àquele proposto por Brooner e outros [4] foi implementado no PDP-11/45, utilizando 64 níveis nas

imagens do LANDSAT com 2 canais, e ocupando 4K bytes de 32 níveis, e 3 canais e utilizando 32 K bytes.

Para a construção da tabela foi empregado o método de classificação por máxima verossimilhança com hipótese gaussiana e usando a transformação de Cholesky. A Figura 3 mostra o espaço de atributos bidimensional resultante, considerando-se 2 classes enquanto que a Figura 4 apresenta o correspondente histograma bidimensional de treinamento.

O tempo de classificação de imagem de 512 x 512 "pixels", em 2 canais e 64 níveis foi de 40 seg. e praticamente independe de número de classes, enquanto que o tempo de construção da tabela para 8 classes foi de 4 seg.

O método de Eppler [5] foi também implementado, permitindo a classificação de 8 categorias e usando indicadores que variavam de 0 a 10.

Na construção da segunda tabela foi, igualmente, usado o método de máxima verossimilhança com hipótese gaussiana. O tempo de construção desta tabela é muito variável, uma vez que ele depende do número de pontos da menor região que engloba todas as classes. Essa região depende, por sua vez, da variância das classes. Para imagens normais, onde não foram efetuadas transformações sobre os diversos canais, os valores típicos variaram de 30 seg. a 3 min. para 8 classes; para uma imagem onde 3 canais haviam sofrido um alongamento ("stretch") do histograma com um conseqüente aumento da variância, a elaboração da tabela consumiu cerca de 10 min. de processamento, também para 8 classes.

A Figura 5 mostra o gráfico do tempo de classificação de parte da imagem com 64 níveis de cinza, como função do número de classes. Considerando que só em entrada e saída são gastos cerca de 20 seg., passando-se de 1 para 8 classes, tem-se apenas o dobro de tempo de processamento de algoritmo.

Para o sistema MAXVER, a Figura 6 apresenta o tempo de classificação da imagem toda (25000 "pixels"), com 256 níveis de cinza e 4 canais, enquanto que a Figura 7 apresenta o mesmo resultado para 2 canais. Observa-se que o aumento do tempo de processamento com o número de classes é maior no algoritmo do MAXVER do que no de EPPLER.

Esses resultados do sistema MAXVER foram obtidos fixando-se o limiar de rejeição igual a 5. Diminuindo-se esse limiar, há uma tendência a cair o tempo de processamento, uma vez que as classes tendem a ser rejeitadas mais rapidamente, caindo, por isso, o número de eixos pesquisados.

Na implementação do algoritmo proposto por Schlien, foram utilizados 128 níveis de cinza, ocupando assim 7 "bits" de memória. O uso de quatro canais exigiu a ocupação de duas palavras de 16 "bits" do PDP - 11/45. O problema de colisões em esquemas de estrutura de dados em "hashing", é crítico e, à medida que a tabela vai sendo ocupada, aumenta a tendência às colisões. Para evitar que o número de colisões se tornasse excessivo, foi construída uma tabela 74,34% cheia. Contou-se o número de tentativas necessárias para localizar um determinado vetor. Mais da metade dos vetores foram acessados sem colisão. Após 5 tentativas, apenas 350 de 5955 vetores não foram encontrados. O maior número de tentativas foi 80. Em média, foram necessárias 2,21 tentativas por vetor.

Utilizando uma tabela com 5100 posições, foi classificada uma imagem da região de Ribeirão Preto, relativamente homogênea, com agricultura e cidade, onde foram definidas 6 classes. O sistema MAXVER precisou de, aproximadamente, 6 min. 00 seg., para classificar a tela toda do sistema I-100, enquanto que o esquema utilizando "hashing" (MAXHAS) necessitou de apenas 2 min. e 30 seg.

AGRADECIMENTOS

Os autores desejam agradecer a Lucila O.C. Prado pela colaboração prestada no desenvolvimento do sistema MAXVER.

REFERÊNCIAS

- [1] DYE, R.H., "Multivariate Categorical Analysis - Bendix Style", BSR 4149, June 1974, Bendix Corporation, Ann Harbor, Michigan.
- [2] EPPLER, W.G., "Canonical Analysis for Increased Classification Speed and Channel Selection", Symposium on Machine Classification of Remotely Sensed Data, Laboratory for Applications of Remote Sensing, Purdue University, West Lafayette, Indiana, June 1975 e IEEE Transactions on Geoscience Electronics, vol. GE - 14, Nº 1, January 1976, pp. 26-33.
- [3] VELASCO, F.R.D., "Proposta de Desenvolvimento de um Classificador de Imagens por Máxima Verossimilhança", Relatório INPE 1126-PPR/027, Setembro de 1977.
- [4] BROONER, W.G.; HARALICH, R.M.; DINSTEIN, I., "Spectral Parameters Affecting Automated Image Interpretation Using Bayesian Probability Techniques", Proc. 7th International Symposium on Remote Sensing of Environment (University of Michigan, Ann Harbor) May 1971, pp. 1929-1949.
- [5] EPPLER, W.G.; HELMKE, C.A.; EVANS, R.M., "Table Look-Up Approach to Pattern Recognition", Ibid pp. 1415-1425.
- [6] EPPLER, W.G., "An Improved Version of the Table Look-Up Algorithm for Pattern Recognition", Proc. 9th International Symposium on Remote Sensing of Environment (University of Michigan, Ann Harbor), April 1974, pp. 793-812.
- [7] SHLIEN, S., "Practical Aspects Related to Automated Classification of LANDSAT Imagery Using Look-Up Tables", Research Report 75-2, Canada Centre for Remote Sensing, Ottawa, Canada, February 1975.

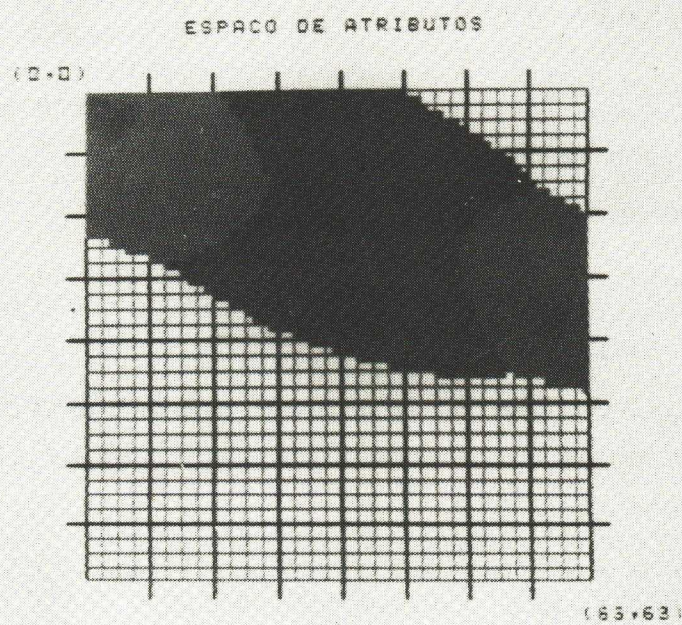


Fig. 1 - Exemplo de Regiões Não Regulares

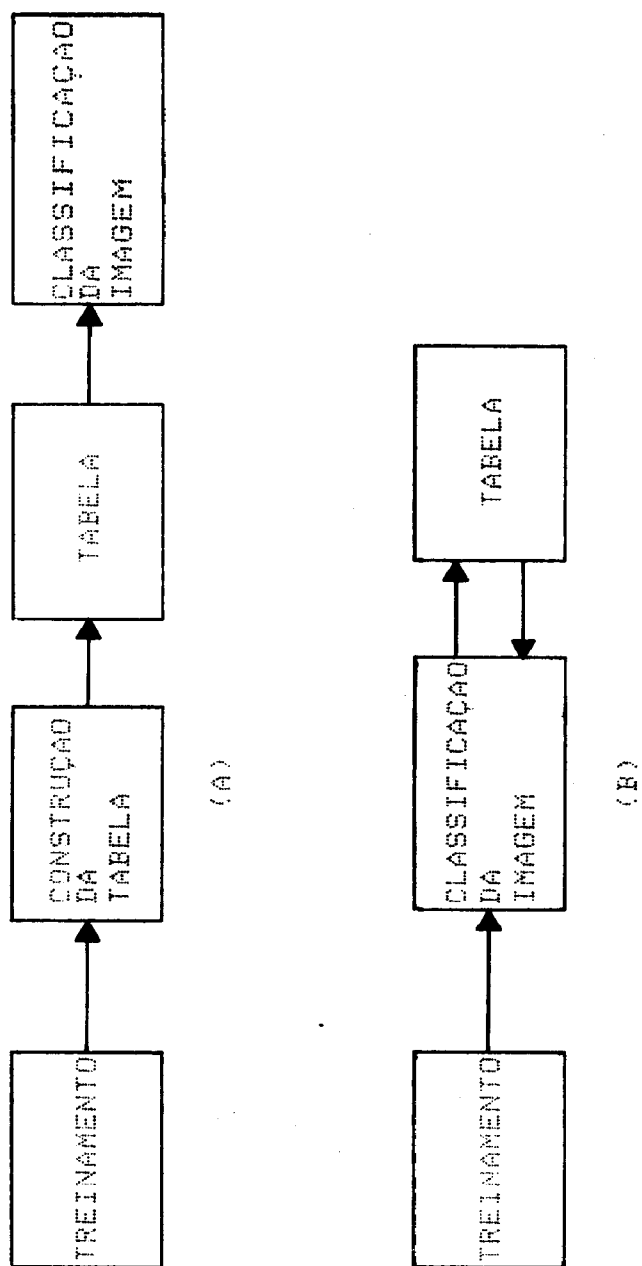


Fig. 2 - Esquemas de Classificação por Tabela

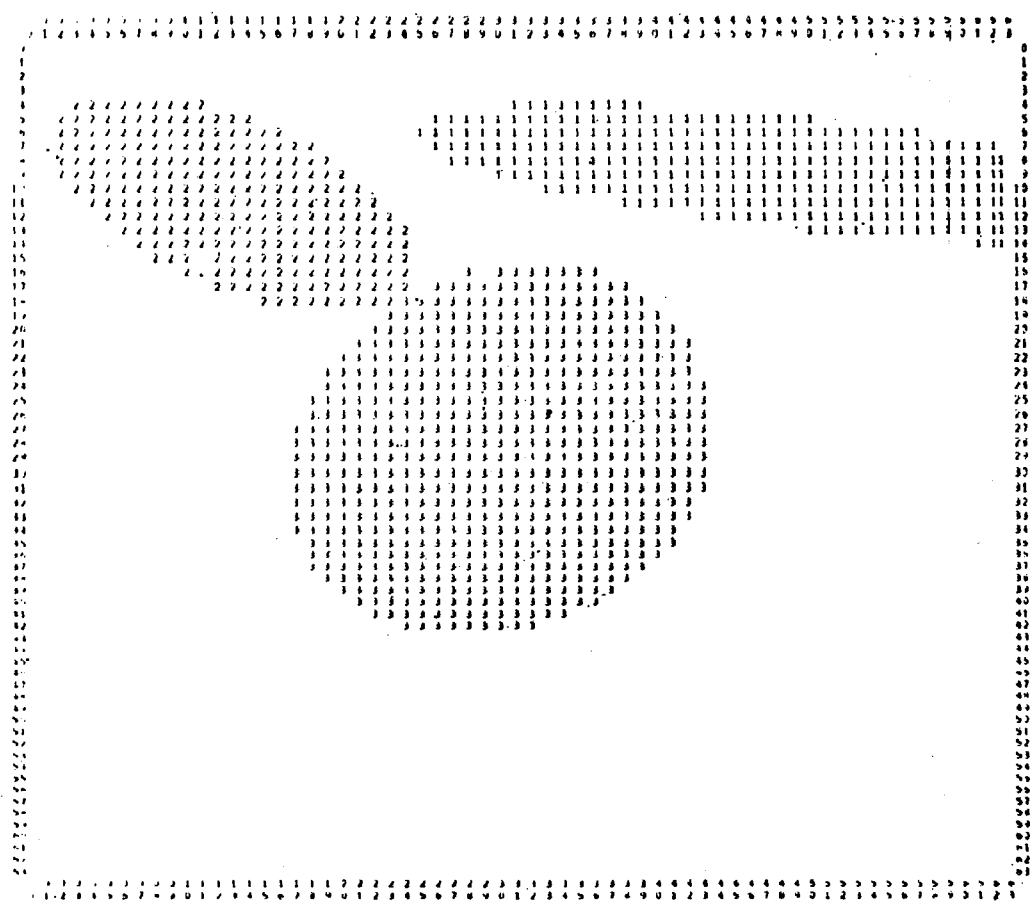


Fig. 3 - Espaço de Atributos Bidimensional

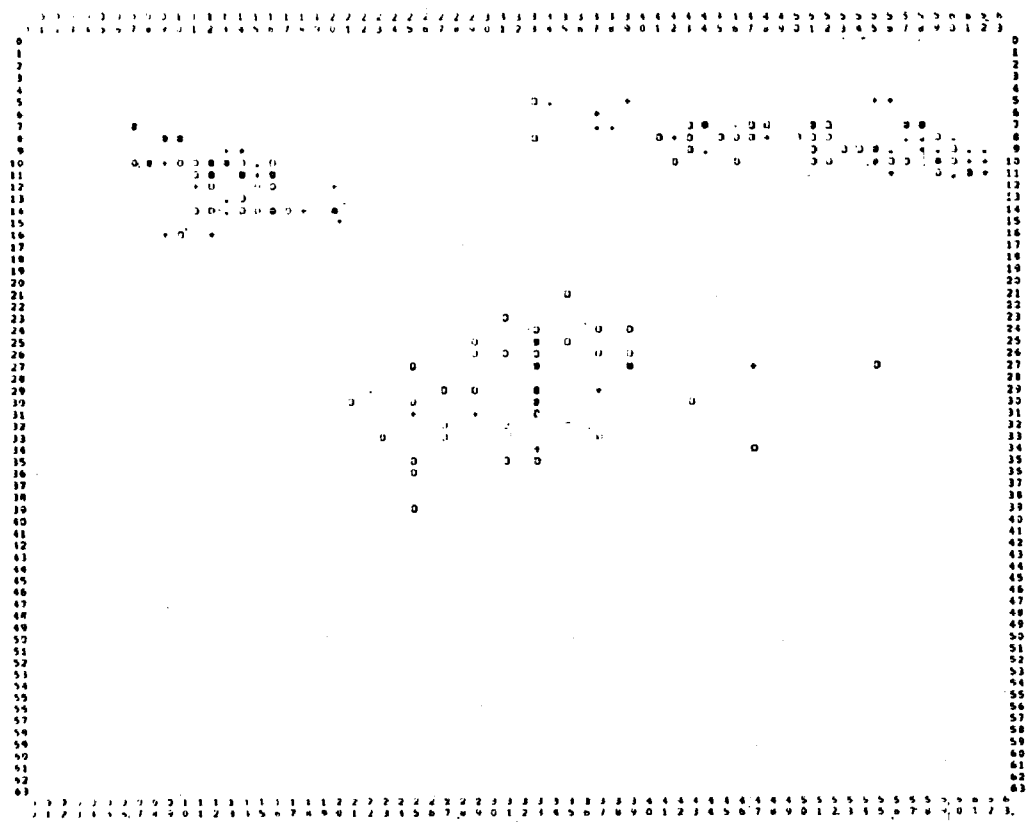


Fig. 4 - Histograma Bidimensional

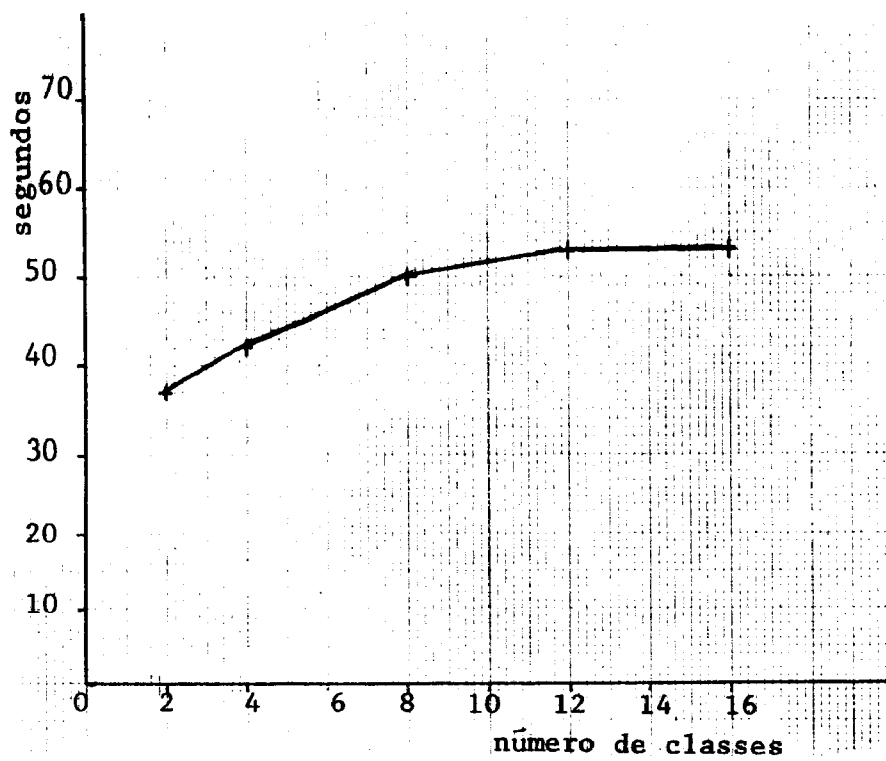


Fig. 5 - Tempo de Classificação de Parte da
Imagem no Esquema de 2 Tabelas

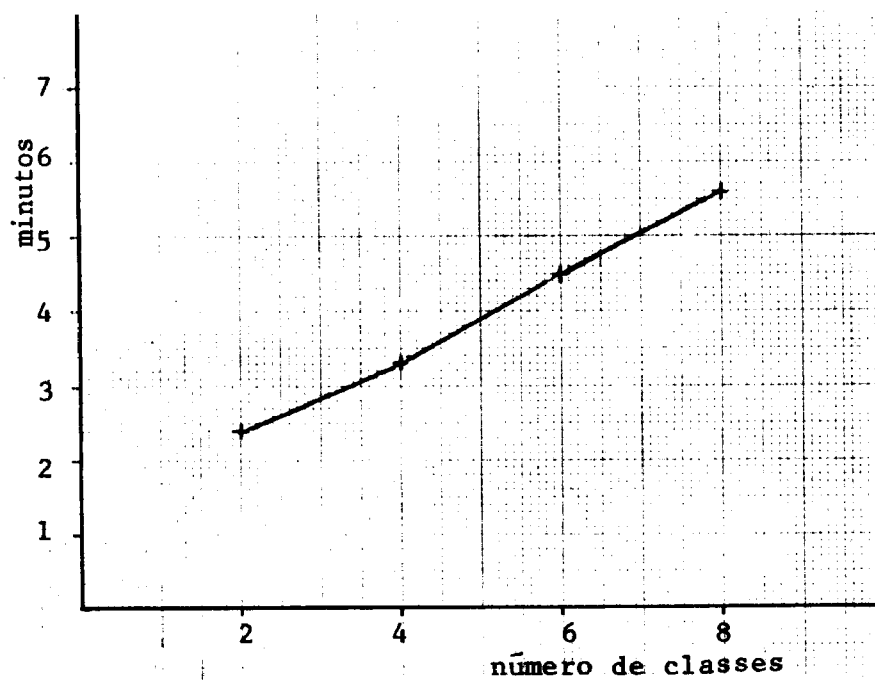


Fig. 6 - Tempo de Classificação de Toda a Imagem
no Sistema MAXVER, com 4 Canais

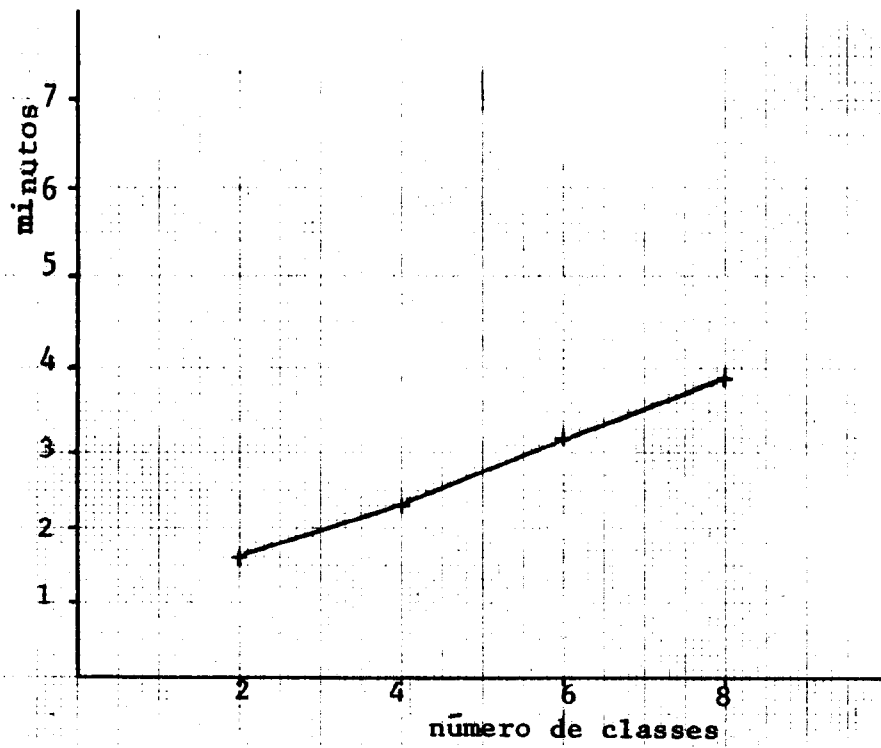


Fig. 7 - Tempo de Classificação de Toda a Imagem
no Sistema MAXVER com 2 Canais