

PI(INPE-1678-RPE/111)

P.I.

-a ed. 1980



11.594

INPE - Bibl.
Central

74944

Ex. 1 TAKAHASHI, W. K.

PARTICIONAMENTO DE IMAGENS MULTIESPECTRAIS DE
RECURSOS NATURAIS DA TERRA.

1. Classificação INPE-CC 681.3.01:519.2:62138SR		
3. Palavras Chaves (selecionadas pelo autor) CLASSIFICAÇÃO DE IMAGENS PARTICIONAMENTO DE IMAGENS ALGORITMOS DE CRESCIMENTO DE REGIÕES SISTEMA GE IMAGE-100		interna ☐ externa ☒
5. Relatório nº INPE-1678-RPE/111	6. Data Fevereiro, 1980	7. Revisado por <i>E. D. Souza</i> Dr. Celso de R. Souza
8. Título e Sub-Título PARTICIONAMENTO DE IMAGENS MULTIESPECTRAIS DE RECURSOS NATURAIS DA TERRA		9. Autorizado por <i>Nelson de J. Parada</i> Nelson de J. Parada Diretor
10. Setor DSE/DIN	Código	11. Nº de cópias 10
12. Autoria Walter Kenkiti Takahashi, Adael Woods de C. Filho Ravindra Kumar		14. Nº de páginas 64
13. Assinatura Responsável <i>Adael Woods de C. Filho</i>		15. Preço
16. Sumário/Notas <i>Desenvolve-se aqui um algoritmo eficiente para o particionamento de imagens multiespectrais de recursos naturais da Terra, através da extração de objetos homogêneos do ponto de vista estatístico; visando a posterior classificação dos mesmos.</i>		
17. Observações		



ÍNDICE

	Página
ABSTRACT	v
LISTA DE FIGURAS	vi
CAPÍTULO I - INTRODUÇÃO	1
CAPÍTULO II - SISTEMAS DE RECONHECIMENTO DE PADRÕES	3
CAPÍTULO III - SENSORIAMENTO REMOTO DE RECURSOS DA TERRA	7
CAPÍTULO IV - CLASSIFICAÇÃO	11
4.1 - Modelo estatístico dos dados do MSS	11
4.2 - Classificação sem memória	14
4.2.1 - Estratégia da máxima probabilidade a posteriori (MAP)	14
4.2.2 - Estratégia da máxima verossimilhança (ML)	15
4.2.3 - Estratégia da máxima verossimilhança generalizada (GML)	15
CAPÍTULO V - PARTICIONAMENTO DE IMAGENS	17
5.1 - A lógica do particionamento	17
5.2 - Modo não supervisionado	18
5.2.1 - Anexação não supervisionada	19
5.2.2 - Seleção não supervisionada de células	21
5.3 - Modo supervisionado	22
5.3.1 - Anexação supervisionada	22
5.3.2 - Seleção supervisionada de célula	26
CAPÍTULO VI - ANÁLISE E DESCRIÇÃO DOS PROGRAMAS UTILIZADOS NO PARTIÇÃOAMENTO DE IMAGENS USANDO A TÉCNICA ECHO	27
6.1 - Programas auxiliares	27
6.1.1 - AUX00	27
6.1.2 - AUX01	27
6.1.2.1 - SUBROTINA DISTR	27

6.1.2.2 - SUB-ROTINA ELEM	28
6.2 - Programa principal: IMAPAR	28
6.3 - Sub-rotinas desenvolvidas para o programa principal IMAPAR	29
6.3.1 - Sub-rotina AINVER	29
6.3.2 - Sub-rotina DETER	29
CAPÍTULO VII - COMENTÁRIOS E CONCLUSÕES	31
BIBLIOGRAFIA	35
APENDICE A - LISTAGEM DOS PROGRAMAS UTILIZADOS AUX00, AUX01, IMAPAR E SUAS SUB-ROTINAS	A.1

ABSTRACT

An efficient algorithm for partitioning multispectral earth resources imagery is described. Statistically homogeneous objects are extracted, for posterior classification.

LISTA DE FIGURAS

	Página
FIGURA VI.1 - Diagrama de blocos simplificado do programa IMAPAR	30
FIGURA VII.1 - Imagem na qual foi testado o programa nos Es tados Unidos	33
FIGURA VII.2 -	34

CAPITULO I

INTRODUÇÃO

O objetivo geral deste trabalho é contribuir para o desenvolvimento do "estado arte" de reconhecimento de padrões na tecnologia de sensoramento remoto.

O método aqui desenvolvido explora a dependência característica entre estados adjacentes da natureza, que é desprezada pela regra de decisão convencional existente. Para tipos gerais de dependência, esta informação incorporada exigiria maior tempo de computação por elemento de resolução do que a regra convencional, mas quando as mesmas ocorrem na forma de redundância, os elementos podem ser classificados coletivamente, ou seja, em grupos, o que reduz o número de classificações. Assim existe um potencial de aumento da eficiência.

Basicamente, o método pode ser interpretado como uma transformação de uma imagem através de um particionamento, pela extração de grupos estatisticamente homogêneos de elementos.

Esta extração no caso será supervisionada, ou seja, serão fornecidas ao computador informações relativas aos grupos a serem extraídos.

O algoritmo constará de dois níveis de testes; o primeiro será um teste de homogeneidade das células (denominação dada a grupos de elementos em que a imagem será subdividida) e posterior teste de comparação entre as mesmas para o particionamento propriamente dito da imagem. (Uma boa parte deste relatório foi baseado no trabalho de Ketig e Landgrebe (1975)). *agrupamento*

Fim da esta etapa, a imagem estará pronta para ser classificada.

CAPITULO II

SISTEMAS DE RECONHECIMENTO DE PADRÕES

A fonte de informação mais abundante de uma cena para o homem, é a energia eletromagnética radiante que dela emana.

Tal informação está incorporada nos padrões espaciais, espectrais e temporais da radiação, que variam conforme o objeto.

O processo geral de extrair informação de padrões é conhecido como reconhecimento de padrões. A forma mais conhecida de reconhecimento de padrões é a "classificação" - a atribuição de um padrão observado para uma das várias categorias pré-especificadas. Isto exige um certo grau de experiência, isto é, o sistema de reconhecimento deve "saber" quais são as classes possíveis e alguma caracterização exclusiva para cada uma.

Tipicamente, esta experiência é obtida de padrões de "treino" representativos, que são fornecidos como referência para cada classe.

No caso mais simples, cada amostra de padrões é uma caracterização completa da classe que ela representa. Logo, a classificação se torna uma espécie de comparação direta. De maneira geral, uma caracterização estatística poderia ser a única aproximação adequada, e os padrões de treinamento poderiam ser usados para estimar quantidades estatísticas. A classificação, neste caso, se torna um problema de teoria de decisão estatística.

Certamente não é sempre possível pre-especificar as categorias à que um padrão pertenceria. Isto é frequente na análise de cenas onde o número de possibilidades pode ser muito grande. Então o reconhecimento de padrões pode tomar a forma de descrição. Em geral, o reconhecimento de padrões pode envolver a ambas, classificação e descrição. Uma cena complexa composta de objetos relativamente simples é fre

quentemente descrita classificando-os e guardando-se suas posições e orientações relativas na cena.

Esta descrição poderia ser considerada como resultado final ou ser utilizada para classificar a própria cena.

Todos os sistemas que extraem informações de uma cena consistem em um sistema de coleta de dados e de um processador de da dos. O propósito do sistema de coleta é reduzir a cena a um número de características manipuláveis sem perder a informação desejada. A redu ção (seleção de características) é frequentemente possível num processo. A escolha de características, obviamente, depende da informação que é desejada mas, por outro lado, a informação que pode ser extraída depen de da escolha de características. A maioria dos sistemas de coleta são similares, sob muitos aspectos, a um olho humano, o qual forma caracte rísticas procurando as dimensões espaciais, espectrais e temporais, e convertendo uma cena em série de pulsos elétricos.

A procura espacial pode ser feita pela formação da ima gem da cena sobre uma disposição de detetores (elétricos ou químicos) ou por exame minucioso da imagem com um detetor elétrico. O elemento de re solução de cada sistema é a reprojeção do detetor através do sistema óptico na cena.

Comumente denomina-se "pixel" (picture element) um ele mento de ~~uma~~ figura. *imagem*

Todos os sistemas de resolução dependem do tamanho do "pixel" e do intervalo entre amostras os quais são, normalmente, iguais.

A amostragem espectral é obtida medindo-se a radiação de cada elemento de resolução com detetores que são sensíveis a diferen tes faixas (canais) espectrais. Um prisma, grade, ou filtro de inter face, é frequentemente usado para separar a energia radiante, espectral mente, antes da detecção.

A amostragem temporal é tomada, simplesmente, tirando-se as amostras espectral e espacial em instantes distintos de tempo.

Dependendo do tipo de informação que é desejada, pode-se enfatizar ou deixar de enfatizar uma dimensão particular, tomando-se a amostra muitas ou poucas vezes. Uma foto em preto e branco, por exemplo, enfatiza a informação espacial, já que ela é criada tomando-se a amostragem temporal e espectral somente uma vez. Já uma foto colorida contém três amostras básicas espectrais e assim enfatiza tanto a informação espectral como a espacial.

A maneira como um padrão é amostrado é enquadrada na categoria de medidas complexas. ?

A sub amostragem resulta numa perda de informação, mas a sobreamostragem num excesso de dados para processar. Tecnicamente, a dimensionalidade dos dados aumenta mais rapidamente do que a dimensionalidade intrínseca.

A dimensionalidade de dados refere-se à dimensão do espaço observado ou medido, na qual uma amostragem de padrões pode ser considerada como uma observação de uma variável aleatória multi-dimensional. A densidade de probabilidade desta variável aleatória é uma função de N variáveis (dimensão), onde N é o número de medidas.

A dimensionalidade intrínseca de uma variável aleatória X é a dimensão mínima que outra variável aleatória Y pode ter se X é relacionado unicamente a Y . Assim, ela é o número mínimo de medidas que poderia ser usado para transmitir a mesma informação que X , se suas relações fossem conhecidas.

A sobreamostragem aumenta a dimensionalidade dos dados, mas as medidas individuais tendem a ser mais altamente correlacionadas, acarretando numa diminuição da informação transmitida por medida. A informação que pode ser extraída de uma imagem é também limitada pela sofisticação do processador que vai manipular os dados.

A mente humana é um processador extremamente bom, particularmente quando a informação é de natureza espacial, e por este motivo, os dados são para ela apresentados na forma de imagem visual, o que é conhecido como um sistema de processamento "orientado para imagens". Por outro lado, num sistema "numericamente orientado", o elemento de decisão é o computador e a imagem visual tem uma pequena ou nenhuma participação.

As vantagens do processamento por computador são a sua alta capacidade de carga, o baixo custo em relação a esta alta carga e a capacidade de manipular medidas altamente complexas.

CAPITULO III

SENSORIAMENTO REMOTO DE RECURSOS DA TERRA

Um assunto importante na engenharia e na comunidade científica, atualmente, é o processamento de cenas, o qual inclui o estudo da superfície da Terra vista de cima.

Uma tal cena típica consiste, principalmente, de regiões regulares e/ou irregulares justapostas, cada qual contendo um tipo de cobertura. Estas regiões homogêneas são denominadas os "objetos" da cena.

A meta básica do processamento é localizar e classificar os objetos, e proceder a uma descrição da cena em função de resultados tabelados e/ou de um mapeamento.

Na aplicação do processamento de imagens, a localização e as características espaciais (tamanho, forma e orientação) dos objetos são reveladas pelas mudanças nas propriedades espectrais médias que ocorrem nas fronteiras. Mas, diferentemente da maioria das outras aplicações, as características espaciais de um objeto frequentemente têm somente uma fraca relação com suas classes.

Pesquisas têm mostrado, contudo, que muitas classes podem ser distinguidas, razoavelmente bem, com base em suas características espectrais, usando-se técnicas estatísticas de classificação de padrões. Atualmente, as pesquisas estão se encaminhando no sentido de usar as características temporais também, mas esse não é o caso nesta investigação.

Nosso interesse está no sistema "numericamente orientado" para processar as cenas. A entrada para o sistema está na forma de dados multi-espectrais digitalizados armazenados numa fita magnética. Um rastreador multi-espectral típico (MSS) fornece a dimensão espectral e uma dimensão espacial. A segunda dimensão é fornecida pelo movimento

da plataforma que carrega o rastreador sobre a região de interesse. A dimensão temporal é fornecida pela passagem do rastreador em diferentes instantes de tempo.

A classificação por computador dos dados do MSS é feita aplicando-se uma regra de decisão "simple-symmetric" para cada pixel. Isto significa que cada pixel é classificada individualmente, com base somente nas medidas espectrais.

A suposição básica desta técnica é a de que os objetos de interesse são grandes, comparados com o tamanho do pixel. Por outro lado, uma grande proporção de pixels seria composta de duas ou mais classes, tornando irrealizável a classificação estatística de padrões, isto é, as categorias pré-especificadas seriam inadequadas para descrever os estados reais da natureza.

Desde que o intervalo de amostragem é usualmente comparável ao tamanho do pixel (para preservar a resolução do sistema), segue-se que cada objeto é representado por um grupo de pixels. Isto sugere uma dependência estatística entre estados consecutivos da natureza, a qual o classificador simples ("simple-symmetric") não explora. Para refletir propriedade, nos referiremos à classificação "simple-symmetric" daqui em diante como classificação sem memória.

Um método para lidar com estados de dependência é o de aplicar o princípio da teoria de decisão composta ou teoria de decisão composta sequencial.

A formulação da decisão composta é uma poderosa aproximação para manipular muitos tipos gerais de dependência. Conclui-se disto que, talvez, fazendo-se uma aproximação mais direta para o problema se pudesse obter resultados similares com considerável simplificação.

Uma característica distinta da dependência espacial nos dados do MSS é a redundância, isto é, a probabilidade de transição do

estado i para o estado j é maior se $j=i$ do que se $j \neq i$, porque o intervalo de amostragem é pequeno comparado com o tamanho do objeto. Isto sugere o uso de uma transformação da imagem por particionamento, a fim de delinear os grupos de pixels estatisticamente similares, antes de classificá-los.

Desde que cada grupo homogêneo representa uma amostra estatística (um conjunto de observações de uma população comum), um classificador de amostras poderia então ser usado para classificar os objetos. Deste modo, a classificação de cada pixel na amostra é um resultado de propriedades espectrais de seus vizinhos assim como de suas próprias. Assim, seu contexto na cena é usado para fornecer uma melhor classificação.

A sigla ECHO (extração e classificação de objetos homogêneos) designa esta abordagem geral.

Uma característica comum à técnica "sem memória" na decisão composta é a de que o número de classificações que precisa ser executado é muito maior do que o número real de objetos na cena. Quando cada classificação exige uma grande quantidade de computação, mesmo o classificador "sem memória" pode ser relativamente lento.

A técnica ECHO reduz substancialmente o número de classificações, resultando no aumento potencial da velocidade de classificação (diminuição no custo). Se este potencial é ou não realizável, depende da eficiência da operação de particionamento.

A meta da corrente investigação é o posterior desenvolvimento do conceito da técnica ECHO.

CAPÍTULO IV

CLASSIFICAÇÃO

A motivação para a extração de objetos é possibilitar a classificação mais rápida e precisa dos pixels dentro do objeto.

4.1 - MODELO ESTATÍSTICO DOS DADOS DO MSS

Como foi dito anteriormente, uma cena típica consiste simplesmente de objetos cujos limites formam uma partição da cena. A partição é geralmente desconhecida, mas nós podemos pelo menos admitir que ela é relativamente grosseira, comparada com o tamanho de um pixel. Cada objeto na cena pertence a alguma classe. Para o propósito de representação, cada classe é dividida em uma ou mais subclasses. Elas são também chamadas de classes espectrais, para indicar que podem ser distinguidas espectralmente embora possa não ser útil fazê-lo.

Seja w_{ij} a j -ésima subclasse da i -ésima classe e F um objeto representado por um grupo de pixels, sendo X um pixel em algum objeto. (A barra em X é usada para indicar uma variável q dimensional, $X \in R^q$, onde q , daqui em diante, representará o número de canais espectrais). Isto significa que $F \in w_{ij}$ denota o evento " F pertence à subclasse w_{ij} ". A probabilidade "a priori" deste evento é representada por $P(F \in w_{ij})$.

De acordo com o item 3, qualquer dependência estatística deste evento com as características espaciais de F é ignorada.

A consequência desta suposição é que $P(X \in w_{ij}) = P(F \in w_{ij})$ e que as quantidades podem ser representadas por $P(w_{ij})$.

Os pixels dentro de um dado objeto de uma dada classe espectral são completamente caracterizados por sua "função distribuição de probabilidade". Para a classificação "sem memória" tal modelo completo é desnecessário; somente a distribuição marginal de cada pixel é exigida.

Posteriormente, os pixels dentro de um único objeto se rão usualmente supostas como tendo uma distribuição marginal comum, a qual se deve à homogeneidade dos tipos de objetos tipicamente encontrados na aplicação de sensoramento remoto.

Embora os dados sejam digitais, é conveniente representar esta distribuição q-variável por uma função densidade de probabilidade (fdp) de um parâmetro contínuo a qual, para a subclasse w_{ij} , é representada por $p(\underline{X}=\underline{x}/\underline{X} \in w_{ij})$ ou simplesmente por $p(\underline{x}/w_{ij})$.

Dois pixels em proximidade espectral são correlacionados incondicionalmente, sendo que os graus de correlação diminuem se a distância entre eles aumenta. Esta correlação é atribuída ao efeito de um estado de dependência, que desejamos explorar.

Para simplificação, ignoramos outras fontes de correlação. Assim, admitimos que os pixels dentro de um mesmo objeto são condicionalmente independente da classe, isto é, que cada objeto é uma amostra simples de uma das populações da classe espectral. Portanto, a fdp pode ser expressa em termos do produto de suas fdp's marginais.

Esta aproximação conduz a um algoritmo de processamento rápido e efetivo (embora sub-ótimo); entretanto, prognóstico baseado neste modelo simplificado deve ser interpretado cautelosamente.

É possível expressar outras características estatísticas em termos das anteriormente definidas. Se w_i representa a i-ésima classe, então:

$$P(\underline{X} \in w_i) = P(\cup_j \underline{X} \in w_{ij}) = \sum_j P(w_{ij}) \quad (\text{IV.1})$$

A fdp de $\underline{X}$ condicional neste evento, é dada por:

$$P(\underline{x}/w_i) = \frac{1}{P(w_i)} \sum_j p(\underline{x}/w_{ij}) P(w_{ij}) \quad (\text{IV.2})$$

Esta equação define a representação de uma classe em termos de suas subclasses.

A função densidade de probabilidade incondicional pode ser escrita de dois modos:

$$p(\underline{x}) = \sum_i \sum_j p(\underline{x}/W_{ij}) P(W_{ij}) = \sum_i p(\underline{x}/W_i) P(W_i) \quad (\text{IV.3})$$

Dentro deste método de trabalho, tudo que é exigido para completar o modelo estatístico para uma dada cena é especificar as classes espectrais que estão presentes e atribuir uma fdp condicional e uma probabilidade "a priori" para cada classe.

Certamente, as distribuições verdadeiras são atributos da natureza e a precisão do modelo depende de quão bem nós podemos estimá-las. Felizmente, somos usualmente capazes de obter estimativas de fdp's condicionais das classes, baseadas em amostras de treinamento tomadas diretamente dos dados. Para isso, contamos com observações da superfície real ou com a interpretação manual da foto, para localizar áreas representando cada classe.

Para o propósito do projeto, supomos que o tamanho de cada amostra é suficientemente grande para que o erro, na correspondente estimativa da distribuição, seja desprezível.

Foi provado que a distribuição normal multi-variável (MVN) é um modelo razoável para os dados do MSS, ou seja:

$$p(\underline{x}/W_{ij}) = N(\underline{M}_{ij}, \underline{C}_{ij}; \underline{x}) \text{ onde:}$$

$$N(\underline{M}, \underline{C}; \underline{x}) = (|2\pi\underline{C}| \exp((\underline{x} - \underline{M})' \underline{C}^{-1} (\underline{x} - \underline{M})))^{-1/2} \quad (\text{IV.4})$$

Segue-se que: se $\underline{X} \in W_{ij}$

$$E(\underline{X}) = \underline{M}_{ij}$$

$E((\underline{X} - \underline{M}_{ij})(\underline{X} - \underline{M}_{ij})') = \underline{C}_{ij}$, onde E representa a esperança estatística. Assim, $\underline{M}_{ij}$ e $\underline{C}_{ij}$ são o vetor média e a matriz co-variância respectivamente da distribuição da subclasse. Notar que, para obter uma estimativa dos parâmetros da distribuição MVN, somente necessitamos estimar os momentos de primeira e segunda ordens, esta sendo a aproximação a ser usada.

4.2 - CLASSIFICAÇÃO SEM MEMÓRIA

Aqui serão introduzidos certos conceitos acerca das técnicas comuns de classificação sem memória, que serão úteis posteriormente.

4.2.1 - ESTRATÉGIA DA MÁXIMA PROBABILIDADE A POSTERIORI (MAP)

Seja $\underline{X}$ um pixel. Sob a hipótese de que $\underline{X} \in W_i$, a fdp de $\underline{X}$ é $p(\underline{X} = \underline{x} / \underline{X} \in W_i)$, que é dada pela equação IV.2.

Supondo que esta função é conhecida exatamente, a hipótese é simples. A meta da classificação é projetar uma estratégia para escolher uma das classes possíveis, baseada em $\underline{x}$, o valor observado de $\underline{X}$, isto é, precisamos especificar uma função $W(\underline{x})$ que mapeia $\underline{x}$ numa das amostras de classes possíveis. Podemos maximizar a probabilidade de uma decisão, para sempre escolher a classe correta W , a qual tem a máxima probabilidade a posteriori $P(\underline{X} \in W_i / \underline{X} = \underline{x})$. Para mostrar isto, simplesmente escrevemos a probabilidade de uma decisão correta da seguinte forma:

$$P(\underline{X} \in W(\underline{x})) = \int_{\underline{x} \in R} P(\underline{X} \in W(\underline{x}) / \underline{X} = \underline{x}) p(\underline{X} = \underline{x}) d\underline{x} \quad (\text{IV.5})$$

É evidente que esta quantidade é maximizada com respeito à função de decisão, adotando-se uma regra de decisão MAP.

Para implementar esta estratégia, usamos a forma mista da regra de Bayes que dá:

$$P(\underline{X} \in W_i / \underline{X} = \underline{x}) = \frac{p(\underline{X} = \underline{x} / \underline{X} \in W_i) P(W_i)}{p(\underline{X} = \underline{x})} \quad (\text{IV.6})$$

O denominador $\bar{e}$ independente de $\underline{i}$, assim, precisamos so mente procurar o $\underline{i}$ que maximiza o numerador. Ou seja, $\underline{x}$ e $W(\underline{x})$ para uma dada observação $\bar{e}$ escolhida tal que:

$$p(\underline{X} = \underline{x} / \underline{X} \in W(\underline{x})) P(W(\underline{x})) = \max_i p(\underline{X} = \underline{x} / \underline{X} \in W_i) P(W_i) \quad (\text{IV.7})$$

4.2.2 - ESTRATÉGIA DA MÁXIMA VEROSSIMILHANÇA (ML)

Quando todas as classes são equiprováveis, a regra de de cisão MAP reduz-se a:

$$p(\underline{X} = \underline{x} / \underline{X} \in W(\underline{x})) = \max_i p(\underline{X} = \underline{x} / \underline{X} \in W_i) \quad (\text{IV.8})$$

A estatística $p(\underline{X} = \underline{x} / \underline{X} \in W_i)$, como uma função de $\underline{i}$, $\bar{e}$ denominada função de verossimilhança, daí esta regra de decisão chamar-se estratégia da máxima verossimilhança. Esta regra $\bar{e}$ uma aproximação razoável mesmo quando as classes não são equiprováveis.

Comparativamente, a MAP tende a discriminar classes cujas probabilidades a "priori" são baixas, ou seja, ela faz com que a probabilidade condicional de erro seja grande quando ocorre uma classe rara, a fim de minimizar a probabilidade de erro geral. Portanto, quando houver necessidade de discriminar uma classe rara, a MAP não $\bar{e}$ uma estratégia conveniente, porque ela distribui equitativamente o erro em todas as classes. Já com o ML tal risco não ocorre, devido ao fato da probabilidade de erro, quando ocorre a classe $\underline{i}$, depender somente do grau de "separação" (ou distância) entre a classe $\underline{i}$ e outras. Ou seja, a decisão independe da probabilidade "a priori" da classe $\underline{i}$.

4.2.3 - ESTRATÉGIA DA MÁXIMA VEROSSIMILHANÇA GENERALIZADA (GML)

Freqüentemente, as probabilidades "a priori" das sub-classes são desconhecidas e, então, a hipótese de que $\underline{X} \in W_i$ se torna

uma hipótese composta, isto é, $p(\underline{x}/w_i) = \sum_j A_{ij} p(\underline{x}/w_{ij})$, onde os coeficientes são desconhecidos. Certamente sabemos que $A_{ij} \geq 0$ e $\sum_j A_{ij} = 1$.

Um modo de contornar tal situação seria o de estimar os parâmetros desconhecidos e resolver, posteriormente, pela estratégia ML.

Este procedimento é designado de estratégia de máxima verossimilhança generalizada.

A regra de decisão resultante pode ser expressa simplesmente da seguinte forma:

$$p(\underline{x}/V(\underline{x})) = \max_i \max_j p(\underline{x}/w_{ij}),$$

onde $V(\underline{x})$, mapeia $\underline{x}$ numa das possíveis classes espectrais existentes.

Portanto, daí conclui-se que a ML e a GML são equivalentes, sendo que quando as classes espectrais são equiprováveis, elas maximizam a probabilidade de classificar a observação na classe correta.

CAPÍTULO V

PARTICIONAMENTO DE IMAGENS

Uma vez que o particionamento é conhecido, poderosas técnicas são disponíveis para classificar os objetos individuais. Assim, quando o mesmo não é conhecido, um algoritmo de particionamento de imagem oferece uma alternativa atrativa em relação à classificação "sem memória. O algoritmo aqui estudado é do tipo de procura conjuntiva de objetos, que nada mais é do que uma junção progressiva de elementos adjacentes que se mostrarem estatisticamente semelhantes. Portanto, este algoritmo consistirá de testes estatísticos aplicados numa sequência lógica.

5.1 - A LÓGICA DO PARTICIONAMENTO

Geralmente não é possível implementar um algoritmo de particionamento livre de erros. Primeiramente, existe uma certa ambiguidade na definição de um particionamento verdadeiro, devido aos efeitos reais, tais como "pixels" que ultrapassam limites físicos. Por outro lado, dois tipos principais de erros de decisão podem ocorrer conduzindo a:

1. limites falsos e
2. limites descontínuos.

Os efeitos combinados destes dois podem ainda produzir limites aproximados, que não formam uma categoria realmente bem definida, devido à ambiguidade do particionamento verdadeiro.

Desde que nem o tamanho, nem a forma do objeto são usados como características para a classificação, os erros do tipo (1) conduzem, menos provavelmente, a uma classificação não real do que os erros do tipo (2). Esta constatação permite certas simplificações na lógica do particionamento.

A abordagem básica adotada consiste de dois níveis de testes. Inicialmente, os "pixels" são divididos por uma grade hipotética

tica em pequenos grupos cujo número será especificado.

No primeiro nível de teste, cada grupo se torna uma unidade denominada célula, desde que passe por um teste de homogeneidade relativamente suave. Os grupos que são rejeitados no teste são geralmente aqueles que se encontram nas fronteiras e seus "pixels" serão classificados individualmente pelo método sem memória. Estes grupos são designados células singulares. Neste nível de teste, é desejável manter uma taxa de rejeição baixa, para refletir a probabilidade a priori relativamente alta de um grupo a ser homogêneo.

A meta deste nível é, essencialmente, detetar tantos "pixels" quanto possível, que se encontram ao longo das fronteiras, sem exigir que sejam detetados contornos fechados ou mesmo contínuos.

No segundo nível de teste, uma célula individual é comparada a um campo adjacente, o qual nada mais é do que um grupo de uma ou mais células que foram previamente agrupadas. Se nesta comparação for detetada certa similaridade estatística, então esta célula será incorporada ao campo. Não mostrando semelhança, esta célula é comparada a outro campo adjacente e, em último caso, torna-se um novo campo. Por ane~~ne~~xação sucessiva, cada campo vai se expandindo até atingir suas fronteiras naturais, onde a taxa de rejeição cresce abruptamente. Neste ponto, o campo é classificado e a classificação transmitida a todos os seus "pixels".

Esta aproximação tem uma vantagem importante que é a possibilidade de implementação sequencial, ou seja, os dados precisam ser acessados somente uma vez e na mesma ordem em que são armazenados.

5.2 - MODO NÃO SUPERVISIONADO

A fim de implementar a aproximação sequencial, precisamos especificar os dois critérios de teste correspondentes aos dois níveis.

Nesta seção, consideramos o modo como sendo não supervisionado, ou seja, o critério de teste é independente do conhecimento específico das distribuições das classes espectrais.

5.2.1 - ANEXAÇÃO NÃO SUPERVISIONADA

Seja $X = (\underline{X}_1, \dots, \underline{X}_n)$, o "pixel" genérico de um grupo de uma ou mais células que foram agrupadas por sucessivas anexações; e seja $Y = (\underline{Y}_1, \dots, \underline{Y}_m)$, o pixel genérico de uma célula adjacente não singular. Desde que X e Y tenham satisfeito certos critérios de homogeneidade, admitiremos que cada qual representa uma amostra de uma população normal multi-variável.

Façamos f e g representarem as funções densidade correspondentes. É desejável que fdg ; isto é uma hipótese composta, desde que f e g não sejam especificadas.

O procedimento de razão de verossimilhança fornece uma estatística efetiva para testar esta hipótese.

$$H_0(x,y) = \{ p(x,y/f,g): g = f, f \in \Omega \}$$

$$H_1(x,y) = \{ p(x,y/f,g): f \in \Omega, g \in \Omega, g \neq f \}$$

onde $p(x,y/f,g)$ é a densidade condicional conjunta de X e Y , $x \in R^{nq}$ e $y \in R^{mq}$ e Ω é um espaço de funções-densidade normais multi-variáveis. A suposição da independência da classe condicional, possibilita-nos expressar a densidade conjunta de pixels como um produto de suas densidades marginais. Assim:

$$p(x,y/f,g) = p(x/f)p(y/g) = \left(\prod_{i=1}^n f(\underline{x}_i) \right) \left(\prod_{j=1}^m g(\underline{y}_j) \right)$$

A razão de verossimilhança generalizada é definida por:

$$\Lambda = \frac{\sup H_0(X,Y)}{\sup H_1(X,Y)} = \frac{\max_{f \in \Omega} p(X/f)p(Y/f)}{\max_{\substack{f \in \Omega \\ g \in \Omega \\ g \neq f}} p(X/f)p(Y/g)} \quad (V.1)$$

Numa aproximação não supervisionada para o particionamento, tomamos como sendo o seguinte o espaço de funções de $\underline{x} \in R^q$:

$$\Omega = \{N(\underline{M}, \underline{C}; \underline{x}) : \underline{M} \in R^q, \underline{C} = \text{simétrico, positivo-definido}\}$$

Desde que para qualquer $f \in \Omega$ existe uma $g \in \Omega$ que é arbitrariamente próxima de f , a condição $g \neq f$ do denominador pode ser desprezada. Daí resulta:

$$\Lambda = \frac{\max N(\underline{M}, \underline{C}; \underline{X}) N(\underline{M}, \underline{C}; \underline{Y})}{\max N(\underline{M}_x, \underline{C}_x; \underline{X}) N(\underline{M}_y, \underline{C}_y; \underline{Y})} \leq 1$$

onde:

$$N(\underline{M}, \underline{C}; \underline{X}) = \prod_{i=1}^n N(\underline{M}, \underline{C}; \underline{X}_i)$$

$$N(\underline{M}, \underline{C}; \underline{Y}) = \prod_{i=1}^m N(\underline{M}, \underline{C}; \underline{Y}_i)$$

Em cada caso, a maximização é com relação ao vetor média e matriz de co-variância.

Λ pode ser escrito ainda como:

$$= \frac{\max N(\underline{M}_x, \underline{C}_x; \underline{X}) N(\underline{M}_y, \underline{C}_y; \underline{Y})}{\max N(\underline{M}_x, \underline{C}_x; \underline{X}) N(\underline{M}_y, \underline{C}_y; \underline{Y})} \cdot \frac{\max N(\underline{M}, \underline{C}; \underline{X}) N(\underline{M}, \underline{C}; \underline{Y})}{\max N(\underline{M}_x, \underline{C}_x; \underline{X}) N(\underline{M}_y, \underline{C}_y; \underline{Y})}$$

$$= \Lambda_2 \cdot \Lambda_1$$

Anderson mostrou que:

$$\Lambda_1 = (|\underline{A}|/|\underline{B}|)^{N/2} \quad (V.2)$$

$$\Lambda_2 = (|\underline{A}_x/n|^{n/2} |\underline{A}_y/m|^{m/2} / |\underline{A}/N|^{N/2})^{0.5} \quad (V.3)$$

sendo que:

$$N = n + m$$

$$\bar{X} = \sum_{i=1}^n \underline{X}_i / n$$

$$\bar{Y} = \sum_{i=1}^m (\underline{Y}_i / m)$$

$$\underline{A}_X = \sum_{i=1}^n (\underline{X}_i - \bar{X})(\underline{X}_i - \bar{X})'$$

$$\underline{A}_Y = \sum_{i=1}^m (\underline{Y}_i - \bar{Y})(\underline{Y}_i - \bar{Y})'$$

Neste ponto, a fim de assegurar que as matrizes sejam não singulares com probabilidade 1, nós introduzimos a restrição $n \neq m$.

$$\underline{A} = \underline{A}_X + \underline{A}_Y$$

$$\underline{M} = (n\bar{X} + m\bar{Y})/N$$

$$\underline{B}_X = \sum_{i=1}^n (\underline{X}_i - \underline{M})(\underline{X}_i - \underline{M})' = \underline{A}_X + n(\bar{X} - \underline{M})(\bar{X} - \underline{M})'$$

$$\underline{B}_Y = \sum_{i=1}^m (\underline{Y}_i - \underline{M})(\underline{Y}_i - \underline{M})' = \underline{A}_Y + m(\bar{Y} - \underline{M})(\bar{Y} - \underline{M})'$$

$$\underline{B} = \underline{B}_X + \underline{B}_Y = \underline{A} + \frac{mn}{N} (\bar{X} - \bar{Y})(\bar{X} - \bar{Y})'$$

Se substituirmos o número de "pixels" em cada amostra pelo número de graus de liberdade, ou seja $\underline{n}$ por $\underline{n}-1$, $\underline{m}$ por $\underline{m}-1$, e $\underline{N}$ por $\underline{N}-2$, nas fórmulas V.2 e V.3, resulta que: $\lambda = \lambda_1 \cdot \lambda_2$, onde λ_1 e λ_2 são obtidos pela modificação de Λ_1 e Λ_2 , como mostrado acima.

De qualquer modo, as estatísticas são invariáveis com relação à transformação linear dos vetores dados. Segue-se portanto, que suas distribuições sob a hipótese do nulo independem da população normal multi-variável real da qual a amostra foi extraída.

O teste segue comparando λ com algum limite $T < 1$, o qual depende geralmente de $\underline{n}$ e $\underline{m}$. A hipótese é aceita se $\lambda \geq T$ e rejeitada se $\lambda < T$:

5.2.2 - SELEÇÃO NÃO SUPERVISIONADA DE CÉLULA

A seleção de célula refere-se ao teste de nível 1, e é usada para detetar células que aparentemente ultrapassam as fronteiras.

Tais células apresentam, frequentemente, variâncias normalmente grandes e assim, um critério possível é o de que a variância das células em cada canal espectral caia abaixo de um limite razoável.

5.3 - MODO SUPERVISIONADO

Neste modo, as distribuições de classes espectral conhecidas são usadas para processar o particionamento sequencial. As aproximações no procedimento de testes da hipótese composta são basicamente as mesmas do caso não supervisionado. O efeito das classes espectrais simplifica grandemente cada hipótese, mas, paradoxalmente, o critério de teste resultante é muito mais complicado. Felizmente, muita computação pode ser efetuada sequencialmente.

5.3.1 - ANEXAÇÃO SUPERVISIONADA

O procedimento é análogo ao de anexação não supervisionada, exceto que, na aproximação para o particionamento, tomamos como:

$$\Omega = \{ p(\underline{x}/V_i): i, 2, \dots, k \}$$

onde k é o número de classes espectrais. Notar que esta condição é consideravelmente mais restritiva do que a anterior.

A razão de verossimilhança generalizada correspondente é:

$$\Lambda = \frac{\max_i (p(X/V_i) p(Y/V_i))}{\max_{\substack{i,j \\ i \neq j}} (p(X/V_i) p(Y/V_j))} \quad (V.4)$$

Esta estatística multi-variável não necessita da restrição $m \leq q$, que era necessária no método não supervisionado. Contudo, os máximos de (V.4) não podem ser expressos numa forma analítica simples como em (V.3): eles precisam ser obtidos por uma pesquisa exaustiva. Ademais, a distribuição de (V.4) é desconhecida, sob quaisquer hipóteses, porque ela depende das classes verdadeiras de X e Y .

Em compensação, ganhamos uma estatística que é muito mais sensível à presença ou ausência de fronteiras. Isto produz um desempenho melhor, além de tornar a especificação do limite de decisão mais crítico.

Somado a isso, ainda temos que os resultados tendem a ser bastante estáveis sob várias ordens de magnitude do limite.

Assim, achamos conveniente representar o limite de decisão como:

$$T = 10^{-t}, t \geq 0$$

Ao contrário do método não supervisionado, o fato de $T=1$, não implica em que a hipótese do nulo seja sempre rejeitada. Mas ele conduz ao mesmo resultado final que quando a GML é usada para classificar os objetos. Isto acontece porque Λ pode exceder 1 somente se X e Y forem classificados juntos pela GML. Consequentemente, os valores práticos de T serão somente aqueles entre 0 e 1.

Na ausência de uma teoria para possibilitar a escolha de um limite conveniente, partimos para uma aproximação empírica para obtê-lo.

O cálculo da razão de verossimilhança generalizada pode ser grandemente simplificado fazendo-se as seguintes considerações:

- 1) Mudar o denominador de Λ para $\max_i (p(X/V_i) p(Y/V_j))$. O único efeito desta mudança é fazer com que o valor de Λ não ultrapasse o limite superior 1, o que não afeta o valor de Λ quando $\Lambda < 1$. Desde que T é sempre menor que 1, a mudança não afeta qualquer decisão. A simplificação que isto fornece permite escrever Λ como segue:

$$\Lambda = \frac{\max_i (p(X/V_i) p(Y/V_i))}{(\max_i p(X/V_i)) (\max_j p(Y/V_j))}$$

o qual é mais simples de computar.

2) Comparar $\ln(\Lambda)$ com $\ln(T)$ ao invés de Λ com T .

$$\ln \Lambda = \max_i (\ln p(X/V_i) + \ln p(Y/V_i)) - \max_i (\ln p(X/V_i)) - \\ - \max_i (\ln p(Y/V_i))$$

Isto converte multiplicações em adições, que são mais fáceis e rápidas de executar.

Para calcular $p(Y/V_i)$:

Seja $Y = (Y_1, \dots, Y_n)$ uma amostra simples de uma população normal multi-variável e seja:

$$\underline{S}_1 = \sum_{i=1}^n \underline{Y}_i$$

$$\underline{S}_2 = \sum_{i=1}^n \underline{Y}_i \underline{Y}_i'$$

Os estimadores de máxima verossimilhança do vetor média e matriz co-variância são:

$$\underline{M} = \underline{S}_1/n$$

$$\underline{C} = (1/n) \sum_{i=1}^n (\underline{Y}_i - \underline{M})(\underline{Y}_i - \underline{M})' = \underline{S}_2/n - \underline{M}\underline{M}'$$

A fim da ML ser computacionalmente competitiva com as outras técnicas, precisamos expressá-la em termos de $\underline{S}_1$ e $\underline{S}_2$, donde se segue que:

$$p(Y/W_i) = \prod_{i=1}^n N(\underline{M}_i, \underline{C}_i; \underline{Y}_i) \\ = (|2\pi \underline{C}_i|)^n \exp\left(\sum_{j=1}^n (\underline{Y}_j - \underline{M}_i)' \underline{C}_i^{-1} (\underline{Y}_j - \underline{M}_i)\right)^{-1/2} \\ = -.5(n \ln |2\pi \underline{C}_i| + \sum_{j=1}^n (\underline{Y}_j' \underline{C}_i^{-1} \underline{Y}_j - 2 \underline{M}_i' \underline{C}_i^{-1} \underline{Y}_j + \underline{M}_i' \underline{C}_i^{-1} \underline{M}_i))$$

$$\sum_{j=1}^n \underline{Y}_j' \underline{C}_i^{-1} \underline{Y}_j = \sum_{j=1}^n \text{tr}(\underline{C}_i^{-1} \underline{Y}_j \underline{Y}_j') = \text{tr}(\underline{C}_i^{-1} \sum_{j=1}^n \underline{Y}_j \underline{Y}_j')$$

Assim:

$$\ln p(Y/W_i) = -.5 \text{tr}(\underline{C}_i^{-1} \underline{S}_2) + \underline{M}_i' \underline{C}_i^{-1} \underline{S}_1 - .5n(\underline{M}_i' \underline{C}_i^{-1} \underline{M}_i + \ln|2\pi \underline{C}_i|)$$

Ou em termos do vetor média e matriz covariância:

$$(1/n) \ln p(Y/W_i) = -.5(\ln |2\pi \underline{C}_i| + \text{tr}(\underline{C}_i^{-1} (\underline{C} + (\underline{M} - \underline{M}_i) (\underline{M} - \underline{M}_i)')))$$

Esta fórmula dá a quantidade $p(Y/V_i)$.

- 3) Supor por um momento que a função log de verossimilhança de X já esteja disponível num vetor G, isto é:

$$G(i) = \ln p(X/V_i), i=1,2,\dots,k$$

$$\text{e } J \text{ é um inteiro tal que } G(J) = \max_i G(i).$$

Sendo assim, o primeiro termo de $\ln(\Lambda)$ pode ser computado utilizando-se um vetor intermediário g como segue:

$$g(i) = G(i) + \ln p(Y/V_i), i = 1,2,\dots,k$$

Se a indicação for para juntar Y com X, a nova ϕ função log de verossimilhança do campo será justamente a soma das funções log de verossimilhança de X e Y. Portanto podemos simplesmente realimentar G e J como segue:

$$G(i) = g(i), i = 1,2,\dots,k$$

$$J = j$$

sendo portanto a suposição preliminar justificada.

Quando é alcançado o ponto em que o campo para de expandir, ele é classificado.

5.3.2 - SELEÇÃO SUPERVISIONADA DE CÉLULA

A estatística para a seleção de células neste caso é:

$$Q_j(Y) = \text{tr}(C_j^{-1} \sum_{i=1}^m Y_i Y_i') - 2M_j' C_j^{-1} \sum_{i=1}^m Y_i + m M_j' C_j^{-1} M_j$$

onde j é tal que:

$$\ln p(Y/V_j) = \max_i \ln p(Y/V_i) = \max_i \{-0.5(m \ln |2\pi C_i| + Q_i(Y))\}$$

A célula é considerada homogênea se $Q_j(Y) < c$, onde c é um limite pré-especificado.

Este teste tem a particular vantagem de rejeitar não somente as células não homogêneas como também as não reconhecíveis. (Células não reconhecíveis são células que representam classes espectrais que o classificador não foi treinado para reconhecer). Outra vantagem é que o uso da função log de verossimilhança torna-se especialmente compatível com o critério de anexação supervisionada, assim como com o classificador GML.

Finalmente, um critério para escolher o limite c é lembrar que a função distribuição $P(Q_j(Y) > c/Y \in V_j)$ é uma qui-quadrada com m_j graus de liberdade.

CAPÍTULO VI

ANÁLISE E DESCRIÇÃO DOS PROGRAMAS UTILIZADOS NO PARTICIONAMENTO DE IMAGENS USANDO A TÉCNICA ECHO

Baseados nas informações teóricas contidas no Capítulo V, daremos uma descrição dos programas implementados para proceder ao particionamento de imagens multiespectrais de recursos naturais da Terra utilizando a técnica ECHO.

6.1 - PROGRAMAS AUXILIARES

6.1.1 - AUX00

Este programa tem por finalidade inicializar os parâmetros independentes ao treinamento (número máximo de classes, número máximo de amostras). Além disso, é nesta etapa que o número de dimensões (canais) é adquirido.

6.1.2 - AUX01

Nesta etapa de processamento efetua-se a aquisição de características (vetor média e matriz autocorrelação) de regiões homogêneas delineadas pelo contorno do cursor. Estas características serão utilizadas para localizar, na imagem, as regiões estatisticamente semelhantes.

Este programa utiliza-se de duas subrotinas, a saber, DISTR e ELEM.

6.1.2.1 - SUBROTINA DISTR

Quando referenciada pelo programa principal AUX01, calcula a média e a matriz autocorrelação da amostra dada pelo contorno do cursor.

Sua chamada é CALL DISTR. Todos os argumentos ocupam uma área COMM1, de modo que os dados se transmitem por meio desta.

6.1.2.2 - SUBROTINA ELEM

Obtém a matriz co-variância da amostra especificada e, através do cálculo do respectivo determinante, determina se a amostra é ou não representativa, para ser utilizada no particionamento.

Sua chamada é CALL ELEM e o princípio de transmissão de dados entre a subrotina e o programa principal é o mesmo que o anteriormente descrito.

Esta subrotina utiliza-se ainda de uma função INTC que faz a conversão de real para inteiro, aproximando para o inteiro mais próximo.

6.2 - PROGRAMA PRINCIPAL: IMAPAR

Com as informações adquiridas pelos programas auxiliares inicia-se a procura de células cujas características (média e covariância) sejam estatisticamente semelhantes.

O número de "pixels" que formam cada célula pode ser especificado conforme a necessidade do usuário. Neste ponto, há o compromisso entre tempo de processamento e o tamanho da célula, ou seja, se a imagem for dividida em células menores, obviamente o número de classificações será maior e consequentemente o tempo será prolongado. O tamanho da célula depende também da característica da imagem; no caso de ser esta formada de poucas regiões homogêneas de geometria considerável, a célula pode ter uma constituição maior, e a classificação não será prejudicada. Por outro lado, no caso de uma imagem de constituição complexa, o tamanho da célula deve ser escolhido de tal modo que se consiga delinear razoavelmente o contorno das regiões homogêneas.

Na figura VI.1 é mostrado o diagrama de blocos simplificado do programa IMAPAR.

6.3 - SUB-ROTINAS DESENVOLVIDAS PARA O PROGRAMA PRINCIPAL IMAPAR

6.3.1 - SUB-ROTINA AINVER

É usada para calcular a inversa de uma matriz. Sua chamada é CALL AINVER (RCVCL, NDIM).

onde:

RCVCL = matriz de entrada, que também retorna com o resultado da inversão.

NDIM = dimensão da matriz.

6.3.2 - SUB-ROTINA DETER

É usada para calcular o determinante de uma matriz. Sua chamada é CALL DETER (RCVCL, NDIM, RDET).

onde:

RCVCL = matriz cujo determinante é desejado

NDIM = dimensão da matriz

RDET = valor do determinante.

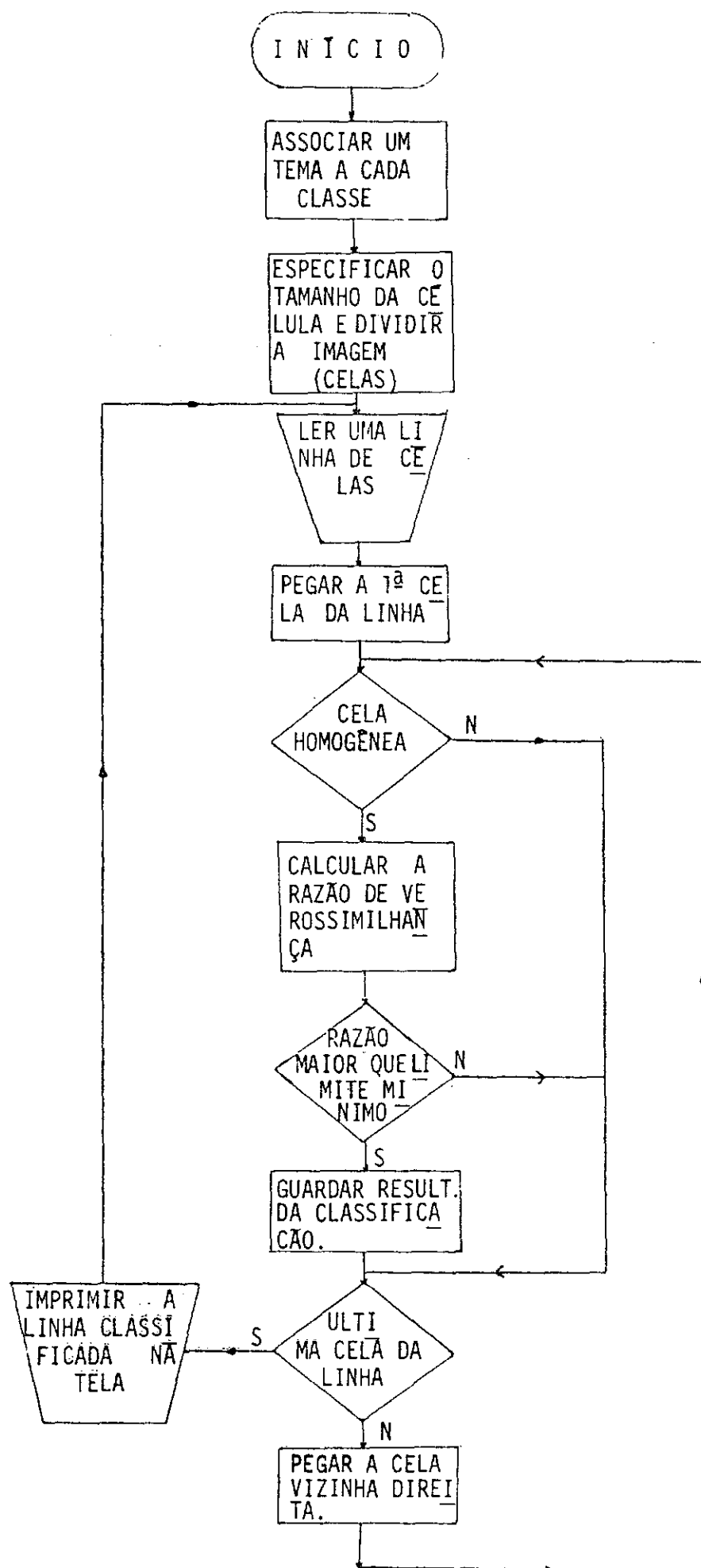


FIG.VI.1 - DIAGRAMA DE BLOCOS SIMPLIFICADO DO PROGRAMA IMAPAR.

CAPÍTULO VII

COMENTÁRIOS E CONCLUSÕES

O presente algoritmo foi inicialmente desenvolvido para o particionamento de imagens, produzidas por satélites LANDSAT e mesmo por aviões, de regiões agrícolas dos Estados Unidos, ou seja, de regiões cuja característica principal é a presença de grandes áreas homogêneas em número razoavelmente pequeno, como se vê na Figura VII.1.

Por outro lado, as imagens para as quais o algoritmo foi implementado, são de constituição complexa, onde a presença de pequenas áreas homogêneas em grande número é a principal característica. Daí, a correlação de estados adjacentes da natureza, que é a justificativa teórica para a divisão da imagem em células e posterior classificação em grupos, é prejudicada, já que existirá uma grande quantidade de células constituindo o contorno e que serão rejeitadas no teste de homogeneidade e, portanto, deixarão de ser classificadas.

Outra consideração a ser feita quanto à utilização deste algoritmo no particionamento das imagens aqui utilizadas (estudo de poluição de água, avanço da erosão, crescimento da área urbana, etc...) é a de que o tamanho das células em que as mesmas serão subdivididas não deve ser muito grande, o que implica num aumento do número de classificações, e, portanto, de tempo de processamento; além disso, um tamanho de célula muito grande não só prejudica a precisão da classificação, como sobrecarrega a memória do computador. Para dar uma idéia, se uma imagem da tela for dividida em células 2x2, serão necessárias 65.536 classificações, havendo necessidade de um arquivo em bytes de (512.4,2) isto é, que seja capaz de armazenar duas linhas da imagem da tela constituída de 512 colunas em cada um dos quatro canais: o que dá aproximadamente 2k de memória. Por outro lado, se tomarmos células 4x4, serão executadas 16384 classificações com o uso de um arquivo de (512.4,4) correspondendo aproximadamente a 4k de memória.

O algoritmo aqui desenvolvido dá opção de escolha de células até 8x8, a critério do operador.

O limite do teste de homogeneidade também se constitui numa variável importante para o particionamento preciso da imagem. Um limite muito baixo pode rejeitar um número muito grande de células e mesmo aquelas que se encontram a uma distância razoável da fronteira, ou seja, uma variação mínima na variância já pode tornar uma célula não homogênea; por outro lado, um limite muito grande pode aceitar células cuja variância seja elevada, como células homogêneas. Portanto há a necessidade de escolher um limite ótimo, o que é conseguido por tentativa e erro, com o valor inicial baseado na informação dada no Capítulo V.

Quanto ao tempo de processamento, ele, em relação ao algoritmo existente, devido ao fator deste ter sido implementado totalmente na linguagem FORTRAN-IV, ainda é grande. Isto se deve também à constituição das imagens em que foram testados os programas, que não permitiu a definição de células maiores.

Para torná-lo competitivo com o sistema existente, há a necessidade de se criarem rotinas em ASSEMBLER, o que poderá diminuir consideravelmente o tempo de processamento.

Na Figura VII.2 é mostrado o resultado obtido com a técnica convencional sem memória e com o ECHO.

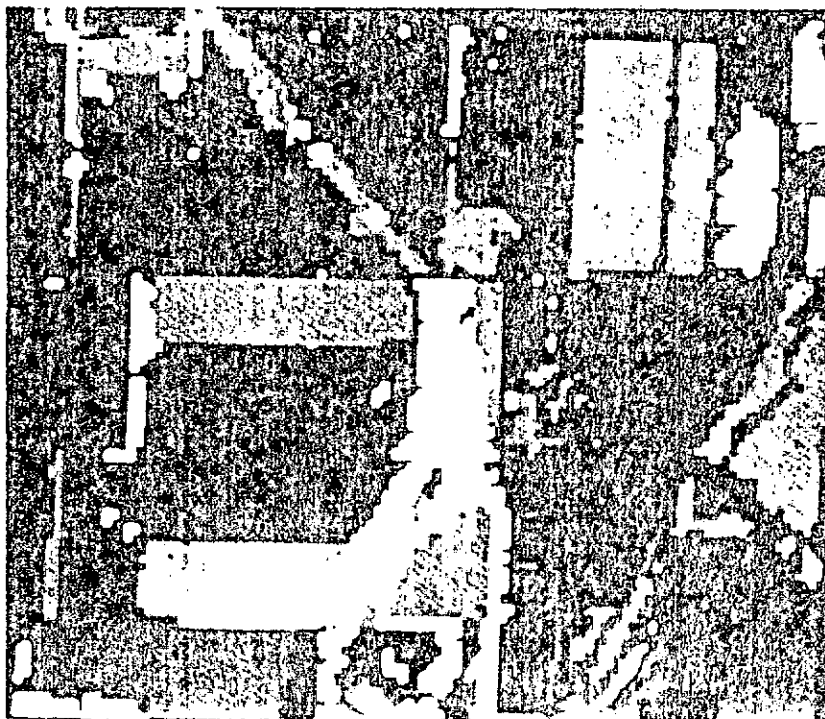
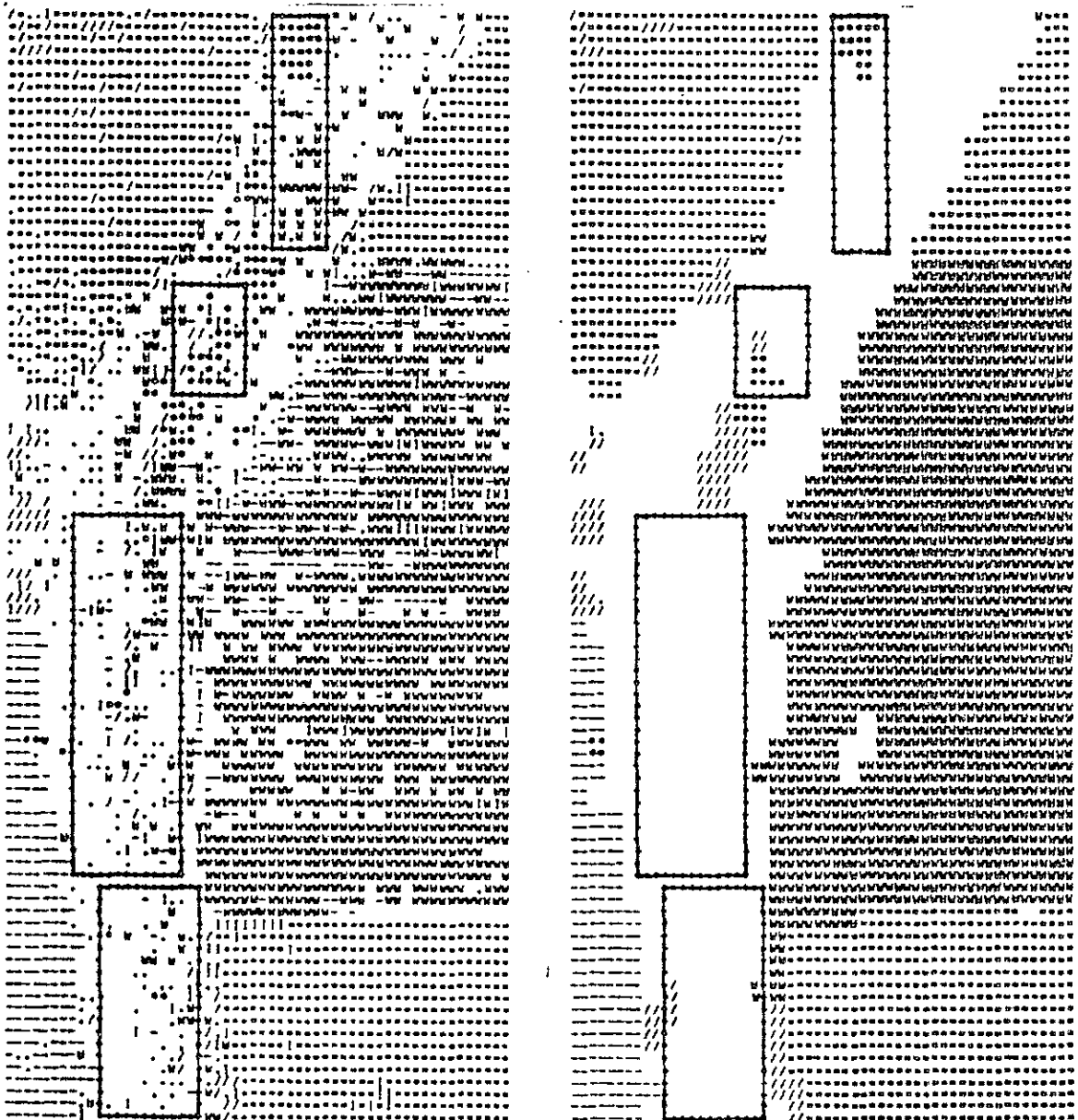


FIG. VII.1 - Imagem na qual foi testado o programa nos Estados Unidos.
(Reproduzido de Ketig & Landgrebe (1975))



CLASSIFYPOINTS

ECHO

SYMBOL	CLASS
W	Wheat
=	Lespedeza
-	Pasture
blank	Wooded Pasture

SYMBOL	CLASS
I	Idle
*	Forest
.	Corn, Soy, Rye, Hay
/	Non-Farm

FIG. VII.2 - A Figura da esquerda mostra o resultado de uma classificação usando a técnica convencional sem memória e da direita usando o método ECHO.

(Reproduzido de Ketig & Landgrebe)

BIBLIOGRAFIA

ANDERSON, T.W. *An introduction to multivariate statistical analysis.*
New York, Wiley & Sons, 1958.

KETTING, R.L.; LANDGREBE, D.A. *Computer classification of remotely
sensed multispectral image data by extraction and classification
of homogeneous objects.* Indiana, ~~University, Purdue~~ 1975.

APENDICE A

Listagem dos programas utilizados AUX00, AUX01, INAPAR e suas sub-rotinas...

6.4.1 Properties NIXOD.

2 INICIALIZACAO

C N U H E : A U X U O

```

C DESCRICAO: INICIALIZA OS PARAMETROS INDEPENDENTES DO TREINA-
C MENTO (NO. MAX. DE CLASSES, DE AMOSTRAS E ESCALA).
C CASO SEJA UMA NOVA SESSAO, ADQUIRE O NO. DE DIMEN-
C SODES.

```

```

IMPLICIT INTEGER(A-Z,S-Z)
LOGICAL*1 LK(74)
DOUBLE PRECISION DOWECL(16)
COMMON/COMMON1/

```

```

1  NDTM, VDTCL, NDMAN,
2  VDMCL, DETERM(15), RAPICL(15), MEDCL(54), RMEDCL(64),
3  TRNSCL(160), PAUTCL(160), AMOSTR(390),
4  NMAXCL, NMAXAM, PSCALE, PAI
   DATA R/6RMAXVER/

```

```
WRITE(6,100)
FORMAT(21X,'*** INICIALIZACAO ***',//)
```

C. PARÂMETROS INDEPENDENTES

```

NMAXCL=16
NMAXAM=78
RSCALE=30.

```

C NOVA SESSÃO

NUMCL=0
NUMAM=0

```

200 WRITE(6,220)
220 FORMAT('SENTE COM O NO. DE DIMENSÕES(CANAIS) >')
    CALL DDINPUT(7)
    READ(6,230)N
230 FORMAT(74A1)
    IF(
    CALL INTEFF(1,N,74,NDIM)
    IF (NDIM.GE.1.AND.NDIM.LE.4)GO TO 800
    WRITE(6,240)
240 FORMAT(1X,'ERRO >>> NO. DE DIMENSÕES INVALIDO')
    GO TO 200

```

```

800      FAT=(NDIM*NDIM+NDIM)/2
        WRITE(6,900)
900      FORMAT(///,21X,' *** INICIALIZACAO TERMINADA *** ')
        CALL OUTPUT(7)
        CALL RESUME(R)
        CALL EXIT
        END

```

6.4.2 Programa AUXOL.

C AQUISIÇÃO DE AMOSTRAS

NAME: AUXOI

3 DESCRICAO: CALCULA A DISTRIBUICAO DA AMOSTRA INDICADA
3 PELO CURSOR, CRIANDO UMA NOVA CLASSE OU ADICIO-
3 NANDO A AMOSTRA A UMA CLASSE JA EXISTENTE.

```

C SUBROUTINAS: DISIR
C               ELE9

```

IMPLICIT INTEGER(A-Z,S-Z)

LOGICAL#1 w(74)

DOUBLE PRECISION BOREC(16)

COMCON/COMINT/

1 NDIV, VEHICLE, NU-4A4.

2 NMEDCL,DETERM(15),RNPICL(15),MEDCL(64),RMEDCL(64),

3 TRANSCL(160),PAUICL(160),A4DSIR(340),

4 - NMAXCL, NMAXAM, RSCALE, FAI

DIMENSION CURSOR(S)

334404/PARAM/RYPTAN, RMEDAN(4), RAUTAN(10)

COMBIV/CDORD/XI, XI, XI, XI

DATA 3/53MAXVFR/

CALL OUTPUT(27,12)

REF(6,50)

```
60      FORMAT(21X,'***AQUISICAO DE PARAMETROS***',///)
```

CALL OUTPUT(7)

IF(VJMAN,LT,NMAXAN) GO TO 73

KRIFE(5,79)

```
70      FORMAT(IX,'ERRO >>> EXCEDEU NO. MAXIMO DE AMOSTRAS')
```

CALL INPUT(7)

GO TO 999

```

C
73      WRITE(6,75)

```

```
75      FORCAT('POSICION DE CURSOR SOBRE A AJUSTA >')
```

CALL OUTPUT(7)

```
READ(5,120)N
```

C LER CURSOR

60 CALL LINK(CURSOR)

TESTAR SE J AGDO ESTA CERTO

```
IF(CURSOR(1).EQ.1)GO TO 95
```

WR 1516, 90)

```
90  FORMAT(IX,'ERRR >>> CURSOR NAU ESTA NO PODE RETANGULO')
```

CALL OUTPUT(7)

GO 13 73

C. CALCULO DOS PARAMETROS DA AMOSTRA

CALCULO DO VO. DE PONTOS


```

15  RC3=CURSQR(3)
    RC5=CURSQR(5)
    RNPIA=15*RC3+RC5+8*(RC3+RC5)+4

    CALCULO DOS LIMITES XI, XF, YI, YF

    XI=CURSQR(2)-2*CURSQR(3)-2
    XF=CURSQR(2)+2*CURSQR(3)+1
    YI=CURSQR(4)-2*CURSQR(5)
    YF=CURSQR(4)+2*CURSQR(5)+1

100  CALL DISIR

    APRESENTACAO DOS PARAMETROS DA AMOSTRA

    WRITE(6,103)RNPIA,(RVEDA*(I),I=1,NDIM)
103  FORMAT(1X,'PARAMETROS DA AMOSTRA',/,
1  1X,'NO. DE PONTOS = ',F7.0,/,
2  1X,'MEDIA = ',AF3.2)
    WRITE(6,107)(RAUT*(I),I=1,FAT)
107  FORMAT(1X,'MATRIZ DE ADICORRELACAO = ',5F8.2,/,24X,5F8.2)

    NOVA CLASSE?

105  WRITE(6,110)
110  FORMAT('ADICIONAR AMOSTRA(A) OU CRIAR NOVA CLASSE(C)? >')
    CALL OUTPUT(7)
    READ(5,120)A
120  FORMAT(74A1)
    CALL FROM1(74,A)
    IF(A(1).EQ."103")GO TO 500
    IF(A(1).EQ."101")GO TO 133
    IF(A(1).EQ."130")GO TO 999
    GO TO 105

    ADICIONAR AMOSTRA

130  IF(NUMCL.GT.0) GO TO 135
    WRITE(6,133)
133  FORMAT(1X,'ERRO >>> NAO HA CLASSES')
    CALL OUTPUT(7)
    GO TO 105

135  CALL OUTPUT(27,12)
    WRITE(6,140)
140  FORMAT(1X,'CLASSES ATE AGORA',/)
    DO 150 I=1,NUMCL
    WRITE(6,160)I,NUMCL(I)
150  CONTINUE

160  FORMAT(1X,12,' ',A8)
170  WRITE(6,175)
175  FORMAT(/,'SQUAL CLASSE(F)? >')
    CALL OUTPUT(7)
    READ(6,120)A
```

```
CALL FROM(74,N)
IF(N(1).EQ."130")GO TO 999

C
I=0
CALL INIFF(1,N,74,CLAT)
IF (CLAT.GE.1.AND.CLAT.LE.NUMCL)GO TO 178
WRITE(6,177)
177 FORMAT(1X,'ERRO >>> CLASSE INVALIDA')
CALL OUTPUT(7)
GO TO 170

C
C ADICIONAR AMOSTRA
C
C
C CALCULO DOS LIMITES DOS PARAMETROS DA CLASSE CLAT:
C 11,12 - RMEDCL
C J1,J2 - RAJICL
C
178 I1=4*(CLAT-1)+1
I2=I1+ND1X-1
J1=10*(CLAT-1)+1
J2=J1+PAT-1

C
C ATUALIZACAO DOS PARAMETROS DA CLASSE
C
RT=RNPICL(CLAT)+RNPTAM
RC=RNPICL(CLAT)/RT
RA=RNPTAM/RT
RNPICL(CLAT)=RT

C
J=0
DO 180 I=I1,I2
J=J+1
RMEDCL(I)=RA*RAEDAN(J)+RC*RMEDCL(I)
180 CONTINUE

C
J=0
DO 200 I=J1,J2
J=J+1
RAJICL(I)=RA*RAUTAX(J)+RC*RAJICL(I)
200 CONTINUE
GO TO 800

C
C CRIAR NOVA CLASSE
C
600 IF(NUMCL.LT.NMAXCL)GO TO 620
WRITE(6,610)NMAXCL
610 FORMAT(1X,'ERRO >>> NO. MAXIMO DE CLASSES EH ',13)
CALL OUTPUT(7)
GO TO 105

C
620 NUMCL=NUMCL+1
CLAT=NUMCL

C
I1=4*(CLAT-1)+1
I2=I1+ND1X-1
```

```

      J1=10*(CLAT-1)+1
      J2=J1+FAI-1
C
C  NOME DA NOVA CLASSE
C
      WRITE(6,630)
630  FORMAT('SENTE COM O NOME DA CLASSE >')
      CALL OUTPUT(7)
      READ(6,640)NOMECL(CLAT)
640  FORMAT(A8)
C
C  PREENCHER DADOS DA CLASSE:
C  RNPICL,RMEDCL,RAUTCL
C
      RNPICL(CLAT)=RNPITAM
C
      J=0
      DO 650 I=11,12
      J=J+1
      RMEDCL(I)=RMEDAM(J)
650  CONTINUE
C
      J=0
      DO 660 I=J1,J2
      J=J+1
      RAUTCL(I)=RAUTAM(J)
660  CONTINUE
C
C  PREENCHER DADOS DA AMOSTRA:
C  CLAT,XI,XF,YI,YF
C
800  NUMAM=NUMAM+1
      AM=5+NUMAM
      AMOSTR(AM-4)=CLAT
      AMOSTR(AM-3)=XI
      AMOSTR(AM-2)=XF
      AMOSTR(AM-1)=YI
      AMOSTR(AM) =YF
C
C  CALCULO DA MATRIZ DE TRANSFORMACAO
C
      CALL ELEM(CLAT)
C
C  DISPLAY DOS PARAMETROS
C
C
D      WRITE(6,830)CLAT,NOMECL(CLAT),RNPICL(CLAT),(RMEDCL(I),I=11,12)
D      WRITE(6,840)(MEDCL(I),I=11,12)
D      WRITE(6,850)(RAUTCL(J),J=J1,J2)
D      WRITE(6,860)(TRVSCL(J),J=J1,J2)
D830  FORMAT(1X,'PARAMETROS DA CLASSE',13,2X,A8,/,
D      1 1X,'NB. DE PONTOS = ',F7.0,/,
D      2 1X,'MEDIA = ',4F8.2)
D840  FORMAT(1X,'MEDCL',416)
D850  FORMAT(1X,'MATRIZ DE AUTOCORRELACAO = ',5F8.2,/,28X,5F8.2)
D860  FORMAT(1X,'TRVSCL',1016)

```

```

C
C
      WRITE(6,870)
870   FORMAT(///,21X,'***AQUISICAO COMPLETADA***')
      CALL DDIPUI(7)
C
999   CALL RESUME(R)
      CALL EXIT
      END

```

PROGRAM SECTIONS

NAME	SIZE	ATTRIBUTES
SCDDSI	002535 687	RW,I,CON,LCL
SPDATA	000020 8	RW,D,CON,LCL
SIDATA	001126 299	RW,D,CON,LCL
SVARS	000174 62	RW,D,CON,LCL
COM41	004474 1182	RW,D,OVR,GBL
PARA4	000074 30	RW,D,OVR,GBL
COORD	000010 4	RW,D,OVR,GBL

TOTAL SPACE ALLOCATED = 010700 2272

AUX01,LP:=OK0:AUX01

6.4.3 - SUB-ROTINA DISTR

C DISTRIBUICAO

C NOME: DISTRI

C DESCRICAO: CALCULA OS SEGUINTEs PARAMETROS: MEDIA (PMEDAM),
C MATRIZ DE AUTOCORRELACAO (RAUTAM) DA AMOSTRA DADA
C PELOS LIMITES XI, XF, YI, YF.

SUBROUTINE DISTR

IMPLICIT INTEGER(A-Z)

LOGICAL*1 BUF1(512),BUF2(512),BUF3(512),BUF4(512)

DOUBLE PRECISION NOMECL(16)

COMMON/COMMON1/

1 NDIM,NUMCL,NUMAM,

2 NOMECL,DETERM(16),RNPICL(16),MEDCL(64),PMEDCL(64),

3 IRNSCL(160),RAUTCL(160),ANDSIR(390),

4 NMAXCL,NMAXAM,RSCALE,FAT

COMMON/PARAM/RNPJAM,PMEDAM(4),RAUTAM(10)

COMMON/COORD/XI,XF,YI,YF

DIMENSION RPTD(4)

C INICIALIZACAO

DO 10 I=1,NDIM

PMEDAM(I)=0.

10 CONTINUE

DO 15 I=1,FAT

RAUTAM(I)=0.

15 CONTINUE

DO 40 LINHA=YI,YF

C LEITURA DAS LINHAS

CALL IRV(1,LINHA,BUF1)

CALL WAIT

CALL IRV(2,LINHA,BUF2)

CALL WAIT

CALL IRV(3,LINHA,BUF3)

CALL WAIT

CALL IRV(4,LINHA,BUF4)

CALL WAIT

DO 40 COLON=XI,XF

C LEITURA DO PONTO

RPIO(1)=IBYTE(0,BUF1(COLON))

RPIO(2)=IBYTE(0,BUF2(COLON))

RPIO(3)=IBYTE(0,BUF3(COLON))

RPIO(4)=IBYTE(0,BUF4(COLON))

C ACUMULAR MEDIA

DO 20 I=1,NDIM

```

      RMEDAM(1)=RMEDAM(1)+RPFD(1)
20    CONTINUE
      C
      C ACUMULAR AUTOCORRELACAO
      C
      K=0
      DO 30 I=1,NDIM
      DO 30 J=1,I
      K=K+1
      RAUTAM(K)=RAUTAM(K)+RPFD(1)*RPFD(J)
30    CONTINUE
40    CONTINUE
      C
      C CALCULO DA MEDIA
      C
      DO 50 I=1,NDIM
      RMEDAM(I)=RMEDAM(I)/K
50    CONTINUE
      C
      C CALCULO DA AUTOCORRELACAO
      C
      DO 60 I=1,FAT
      RAUTAM(I)=RAUTAM(I)/K
60    CONTINUE
      C
      RETURN
      END

```

PROGRAM SECTIONS

NAME	SIZE	ATTRIBUTES
SCODE1	001054 278	RW,1,CON,LCL
SPDATA	000024 10	RW,0,CON,LCL
SIDATA	000050 20	RW,0,CON,LCL
SVAPS	004032 1037	RW,0,CON,LCL
STEPS	000004 2	RW,0,CON,LCL
CDRM1	004474 1182	RW,0,OVR,GBL
PAR44	000074 30	RW,0,OVR,GBL
CDRMD	000010 4	RW,0,OVR,GBL

TOTAL SPACE ALLOCATED = 012006 2563

DISTR,LP:=DOKO:DISTR

6.4.4 Subrotina ELEM.

C PARAMETROS DA CLASSIFICACAO

C NOME: ELEM

C DESCRICAO: CALCULA, A PARTIR DA AUTOCORRELACAO DA CLASSE CLAT
(RAUTCL(J1,J2)) E DA MEDIA (RMEDCL(11,12)), A MATRIZ
DE TRANSFORMACAO (TRNSCL(J1,J2)) DA CLASSE CLAT, A MEDIA
(MEDCL(11,12)) E O LOGARITMO DO DETERMINANTE
(DETERM(CLAT)).

C SUBROTINAS: INTC

SUBROUTINE ELEM(CLAT)

IMPLICIT INTEGER(A-Z)

REAL ALOS

DOUBLE PRECISION RMEDCL(16)

COMMON/COMMON1/

1 NDIM,NUMCL,NUMAN,

2 RMEDCL,DETERM(16),RSPICL(16),RMEDCL(64),RMEDCL(64),

3 TRNSCL(160),RAUTCL(160),AOSIR(390),

4 VMAXCL,VMAXAN,RSCALA,FAT

DIMENSION RCOVAR(10),RELE(10),RELEM1(10),RSAI(4)

I1=4*(CLAT-1)+1

I2=I1+NDIM-1

J1=10*(CLAT-1)+1

J2=J1+FAT-1

C CALCULO DA MATRIZ DE COVARIANCIA

C CALCULO DE M*MT

K=0

DO 10 I=I1,I2

DO 10 J=I1,I

K=K+1

RCOVAR(K)=RMEDCL(I)*RMEDCL(J)

10 CONTINUE

J=0

DO 20 I=J1,J2

J=J+1

RCOVAR(J)=RAUTCL(I)-RCOVAR(J)

20 CONTINUE

0 WRITE(6,25)(RCOVAR(J),J=1,FAT)

25 FORMAT(1X,'COV',/,1X,10F7.2)

C CALCULO DE L

IF(RCOVAR(1).LE.0)GO TO 200

RELE(1)=SQRT(RCOVAR(1))

DO 30 I=1,NDIM

J=(I+1-1)/2+1

RELE(J)=RCOVAR(J)/RELE(1)

30 CONTINUE

```

C
C
DO 40 I=2,NDIM
  L=(I+I-1)/2
DO 40 J=2,I
  M=(J+J-J)/2
  RF=0.
  K2=J-1
DO 35 K=1,K2
  RF=RF+RELE(L+K)*RELE(M+K)
35 CONTINUE
  IF(I.EQ.J) GO TO 37
  IF(RELE(M+J).EQ.0) GO TO 200
  RELE(L+J)=(RCOVAR(L+J)-RF)/RELE(M+J)
  GO TO 40
37 RF=RCOVAR(L+J)-RF
  IF(RF.LT.0)GO TO 200
  RELE(L+J)=SORT(RF)
40 CONTINUE
D WRITE(6,41)(RELE(J),J=1,FAT)
041 FORMAT(1X,'ELE',/,1X,10F7.2)
C

```

C CALCULO DE DETERM

```

C
RDET=1.
DO 50 I=1,NDIM
  J=(I+I+1)/2
  RDET=RDET*RELE(J)
50 CONTINUE
D WRITE(6,55)RDET
055 FORMAT(1X,'RDET=',F7.2)
  RDET0=RDET*RDET
  IF(RDET0.EQ.0) GO TO 200
  RDOL=ABS(RDET0)
D WRITE(6,50)RDOL
D60 FORMAT(1X,'RDOL=',F7.2)
  RIEMP=RSCALE*HSCALE*RDOL
  DETERM(CLAT)=INTC(RIEMP)
D WRITE(6,65)DETERM(CLAT)
D65 FORMAT(1X,'DETERM=',1X)
C

```

C INVERSAO DE L

```

C
DO 130 I=1,NDIM
DO 130 JJ=1,I
  J=I+I-JJ
  L=(I+I-1)/2+J
  IF(I.EQ.JJ)GO TO 145
  RF=0.
  K2=I-1
DO 143 K=J,K2
  M=(K+K-K)/2+J
  N=(I+I-1)/2+K
  RF=RF+RELE(N)*RELE*1(K)
143 CONTINUE
  H=(I+I+1)/2

```



```

      RELEM1(L)=-RF*RELEM1(H)
      GO TO 150
145    IF (RELE(L).EQ.0) GO TO 200
      RELEM1(L)=1./RELE(L)
150    CONTINUE
C
D      WRITE(6,155)(RELEM1(J),J=1,FAT)
D155   FORMAT(1X,'ELEM1',/,1X,10F7.2)
C
C  ESCALA E GUARDE A INVERSA DE L
C
      J=0
      DO 160 I=J1,J2
      J=J+1
      RTEMP=RSCALA*RELEM1(J)
      TRNSCL(I)=INIC(RTEMP)
160    CONTINUE
C
C  GUARDAR MEDIAS DAS CLASSES
C
      DO 170 I=11,12
      MEDCL(I)=INTC(RMEDCL(I))
170    CONTINUE
C
C  APRESENTACAO DOS NOVOS PARAMETROS
C
      WRITE(6,172)CLAT,NOMECL(CLAT)
172    FORMAT(/,1X,'PARAMETROS DA CLASSE',13,2X,A6,/)
      WRITE(6,174)RNPICL(CLAT),(RMEDCL(I),I=11,12)
174    FORMAT(1X,'NO. DE PONTOS = ',F7.0,/,
1 1X,'MEDIA = ',4F8.2)
      WRITE(6,176)
176    FORMAT(1X,'MATRIZ DE COVARIANCIA:')
C
      DO 185 I=1,NDIM
      DO 182 J=1,NDIM
      IF (I.LT.J) GO TO 178
      K=(I+J-1)/2+J
      GO TO 180
178    K=(J+J-J)/2+1
180    RSAI(J)=RCOVAR(K)
182    CONTINUE
      WRITE(6,184)(RSAI(J),J=1,NDIM)
184    FORMAT(23X,4F8.2)
186    CONTINUE
C
C
      RETURN
C
200   WRITE(6,300)
300   FORMAT(1X,'AVISO >>> AMOSTRA INSUFICIENTE PARA ',
1 1X,'CLASSIFICACAO')
      CALL OUTPUT(7)
C
      RETURN
      END

```

6.4.5 Function INTG

```
C CONVERSÃO COM APROXIMAÇÃO
C
C NOME: INTG
C
C DESCRIÇÃO: FAZ A CONVERSÃO DE REAL PARA INTEIRO APROXIMANDO
C             PARA O INTEIRO MAIS PROXIMO.
C
C
C     FUNCTION INTG(RX)
C     IMPLICIT INTEGER(A-Z)
C     REAL AINT
C     RY=AINI(RX)
C     I=INT(RX)
C     IF (RX.LT.0) GO TO 10
C
C     CASO POSITIVO
C
C         RZ=RX-RY
C         IF (RZ.GT.0.5) I=I+1
C         GO TO 20
C
C     CASO NEGATIVO
C
C 10      RZ=RY-RX
C         IF (RZ.GT.0.5) I=I-1
C
C 20      INTG=I
C
C     RETURN
C     END
```

[illegible][illegible]

THE UNIVERSITY OF CHICAGO

[illegible][illegible][illegible]

() () () () () ()

, , r s ()

[illegible][illegible][illegible]

20

```
30  WRITE(6,40)1
40  FORMAT('SCLASSE>',12,2X,'TEMA>')
    CALL OUTPUT(7)
    READ(5,80)A
    CALL FRONT(W,74)
    J=0
    CALL INTFF(J,A,74,TE)
    IF(TE.GE.0.AND.TE.LE.8)GO TO 60
    WRITE(6,50)
50  FORMAT(1X,'ERRO>>>TEMA INVALIDO')
    GO TO 30
60  CSCL(I+1)=TEMA(TE+1)
70  CONTINUE
80  FORMAT(74A1)
C
C
C
C
C
C
90  WRITE(6,100)
100 FORMAT(7,1X,'CONSIDERAR CLASSES EQUIPROVAVEIS(C) DO
*  ENTRAR COM PROBABILIDADES(R)>')
    CALL OUTPUT(7)
    READ(6,110)R
110  FORMAT(74A1)
    CALL FRONT(74,R)
    IF(R(1).EQ."103")GO TO 150
    IF(R(1).EQ."105")GO TO 120
    GO TO 90
C
C
C
C
C
C
    ENTRAR COM AS PROBABILIDADES A PRIORI
120  DO 140 I=1,NUMCL
    WRITE(6,130)1
130  FORMAT(7,1X,'PROBAB. A PRIORI DA CLASSE',13,'=')
    READ(5,110)R
    CALL FRONT(R,74)
    J=0
    CALL REALFF(J,R,74,RLGGC(1))
140  CONTINUE
    GO TO 170
C
C
C
C
C
C
    CONSIDERAR AS CLASSES EQUIPROVAVEIS
150  DO 160 I=1,NUMCL
    RLGGC(1)=1./NUMCL
160  CONTINUE
170  DO 190 I=1,NUMCL
    WRITE(6,180)I,RLGGC(1)
180  FORMAT(7,1X,'CLASSE=',13,2X,'PROB. A PRIORI =',15.3)
    RLGGC(1)=ALOG(RLGGC(1))
190  CONTINUE
```

CCCCC

INICIALIZACAO DO NO. DE PONTOS DA AMOSTRA(RNPTAM)

200

RNPTAM=ICELA*ICELA
DO 200 I=1,NUMCL
RNPT(I)=0.
CONTINUE
ICD+I=0

CCCCC

PREPARACAO DOS PARAMETROS DAS CLASSES PROPOSTAS

CCCCC

DO 200 CLAT=1, NUMCL
I1=4*(CLAT-1)+1
I2=I1+NDIA-1
J1=10*(CLAT-1)+1
J2=J1+FAT-1

CALCULO DA MATRIZ COVARIANCA DA CLASSE DEFINIDA PELOS LIMITES I1,I2 E J1,J2

210

220

230

240

250

K=J1
DO 210 IC=I1,I2
DO 210 J=I1,IC
RCOVAR(K)=R*EDCL(IC)*R*EDCL(J)
K=K+1
CONTINUE
DO 220 IC=J1,J2
RCOVAR(IC)=RAUECL(IC)-RCOVAR(IC)
CONTINUE
DO 230 IC=1,NDIA
DO 230 J=1,NDIA
IF(IC.LT.J)GO TO 230
K=(J+J*J)/2+IC+J1-1
GO TO 240
K=(J+J*J)/2+IC+J1-1
RCOVCL(IC,J)=RCOVAR(K)
RAUX(IC,J)=2.*3.142*RCOVCL(IC,J)
CONTINUE

CCCCC

APRESENTACAO DAS CARACTERISTICAS DAS CLASSES

260

WRITE(5,260)CLAT,(RMEUCL(IC),IC=I1,I2),((RCOVCL(IC,J),
*J=1,NDIA),J=1,NDIA)
FOR=4(7,1X,'CLASSE',13,7,1X,'VEIOR MEDIA',4F14.2,/,
*1X,'MATRIZ COVAR.',7,1X,4(7,1X,4F14.2))

CCCCC

CALCULO DO LOG(P(Y/V1))

```

CALL DETER(RAUX,NDIM,RDET)
WRITE(5,270)RDET
270  FORMAT(7,1X,'DETERMINANTE=',F14.2)
      RALD(CLAT)=ALOG(ABS(RDET))
      CALL ATINVER(RCIVCL,NDIM)
      IC=15*(CLAT-1)+1
      K=IC
      DO 280 I=1,NDIM
      DO 290 L=1,NDIM
      RVEIUR(K)=RCIVCL(I,L)
      K=K+1
280  CONTINUE
290  CONTINUE

      DEFINICAO DA REGIAO A SER PARTICIONADA

300  WRITE(6,310)
310  FORMAT(7,1X,'TODA TELA ?(S/N)>')
      CALL OUTPUT(7)
      READ(5,320)*
320  FORMAT(74A1)
      CALL PRJNT(74,*)
      IF(N(1).EQ."123")GO TO 330
      IF(N(1).EQ."116")GO TO 340
      GO TO 300

      TODA A TELA

330  SINAD=0
      XI=1
      XF=512
      YI=1
      YF=512
      RNP=512*512
      GO TO 380

      SO A AREA DO CURSOR

340  SINAD=1
      WRITE(6,350)
350  FORMAT(1X,'POSICIONE O CURSOR>')
      READ(5,320)*

      LER O CURSOR

      CALL IRS(CURSOR)

```

TESTAR SE O CURSOR ESTA NO MODO RETANGULO

IF(CURSOR(1).EQ.1)GO TO 370

WRITE(6,360)

360 FORMAT(1X,'ERRO>>>CURSOR NAO ESTA NO MODO RETANGULO')
GO TO 340

CALCULO DAS COORDENADAS DO CURSOR

370 X1=CURSOR(2)-2*CURSOR(3)-2

X2=CURSOR(2)+2*CURSOR(3)+1

Y1=CURSOR(4)-2*CURSOR(5)

Y2=CURSOR(4)+2*CURSOR(5)+1

CALCULO DA AREA DO CURSOR

RC3=CURSOR(3)

RC5=CURSOR(5)

RNP=16*RC3*RC5+8*(RC3+RC5)+4

ENTRAR COM O LIMITE PARA TESTAR A HOMOGENEIDADE
DAS CELAS

380 WRITE(6,385)

385 FORMAT(1X,'LIMITE DE HOMOGENEIDADE>')

READ(8,80)A

CALL FRONT(A,74)

J=0

CALL READFF(J,A,74,RLIM)

ENTRAR COM LIMITE PARA TESTAR ACITACAO DA CELA
NA CLASSE

WRITE(6,386)

386 FORMAT(1X,'LIMITE PARA FAZAO DE VEROSSIM.>')

READ(8,80)B

CALL FRONT(B,74)

J=0

CALL READFF(J,B,74,RI)

RI=-ALOG(10.)*RI

DIVISAO DA IMAGEM EM CELAS

```

DO 720 I=YI-1,YF-ICELA,ICELA
DO 390 J=0,TECLA-1
DO 390 L=1,NDIM
CALL IRV(L,I+J+1,TELA(1,L,J+1))
CALL WAIT
CONTINUE
IF(SIGNAL.EQ.0)GO TO 420

PREPARACAO DAS LINHAS DE IMPRESSAO

DO 400 IK=1,ICELA
DO 400 JK=1,XI
TELA(JK,1,IK)=CSCL(1)
CONTINUE
DO 410 IK=1,ICELA
DO 410 JK=XF+1,512
TELA(JK,1,IK)=CSCL(1)
CONTINUE
DO 700 ICOL=XI-1,XF-ICELA,ICELA

INICIALIZACAO DOS PARAMETROS DA CELA

DO 430 IV=1,NDIM
RAMEO(IV)=0.
CONTINUE
DO 440 IW=1,FATR
RAOTIA4(IV)=0.
CONTINUE

DO 470 ILIN=0,ICELA-1
DO 470 II=0,TECLA-1
DO 450 IC=1,NDIM
RPID(1C)=1BYTE(0,TELA(ICOL+II+1,IC,ILIN+1))
CONTINUE
DO 450 IC=1,NDIM
RAMEO(1C)=RAMEO(1C)+RPID(1C)
CONTINUE
K=0
DO 470 IC=1,NDIM
DO 470 J=1,IC
K=K+1
RAOTIA4(K)=RAOTIA4(K)+RPID(1C)+RPID(J)
CONTINUE

CALCULO DA MEDIA DA CELA

DO 480 IC=1,NDIM
RAMEO(1C)=RAMEO(1C)/RNPTRAS
CONTINUE

```


CALCULO DA AUTOCORRELACAO DA CEGA

```
DO 490 J=1,FAT
  RAUTAM(J)=PAUTAM(J)/RDEPTAM
CONTINUE
```

CALCULO DA COVARIANCA(RACOV) A PARTIR DA AUTOCORRELACAO

```
K=0
DO 500 IC=1,NDIM
  DO 500 J=1,IC
    K=K+1
    RACOV(K)=RAMEO(IC)*RAMEO(J)
  CONTINUE
DO 510 II=1,FAT
  RACOV(II)=RAUTAM(II)-RACOV(II)
CONTINUE
```

APRESENTACAO DO VETOR ACOVAR NA FORMA MATRICIAL PCOV(1,

```
DO 540 IC=1,NDIM
  DO 540 II=1,NDIM
    IF(IC.GT.II)GO TO 520
    K=(IC+IC-IC)/2+II
    GO TO 530
  520 K=(II+II-II)/2+IC
  530 PCOV(IC,II)=RACOV(K)
  540 CONTINUE
```

```
DO 610 CLAT=1,NUNCL
  II=4*(CLAT-1)+1
  I2=II+NDIM-1
  J1=10*(CLAT-1)+1
  J2=J1+FAT-1
  RTRABO=0.
```

CALCULO DE RAMEO-RAMEOCL,(K=1)

```
K=1
DO 550 J=II,12
  RAME(K)=RAMEO(K)-RAMEOCL(J)
  K=K+1
CONTINUE
```

CALCULO DE (K-31)*(K-41)

C

```
DO 560 IC=1,NDIM
DO 560 J=1,NDIM
RAUX(IC,J)=RAME(IC)*RAYE(J)
560 CONTINUE
```

560

CCCCCCCC

CALCULO DE $C+(M-MI)*(M-MI)'$

```
DO 570 IC=1,NDIM
DO 570 J=1,NDIM
RCOV(IC,J)=RCOV(IC,J)+PAUX(IC,J)
570 CONTINUE
```

570

CCCCCCCC

COLOCACAO DA INVERSA DA MATRIZ COVARIANCA(RVETOR),
NA FORMA MATRICIAL PADRAO

```
IC=15*(CLAT-1)+1
J=IC
DO 580 K1=1,NDIM
DO 580 K2=1,NDIM
RCOVCL(K1,K2)=RVETOR(J)
J=J+1
580 CONTINUE
```

580

CCCCCCCC

-1
CALCULO DO TRACO DE $C1+(C+(M-MI)*(M-MI))'$

```
DO 590 IC=1,NDIM
DO 590 J=1,NDIM
DO 590 K=1,NDIM
IF(IC.EQ.J)RTRACO = RTRACO + RCOVCL(IC,K)*RCOV(K,J)
590 CONTINUE
```

590

CCCCCCCC

CALCULO DE $\text{LOG}(P(Y/YI))$

$\text{RLOG}(CLAT)=-.5*RNPTAM*(RLOG(CLAT)+RTRACO)$

CALCULO DO MAXIMO DE $\text{LOG}(P(Y/YI))$

```
IF(CLAT.EQ.1)GO TO 600
IF((RMAAA=RLOG(CLAT)).31.0)GO TO 610
RMAAA=RLOG(CLAT)
RIR=RTRACO
610 CONTINUE
RJJ=RNPTAM*RIR
```

600

610

CCCC

TESTE DE HOMOGENEIDADE DA CELA

IF(RJJ.GT.RLIM)GO TO 680
IF(ICONT.EQ.1)GO TO 640

CALCULO DO MAXIMO DE $\text{LDG}(P(X/V1))$

DO 630 CLAT=1,NUNCL
IF(CLAT.EQ.1)GO TO 620
IF((RMAXC-RLOGC(CLAT)).GT.0)GO TO 630
RMAXC=RLOGC(CLAT)
CONTINUE

CALCULO DO MAXIMO DE $(\text{LDG}(P(X/V1)) + \text{LDG}(P(Y/V1)))$

DO 650 CLAT=1,NUNCL
RLOGS(CLAT)=RLOGC(CLAT)+RLOG(CLAT)
IF(CLAT.EQ.1)GO TO 650
IF((RMAXS-RLOGS(CLAT)).GT.0)GO TO 660
RMAXS=RLOGS(CLAT)
CLAGR=CLAT
CONTINUE

CALCULO DO LOGARITMO DA RAZAO DE VEROSSIMILHANCA

$RLV=RMAXS-RMAXC-PNAXA$

TESTE ESTATISTICO ENTRE CELA E CLASSE

IF(RLV.LT.RT)GO TO 680
ICONT=1

NO CANAL I DO ARQUIVO "TELA", ARMAZENA-SE
O RESULTADO DA JUNCAO

RNPT(CLAGR)=RNPT(CLAGR)+RNPTAM
DO 670 ILIN=0,ICELA-1
DO 670 J=0,ICELA-1
TELA(ICOL+J+1,1,ILIN+1)=CSCL(CLAGR+1)
CONTINUE
GO TO 700
DO 690 ILIN=0,ICELA-1
DO 690 J=0,ICELA-1
TELA(ICOL+J+1,1,ILIN+1)=CSCL(1)
CONTINUE

```
700 CONTINUE
    DO 710 J=0, ICELA-1
    CALL JAV(5, I+J+1, TELA(1,1,J+1))
    CALL EXIT
710 CONTINUE
720 CONTINUE
C
C
C
C
C
    APRESENTACAO DOS PARAMETROS FINAIS DAS CLASSES

    CALL OUTPUT(27,12)
    WRITE(6,730)
730  FORMAT(/,1X,'PARAMETROS FINAIS DAS CLASSES',/)
    RNPTOT=0.
    DO 750 CLAT=1, NUMCL
    I1=4*(CLAT-1)+1
    I2=I1+NDIM-1
    J1=10*(CLAT-1)+1
    J2=J1+FAI-1
    RPERC=100.*RNPT(CLAT)/RNP
    WRITE(6,740)CLAT,NDIMCL(CLAT),RPERC
740  FORMAT(1X,'CLASSE',13,2X,A8,2X,'PERCENTAGEM='
    *,F5.0,' %',/)
    RNPTOT = RNPTOT + RNPT(CLAT)
750  CONTINUE
    RNPTOT=100.*RNPTOT/RNP
    WRITE(6,760)RNPTOT
760  FORMAT(/,1X,'PERCENTAGEM CLASSIF. =',F5.0,' %',/)
    WRITE(6,770)
770  FORMAT(/,19X,'***IMAGE PARTITIONING TERMINATED***')
    CALL OUTPUT(7)
    CALL EXIT
    END
```

6.4.7 - SUB-ROTINA DETER

```

SUBROUTINE DETER(A,M,DET)
  DIMENSION TXX(4,4),LV(4),ALFA(4),A(4,4)
  IF(M.EQ.1) GO TO 101
  DO 131 I=1,M
    DO 131 J=1,M
131  TXX(I,J)=A(I,J)
    DO 55 I=1,M
      55  LV(I)=1
      DET=1.
      DO 4 I=1,M
        SUP=0.
        DO 3 J=1,M
          IF(ABS(TXX(LV(I),J))-SUP) 3,3,2
2        SUP=ABS(TXX(LV(I),J))
3        CONTINUE
          IF(SUP.EQ.0.) GO TO 9
4        ALFA(LV(I))=1./SUP
          LCON=1
1        SUPC=ABS(TXX(LV(LCON),LCON)*ALFA(LV(LCON)))
          PIVD=TXX(LV(LCON),LCON)
          LPI=LCON
          ISER=0
          DO 5 J=LCON,M
            AMUL=TXX(LV(J),LCON)*ALFA(LV(J))
            IF(ABS(AMUL).LE.SUPC) GO TO 5
            SUPC=ABS(AMUL)
            PIVD=TXX(LV(J),LCON)
            LPI=J
            ISER=ISER+1
            IF(ISER.GT.1) GO TO 5
            DET=DET*(-1.)
5          CONTINUE
            IF(PIVD.EQ.0.) GO TO 9
            IAXX=LV(LCON)
            LV(LCON)=LV(LPI)
            LV(LPI)=IAXX
            LCON1=LCON+1
            DO 6 I=LCON1,M
              TXX(LV(I),LCON)=TXX(LV(I),LCON)/PIVD
            DO 6 J=LCON1,M
              TXX(LV(I),J)=TXX(LV(I),J)-(TXX(LV(LCON),J)*TXX(LV(I),LCON))
6            CONTINUE
            LCON=LCON+1
            IF(LCON=M) 7,7,7
7          DO 8 I=1,M
8            DET=DET*TXX(LV(I),I)
          RETURN
          *
          RETURN
101  DET=A(1,1)
      RETURN
      END

```

6.4.8 - SUB-ROTINA DETER

```
SUBROUTINE AINVER(A,N)
  DIMENSION A(4,4),U(4,8)
  IF(N.EQ.1) GO TO 5
  A(1,1)=1./A(1,1)
  GO TO 100
5  DO 10 I=1,N
  DO 15 J=1,N
15 U(I,J)=A(I,J)
  N1=N+1
  N2=N+2
  DO 10 J=N1,N2
  U(I,J)=0.
  IF(J-1.EQ.N) U(I,J)=1.
10 CONTINUE
  N1=N-1
  DO 60 K=1,N
  IP=K
  IF(U(K,K).EQ.0.) GO TO 17
  N1=N-1
  DO 30 J=1,N1
  IP=IP+1
  IF(IP.GT.N) IP=1
  IF(U(IP,K).EQ.0.) GO TO 30
  DO 40 I=1,N2
40 U(K,I)=U(K,I)+U(IP,I)
  GO TO 17
30 CONTINUE
  WRITE(5,2)
  2 FORMAT(1X,'NAO EXISTE INVERSA P/ ESTA MATRIZ')
  GO TO 100
17 AUX=U(K,K)
  DO 50 I=1,N2
50 U(K,I)=U(K,I)/AUX
  IP=K
  DO 60 J=1,N1
  IP=IP+1
  IF(IP.GT.N) IP=1
  AUX=-U(IP,K)
  DO 60 I=1,N2
60 U(IP,I)=AUX+U(K,I)+U(IP,I)
  DO 20 I=1,N
  N11=N+1
  DO 20 J=N11,N2
  K=J-N
  20 A(I,K)=B(I,J)
100 RETURN
END
```