

5 Implementação

O capítulo seguinte apresenta uma bateria de experimentos práticos realizados com o objetivo de validar o método proposto neste trabalho. O método envolve, contudo, alguns passos que podem ser implementados de maneiras distintas, aplicando diferentes abordagens disponíveis na literatura específica. Para validar experimentalmente a proposta, foi necessário escolher uma implementação específica para alguns dos passos indicados na figura 5. Este capítulo descreve como foram realizados neste trabalho os procedimentos de segmentação; seleção automática dos conjuntos de treinamento (produção automática das assinaturas / conhecimento espectral; método de fuzzificação espectral); representação do conhecimento multitemporal.

5.1. Procedimento de segmentação

Na literatura diversos métodos de segmentação de regiões contíguas homogêneas podem ser encontrados (Cocquerez, 1995; Gonzalez, 1993; Jain, 1989; Jähne 1991; Sonka, 1999; Priese, 2003) . O procedimento de segmentação produzido para este trabalho se baseia no método *watersheds* (divisor de águas) (Vincent, 1991; Gonzalez, 1993). Originalmente, este método baseia-se numa visão topográfica de uma imagem em tons de cinza, isto é, como uma superfície em três dimensões, onde a terceira dimensão é fornecida pela intensidade do pixel. Nesta superfície podem ser observados três tipos de pontos: **Tipo 1)** Pontos pertencentes a um mínimo local; **Tipo 2)** Pontos nos quais uma gota d'água ao cair escorreria para o mínimo local; **Tipo 3)** Pontos nos quais uma gota d'água ao cair teria a mesma probabilidade de escorrer para um ou outro mínimo local.

Para cada mínimo local (**Tipo 1**) o conjunto dos respectivos pontos do **Tipo 2** juntamente com o próprio mínimo local formam a **bacia de captação** daquele

mínimo. Os pontos do **Tipo 3** formam as chamadas **linhas de crista**. O objetivo do algoritmo *watersheds* é delinear as **linhas de crista**, para isso se supõe que os pontos do **Tipo 1** têm um furo por onde é injetada “água” numa taxa “uniforme”. Quando a água de duas bacias distintas está para se juntar, começa a ser construída uma **barreira**. Este processo prossegue até que todas as bacias estejam completamente “inundadas”. Assim, o conjunto das barreiras produzidas delimita e separa as regiões segmentadas; sem pertencerem a nenhuma das regiões por ela delimitadas.

O algoritmo baseia-se numa imagem contendo uma única intensidade. Como se está trabalhando com imagens multiespectrais, algumas extensões em relação ao algoritmo original são necessárias. Estas extensões são sumarizadas na seguinte sequência de passos:

- a) Inicialmente calcula-se o módulo do gradiente sobre cada uma das bandas das duas imagens utilizadas na análise.
- b) Aplica-se em seguida um filtro passa-baixa com o objetivo de reduzir o efeito de ruído na imagem.
- c) Calcula-se então para cada coordenada espacial da imagem o máximo entre os valores do módulo do gradiente computado sobre todas as bandas.
- d) Antes de se aplicar o algoritmo *watersheds* propriamente dito, faz-se uma limiarização sobre a matriz com os máximos dos gradientes, de modo que seus elementos com valores inferiores a dado limiar sejam igualados a zero, sem alterar os demais elementos. Esta medida tem como objetivo evitar o fenômeno comum ao algoritmo *watersheds*, conhecido como super-segmentação, em que surgem muitos segmentos muito pequenos.
- e) A matriz assim resultante é submetida ao *watersheds*.

A figura 7 apresenta o resultado obtido após a aplicação do algoritmo watershed, etapa e).

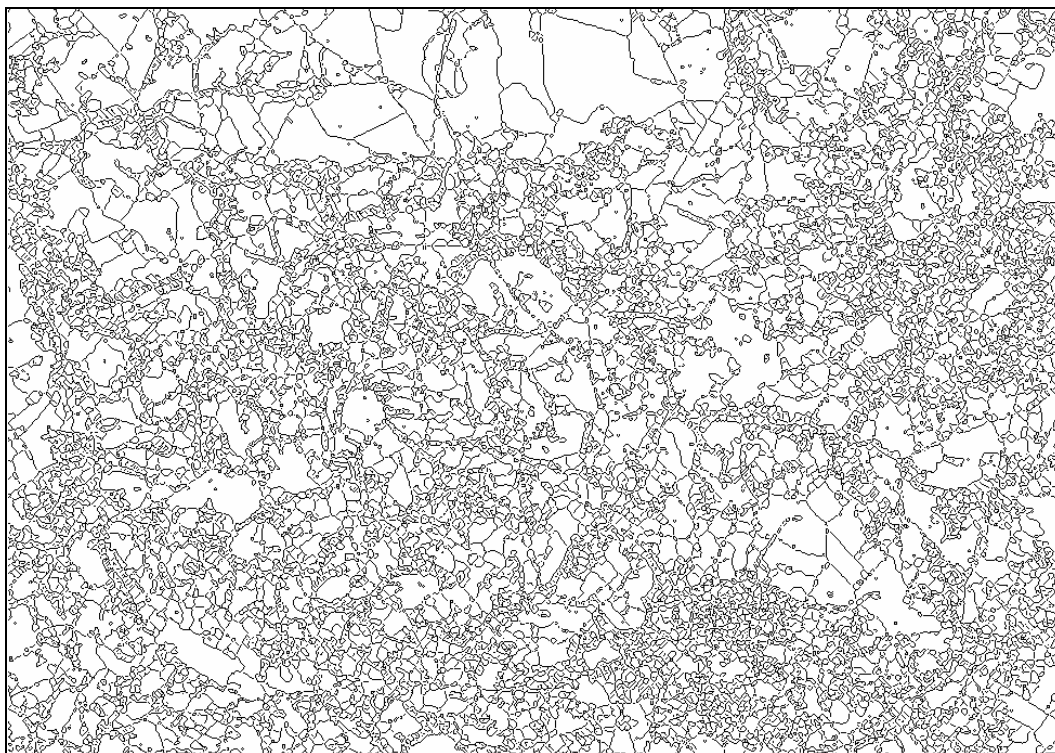


Figura 7 - Resultado obtido após a aplicação do algoritmo watershed, etapa e).

Conforme mencionado anteriormente, no resultado originalmente produzido pelo algoritmo *watersheds*, as linhas de separação não pertencem a qualquer região. Assim sendo, foi incluído no processo de segmentação ainda um passo final que atribui cada pixel da região de crista à mesma região do pixel adjacente com resposta espectral mais próxima considerando todas as bandas de todas as imagens simultaneamente. A figura 8 apresenta o resultado final do procedimento de segmentação.

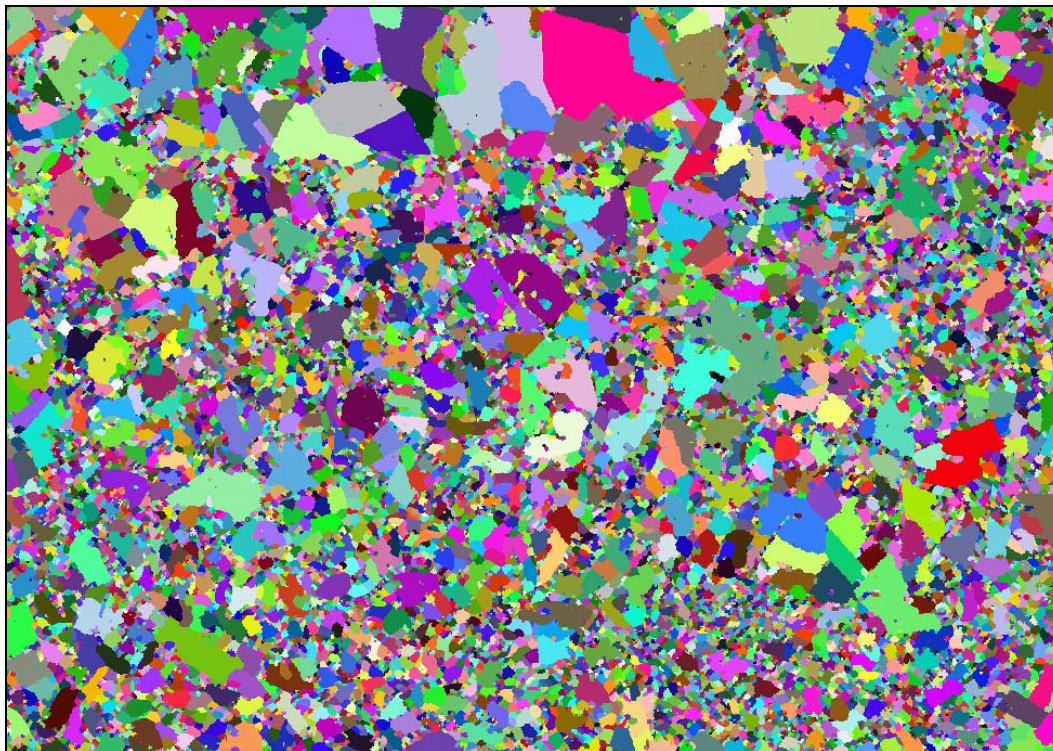


Figura 8 – Resultado do procedimento de segmentação

5.2.

Seleção automática dos conjuntos de treinamento

Com o objetivo de definir qual o método de detecção automática de mudanças na cobertura do solo melhor se adapta às características das imagens empregadas nos experimentos para a validação do método proposto, na presente seção, são comparados diversos algoritmos não supervisionados de detecção automática de mudanças em imagens bi-temporais.

Cazes e Feitosa (2004) comparam, com base nos mesmos resultados de referência empregados nesta tese, os desempenhos dos seguintes métodos de detecção de mudanças: subtração banda a banda (Coppin, 2001), análise de componentes principais (PCA) (Fung, 1987), regressão linear (Feitosa, 2001) e rede neural RBF (Duarte, 2003). Os resultados obtidos indicaram que o método baseado na regressão linear apresentou melhor desempenho. Por isso, este método será implementado nos experimentos que se seguem.

Basicamente, o método da regressão linear (Feitosa, 2001) considera que existe uma função linear que prevê a resposta espectral de um segmento em t em

função da sua resposta espectral em $t-1$. O modelo de regressão linear para p bandas é apresentado na Eq. (6).

$$y_{ij} = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip} + \varepsilon_{ij} \quad \text{Eq. (6)}$$

onde, y_{ij} corresponde ao valor (medido) da resposta espectral do segmento i na banda j em t ; os β 's são os coeficientes do modelo linear associados a cada uma das p bandas, x_{ij} , associadas ao segmento i em $t-1$; ε_{ij} corresponde ao erro entre os valores previsto ($\beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip}$) e medido.

Assumindo que os resíduos possuem distribuições normais com média zero e variância constante, é comum realizar a normalização da variância dos resíduos. Neste caso, os resíduos normalizados r_i^* são dados por

$$r_i^* = \frac{r_i}{\sqrt{\text{var}(r_i)}} \quad \text{Eq. (7)}$$

onde, $\text{var}(r_i)$ é a variância estimada do erro. Na eq. (8), os resíduos normalizados possuem distribuição t de Student com $(n-p-1)$ graus de liberdade, onde n corresponde ao número total de pontos de dados e p ao número de variáveis empregadas na previsão. O intervalo de confiança para a média de cada erro é dado por:

$$c_i = r_i \left[1 \pm t_{(n-p-1)}(1-\alpha/2) \sqrt{\text{var}(r_i)} \right] \quad \text{Eq. (8)}$$

onde $t_{(n-p-1)}(1-\alpha/2)$ é o $100(1-\alpha/2)$ ésimo percentil da distribuição t de Student com $(n-p-1)$ graus de liberdade. É Usual considerar que, quando o intervalo de confiança dado pela eq. (8) não inclui o valor (zero), há uma forte evidência que o padrão em questão é um *outlier*. Nesta implementação os segmentos indicados como *outliers* são considerados alterados, isto é, passaram de uma classe em $t-1$ para outra classe em t .

O montante de erro aceitável é inversamente proporcional ao parâmetro α que define o intervalo de confiança do erro. Quanto menor o valor de α , mais

candidatos são aceitos, porém mais falsos negativos (regiões cuja cobertura do solo mudou, mas o método considerou inalteradas) são aceitos. Por outro lado, quanto maior α menos falsos negativos estarão no conjunto de treinamento, porém mais falsos positivos (candidatos que não sofreram mudança na cobertura, mas o método considerou alterados) são eliminados do conjunto de treinamento. Assim sendo, o ajuste do parâmetro α tem grande influência na representatividade do conjunto de treinamento resultante. Devido a ausência na literatura de um critério para a escolha deste parâmetro, o valor ótimo de α será escolhido a partir de resultados empíricos.

5.3.

Procedimento automático de aprendizado das assinaturas

Os conjuntos nebulosos associados aos rótulos lingüísticos da variável resposta espectral podem ser definidos de forma automática, como descrito na seção anterior. Este processo considera como conjunto de treinamento as regiões que não mudaram entre $t-1$ e t , resultado fornecido pelo procedimento de detecção de mudanças. Vale mencionar novamente que uma alternativa não automática seria pedir a um foto-intérprete que selecionasse manualmente os conjuntos de treinamento.

O conjunto de treinamento é formado colocando-se a resposta espectral média em t (bandas 5, 4 e 3) de cada um dos segmentos nas linhas de uma matriz $\mathbf{X}_i$ de dimensões $s \times 3$, onde s corresponde ao número de segmentos das classes agrupadas pelo rótulo $\mathbf{A}_{1,i}$, formando o respectivo conjunto de treinamento. No presente texto, o conjunto $\mathbf{X}_i$ se refere ao rótulo $\mathbf{A}_{1,i}$.

Admite-se, neste trabalho, que para cada rótulo lingüístico a resposta espectral média dos segmentos pode ser adequadamente modelada por uma Normal $N(\mathbf{m}_i, \Sigma_i)$, onde $\mathbf{m}_i$ e Σ_i são o centróide e a matriz de covariâncias correspondentes ao rótulo $\mathbf{A}_{1,i}$.

De posse dos conjuntos de treinamento, estimam-se $\mathbf{m}_i$ e Σ_i respectivamente como a média do conjunto de treinamento, $\bar{\mathbf{x}}_i$, (vetor coluna) e a matriz de covariância amostral do conjunto de treinamento, $\mathbf{S}_i$, dados pela Eq. (9) e pela Eq. (10).

$$\mathbf{m}_i = \frac{1}{N_i} \sum_{\mathbf{x} \in A_{1,i}} \mathbf{x} \quad \text{Eq. (9)}$$

$$\mathbf{S}_i = \frac{1}{N_i - 1} \sum_{\mathbf{x} \in A_{1,i}} (\mathbf{x} - \bar{\mathbf{x}}_i)(\mathbf{x} - \bar{\mathbf{x}}_i)^T \quad \text{Eq. (10)}$$

onde N_i é o número de segmentos de treinamento elementos da classe $\mathbf{A}_{1,i}$.

5.4. Fuzzyficação das assinaturas espectrais

Para calcular da pertinência do valor da variável resposta espectral ao modelo do rótulo $\mathbf{A}_{1,i}$ explora-se uma propriedade das distribuições normais, segundo a qual a distância de mahalanobis ao centróide de uma população com distribuição $N(\mathbf{m}_i, \Sigma_i)$ é uma variável aleatória com distribuição Chi-quadrado com p graus de liberdade, onde p é a dimensão do espaço de atributos. Decorre daí que a probabilidade de que um padrão $\mathbf{x}$ pertencente à população esteja a uma distância

$$(\mathbf{x} - \mathbf{m}_i)^T \Sigma_i^{-1} (\mathbf{x} - \mathbf{m}_i) \leq \chi_p^2(\gamma) \quad \text{Eq. (11)}$$

é igual a $(1-\gamma)$ onde, $\chi_p^2(\gamma)$ é o percentil (100γ) da função Chi-quadrado com p graus de liberdade (Duda, .2001).

Assim, adotou-se como grau de pertinência do padrão $\mathbf{x}$, constituído neste caso pela resposta espectral média de um segmento nas três bandas, em relação à classe $\mathbf{A}_{1,i}$, o valor de γ que satisfaz a eq. (12) :

$$(\mathbf{x} - \bar{\mathbf{x}}_i)^T \mathbf{S}_i^{-1} (\mathbf{x} - \bar{\mathbf{x}}_i) = \chi_3^2(\gamma) \quad \text{Eq. (12)}$$

onde $\bar{\mathbf{x}}_i$ e $\bar{\mathbf{S}}_i$ são respectivamente a média e a matriz de covariância amostral da classe correspondente.

Cabe notar que, baseado nesta definição, um classificador que associa um padrão à classe de maior pertinência é equivalente ao classificador de distância de Mahalanobis, cuja definição corresponde ao lado esquerdo da igualdade na Eq. (12).

5.5.
Representação do conhecimento multitemporal

Um diagrama de transição de estados pode ser definido como um grafo, cujos nós representam os estados possíveis e os caminhos as possíveis transições (Word IQ, 2004). Este modelo de diagrama foi empregado anteriormente para a representação de conhecimento multitemporal (Bückner, 1999; Pakzad, 2001; Growe, 2000, 2001). A figura 9 apresenta um exemplo de diagrama de transição de estados.

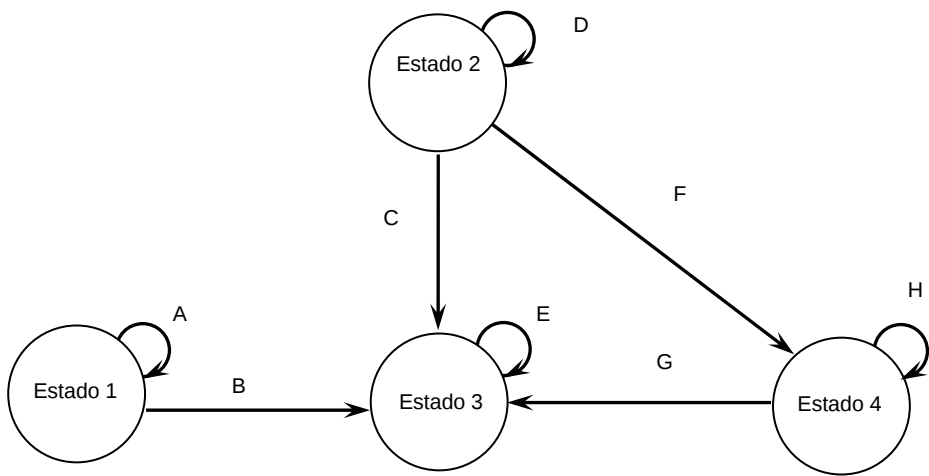


Figura 9 – Exemplo de diagrama de transição de estados.

No diagrama de transição de estados anterior, estão representados 4 estados: estado 1, estado 2, estados 3 e estado 4. Os estados anteriores possíveis em função do estado atual representados no diagrama anterior são apresentadas na tabela 1.

Estado atual	Possíveis estados anteriores
Estado 1	Estado 1
Estado 2	Estado 2
Estado 3	Estado 1; Estado 2; Estado 3; Estado 4
Estado 4	Estado 2; Estado 4

Tabela 1 – Estados anteriores que podem ocasionar cada um dos estados atuais.

Neste trabalho, cada transição está associada a um valor correspondente à possibilidade de sua ocorrência. No diagrama de transição de estados apresentado na figura 9, A, B, C, D, E, F, G e H representam os valores das possibilidades.

A presente abordagem propõe a utilização de conjuntos nebulosos para a representação das diversas modalidades de conhecimento. No caso particular do conhecimento multitemporal, são empregados conjuntos nebulosos discretos. Conjuntos nebulosos discretos podem ser representados (Mendel, 1995) por uma sequência de pares, formados por um valor de pertinência e pelo valor discreto da variável lingüística, respectivamente, a possibilidade e o estado do segmento no instante anterior. Os diversos pares são separados pelo sinal +, que nesta notação é um mero separador. Os conjuntos nebulosos derivados do diagrama de transição de estados anterior, figura 9, são apresentados na tabela 2.

Estado Atual	Conjunto nebuloso
Estado 1	A/Estado1 + 0/Estado2 + 0/Estado3 + 0/Estado4
Estado 2	0/Estado1 + D/Estado2 + 0/Estado3 + 0/Estado4
Estado 3	B/Estado1 + C/Estado2 + E/Estado3 + G/Estado4
Estado 4	0/Estado1 + F/Estado2 + 0/Estado3 + H/Estado4

Tabela 2 – Conjuntos nebulosos que representam o diagrama de transição de estados apresentado na figura 9.