

STATISTIQUE ET ANALYSE DES DONNÉES

NICOLE MARHIC

PASCALE MASSON

WOJCIECH PIECZYNSKI

**Mélange de lois et segmentation non supervisée
des données SPOT**

Statistique et analyse des données, tome 16, n° 2 (1991), p. 59-79

<http://www.numdam.org/item?id=SAD_1991__16_2_59_0>

© Association pour la statistique et ses utilisations, 1991, tous droits réservés.

L'accès aux archives de la revue « Statistique et analyse des données » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques

<http://www.numdam.org/>

MELANGE DE LOIS ET SEGMENTATION NON SUPERVISEE DES DONNEES SPOT

Nicole MARHIC* Pascale MASSON* Wojciech PIECZYNSKI**

* Groupe Traitement d'Image

MSC, ENST Bretagne, B.P. 832, 29285 BREST Cedex

**Groupe Image

DSR, INT- 9, rue C. Fourier, 91011 EVRY Cedex

Résumé

Cet article traite de segmentation Bayésienne non supervisée d'images. Le problème, dans le cas de la segmentation dite "contextuelle", réside dans l'estimation préalable des paramètres du mélange de distributions correspondant. La méthode d'estimation utilisée est une variante originale de l'algorithme EM : la reestimation de la loi a priori par la moyenne des lois a posteriori est remplacée par l'estimateur de Tilton. Les performances de l'algorithme ainsi obtenu sont comparées à celles de deux autres méthodes de segmentation non supervisée : une méthode contextuelle utilisant le SEM pour l'estimation des paramètres et une méthode globale fondée sur le MPM, laquelle requiert l'utilisation d'un modèle Markovien.

Mots clés : mélange de distributions, estimation, classification Bayésienne, segmentation contextuelle non supervisée, champs aléatoires de Markov.

Classification AMS : 62 A 15, 62 H 30.

Classification STMA : 04 070, 04 900.

Abstract

This paper deals with Bayesian unsupervised segmentation of images. The problem, in the case of the so-called "contextual" segmentation, lies in the estimation of the corresponding distribution mixture. The estimation method used is an original variation of the EM algorithm : reestimation of the prior distribution by the mean value of the posterior distributions is replaced by Tilton's estimator. The way this algorithm performs is compared with that of two other unsupervised segmentation methods : a contextual one using the SEM to carry out parameter estimation and a global one, based on the MPM algorithm, which requires a Markov fields model.

Keywords : distribution mixture, estimation, Bayesian classification, unsupervised contextual segmentation, Markov random fields.

1. INTRODUCTION

1.1 Segmentation Bayesienne

Ce travail porte sur une approche statistique de la segmentation non supervisée d'images. Le problème de la segmentation Bayesienne peut s'écrire de la façon suivante. S désignant un ensemble de pixels (ou sites), une image est modélisée par un couple $(\zeta=(\zeta_s)_{s \in S}, X=(X_s)_{s \in S})$ de champs aléatoires. Chaque ζ_s est à valeurs dans un ensemble fini de classes $\Omega=\{\omega_1, \dots, \omega_k\}$ et chaque X_s est une variable aléatoire réelle. L'image que l'on observe, généralement qualifiée de "réelle", est considérée comme une réalisation du champ X et le problème est celui de l'estimation de la réalisation invisible de ζ à partir des données $X=x$. Dans le cas de données satellite, ζ représente ce que l'on appelle généralement "la vérité terrain".

Pour A et B tels que $S \supseteq B \supseteq A$, désignons par ζ_A (resp. X_B) la restriction de ζ (resp. X) à A (resp. B). De manière la plus générale, la segmentation Bayesienne induite par la fonction de perte 0-1 consiste à estimer ζ_A par la configuration ζ_A^* pour laquelle la probabilité conditionnelle à $X_B=x_B$ (la *probabilité a posteriori*) est maximale. Ainsi, selon le choix que l'on fait pour A et B , on dispose d'un large éventail de méthodes de segmentation Bayesienne.

Parmi celles-ci on peut distinguer deux groupes selon le choix fait pour B : celui des méthodes dites locales ou contextuelles et celui des méthodes globales. Dans le cas des méthodes locales A est pris égal à $\{s\}$ et B est un ensemble, V_s , contenant s et un petit nombre de pixels qui lui sont voisins. Les méthodes globales sont celles pour lesquelles B coïncide avec l'ensemble S de tous les pixels de l'image. Dans ce cas, si A est pris égal à $\{s\}$, on obtient le MPM [Marroquin, Mitter, Poggio (1987)] et, s'il est égal à S , la méthode obtenue est celle du MAP décrite dans [Geman, Geman (1984)].

1.2 Estimation et robustesse des méthodes de segmentation

La mise en œuvre de ces méthodes implique l'utilisation de techniques différentes selon le groupe auquel elles appartiennent. Par ailleurs, une méthode, selon qu'elle appartienne à l'un ou l'autre groupe, requiert la connaissance de paramètres spécifiques parmi ceux définissant la loi de probabilité du couple (ζ, X) et en ignore certains. Ainsi, l'utilisation d'une méthode donnée implique une simplification implicite du modèle général.

De nombreuses études ont montré l'efficacité des méthodes Bayésiennes dans les cas où le modèle adopté correspond au modèle sous-jacent. De ce fait, il se pose le problème de la "robustesse" des méthodes Bayésiennes par rapport à l'évolution des paramètres qu'elles ignorent. Dans la plupart des cas, la distribution du couple (ζ, X) est inconnue et on doit alors, lors d'une phase préalable ou simultanément, estimer les paramètres requis par la méthode de segmentation retenue. Désignons par ASNS, Algorithme de Segmentation Non Supervisée, l'ensemble de la démarche Estimation-Segmentation. Nous évaluerons la qualité d'un ASNS en termes de taux de pixels bien classés; elle dépend, en dehors du problème de la robustesse par rapport aux paramètres ignorés mentionné ci-dessus, de deux facteurs, lesquels sont la qualité des estimateurs utilisés et la robustesse, au sens statistique, de la méthode de segmentation choisie, robustesse par rapport aux paramètres utiles. Ainsi, lorsqu'une méthode est robuste par rapport à un paramètre donné, c'est-à-dire quand son comportement varie peu quand la valeur du paramètre s'écarte de sa vraie valeur, on peut se contenter d'un estimateur rapide et pas nécessairement très performant; par contre, il est important de disposer d'un estimateur performant lorsque la variation du paramètre correspondant influe de manière significative sur le comportement de la méthode considérée. Finalement, l'efficacité d'un tel algorithme face à un problème concret dépend de trois facteurs : l'adéquation à la réalité du modèle sous-jacent à la méthode de segmentation, la qualité des estimateurs et, enfin, la robustesse de la méthode de segmentation par rapport aux paramètres utiles.

Dans ce travail, le centre d'intérêt n'est pas la qualité des estimateurs seule mais la qualité "conjointe" du couple formé de l'estimation des paramètres et de la segmentation sur la base des paramètres estimés. Ainsi l'étude d'un ASNS portera sur deux facteurs F1 et F2 où F1 désigne la robustesse par rapport à l'évolution de certains paramètres "ignorés" et F2 représente l'ensemble des facteurs "qualité des estimateurs" et "robustesse de la méthode par rapport aux paramètres utiles". Nous présentons un nouvel ASNS et l'objectif de notre travail est de comparer son efficacité avec celle de deux ASNS déjà existants, un local [Masson, Pieczynski (1990)] et un global [Chalmond (1989)]. Nous visons ainsi une double comparaison : une comparaison entre deux ASNS locaux où seul le facteur F2 intervient et une comparaison entre les ASNS locaux et un ASNS global où interviennent les deux facteurs F1 et F2. Pour ce faire, les algorithmes seront, dans un premier temps, appliqués à des images modélisant les images SPOT; ces images résultent de l'addition d'un bruit Gaussien à des images obtenues par simulation.

1.3 Organisation de l'article

Dans le paragraphe suivant, nous décrivons le modèle retenu et donnons une brève interprétation des notions introduites en considérant une image SPOT réelle. Le troisième paragraphe est consacré à la description d'un nouvel ASNS et nous y rappelons la démarche des deux algorithmes auxquels il sera comparé. Nous y précisons également la signification des facteurs F1 et F2 dans le cadre retenu.

Les résultats des segmentations, accompagnés de commentaires, sont présentés dans le paragraphe 4. Avant de conclure, nous donnons un exemple de segmentation d'une image SPOT réelle.

2. MODELE HIERARCHIQUE

2.1 Description du modèle

La loi $P_{\zeta, X}$ du couple (ζ, X) est donnée par la loi P_{ζ} de ζ et la famille des lois conditionnelles, $\{P_{X^e}, e \in \Omega^{\text{card} S}\}$, de X sachant $\zeta=e$. ζ sera supposé Markovien : la loi de ζ_s conditionnelle à $(\zeta_t)_{t \neq s}$ est égale à la loi de ζ_s conditionnelle à $(\zeta_v)_{v \in V}$ où V est un voisinage de s . Supposons, de plus, toutes les réalisations de ξ possibles, la loi de ζ est alors une distribution de Gibbs, celle-ci pouvant se mettre sous la forme (voir par exemple [Dubes, Jain (1989)]) :

$$P_z[\epsilon] = K e^{-U[\epsilon]} \quad (1)$$

où la forme de l'énergie U est liée à la forme de V , K étant la constante normalisatrice.

Dans le cas des images SPOT, Ω représente les classes de nature de terrain ("forêt", "désert", "zone urbainisée", ...) et la Markovianité de ζ ne paraît pas être une hypothèse forte.

Nous supposerons que X est la somme de deux champs Y et B . Y modélise la "variabilité naturelle" ou "texture" des classes (forêt plus ou moins dense, présence de roches dans le désert, ...) et B est le bruit de transmission.

$$X = Y + B \quad (2)$$

La loi de Y conditionnellement à ζ est définie de la manière suivante. On se donne k distributions $P_1, P_2, \dots, P_k$ de Y conditionnelles aux réalisations uniformes de ζ : $\zeta_s = \omega_1$ pour tout $s \in S, \dots, \zeta_s = \omega_k$ pour tout $s \in S$, respectivement.

Pour une réalisation ε quelconque de ζ , soit alors $C_1, C_2, \dots, C_k$ la partition de S définie par $(s \in C_i) \Leftrightarrow (\varepsilon_s = \omega_i)$ et $Y_1, Y_2, \dots, Y_k$ les restrictions de Y à $C_1, C_2, \dots, C_k$ respectivement. La distribution de chaque Y_i est définie en tant que distribution marginale de P_i . On définit alors la loi de Y (conditionnelle à $\zeta = \varepsilon$) en supposant l'indépendance des variables $Y_1, Y_2, \dots, Y_k$. Cela définit la loi de (ζ, Y) ; en supposant l'indépendance de (ζ, Y) et B on définit la loi de $(\zeta, X) = (\zeta, Y+B)$.

2.2 Simulation

Le modèle hiérarchique ainsi défini est relativement général et la simulation des réalisations de (ζ, X) est possible dès que l'on dispose d'une méthode de simulation pour $P_\zeta, P_1, P_2, \dots, P_k$ et P_B .

P_ζ peut être simulé par l'échantillonneur de Gibbs et nous supposons $P_1, P_2, \dots, P_k$ et P_B Gaussiennes; il existe alors des méthodes pour les simuler. En particulier, si les champs correspondants sont Markoviens, on peut encore utiliser l'échantillonneur de Gibbs.

La procédure de simulation est la suivante :

(i) on simule une réalisation ε de ζ .

(ii) on simule k réalisations $y_1, y_2, \dots, y_k$ de Y selon les lois $P_1, P_2, \dots, P_k$, respectivement; y donné par :

$$y_s = \sum_{i=1}^k \chi_{[\varepsilon_s = \omega_i]} y_s^i \quad (3)$$

où χ désigne la fonction indicatrice, est alors une réalisation de Y .

3. SEGMENTATION BAYESIENNE NON SUPERVISEE

3.1 Méthodes contextuelles

Lorsque l'on choisit d'utiliser une méthode contextuelle pour la segmentation, il n'est pas nécessaire de supposer ζ Markovien et on peut utiliser directement des fonctions discriminantes. Ces fonctions sont déduites de la distribution de (ζ_V, X_V) laquelle est supposée indépendante de V .

Pour tout ω_i dans Ω :

$$FD(\omega_i, x_v) = \sum_{\omega \in \Omega^{|V|-1}} P\{\zeta_s = \omega_i, \zeta_{v-\{s\}} = \omega, X_v = x_v\} \quad (4)$$

où $|V|$ désigne le cardinal de V et Ω est l'ensemble des classes où ζ_s prend ses valeurs.

Alors :

$$\zeta_s^* = \omega_i \Leftrightarrow FD(\omega_i, x_v) = \max_{\omega_j \in \Omega} FD(\omega_j, x_v) \quad (5)$$

Dans le cas général, la distribution de (ζ_V, X_V) , nécessaire au calcul des fonctions discriminantes, est inconnue et doit être estimée à partir de l'observation $X=x$.

Soit F_i la distribution de X_V sachant $\zeta_V = \epsilon_i$ et supposons que ζ_V est égal à ϵ_i avec la probabilité α_i . La loi de X_V peut alors s'écrire :

$$F = \sum_{i=1}^N \alpha_i F_i \quad (6)$$

où $N = (\text{card} \Omega)^{|V|}$.

Ainsi, le problème est celui de l'estimation des composants d'un mélange de distributions.

Dans la suite, X est supposé Gaussien conditionnellement à ζ . Dans ce cas, F_i dépend d'un paramètre β_i composé d'un vecteur moyenne μ_i et d'une matrice de covariance Γ_i et le problème réside dans l'estimation de $\theta = \{\alpha = (\alpha_1, \dots, \alpha_N); \beta = (\beta_1, \dots, \beta_N)\}$ à partir d'une réalisation x de X fournissant un échantillon de X_V .

3.2 L'algorithme EM Modifié

3.2.1 L'algorithme EM Classique

Considérons une suite $V_1, V_2, \dots, V_n$ de voisinages dans S de forme V et notons ζ_i et X_i les restrictions respectives de ζ et X à V_i .

Le problème de l'estimation de $\theta = (\alpha, \beta)$ à partir de $X_1, X_2, \dots, X_n$ étant celui de l'estimation d'un mélange de lois, on peut le résoudre en utilisant l'algorithme EM classique décrit dans [Dempster, Laird, Rubin (1977)] et [Redner, Walker (1984)].

Le principe de cet algorithme consiste à maximiser par rapport à θ la vraisemblance en $x_1, x_2, \dots, x_n$. Il définit, à partir d'une valeur initiale $\theta^{(0)}$, une suite $(\theta^{(q)})_q$ de paramètres qui, sous certaines conditions, converge vers θ .

Pour $i \in \{1, \dots, n\}$ et ϵ_j dans $\Omega^{|V|}$, désignons par $p_{ij}^{(q)}$ la probabilité a posteriori, calculée sur la base de $\theta^{(q)} = (\alpha^{(q)}, \beta^{(q)})$, d'avoir la configuration ϵ_j sur V_i :

$$p_{ij}^{(q)} = p_{\theta^{(q)}} \{ \zeta_i = \epsilon_j \mid X_i = x_i \} \quad (7)$$

Alors, les nouvelles valeurs $\alpha^{(q+1)}$ et $\beta^{(q+1)}$ données par l'algorithme EM satisfont :

$$\alpha_j^{(q+1)} = \frac{1}{n} \sum_{i=1}^n p_{ij}^{(q)} \quad (8)$$

et

$$\mu_j^{(q+1)} = \frac{\sum_{i=1}^n x_i p_{ij}^{(q)}}{\sum_{i=1}^n p_{ij}^{(q)}} \quad (9.1)$$

$$\Gamma_j^{(q+1)} = \frac{\sum_{i=1}^n (x_i - \mu_j^{(q+1)}) (x_i - \mu_j^{(q+1)})^t p_{ij}^{(q)}}{\sum_{i=1}^n p_{ij}^{(q)}} \quad (9.2)$$

On peut remarquer que les probabilités a priori sont reestimées par la moyenne des probabilités a posteriori.

3.2.2 L'estimateur de Tilton

L'algorithme que nous appelons EM Modifié (EMM) est obtenu en utilisant l'estimateur décrit par Tilton *et al* dans [Tilton, Vardeman, Swain (1982)], une étude théorique de celui-ci étant proposée dans [Pieczynski (1989)].

En effet, dans l'hypothèse où β est connu, Tilton propose un estimateur, que nous noterons α^* , pour le paramètre α . Cet estimateur est construit de la façon suivante :

Soient $f_1, f_2, \dots, f_N$ les densités des lois de probabilités de paramètres $\beta_1, \beta_2, \dots, \beta_N$ respectivement. Posons, pour tout $j \in \{1, \dots, N\}$:

$$Y_j = \frac{1}{n} \sum_{i=1}^n f_j(x_i) \quad (10)$$

et considérons la matrice $A(N \times N)$ de terme général :

$$a_{ij} = \int f_i f_j \quad (11)$$

Alors α^* est donné par :

$$\alpha^* = A^{-1} Y \quad (12)$$

où $Y = {}^t(Y_1, \dots, Y_N)$.

Les formules de reestimation de l'algorithme EM Modifié sont obtenues directement en remplaçant (8) par :

$$\alpha^{(q+1)} = (A^{(q)})^{-1} Y^{(q)} \quad (13)$$

où $Y^{(q)}$ et $A^{(q)}$ sont calculées comme dans (10) et (11), respectivement, en utilisant la valeur courante $\beta^{(q)}$ du paramètre β .

Cette modification peut être justifiée par les deux remarques suivantes :

(i) La valeur de $\alpha^{(q)}$ n'intervient pas dans le calcul de la nouvelle valeur $\alpha^{(q+1)}$; on évite ainsi les problèmes de robustesse de (8) par rapport à α .

(ii) L'étude numérique d'un cas simple de mélange de deux Gaussiennes univariées de variance 1 montre que l'erreur quadratique moyenne de l'estimation donnée

par (13) est moindre (lorsque $\beta^{(q)}$ tend vers β) que celle de l'estimation donnée par (8), et ce, même dans le cas qui lui est le plus favorable, à savoir $\alpha^{(q)} = \alpha$.

3.3 L'algorithme SEM

L'algorithme SEM, dont les avantages par rapport à l'algorithme EM sont exposés dans [Celeux, Diebolt (1986)], est une procédure permettant de définir une suite $(\theta^{(q)})_q$ en utilisant des tirages stochastiques. On se place dans les mêmes conditions que ci-dessus et on adopte les mêmes notations. Le calcul de $\theta^{(q+1)}$ à partir de $\theta^{(q)}$ et $x = (x_1, x_2, \dots, x_n)$ s'effectue en trois étapes :

(i) Calcul, pour chaque x_i , des probabilités a posteriori $p_{ij}^{(q)}$, $j \in \{1, \dots, N\}$, sur $\Omega^{|V|}$.

(ii) Affectation de chacun des x_i à une "classe", qui est ici un élément de $\Omega^{|V|}$, tirée dans $\Omega^{|V|}$ selon la loi de probabilité :

$$p_i^{(q)} = \{P_{ij}^{(q)} : j \in \{1, N\}\} \quad (14)$$

On obtient ainsi une partition $Q_1, Q_2, \dots, Q_N$ de l'échantillon $x_1, x_2, \dots, x_n$.

(iii) On pose alors :

$$\alpha_j^{(q+1)} = \frac{1}{n} \text{Card}(Q_j) \quad (15)$$

et chaque $\beta_j^{(q+1)} = (\mu_j^{(q+1)}, \Gamma_j^{(q+1)})$ est donné par l'application des estimateurs classiques, moyenne et covariance empiriques, sur le sous-échantillon Q_j . χ désignant la fonction indicatrice, on a :

$$\mu_j^{(q+1)} = \frac{\sum_{i=1}^n x_i \chi_{Q_j}(x_i)}{\sum_{i=1}^n \chi_{Q_j}(x_i)} \quad (16.1)$$

$$\Gamma_j^{(q+1)} = \frac{\sum_{i=1}^n (x_i - \mu_j^{(q+1)})(x_i - \mu_j^{(q+1)})^t \chi_{Q_j}(x_i)}{\sum_{i=1}^n \chi_{Q_j}(x_i)} \quad (16.2)$$

L'utilisation du SEM comme préalable à la segmentation contextuelle des images SPOT est décrite dans [Masson, Pieczynski (1990)]. Son principal intérêt réside dans le fait que la connaissance a priori du nombre de classes n'est pas nécessaire; en effet, s'agissant d'images réelles, ce nombre est parfois difficile à déterminer.

3.4 Une méthode globale : l'algorithme EM Gibbsien

Cet algorithme global, développé dans [Chalmond (1989)] combine l'algorithme MPM, pour la segmentation, avec l'algorithme EM, ce dernier permettant la reestimation des paramètres. La loi de ζ est une distribution de Gibbs relativement à un système de voisinages $\{V_s\}$ et dont la forme générale est donnée par (1). Un voisinage de forme V_s a $m=k \text{ Card}(V_s)$ configurations possibles dans le cas où Ω , l'ensemble des classes, contient k éléments. L'auteur suppose qu'à chacune de ces configurations correspond un "label"; alors le paramètre α caractérisant la distribution a priori est donné par les probabilités conditionnelles :

$$P_{ij} = P(\epsilon_s = \omega_i \mid V_s \text{ de type } j) \quad (17)$$

Par ailleurs le champ X modélisant l'image observée est supposé Gaussien conditionnellement à ζ et les $(X_s)_{s \in S}$ sont indépendantes conditionnellement à ζ :

$$P(X|\epsilon) = \prod_s P(X_s \mid \epsilon_s) \quad (18)$$

Le paramètre caractérisant le bruit est $\beta = (\{\mu_i\}, \sigma)$ si $X_s/\epsilon_s = \omega_i$ suit une loi Gaussienne de moyenne μ_i et d'écart-type σ . On doit donc estimer $\theta = (\alpha, \beta)$. Pour ce faire, on considère la pseudo-vraisemblance f_θ , définie à partir des densités conditionnelles et des distributions a priori par :

$$f_\theta(\zeta = \epsilon, X) = P_\theta(X \mid \zeta = \epsilon) P_\theta(\epsilon) \quad (19)$$

avec :

$$P_{\theta}(\epsilon) = \prod_s P_{\theta}(\epsilon_s | (\epsilon_t)_{t \in V_s}) \quad (20)$$

L'algorithme EM définit alors $\theta^{(q+1)}$ par :

$$\text{MAX}_{\theta} E[\text{Log } f_{\theta}(\zeta, X) | X = x, \theta^{(q)}] \quad (21)$$

On obtient alors les formules de reestimation suivantes :

$$p_{ij}^{(q+1)} = \frac{E[n_{ij} | x, \theta^{(q)}]}{\sum_i E[n_{ij} | x, \theta^{(q)}]} \quad (22)$$

où n_{ij} est le nombre d'occurrences dans ϵ telles que $\epsilon_s = i$ et V_s de type j .

$$\mu_i^{(q+1)} = \frac{\sum_{s \in S} \gamma_s^{(q)}(i) x_s}{\sum_{s \in S} \gamma_s^{(q)}(i)} \quad (23.1)$$

et :

$$\sigma_i^{(q+1)} = \frac{\sum_i \sum_{s \in S} \gamma_s^{(q)}(i) (x_s - \mu_i^{(q+1)})^2}{\sum_i \sum_{s \in S} \gamma_s^{(q)}(i)} \quad (23.2)$$

où :

$$\gamma_s^{(q)}(i) = P_{\theta^{(q)}}(\omega_s = i | x_s) \quad (24)$$

Sous l'hypothèse d'indépendance conditionnelle des variables X_s la distribution a posteriori est aussi une distribution de Gibbs et on peut en simuler des réalisations à l'aide de l'échantillonneur du même nom. En utilisant une méthode de Monte-Carlo, on peut alors calculer des approximations pour $E[n_{ij} | x, \theta^{(q)}]$ et pour (24), ces quantités ne pouvant être évaluées de façon directe.

4. OBJECTIFS ET RESULTATS EXPERIMENTAUX

Dans la suite, nous désignerons respectivement par A1 et A2 les ASNS locaux fondés sur EMM et SEM, respectivement, et par A3 l'ASNS global de B. Chalmond.

4.1 Objectifs

Notre étude expérimentale vise à mesurer l'influence, sur le comportement des trois algorithmes A1, A2, A3, de deux facteurs qui sont l'homogénéité de l'image, d'une part, et la corrélation du bruit, d'autre part, lesquels peuvent être considérés comme étant objectifs. Plus précisément, nous nous intéressons aux deux problèmes suivants :

4.1.1 *Importance du caractère stochastique dans les algorithmes A1 et A2*

La différence entre les algorithmes A1 et A2 se situe au niveau de la phase estimation, basée sur les algorithmes EMM et SEM, respectivement. Ces derniers peuvent être interprétés comme étant des cas particuliers d'une procédure générale, appelée ICE (Iterative Conditional Estimation) et décrite dans [Pieczynski (1990)]. Cette procédure engendre plusieurs méthodes d'estimation dans le cas de mélanges Gaussiens, le caractère stochastique de celles-ci pouvant être plus ou moins important [Marhic, Pieczynski (1991)]. Les algorithmes EMM et SEM apparaissent, respectivement, comme la plus "déterministe" et la plus "stochastique" parmi toutes ces méthodes. Ainsi, la seule comparaison de A1 et A2 peut apporter des renseignements sur les interactions éventuelles entre les facteurs "objectifs" cités ci-dessus et l'importance du caractère stochastique de la méthode retenue. Plus précisément, nous allons regarder sur des simulations l'influence du caractère stochastique et des caractéristiques des données sur le comportement des algorithmes A1 et A2, le facteur auquel nous nous intéressons étant le facteur F2.

4.1.2 *Importance de la corrélation spatiale du bruit*

Les méthodes globales ne pouvant tenir compte de la corrélation spatiale du bruit (en effet, lorsque le bruit est corrélé le champ a posteriori n'est plus Markovien, ce qui rend sa simulation impossible), nous nous intéressons à l'influence de celle-ci sur le comportement des algorithmes A1 et A2, d'une part, et A3, d'autre part. Concernant l'algorithme A3, la corrélation du bruit est un facteur de type F1, lequel représente la robustesse de l'algorithme par rapport à l'évolution des paramètres ignorés.

4.2 Simulations

Nous avons, dans un premier temps, généré, à l'aide de l'échantillonneur de Gibbs, des réalisations de champs Markoviens aux quatre plus proches voisins, le modèle retenu étant celui de Ising. L'énergie U sur le réseau S est de la forme :

$$U(\epsilon_s) = \gamma_1 U_1 + \gamma_2 U_2 \quad (25)$$

où :

$$\begin{cases} U_1 = \sum_{x \in S} \omega_x \\ U_2 = \sum_{s, t \in v_s} \omega_s \omega_t \end{cases}$$

V_S étant l'ensemble des couples (s, t) tels que s et t voisins.

Des valeurs de γ_1 et γ_2 dépend l'homogénéité des images (binaires) obtenues; pour $(\gamma_1, \gamma_2) = (+5.0, -2.5)$ (resp. $(+2.0, -1.0)$) on obtient celle désignée par IM1 (resp. IM2) et présentée figure 1 (resp. figure 4).

En associant respectivement 0 et 2 aux deux classes nous avons dégradé ces images par addition de bruit Gaussien centré, de variance 1. De l'addition d'un bruit blanc (BB) résultent les images présentées figures 2 et 5. Les images présentées figures 3 et 6 sont obtenues par addition d'un bruit spatialement corrélé (BC), généré par moyenne mobile aux vingt quatre plus proches voisins, ce qui implique :

$$\text{COV}(BC(s), BC(t)) = 0.8 \quad (26)$$

si s et t sont deux sites adjacents.



Figure 1: IM1



Figure 2: IM1+BB



Figure 3: IM1+BC



Figure 4: IM2

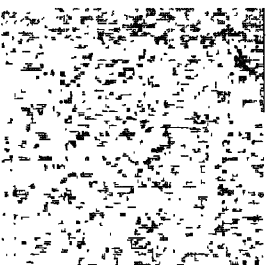


Figure 5: IM2+BB

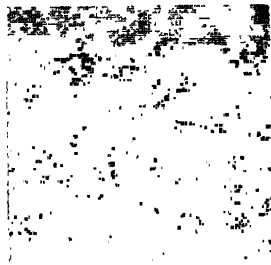


Figure 6: IM2+BC

La méthode d'initialisation retenue est celle de l'histogramme cumulé qui nous donne une première estimation des paramètres du bruit et permet d'obtenir une première segmentation par classification aveugle, laquelle consiste à ne tenir compte que de l'observation faite sur le pixel à classer, i.e. à considérer $V_s = \{s\}$.

méthode image	A4	A1	A2	ERT2	A3	ERT5
IM1+BB	16.1	10.7	11.0	10.4	03.2	03.8
IM1+BC	15.3	14.9	15.6	15.0	10.5	13.8
IM2+BB	16.0	16.1	17.1	14.5	13.6	12.7
IM2+BC	16.2	13.3	13.6	10.6	31.7	03.4

Tableau 1: erreur théorique et pourcentages de pixels mal classés

Les pourcentages de pixels mal classés sont donnés dans le tableau 1 ci-dessus, A4 désignant la méthode "histogramme cumulé-classification aveugle". Concernant les méthodes A1 et A2, les résultats présentés sont ceux obtenus en considérant $V_s = \{(s, t) \text{ tels que } t \text{ adjacent à } s\}$.



Figure 7:
A1(IM1+BB)



Figure 8:
A2(IM1+BB)



Figure 9:
A3(IM1+BB)



Figure 10:
A1(IM1+BC)



Figure 11:
A2(IM1+BC)



Figure 12:
A3(IM1+BC)



Figure 13:
A1(IM2+BB)

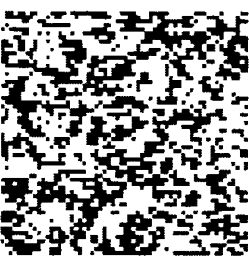


Figure 14:
A2(IM2+BB)

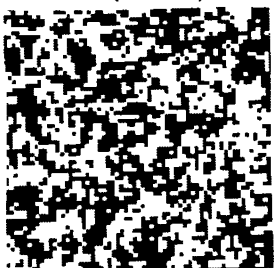


Figure 15:
A3(IM2+BB)

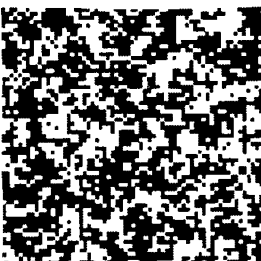


Figure 16:
A1(IM2+BC)



Figure 17:
A2(IM2+BC)



Figure 18:
A3(IM2+BC)

Les segmentations obtenues par les méthodes A1, A2 et A3 sont présentées dans les figures précédentes (7 à 15), $A_j(IM_i+B^*)$ désignant la segmentation de IM_i+B^* par la méthode A_j , et ce, pour $i=1, 2, j=1, 2, 3$ et $*$ =B ou C.

Dans le tableau 1 nous donnons également, pour chacun des cas, une estimation de l'erreur théorique de la méthode contextuelle; cette estimation est obtenue en simulant 10000 données indépendantes de (ζ_V, X_V) et en calculant le taux de données mal classées. ERT2 désigne l'erreur théorique relative à la méthode contextuelle pour laquelle V_s est de la forme ci-dessus et ERT5 celle relative à la méthode prenant en compte les quatre plus proches voisins.

4.3 Interprétation des résultats

(i) L'algorithme A4, classique et facile à mettre en œuvre, s'avère, dans certains cas, plus efficace que les algorithmes A1, A2 ou A3.

(ii) Les performances des algorithmes A1(EMM) et A2(SEM) sont comparables, A1 manifestant toutefois une légère supériorité. Ainsi, la méthode contextuelle semble être robuste par rapport à l'évolution du caractère stochastique de l'estimateur des composantes du mélange à considérer.

(iii) On peut remarquer que les pourcentages de pixels mal classés relatifs à A1 et A2 évoluent comme l'erreur théorique ERT2 en fonction de l'homogénéité de l'image et de la corrélation spatiale du bruit. Cela semble indiquer une stabilité des ASNS locaux par rapport à ces deux facteurs.

(iv) En revanche, l'algorithme A3(EM Gibbsien) est très sensible à l'évolution de ces deux facteurs. Cet algorithme est plus performant que les algorithmes A1 et A2 dans trois des cas présentés mais ses performances se dégradent de manière significative dans le dernier cas (IM2+BC). Ainsi, A3 n'est pas robuste par rapport à la corrélation spatiale du bruit, laquelle est dans ce cas un facteur de type F1. --

5. SEGMENTATION D'UNE IMAGE SPOT REELLE

Nous avons appliqué les trois algorithmes A1, A2 et A3, les paramètres étant initialisés, comme précédemment, à l'aide de l'algorithme A4, sur une image SPOT d'une région de Guinée-Bissau (figure19). La segmentation obtenue par A4 à partir de laquelle sont initialisés les paramètres à estimer est présentée figure 20 et les segmentations obtenues par A1 et A2 sont présentées figures 21 et 22, respectivement.



Figure 19: Image observée

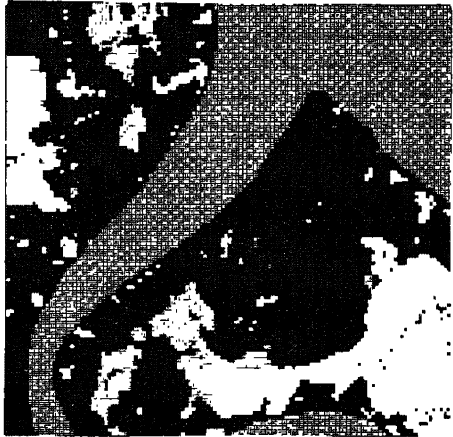


Figure 20: Initialisation par A4

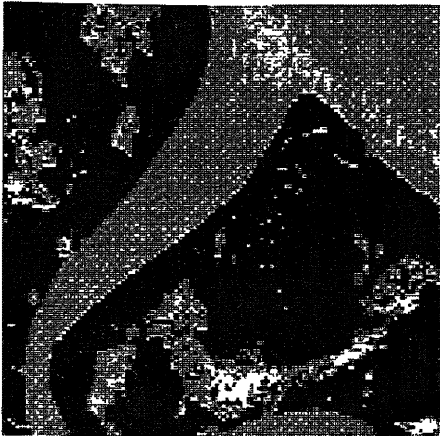


Figure 21: Segmentation par A1

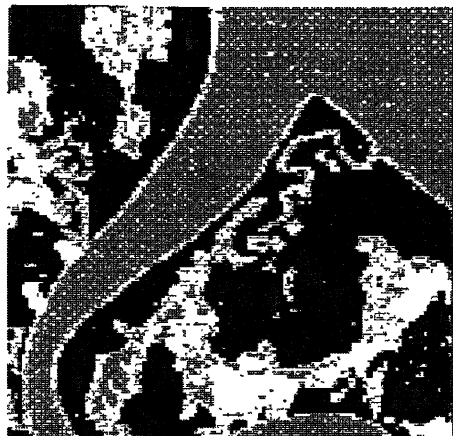


Figure 22: Segmentation par A2

classe	1	2	3	4
moyenne	57.9	62.9	68.3	68.7
variance	1.9	14.1	1.1	45.7
covariance	1.4	9.5	0.8	35.4

Tableau 2: estimations obtenues par SEM

Le nombre de classes estimé par le SEM est de 4; le nombre de voisins ayant été limité à un, les paramètres estimés par le SEM sont donnés dans le tableau 2; les covariances données dans ce tableau sont les covariances intra-classe. Les estimations obtenues par EMM, à savoir les vecteurs moyenne et les matrices de covariance, sont présentées dans le tableau 3.

j	vecteur moyenne	matrice des covariances
1	71.1	25.9 26.4
	71.4	26.4 28.4
6	68.8	1.7 1.4
	68.9	1.4 2.0
11	58.2	3.2 2.7
	58.3	2.7 3.3
16	68.1	0.9 0.9
	68.1	0.9 0.9

Tableau 3: estimations obtenues par EMM

Dans ce tableau j désigne le type de la configuration; il peut prendre 16 valeurs puisque ici $\text{card}(\Omega)=4$ et $|V|=2$, mais, dans ce tableau, nous ne donnons que les estimations relatives aux configurations telles que $\zeta_s = \zeta_t$.

Si l'on considère les tableaux 2 et 3, on constate que les variances estimées par l'une ou l'autre des méthodes EMM ou SEM sont différentes selon la classe considérée. C'est la raison pour laquelle nous n'avons pas appliqué l'algorithme EM Gibbsien à l'image SPOT; en effet, dans [Chalmond (1989)], l'auteur précise que, si la variance n'est pas constante sur toute l'image, alors les résultats obtenus par son algorithme ne sont pas très bons.

Par ailleurs, il est difficile d'établir des conclusions à partir des segmentations obtenues dans le cas de l'image SPOT; en effet, concernant celle-ci nous ne disposons pas de la vérité terrain, laquelle nous aurait permis de comparer de manière "objective" la qualité respective de ces segmentations.

6. CONCLUSIONS

L'étude que nous avons faite sur des simulations permet d'avancer trois conclusions générales :

1-Stabilité des méthodes locales : La méthode contextuelle est, d'une part, peu sensible au choix des estimateurs pour l'estimation des composantes du mélange et, d'autre part, elle est stable par rapport à l'évolution de l'homogénéité de l'image et de la corrélation du bruit; en effet, les résultats relatifs à A1 et A2 montrent que les erreurs commises par l'un ou l'autre de ces algorithmes sont relativement indépendantes de l'image considérée et évoluent comme l'erreur théorique ERT2.

2- Instabilité de la méthode globale étudiée : Les performances de cette méthode sont excellentes dans le cas IM1+BB, très bonnes dans le cas IM1+BC, acceptables dans le cas IM2+BB et très mauvaises dans le cas IM2+BC. Concernant cette méthode, l'erreur, évaluée en termes de pourcentage de pixels mal classés, ne se comporte pas de la même façon que ERT5 qui est un majorant de l'erreur théorique de la méthode MPM. L'algorithme A3 apparaît donc comme étant très sensible à l'évolution de l'homogénéité de l'image et de la corrélation du bruit. Ce dernier facteur étant de type F1 pour l'ensemble des ASNS globaux, il serait intéressant de savoir si cette instabilité est partagée par d'autres algorithmes faisant partie de cette famille, en particulier ceux faisant appel au MAP dont l'algorithme de recuit simulé.

3- Non universalité des algorithmes : L'efficacité des algorithmes est relative au type d'image considéré; de plus, dans certains cas, les performances de l'algorithme A4, très simple et rapide, sont comparables à celles des algorithmes A1, A2 et A3. Ainsi, il est probable qu'il existe toute une classe d'images réelles dont la segmentation ne nécessite pas l'utilisation d'ASNS statistiques relativement complexes.

Notre étude apporte ainsi quelques éléments de réponse au problème du choix, face à une image réelle donnée, de la meilleure méthode de segmentation. Dans le cas simple d'une image à deux classes les méthodes locales semblent assez stables, le cas de type IM2+BC apparaissant comme le plus propice à leur utilisation. La méthode globale A3 est très instable; elle est, par ailleurs, très performante dans les cas de type IM1+BB. Ainsi, s'il s'agit de traiter une image réelle à deux classes, on peut, dans un premier temps, estimer par EMM ou SEM la loi de $(\zeta_t, \zeta_s, X_t, X_s)$ pour s et t voisins. On peut ensuite décider, à partir des indications quant à l'homogénéité de l'image et à la corrélation du bruit contenues dans la loi estimée, lequel parmi les quatre cas présentés est

le plus proche de l'image considérée. On choisit alors, en considérant le tableau 1, la méthode la plus intéressante parmi celles qui y sont étudiées.

Notre étude appelle d'autres questions parmi lesquelles :

- (i) Les propriétés constatées se conservent-elles lorsque le nombre de classes augmente?
- (ii) L'instabilité de A3 est-elle commune à tous les ASNS globaux?
- (iii) Quel est le comportement des ASNS globaux et locaux lorsque la variance du bruit dépend de la classe?

Références

Besag, J. (1986) On the statistical analysis of dirty pictures. *Journ. of the Royal Stat. Soc., Series B*, 48, pp. 259-302.

Celeux, G. et Diebolt, J.(1986) L'algorithme SEM : un algorithme d'apprentissage probabiliste pour la reconnaissance de mélanges de densités. *Rev. de Stat. Appl.*, Vol. 34, No.2.

Chalmond, B. (1989) An iterative Gibbsian technique for reconstruction of m-ary images. *Pattern Recognition*, Vol. 22, No.6, pp. 747-761.

Dempster, M.M. Laird, N.M. and Rubin, D.B. (1977) Maximum likelihood from incomplete data via the EM algorithm. *Journ.of the Roy.Stat.Soc., Series B*, 39, pp. 1-38.

Devijver, P.A. and Dekesel, M.M. (1987) Learning the parameters of a hidden Markov random field image model : a simple example. *Pattern Recog. Theory and Appl.*, NATO ASI Series, Vol. F30, pp 141-163.

Dubes, R.C. and Jain, A.K. (1989) Random field models in image analysis. *Journ. of Appl. Stat.*, Vol. 6, No.2.

Geman, S. and Geman, D. (1984) Stochastic relaxation, Gibbs distributions and the Bayesian restoration of images. *IEEE Trans. on PAMI, PAMI-6*, pp. 721-741.

Guyon, X. et Yao, J.F. (1987) Analyse discriminante contextuelle. *Fifth Intern.Symp.of Data Anal.and Informatics*, (INRIA, B.P. 105,78153 Le Chesnay Cedex).

Marhic, N. and Pieczynski, W. (Juin 1991) Estimation of mixture and unsupervised segmentation of images. *Proceedings of IGARSS*, Helsinki, Finlande.

Marroquin, J. Mitter, S. and Poggio, T. (1987) Probabilistic solution of ill-posed problems in computational vision. *Journ. of the Am.Stat. Ass.*, 82, pp. 76-89.

Masson, P. and Pieczynski, W. (Sept. 1990) Segmentation of SPOT images by contextual SEM. *Proceedings of EUSIPCO*, Barcelone, Espagne.

Pieczynski, W. (1989) Estimation of context in random fields. *Jour. of Appl. Stat.*, Vol. 6, No. 2.

Pieczynski, W. (Nov. 1990) Mixture of distributions, Markov random fields and unsupervised segmentation of images. Rapport techn., No.122, LSTA, Univ. de Paris 6.

Redner, R.A. and Walker, H.F. (Apr. 1984) Mixture densities, maximum likelihood and the EM algorithm. *Soc. for Ind. and Appl.Math. Rev.*, Vol. 26, No. 2, pp. 195-239.

Tilton, J. Vardeman, S. and Swain, P. (Oct. 1982) Estimation of context for statistical classification of multispectral image data. *IEEE Trans. on GRS, GE-20*, No. 4, pp. 445-452.

Younes, L. (1988) Estimation and annealing for Gibbsian fields. *Annales de l'Institut Henri Poincaré*, 2.