

Segmentation of image hyperspectral

1

Serge Aïme / Gilles Celeux / Cécile Maudin

⇒ A multiscale approach for statistical characterization of functional images

Alexis Antoniadis, Pierre Bickel et Pierre W. Jones
Journal of Computational and Graphical Statistics (2008)

+ results • Kolarczyk, Xu et Gopal / Multiscale, multiparameter statistical image segmentation
Journal of the AMS (2005)

• Kolarczyk et Nanki / Multiscale likelihood analysis and complexity penalize estimators
Annals of Stat (2004)

• Barron, Huang, Li et Xiao / The MML principle, penalized likelihoods and statistical risk, Festschrift for J. Kiefer (2008)

• Aïme, Elad et Bruckstein

K-SVD: An algorithm for designing overcomplete dictionaries for sparse representation.

IEEE Transactions on Signal Processing (2006)

I Cadre et, objectif et solution clinique

- Projet de segmentation d'image hyperspectrale
- Collaboration INRIA Soelun / Université Paris Sud / SOLIEL
SERCCY IPANEMA

IPANEMA (Institut Méricaux Amiens)

- ⇒ ~~Mesure~~ de Mesure sur des objets vivants
oculiste, pédiatre, rhinome cutané
syndromes: zona de "lumière" per
- ⇒ Mesure de spectre très précis
d'absorption / fluorescence
- ⇒ Acquisition de beaucoup de spectres
⇒ traitement automatique → classification
sur expériences précédentes
- ⇒ Spécialité : ~~acquisition~~ organisation spatiale des spectres
⇒ 1 img
- ⇒ Imagerie hyperspectrale
- ⇒ Imagerie satellite, Imagerie médicale
évolution au cours du temps
(Bigot...)

Ex

Notations :

13

pixel: $p \in [1, N]^2 = \mathcal{P}$

$N \approx 100$

Spectre: $S(p) \in \mathbb{R}^m \leftarrow$ résolution du spectre
 $m \approx 1000$

Modélisation stochastique des spectres

$S \sim \mathcal{N}\left(\frac{\mu}{N}, \Gamma\right)$ μ moyen
 $= \mu_0$ Γ covariance de Γ^{-1} matrice de précision

Hypothèse: $\exists K$ classes différentes

de loi $\mathcal{N}(\mu_q, \Gamma_q) \neq \mu = \mu_0$
+ indépendance

$\Rightarrow$ Modèle de mélange

$$S \sim \sum_{q=1}^K \pi_q P_{\theta_q}$$

$$P(S) = \sum_{q=1}^K \pi_q P_{\theta_q}(S)$$

~~$P(S) = \sum_{q=1}^K \pi_q P_{\theta_q}(S)$~~

$$P_{\pi, \theta}(S) = \sum_{q=1}^K \pi_q P_{\theta_q}(S)$$

$\Rightarrow$ Modèle compliqué (avec variable cachée indiquant la classe)

$$P_{\pi, \theta}(S, c=q) = \pi_q P_{\theta_q}(S)$$

$$P_{\pi, \theta}(c=q) = \pi_q$$

Estimation de la "dose" de chaque pixel par M.V.

4

- Estimation de π et θ en $\hat{\pi}$ et $\hat{\theta}$
 $\hat{\pi}, \hat{\theta} = \underset{p \in P}{\operatorname{argmax}} \prod P_{\pi, \theta}(s(p)) = \underset{p \in P}{\operatorname{argmax}} P_{\pi, \theta}(D)$
 $\hat{c}(p) = \underset{c}{\operatorname{argmax}} \prod_{q \in Q} \hat{\pi}_q P_{\hat{\theta}_q}(s)$

- Estimation ^{directe} de ξ , π et θ en $\hat{\xi}$, $\hat{\pi}$ et $\hat{\theta}$

$$\hat{\xi}, \hat{\pi}, \hat{\theta} = \underset{p \in P}{\operatorname{argmax}} \prod P_{\pi, \theta}(s(p), c(p))$$

III

Algorithme E.M.

Algo itératif pour estimer le M.V.

Basé sur un principe de forçage de substitution
(surrogate function)

$$\begin{aligned} \log P_{\pi, \theta}(s) &= \int q(c) \log P_{\pi, \theta}(s, c) dc \\ &= \int q(c) \log \left[\frac{P_{\pi, \theta}(s, c)}{P_{\pi, \theta}(c | s)} \frac{q(c)}{q(c)} \right] dc \\ &= - \int q(c) \log \left(\frac{q(c)}{P_{\pi, \theta}(s, c)} \right) dc \\ &\quad + \int q(c) \log \left(\frac{q(c)}{P_{\pi, \theta}(c | s)} \right) dc \end{aligned}$$

$$\log P_{\pi, \theta}(s) = -K(q, P_{\pi, \theta}(s, c)) + K(q, P_{\pi, \theta}(c | s))$$

Valeur courante $\hat{\pi}^e, \hat{\theta}^e \Rightarrow$ optimisation de la vraisemblance

5

$$q = P_{\hat{\pi}^e, \hat{\theta}^e}(s)$$

$$\Rightarrow \log P_{\pi, \theta}(s) = -KL(P_{\hat{\pi}^e, \hat{\theta}^e}(c|s), P_{\pi, \theta}(s, c)) \\ + \underbrace{KL(P_{\hat{\pi}^e, \hat{\theta}^e}(c|s), P_{\pi, \theta}(c|s))}_{\geq 0}$$

$$Q^e(\pi, \theta) = -KL(P_{\hat{\pi}^e, \hat{\theta}^e}(c|s), P_{\pi, \theta}(s, c))$$

Prop $Q^e(\pi, \theta) \leq \log P_{\pi, \theta}(s)$

$$Q^e(\hat{\pi}^e, \hat{\theta}^e) = \log \hat{\pi}^e, \hat{\theta}^e(s)$$

$\Rightarrow$ Améliorer Q^e par rapport à $Q^e(\hat{\pi}^e, \hat{\theta}^e)$
implique une optimisation sur $\log P_{\pi, \theta}(s)$

Prop $Q^e(\pi, \theta) = -KL(P_{\hat{\pi}^e, \hat{\theta}^e}(c|s), P_{\pi, \theta}(s, c))$
 $= \int P_{\hat{\pi}^e, \hat{\theta}^e}(c|s) \log P_{\pi, \theta}(s, c) dc$
 $- \int \hat{P}_{\pi, \theta}(c|s) \log P_{\hat{\pi}^e, \hat{\theta}^e}(c|s) dc$

$\Rightarrow$ Il suffit d'augmenter $\int P_{\hat{\pi}^e, \hat{\theta}^e}(c|s) \log P_{\pi, \theta}(s, c) dc$

Def $P_{\hat{\pi}^e, \hat{\theta}^e}(c|s) = \frac{P_{\hat{\pi}^e, \hat{\theta}^e}(c, s)}{\int P_{\hat{\pi}^e, \hat{\theta}^e}(c, s) dc}$

$\Rightarrow$ Calculer $\int P_{\hat{\pi}^e, \hat{\theta}^e}(c|s) \log P_{\pi, \theta}(s, c) dc$

$\Rightarrow$ Etape E

$\Rightarrow$ Maximisation (ou augmentation) on π, θ

$\Rightarrow$ Etape M ou GM

La : • Algo CEM immer au stage de "qualification" → acquisition de l'éditorialité

16

- Variante stochastique pour l'échapper des minima locaux.
- Logiciel MIXMOD

Pratique en résumé,

- Nombre de classes K , structure de la matrice de covariance → pb beaucoup étudié dans le cadre de SELECT
 - ⇒ solution par pénalisation
 - ⇒ Code de la sélection de modèles!
 - ⇒ & autres exos...
- Sélection de variables = réduction de dimension
= sélection de modèles "cliniques" mais pb de complexité
⇒ Besoin de réduire la dimension des spectres!
 - Projection sur un sous-espace
 - ⇒ ~~ACP~~ ACP ⇒ norme "parcimonieuse"
 - ⇒ image / noyau linéaire ⇒ norme asymptotique
 - ⇒ Descripteurs globalement différents
- Prise en compte de la structure spatiale
 - ⇒ Pour l'instant absence totale!
 - Régularisation de la carte de décision
 - sur les étiquettes ⇒ Champ Markovien (SOPELEC)
 - à travers les poids de mélange ⇒ Koleszyk (structure multicellulaire) Ambroise

II Vraisemblance multi-échelle, segmentation hiérarchique et estimateurs pénalisés.

7

Idée de Kolaczyk ~~et al.~~ Ju et Gopal.

Pour de
$$P(s(p)) = \sum \pi_e P_{\theta_e}(s(p))$$
 (iid)

où
$$P(s(p)) = \sum \pi_e(p) P_{\theta_e}(s(p))$$
 (indépendance)

⇒ Prise en compte de la régularité spatiale à travers les poids.

⇒ Contraintes nécessaires pour rendre le modèle identifiable

→ θ_e fixe !

→ π_e constant par morceaux

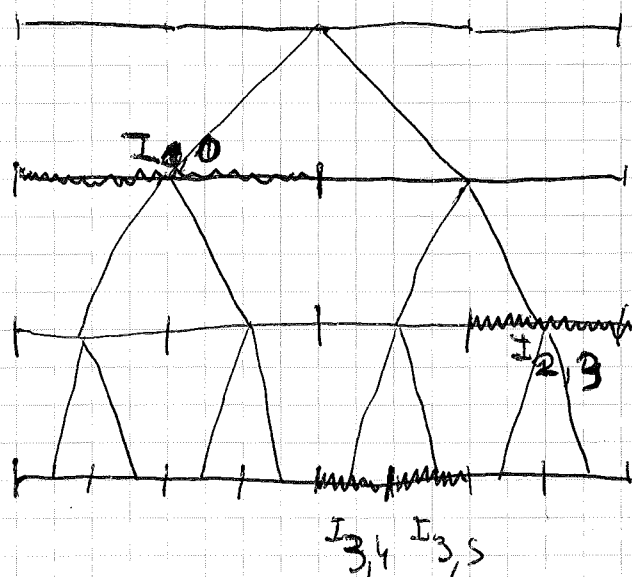
→ Partition structurée. (hiérarchique)

⇒ Cas particuliers des modèles de vraisemblances multi-échelles de Kolaczyk et Nowak

→ Algorithme et théorie !

⇒ Restriction au cas des partitions dyadiques récursives plus facile à présenter

1D



Structure
d'arbre

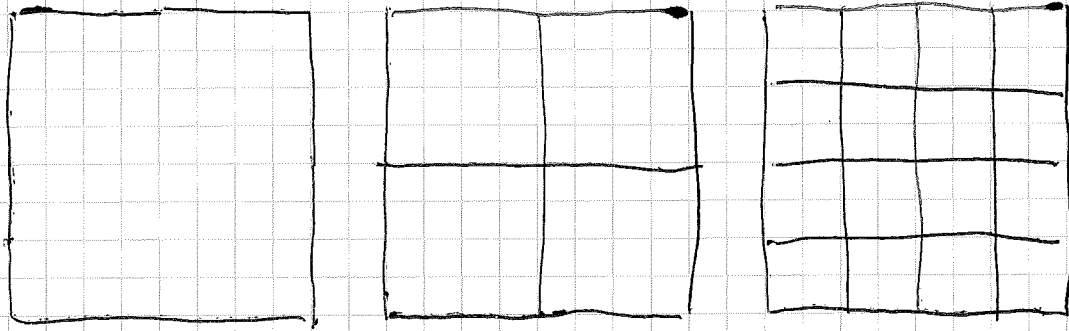
Partition $\mathcal{P} = \{I\}$

arbre partition

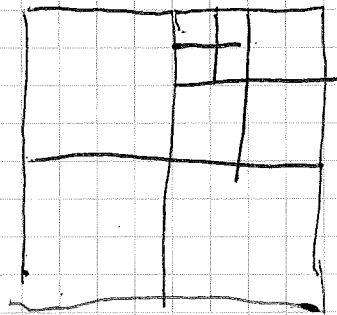
cf. ordonnées

2D Même chose avec des carrés

8



⇒ Exemple de récursion.



Modèle : $(P, (\pi_I))$ (θ_2 est fixe!)

Appelle MV

⇒ recherche de (π_I) optimal à P fixé facile

⇒ Algo "EM" sur chaque carré (calcul direct on peut)

⇒ recherche de P optimal → pénalisation nécessaire
sur c'est toujours la partition la plus simple qui gagne!

$$\text{si } \text{pen}(P) = \sum_{I \in P} \text{pen}(I), \quad \log P_{P, \pi, \theta}(D) - \text{pen}(P)$$

$$\text{ce revient à } \sum_{I \in P} \log P_{\pi_I, \theta}(S_I) - \text{pen}(I)$$

⇒ Séparation du problème sur chaque feuille / optimisation du critère

⇒ Algorithme efficace d'optimisation de la partition
Programmation dynamique / Bottom-up / CART (Breiman)

⇒ Clé: meilleure partition sur un carré est

- soit la partition triviale
- soit l'unin des meilleurs partitions sur les 4 carrés qui le compose.

⇒ Algo qui commence par la partition "triviale" et qui remonte jusqu'à tester la partition "triviale" correspondant au modèle de mélange gaussien classique.

⇒ Choix de la pénalisation ?

Outil théorique: résultats de Barron et al. ...

Cadre général : . $D \sim P_{\theta_0}$

. collection Π d'observables de $\mathcal{X}$

. $\hat{\gamma} = \argmax [\log P_{\gamma}(D) - \text{pen}(\gamma)]$

TH Si $\sum_{\gamma \in \Pi} e^{-\text{pen}(\gamma)/2} \leq 1$ alors

$$E(d(P_{\theta_0}, P_{\hat{\gamma}})) \leq \inf_{\gamma} [KL(P_{\theta_0}, P_{\gamma}) + \text{pen}(\gamma)]$$

Explication

$$KL(P_{\theta_0}, P_{\gamma}) = E_{P_{\theta_0}} \left[-\log \frac{dP_{\gamma}}{dP_{\theta_0}} \right] \quad \text{Kullback}$$

$$A(P_{\theta_0}, P_{\gamma}) = E_{P_{\theta_0}} \left[\left(\frac{dP_{\gamma}}{dP_{\theta_0}} \right)^{1/2} \right] = \int \left(\frac{dP_{\gamma}}{d\lambda} \right)^{1/2} \left(\frac{dP_{\theta_0}}{d\lambda} \right)^{1/2} d\lambda$$

$$h^2(P_{\theta_0}, P_{\gamma}) = \int \left(\left(\frac{dP_{\theta_0}}{d\lambda} \right)^{1/2} - \left(\frac{dP_{\gamma}}{d\lambda} \right)^{1/2} \right)^2 d\lambda = 2 - 2A(P_{\theta_0}, P_{\gamma})$$

$$d(P_{\theta_0}, P_{\gamma}) = -2 \log A(P_{\theta_0}, P_{\gamma}) \quad (A = 1 - \frac{1}{2} h^2)$$

On a

$$h^2(P_{\theta_0}, P_{\gamma}) \leq d(P_{\theta_0}, P_{\gamma}) \leq KL(P_{\theta_0}, P_{\gamma})$$

2 égalité ouale

Pr: si $\| \log \frac{dP_{\tilde{\gamma}}(D)}{dP_{\gamma_0}(D)} \|_{\infty} \leq B$ alors

$$d(P_{\gamma_0}, P_{\tilde{\gamma}}) \leq KL(P_{\gamma_0}, P_{\tilde{\gamma}}) \leq (2+B) d(P_{\gamma_0}, P_{\tilde{\gamma}})$$

Preuve:

$$d(P_{\gamma_0}, P_{\tilde{\gamma}}) = -2 \log A(P_{\gamma_0}, P_{\tilde{\gamma}})$$

$$= -2 \log A(P_{\gamma_0}, P_{\tilde{\gamma}}) + \log \frac{dP_{\gamma_0}(D)}{dP_{\tilde{\gamma}}(D)} + \text{pen}(\tilde{\gamma})$$

$$+ 2 \log \left(\frac{dP_{\tilde{\gamma}}}{dP_{\gamma_0}} \right)^{1/2} + 2 \log \left(e^{-\frac{\text{pen}(\tilde{\gamma})}{2}} \right)$$

$$= 2 \log \left[\frac{\left(\frac{dP_{\tilde{\gamma}}}{dP_{\gamma_0}} \right)^{1/2} e^{-\frac{\text{pen}(\tilde{\gamma})}{2}}}{E_{P_{\gamma_0}} \left(\left(\frac{dP_{\tilde{\gamma}}}{dP_{\gamma_0}} \right)^{1/2} \right)} \right] + \log \frac{dP_{\gamma_0}(D)}{dP_{\tilde{\gamma}}(D)} + \text{pen}(\tilde{\gamma})$$

$$\leq 2 \log \left[\frac{\sum_{\gamma \in \Pi} \left(\frac{dP_{\gamma}}{dP_{\gamma_0}} \right)^{1/2} e^{-\text{pen}(\gamma)}}{E_{P_{\gamma_0}} \left(\left(\frac{dP_{\gamma}}{dP_{\gamma_0}} \right)^{1/2} \right)} \right] + \min_{\gamma \in \Pi} \left[-\log \frac{dP_{\gamma_0}}{dP_{\gamma}} + \text{pen}(\gamma) \right]$$

$$E(d(P_{\gamma_0}, P_{\tilde{\gamma}})) \leq 2 E \log \left[\sum_{\gamma} \right] + E \min_{\gamma \in \Pi}$$

$$\leq 2 \log E \left(\sum_{\gamma} [] \right) + \min_{\gamma \in \Pi} E ()$$

$$\leq \underbrace{2 \log \left(\sum_{\gamma} e^{-\text{pen}(\gamma)} \right)}_{\leq 0} + \min_{\gamma \in \Pi} (KL(P_{\gamma_0}, P_{\gamma}) + \text{pen}(\gamma))$$

Ici : il faut discrétiser le problème

11

$\Rightarrow$ partition discrète par nature

$\Rightarrow \pi_I$ restreint à une grille de précision $N^{-\beta}$

Rq $\sum_x e^{-\frac{\text{pen}(x)}{2}} \leq 1 \Rightarrow$ Kraft en théorie de l'information

$\Rightarrow \frac{\text{pen}(x)}{2} \approx$ coût de codage de la loi

Ici : $\mathcal{P}$ partition à $|\mathcal{P}|$ éléments.

$\Rightarrow$ codage à travers l'arbre

$\Rightarrow |\mathcal{P}|$ feuilles $\Rightarrow \frac{|\mathcal{P}|-1}{3}$ nœuds intérieurs
(chaque nœud intérieur a 4 fils...)

$\Rightarrow$ Codage rec $|\mathcal{P}| - 1$ et $\frac{|\mathcal{P}|-1}{3}$ 0

$\Rightarrow$ Coût $\left[\frac{4}{3} |\mathcal{P}| - \frac{1}{3} \right] \log 2$

• Dans chaque carré, spécification de π_I : $k-1$ échantillons dans $[0,1]$ rec une précision $N^{-\beta}$

$\Rightarrow (k-1) \beta \log N$ mots ($\leftarrow$ bits en base e)

$$\Rightarrow \text{pen}(\mathcal{P}, \pi_I) = 2 \times \left[\frac{4}{3} |\mathcal{P}| - \frac{1}{3} \right] \log 2 + |\mathcal{P}| (k-1) \beta \log N$$

$$= \left[\frac{8}{3} \log 2 + 2(k-1) \beta \log N \right] |\mathcal{P}| - \frac{2}{3} \log 2$$

$$= \sum_{I \in \mathcal{P}} \underbrace{\left(\frac{8}{3} \log 2 + 2(k-1) \beta \log N \right)}_{\text{pen}(I)} - \frac{2}{3} \log 2$$

Rq : raffinement possible si on fait varier le nombre de classe

Contrôle de $KL(P_{\gamma_0}, P_{\gamma}) + \rho_n(\gamma)$

12

Ici $P_{\gamma_0} : P(s(p)) = \sum \pi_e^0(p) P_{\theta_e^0}(s(p))$

avec $P_{\gamma} : P(s(p)) = \sum_e \pi_{I(p)} P_{\theta_e^0}(s(p))$

$$\Rightarrow KL(P_{\gamma_0}, P_{\gamma}) \leq 2 \sum_p \sum_e |\pi_e^0(p) - \pi_{I(p)}|$$

$$\min_{\pi} KL(P_{\gamma_0}, P_{\gamma}) \leq 2 \|\pi^0(p) - \pi_{\mathcal{P}}^0(p)\|_1 + 2KN^2 \cdot N^{-\beta}$$

Théorie de l'approximation par π^0

$$\min_{|P| \leq M} \frac{1}{N^2} \|\pi^0(p) - \pi_P^0(p)\|_1 \leq CKM^{-\alpha}$$

$$\frac{1}{N^2} \left(\min_P \min_{\pi} KL(P_{\gamma_0}, P_{\gamma}) + \rho_n(\gamma) \right) \leq 2CKM^{-\alpha} + 2KN^{-\beta} + (C_1 + C_2 \beta \log N) \frac{M}{N^2}$$

$\Rightarrow$ optimisation en M et $\beta = 1$

$$\leq C \left(\frac{\log N}{N^2} \right)^{\frac{\alpha}{\alpha+1}}$$

$\Rightarrow$ Vitesse "béliakuelle" en dimension 2 !

$\Rightarrow$ Extension possible pour éviter la discrétisation ?

$\leadsto$ passage par une ligne généralisée

$\leadsto$ argument d'entropie (cf preuve des sélections de modèles)

$\Rightarrow$ Bouclage possible pour réestimer les θ_e !

III Réduction de dimension initiale

- Algo automatique limité par la dimension du problème
- Penalisation L^1 sur la matrice de précision $\Rightarrow$ Christophe...
- $\Rightarrow$ Méthode ad hoc de réduction de dimension

- Antoniadis, Bérut, van de Sijde : même méthode basée sur les ondelettes
 - $\Rightarrow$ seuillage des coeffs d'ondelettes
 - $\Rightarrow$ Conservation de l'ensemble des coefficients significatifs sur les courbes
 - $\Rightarrow$ Élimination des coefficients qui semblent constants...
- Justification théorique!

- Méthode basée sur l'apprentissage de la représentation en k -SVD de Elad et al.
 - $\Rightarrow$ Recherche de $S_1, \dots, S_d$ tel que toutes les courbes observées se représentent correctement de manière parcimonieuse

$$\text{over } S_1, \dots, S_d \quad \sum_P \min_{\|\alpha_i\|_0 \leq c} \|S[P] - \sum \alpha_i S_i\|^2$$

Cas $c=1 \Rightarrow k$ -Means = clustering
 $\approx$ Recherche d'espace affine adapté à chaque cluster

clustering sur les p.s. $c: |p| < S[p], S_d >$

$\Rightarrow$ résultats préliminaires encourageants!

- Comment combiner cette recherche de représentation avec l'objectif de clustering??

Seize fait une domif supervisée mais échoue à en superviser une linéaire !!!