

TRAVAIL PRATIQUE N° 5

SEGMENTATION D'IMAGE TEXTURÉE ET STRATÉGIE DE RECONNAISSANCE DE FORMES

Max Mignotte

DIRO, Département d'Informatique et de Recherche Opérationnelle.

[http : //www.iro.umontreal.ca/~mignotte/ift6150](http://www.iro.umontreal.ca/~mignotte/ift6150)

e-mail : mignotte@iro.umontreal.ca

A. INTRODUCTION

L'image **texture.pgm** (cf. Figure 1a) est une image texturée dans laquelle se cachent plusieurs régions homogènes (homogène au sens de la texture). Le cerveau humain peut aisément distinguer quatre zones de région homogène et géométriques dans cette image, i.e., une ellipse dans la moitié droite de l'image ainsi qu'un carré et un triangle de texture différente sur un fond texturé dans l'autre moitié de l'image. Dans la suite, on supposera connu ce nombre de régions homogènes¹.

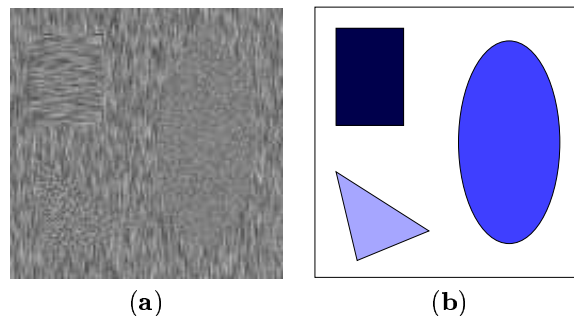


FIG. 1 – (a) Image originale **texture.pgm**. (b) Image segmentée. Le résultat présenté est une segmentation parfaite (difficilement atteignable en pratique ...) et présenté à titre indicatif seulement.

B. STRATÉGIE DE SEGMENTATION D'IMAGE TEXTURÉE

L'objet de ce TP est de retrouver (i.e., segmenter) par une procédure automatique ces quatre régions à partir de leurs caractéristiques (texturales) propres (cf. Figure 1). La stratégie que nous allons utiliser est une stratégie classique utilisée en reconnaissance des formes². Elle suppose que l'on connaît le nombre exacte de classes existante. Dans notre application, elle suppose que l'on connaît le nombre de régions homogènes présent dans l'image. Elle consiste en deux étapes.

¹En pratique, on doit préalablement estimer ce nombre de région avant de réaliser toute segmentation d'image texturée.

²On peut utiliser cette stratégie pour classer ou reconnaître des formes selon leur caractéristique géométrique, spectrales,

- Chaque élément que l'on doit classer (dans notre application, chaque pixel) en $K = 4$ classes est caractérisé par un ensemble de n mesures de textures. Pour notre application, on calcule ces mesures en utilisant les niveaux de gris de l'image texturée contenu dans une imagerie carrée centrée sur le pixel à classer. Chaque imagerie ou pixel à classer est ainsi caractérisée par un vecteur à n composantes (appelé "échantillon" ou $\mathbf{x}$ par la suite). Si les mesures texturales utilisées sont discriminantes, l'ensemble des "échantillons" associé à cette image représenté dans un espace à n dimension (une dimension pour chaque mesure), vont se regrouper en $K = 4$ "clusters" ou groupement³.
- On utilise ensuite un algorithme de classification supervisé ou non qui cherchera les K groupement les plus probables (au sens d'un certain critère).

Vous pouvez utiliser pour ce TP, les mesures de textures que vous jugerez discriminantes pour résoudre ce problème. Une façon efficace de savoir si les mesures que vous déciderez d'utiliser sont discriminantes consistent à les tester individuellement. Associer à chaque site de coordonnées (x, y) la valeur obtenue par la mesure texturale. Re-calibrer ensuite les niveaux de gris de l'image obtenue entre 0 et 255 et visualiser la. Si l'image obtenue permet de distinguer une ou plusieurs régions texturées, alors la mesure est discriminante. Vous pouvez utiliser l'algorithme de classification ou la stratégie de classification qui vous plaira ou encore utiliser l'algorithme des K-Moyennes décrit dans la suite de cet énoncé.

C. CONSEILS

(a) Utiliser des fenêtres (ou images) qui se chevauchent de taille par exemple 11×11 . Vous pouvez utiliser, entre autre, les mesures texturales caractéristiques suivantes : (1) la différence entre le niveau de gris maximal moins le niveau de gris minimal contenu dans la fenêtre. Vous pouvez ensuite seuiller l'image originale en deux classes (i.e., binariser l'image) avec un seuil compris à mi-chemin entre le niveau de gris minimal et maximal de cette image et calculer ensuite ; (2) la moyenne des longueurs de plage orientée horizontalement, (3) la moyenne des longueurs de plage orientée verticalement, etc. Je vous conseille ensuite de Re-calibrer chaque mesure (par exemple entre 0 et 255) pour leur donner une importance égale dans la procédure de classification.

(b) Vous pouvez utiliser l'algorithme des K moyennes pour rassembler ces échantillons en $K = 4$ groupements en fonction d'un critère de "ressemblance". Dans ce cas, la répartition des différents échantillons $\mathbf{x}_l$ dans les différents groupements est effectuée de façon à minimiser un indice de dispersion. Cet indice de dispersion s'écrit,

$$J = \sum_{i=1}^K \sum_{\mathbf{x}_l \in C_i} \|\mathbf{x}_l - c_i\|^2,$$

où la seconde sommation se fait sur tous les échantillons $\mathbf{x}_l$ appartenant au groupement C_i dont le centre est c_i . On peut montrer facilement que, pour une partition $\{C_i, i = 1, \dots, K\}$ donnée, l'indice de dispersion J est minimal si le centre de chaque groupement C_i est donné par $c_i = 1/N_i \sum_{\mathbf{x}_l \in C_i} \mathbf{x}$. avec N_i le nombre d'échantillons appartenant à C_i . L'algorithme des K moyennes utilise cette propriété et peut se résumer dans le Tableau 1.

etc.

³En pratique lorsque le vecteur de caractéristique est très grand et les mesures caractéristiques utilisées sont redondantes, on doit utiliser des techniques comme l'ACP (Analyse en Composantes principales) ou AFD (Analyse Factorielle Discriminante) pour éliminer cette redondance et réduire la dimension de ce vecteur.

□ **Algorithme K-Moy.**

1. On choisit les K premiers échantillons comme étant les K centres $(c_1^{[1]}, \dots, c_K^{[1]})$ des K groupements,

$$c_i^{[1]} = \mathbf{x}_i, \quad 1 \leq i \leq K.$$

2. A l'itération p , on associe l'échantillon $\mathbf{x}_l$, ($1 \leq l \leq M$) au groupement C_i (noté $C_i^{[p]}$) si,

$$\|\mathbf{x}_l - c_i^{[p]}\| < \|\mathbf{x}_l - c_j^{[p]}\|, \quad \forall j \neq i.$$

En fait, on associe l'échantillon $\mathbf{x}_l$ au groupement dont le centre lui est le plus proche. En cas d'égalité, on associe $\mathbf{x}_l$ arbitrairement au groupement i ou j .

3. $C_i^{[p]}$ correspond au groupement i constitué à l'itération p après l'étape 2. On détermine le nouveau centre de chaque groupement par,

$$c_i^{[p+1]} = \frac{1}{N_i} \sum_{\mathbf{x} \in C_i^{[p]}} \mathbf{x}, \quad \forall i.$$

avec N_i le nombre d'échantillons de $C_i^{[p]}$.

4. Si les centres des groupements à l'itération p sont les mêmes que ceux obtenus à l'itération précédente, alors les groupements constitués à l'itération p et $(p-1)$ sont identiques et l'algorithme a convergé, sinon on retourne à l'étape 2.

si $c_i^{[p+1]} = c_i^{[p]}, \quad \forall i,$ fin de l'algorithme sinon on retourne en 2.

TAB. 1 – Algorithme des K-moyennes.