

République Algérienne Démocratique et Populaire
Ministère de l'Enseignement Supérieur et de la Recherche Scientifique
Université Mouloud MAMMERI, Tizi-Ouzou



Faculté de Génie Electrique et d'Informatique
Département d'Electronique

Mémoire de Fin d'Etudes

En vue de l'obtention du diplôme

D'Ingénieur d'Etat en Electronique
Option : Instrumentation

Thème

Traitement d'une image médicale
Application à la segmentation

Proposé et dirigé par : **M^r Y. ATTAF**

Présenté par :
M^{elle} KHEROUBI Ghania
M^{elle} ARAR Cherifa

Promotion 2011

REMERCIEMENTS

Remerciements

Nous exprimons notre profonde reconnaissance envers monsieur Youcef ATTAF pour son parfait encadrement, ses précieux conseils, sa disponibilité ainsi que pour ses qualités humaines et son bonne humeur.

Que tous les membres du jury trouvent ici l'expression de nos profonds respects pour avoir pris la peine d'examiner ce manuscrit.

Nos remerciements et notre gratitude vont aux professeurs et enseignants de l'UMMTO ainsi que son personnel côtoyés tout au long de notre cursus universitaire.

SOMMAIRE

Sommaire

Introduction générale.....	1
Chapitre 1: Généralités sur le traitement d'images	
1. Introduction	3
2. Notion d'image	3
3. Image et propriétés	4
3.1. Définition d'une image	4
3.2. Définition d'une image numérique	4
4. Différents types d'une image numérique	4
4.1. Image binaire.....	4
4.2. Image en niveau du gris	5
4.3. Image couleur.....	5
5. Caractéristiques d'une image numérique	6
5.1. Pixel.....	6
5.2. Dynamique de l'image.....	7
5.3. La dimension.....	7
5.4. La résolution.....	7
5.5. Le bruit	7
5.6. La luminance.....	7
5.7. Le contraste	7
5.8. Le contour	8
5.9. La région	8
5.10. Distance entre deux couleurs	8
5.11. Histogramme d'une image	9
6. Architecture générale d'un système de traitement d'images	9
6.1. Acquisition	9
6.2. Prétraitement	10
6.3. Le traitement numérique	11
6.4. Le post-traitement	11
6.5. Visualisation, transmission et stockage.....	11
7. Notion sur l'imagerie médicale.....	12
8. Exemple d'images médicales.....	13
9. Conclusion.....	14
Chapitre 2: La segmentation	
1. Introduction	15
2. Définition de la segmentation.....	15

3. Méthodes de segmentation	16
3.1.Segmentation par approche contour	16
3.2.Segmentation par approche région	16
3.2.1. Segmentation par division de région	16
3.2.2. Segmentation par croissance de région.....	16
3.2.3. Segmentation par division-fusion de région	17
3.3.Segmentation par classification	18
4. Texture	18
4.1.Analyse de texture	19
4.1.1. Les méthodes statistiques.....	19
4.1.2. Les méthodes structurales.....	19
4.2.Les attributs statistiques de textures	19
4.2.1. attributs statistiques de premier ordre	20
4.2.1.1.la moyenne	20
4.2.1.2. La variance	20
4.2.1.3.Le moment d'ordre 3	21
4.2.1.4. Le moment d'ordre 4	22
4.2.2. attributs statistiques de second ordre	23
4.2.3. attributs statistiques d'ordre supérieur.....	25
5. conclusion.....	27

Chapitre 3: La classification

1. Introduction	28
2. Définition d'une classe	28
3. Définition de la classification.....	28
4. Les techniques de classification	30
4.1. Classification supervisée (semi automatique)	30
4.1.1. Méthode métrique	30
4.1.2. Méthode statique.....	31
4.2. Classification non supervisée (automatique)	31
4.2.1. Méthode statique.....	32
4.2.1.1. Méthode non paramétrique	32
4.2.1.2. Méthode paramétrique	32
4.2.2. Méthode métrique.....	33
4.2.2.1. Classification hiérarchique	34
4.2.2.2. Classification non hiérarchique	34
5. Quelques méthodes de classification.....	35
5.1. Méthode des k plus proches voisins.....	35
5.2. L'approche densité des informations.....	36
5.3. Méthode de k-means	37
5.4. Méthode Fuzzy.C.means.....	39
6. Conclusion.....	41

Chapitre 4: Tests et résultats

1. Introduction	42
2. Etapes de segmentation	42
1. Télécharger et lire une image	42
3. Les différentes classes obtenues par les résultats de k-means.....	43
4. Discussion des résultats et conclusion	52
Conclusion générale	53

Annexes

Bibliographie

INTRODUCTION GÉNÉRALE

Introduction générale

L'image est devenue de nos jours le support d'informations le plus répandu et le plus performant vu son universalité, son accessibilité, son riche contenu et ses multiples applications dans tous les domaines de la science.

Avec le développement technologique et l'apparition de l'image numérique, la discipline traitement d'image a vu le jour. Elle a donc des objectifs variés, dont le but est à la fois simple dans son concept et difficile dans sa réalisation. Simple en effet, puisqu'il s'agit de reconnaître des objets que notre système visuel aperçoit rapidement, difficile cependant car dans la grande quantité d'informations contenue dans l'image, il faut extraire des éléments pertinents pour l'application visée et ceci indépendamment de la qualité de l'image. Le traitement d'image est alors doté d'outils et de méthodes puissants tels que les mathématiques et l'informatique. L'analyse d'images a pour but l'extraction de l'information caractéristique contenue dans une image.

Le résultat d'une telle analyse s'appelle très souvent la description structurelle. Celle-ci peut prendre la forme d'une image ou de toute structure de données permettant une description des entités contenues dans l'image. Par opposition avec la phase d'interprétation, l'analyse tente, dans la mesure du possible, de ne pas prendre en compte le contexte (*i.e.* l'application).

Essentiellement, l'analyse de l'image fait appel à la segmentation où l'on va tenter d'associer à chaque pixel de l'image un label en s'appuyant sur l'information portée (niveaux de gris ou couleur), sa distribution spatiale sur le support image, des modèles simples (le plus souvent des modèles géométriques).

La segmentation d'images ainsi définie est un domaine vaste où l'on retrouve de très nombreuses approches.

Toutes ces approches visent à l'extraction des indices visuels. Après de nombreuses années passées à rechercher la méthode optimale, les chercheurs ont compris que la segmentation idéale n'existait pas. On peut même montrer que le problème de la segmentation est le plus souvent un problème mal posé. Etant donnée une image, il existe toujours plusieurs segmentations possibles. Une bonne méthode de segmentation sera donc celle qui permettra d'arriver à une bonne interprétation.

Le cadre général dans lequel s'inscrit ce mémoire est celui de la segmentation d'image médicale qui donne une représentation du corps humain ; elle peut être obtenue à partir de différents phénomènes physiques tels que l'absorption de rayons X, la résonance magnétique nucléaire, la réflexion d'ondes ultrasons ou la radioactivité et cela par la méthode de classification non supervisée k-Means en changeant à chaque fois le nombre de classes et comparer les résultats obtenus.

Ce mémoire est structuré en quatre chapitres:

- Dans le premier chapitre, nous donnons un aperçu général sur l'image
- le second chapitre a pour contenu les différentes méthodes de segmentation, les principales méthodes d'analyse de textures et la définition des attributs statistiques de texture.
- Dans le troisième chapitre, nous expliquons les principales méthodes de classification et nous présentons la méthode de classification non supervisée utilisée à savoir la méthode des k-means.
- Dans le dernier chapitre nous allons discuter les résultats issus des différents tests effectués en appliquant la méthode choisie.
- Nous terminons par une conclusion générale, la bibliographie et les annexes.

CHAPITRE I : GÉNÉRALITÉS SUR LE TRAITEMENT D'IMAGES

1. Introduction

Avec la parole, l'image constitue l'un des moyens les plus importants qu'utilise l'homme pour communiquer avec autrui. C'est un moyen de communication universel dont la richesse du contenu permet aux êtres humains de tout âge et de toute culture de se comprendre.

C'est aussi le moyen le plus efficace pour communiquer, chacun peut analyser l'image à sa manière, pour en dégager une impression et d'en extraire des informations précises.

De ce fait, le traitement d'images est l'ensemble des méthodes et techniques opérant sur celles-ci, en corrigeant les dégradations subies lors de son acquisition ou en extrayant les informations permettant une interprétation visuelle ou automatique, pour mieux exploiter les informations qu'elle contient.

2. La notion d'image

Avant d'entrer dans le sujet proprement dit, il est important de comprendre la nature des objets que nous allons manipuler. Intéressons-nous à la notion d'image, qu'est-ce qu'une image ?

La définition du dictionnaire Larousse donne : «image : Représentation imprimée d'un sujet quelconque».

Une image est une représentation imprimée d'un sujet quelconque. Cela signifie qu'une image nécessite un support sur lequel elle sera imprimée. Ainsi une photographie papier et une peinture sont des exemples d'images au même titre qu'une image numérique affichée sur un écran d'ordinateur.

Informatiquement, une image sera une représentation numérique en mémoire d'un sujet imprimé sur une rétine artificielle (matricielle comme le capteur d'un appareil photographique numérique ou la scène virtuelle d'une image de synthèse ou bien linéaire comme le capteur optique du télécopieur, du photocopieur ou du scanner). Nous allons donc travailler sur des ensembles de nombres numériques codés sur un ordinateur.

Les premières images sont nées il y a de nombreuses années sous une forme primitive, certes, mais déjà très fidèle à la réalité du monde.

3. Image et propriétés

3.1. Définition d'une image

Une image peut avoir d'autres définitions. En traitement du signal, on définit une image comme étant un signal bidimensionnel. Mathématiquement parlant, une image est une application d'un sous ensemble $M \times N$ de $R \times R$ vers l'ensemble des réels R , qui à chaque couple de réels (x, y) associer le réel $f(x, y)$:

$$\left\{ \begin{array}{l} f(M, N) \longrightarrow R \\ (x, y) \longrightarrow f(x, y) \end{array} \right.$$

3.2 Définition d'une image numérique [19]

Une image numérique, encore appelée image BITMAP est une image dont la surface est divisée en éléments de tailles fixes appelés cellules ou pixels, ayant chacun comme caractéristique un niveau de gris ou de couleur prélevé à l'emplacement correspondant dans l'image réelle, ou calculé à partir d'une description interne de scène à représenter.

Il est très important de mentionner qu'une image numérique est obtenue par une conversion de celle-ci de son état analogique (signal électrique) vers un état numérique (binaire) pour qu'on puisse le traiter avec un système informatique.

Une image est une représentation bidimensionnelle d'objets tridimensionnelles de natures diverses. Elle peut être considérée comme un signal continu bidimensionnel ou présentée sous forme d'un ensemble de points répartis dans une surface donnée. Elle peut être définie comme étant une matrice rectangulaire, où chaque élément appelé pixel (extrait de l'expression anglaise " Picture élément") a comme caractéristique un niveau de gris ou de couleur, prélevé à l'emplacement correspondant dans chaque image réelle, ou calculé à partir d'une description interne de la scène à représenter.

4. Différents types d'une image numérique [2]

Principalement, il existe trois types d'images numériques qui sont des images : binaires, en niveau du gris et en couleurs

4.1. Image binaire

Une image binaire est une matrice rectangulaire dont le nombre de niveaux de gris est réduit en deux éléments 1 et 0, où le niveau 0 représente le noir et le niveau 1 représente le blanc.



Figure I.1: Exemple d'image binaire.

4.2. Image en niveau de gris

Le niveau de gris est la valeur de l'intensité lumineuse en un point. Le pixel peut prendre des valeurs allant du noir au blanc, en passant par un nombre fini de niveaux intermédiaires. Donc pour représenter des images en niveau de gris, on peut attribuer à chaque pixel de l'image une valeur correspondante à la qualité de la lumière renvoyée. Cette valeur est souvent comprise entre 0 et 255, chaque pixel est donc codé par 8 bits (un octet).



Figure 1.2: Exemple d'image en niveau de gris.

4.3. Image couleur

Ces images sont en général codées, en utilisant le codage des trois couleurs fondamentales (rouge, vert, bleu) ; on parle alors d'image RVB. Une couleur contient donc trois plans couleurs : rouge, vert et bleu (RVB). Chaque plan est codé comme une image en niveau de gris avec des valeurs allant de 0 à 255. Lorsque $R=V=B$, la valeur associée est un niveau de gris. D'autre part, pour aller d'une image couleur à une image en niveau de gris, on réalise la fonction suivante:

$$I(i, j) = \frac{R(i, j) + V(i, j) + B(i, j)}{3} \quad (I.2)$$

Avec:

$I(i, j)$: niveau de gris du pixel situé à la ligne i et à la colonne j .

$R(i, j)$: intensité de la couleur rouge du pixel (i, j) .

$V(i, j)$: intensité de la couleur verte du pixel (i, j) .

$B(i, j)$: intensité de la couleur bleue du pixel (i, j) .



Figure I.3: Exemple d'image couleur.

5. Caractéristique d'une image numérique [3]

Les principales caractéristiques d'une image numérique sont:

5.1. Le pixel

Contraction de l'expression anglaise " picture éléments ": éléments d'image. Le pixel est le plus petit point de l'image, c'est une entité calculable qui peut recevoir une structure et une quantification. Si le bit est la plus petite unité d'information que peut traiter un ordinateur, le pixel est le plus petit élément que peuvent manipuler les matériels et logiciels d'affichage ou d'impression. La lettre A, par exemple, peut être affichée comme un groupe de pixels dans la figure ci-dessous :



Figure I.4: un groupe de pixels

La quantité d'information que véhicule chaque pixel donne des nuances entre images monochromes et images couleurs. Dans le cas d'une image monochrome, chaque pixel est codé sur un octet, et la taille mémoire nécessaire pour afficher une telle image est directement liée à la taille de l'image.

Dans une image couleur (R.V.B.), un pixel peut être représenté sur trois octets : un octet pour chacune des couleurs : rouge (R), vert (V) et bleu (B).

5.2. Dynamique de l'image

C'est le nombre de niveaux de gris dans une image (c'est le nombre de niveaux de quantification).

5.3. La dimension

C'est la taille de l'image. Cette dernière se présente sous forme de matrice dont les éléments sont des valeurs numériques représentatives des intensités lumineuses. Le nombre de lignes de cette matrice multiplié par le nombre de colonnes nous donne le nombre total de pixels dans une image.

5.4. La résolution

La résolution est exprimée en nombre de pixels par unité de mesure (pouce ou centimètre). On utilise aussi le mot résolution pour désigner le nombre total de pixels affichables horizontalement ou verticalement sur un moniteur; plus grand est ce nombre, meilleure est la résolution.

5.5. Le bruit

Un bruit (parasite) dans une image est considéré comme un phénomène de brusque variation de l'intensité d'un pixel par rapport à ses voisins.

5.6. La luminance

C'est le degré de luminosité des points de l'image. Elle est définie aussi comme étant le quotient de l'intensité lumineuse d'une surface par l'aire apparente de cette surface. Pour un observateur lointain, le mot luminance est substitué au mot brillance, qui correspond à l'éclat d'un objet.

5.7. Le contraste

C'est l'opposition marquée entre deux régions d'une image, plus précisément entre les régions sombres et les régions claires de cette image. Le contraste est défini en fonction des luminances de deux zones d'images. Si L_1 et L_2 sont les degrés de luminosité respectivement de deux zones voisines A_1 et A_2 d'une image, le contraste C est défini par le rapport:

$$C = \frac{L1-L2}{L1+L2} \quad (1.3)$$

5.8. Le contour

Les contours représentent les frontières entre les objets de l'image, ou la limite entre deux pixels dont les niveaux de gris représentent une différence significative.

5.9. La région

C'est un ensemble connexe de pixels ayant une ou plusieurs propriétés communes; on parle de zone homogène de l'image.

Notion de voisinage:

Chaque pixel P (x, y) a quatre voisins horizontaux- verticaux et quatre voisins diagonaux

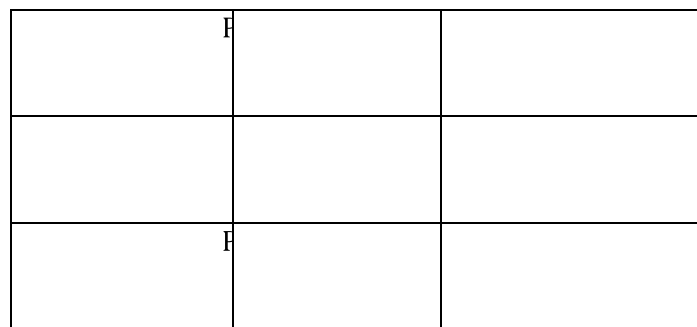


Figure 1.5: les huit voisins de P(x,y).

5.10. Distance entre deux couleurs:

Pour les images en niveau de gris, la distance entre deux couleurs est la différence des niveaux de gris. Par contre, pour les images couleur, cette distance peut être calculée différemment. Elle se traduit généralement par une distance euclidienne dans l'espace colorimétrique utilisé.

La distance euclidienne dans l'espace RVB est donnée par la formule suivante:

$$\Delta d = \sqrt{\Delta R^2 + \Delta V^2 + \Delta B^2} \quad (1.4)$$

5.11. Histogramme d'une image [6]

L'histogramme d'une image représente le nombre (ou la proportion) de pixels en fonction de niveau de gris dans l'image, c'est-à-dire la distribution des valeurs de niveaux de gris dans une image.

Pour les images en couleur, plusieurs histogrammes sont nécessaires. Par exemple pour une image codée en RVB, trois histogrammes représentant respectivement la distribution des valeurs respectives des composantes : rouges, bleues et vertes sont nécessaires.

6. Architecture générale d'un système de traitement d'images [2]

Le traitement d'image désigne une discipline de l'informatique et de mathématique appliquées, qui étudie les images numériques et leurs transformations dans le but d'améliorer leur qualité ou d'extraire des informations pertinentes.

Un système de traitement d'image est composé de plusieurs parties, la figure (I.6) montre les différentes étapes de traitement

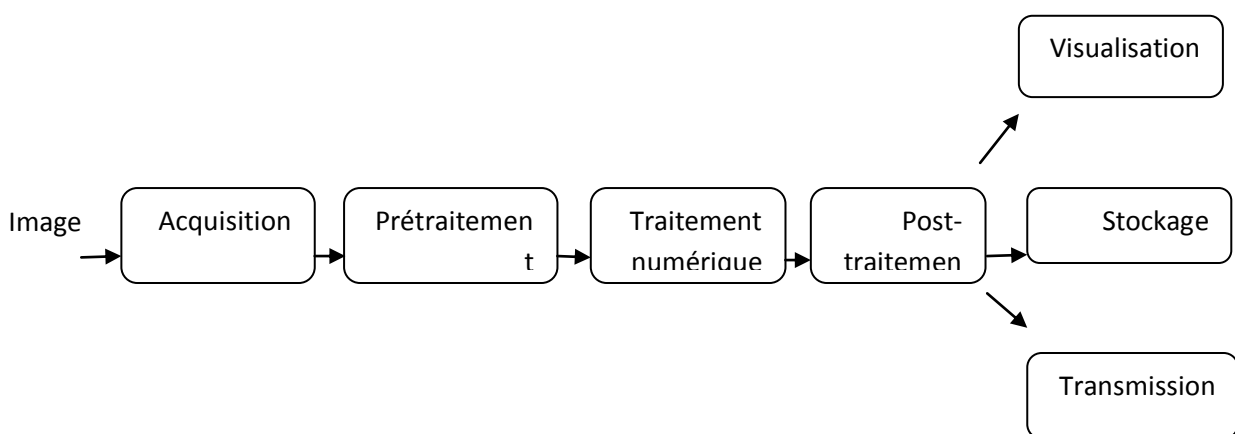


Figure I.6: Schéma synoptique d'un système de traitement d'image.

6.1. Acquisition

L'acquisition d'une image est l'opération qui permet le passage de l'information réelle à une représentation numérique. Elle est réalisée en 3 étapes :

- La transformation du signal optique en un signal analogique (électrique).
- l'échantillonnage qui consiste à multiplier le signal analogique par une série d'impulsion unité dans le but d'avoir des échantillons du signal.
- la quantification qui est une traduction des échantillons en valeurs numériques selon une règle de codage.

6.2. Prétraitement

En général, cette étape consiste à filtrer l'image. Deux méthodes seront abordées :

- Le filtre médian
- L'algorithme de Nagao et Al.

➤ Filtre Médian

L'algorithme est relativement simple. Le filtre remplace chaque pixel par la valeur médiane du pixel et de ses voisins. Le nombre de voisins peut être déterminé à l'aide d'un paramètre. Cette variable indique la taille du carré dans lequel seront effectués les calculs. L'algorithme est appliqué sur chaque composante RGB indépendamment. Voici un exemple d'une image choisie, de gauche à droite, l'originale aerial.jpg, filtrée avec une taille de 3 et filtrée avec une taille de 5 :

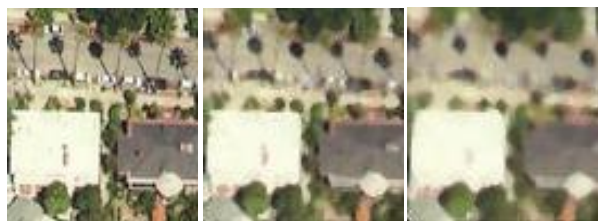


Figure 1.7 : Exemple du filtre médian.

➤ Algorithme de Nagao et Al

L'algorithme de Nagao est un peu plus complexe. Il permet d'adoucir l'image sans pour autant rendre les bords flous. Dans un carré de 5x5 entourant le pixel à traiter, la variance de 9 différentes régions est calculée. La valeur du pixel central est remplacée par la valeur moyenne de la région qui a la plus faible variance.

L'implémentation est la suivante :

Pour chaque pixel de l'image :

- Filtrer l'image par chacune des 9 matrices de Nagao.
- Calculer la variance des 9 résultats ci-dessus.
- Rechercher l'ensemble qui a la plus faible variance.
- Remplacer la valeur du pixel par la moyenne de l'ensemble ci-dessus.

Comme dans le cas du filtre médian, le filtre de Nagao est appliqué sur les trois composantes RGB indépendamment.

Ceci paraît facile, mais les boucles imbriquées sont très lentes. Par conséquent, cet algorithme fonctionne relativement lentement sous Matlab. Ceci peut devenir un inconvénient majeur si le nombre d'images à traiter est important. Voici un exemple du filtre de Nagao sur la même image que ci-dessus :

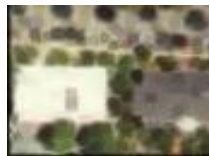


Figure 1.8 : Exemple du filtre de Nagao.

6.3. Le traitement numérique

Le traitement numérique d'image est l'ensemble de méthodes qui permettent de décrire quantitativement le contenu d'une image.

On distingue deux types de traitement : les traitements de bas niveau et les traitements de haut niveau. Cette distinction est liée au contenu sémantique des entités traitées et extraites de l'image.

Les traitements de bas niveau opèrent en général sur les grandeurs calculées à partir des valeurs attachées à chaque point de l'image sans faire la liaison avec la réalité. Ils opèrent plutôt sur des données de nature numérique.

Les traitements de haut niveau s'appliquent sur des entités de natures symboliques associées à une représentation de la réalité extraite de l'image ; ils sont réactifs à l'interprétation.

6.4. Le post-traitement

Le traitement numérique n'est en général pas parfait, le post-traitement est donc censé améliorer le résultat.

1. Dans tous les cas, les frontières ne font partie d'aucune région. Il est alors possible de les affecter à une région proche selon un critère, par exemple la différence de couleur.
2. Dans les cas de sur-segmentations, les petites régions peuvent être regroupées à d'autres.
3. Il serait aussi possible d'utiliser plusieurs méthodes de segmentation et de les unir afin d'obtenir un résultat plus satisfaisant.

6.5. Visualisation, transmission et stockage

La visualisation est une opération qui permet de transformer le signal numérique qui est la matrice image en un signal analogique visible par l'œil de l'observateur à travers un dispositif d'affichage.

L'image peut être transmise vers une station d'archivage ou de traitement qui peut être différente de la station d'acquisition.

Les images sont archivées sur des supports destinés à cet effet comme par exemple les bandes magnétiques, les disques durs,...

7. Notions sur l'imagerie médicale [3]

L'imagerie médicale regroupe un ensemble de méthodes permettant de visualiser des processus biologiques, au sein même des organismes vivants sans procéder à une opération.

Elle est essentielle à la compréhension de leur physiologie et de leurs pathologies ; afin de mieux les diagnostiquer, les pronostiquer, et les traiter. L'imagerie médicale regroupe les moyens d'acquisition et de restitution d'images à partir de différents phénomènes physiques ou chimiques tels que :

- **Les rayons X** : sont des rayonnements électromagnétiques de très courte longueur d'onde et donc très pénétrants. Ils sont utilisés en médecine pour faire les différents tests médicaux tels que la radiographie et le scanner.
- **Les ultrasons** : sont des ondes sonores imperceptibles à l'oreille humaine. Les ultrasons sont absorbés ou réfléchis par les substances qu'ils rencontrent. Les ultrasons sont utilisés dans l'imagerie médicale et principalement dans le domaine de l'échographie et du doppler. L'imagerie par ultrasons permet d'explorer toutes sortes d'organes qu'ils soient superficiels (thyroïde, sein, muscle, articulation) ou profonds (foie, vésicule biliaire, pancréas, rein, vessie, gros vaisseaux utérus).
- **Les rayons gamma**: Les rayons gamma sont sur le spectre lumineux des rayonnements ayant la plus grande fréquence et la plus grande longueur d'onde, ce qui implique que les rayons gamma soient si énergétique car il faut savoir que plus un photon a une fréquence élevée plus il est énergétique.

Deux techniques d'imagerie médicale font appel au phénomène radioactivité : la scintigraphie et la tomographie par émissions de positions.

- **L'imagerie par résonance magnétique** : Elle est utilisée pour faire un diagnostic qui se fonde sur les principes de la résonance magnétique nucléaire (RMN). L'IRM est la méthode de diagnostic la plus puissante et la plus sensible disponible actuellement. Cet outil permet d'obtenir des images plus précises de tissus à l'intérieur du corps humain, que celles obtenues par un scanner ou par ultrasons.
- **L'endoscopie** : C'est une technique d'imagerie médicale très précise puisqu'elle offre des indications de texture du tissu, de volume ou de coloration. Elle permet un examen "en direct" de certains organes. L'endoscopie consiste à introduire après une incision cutanée un tube, l'endoscope, muni d'une optique (mini camera couleur), d'un système d'éclairage, d'une pince à biopsie pour la réalisation d'un prélèvement et d'une pince à corps étranger, dans l'organisme. L'endoscopie peut être utilisée pour explorer le thorax (thoracoscopie), l'abdomen (coelioscopie) ou dans une articulation (arthroscopie), ou encore par les voies naturelles, par le méat urinaire pour les uretères, la vessie et l'urètre, par l'anus pour le rectum et le côlon (coloscopie).

8. Exemples d'images médicales



Figure 1.9 : Image radiographique.



Figure I.10 : Image obtenue par scanner.



Figure I.11 : Image endoscopique.

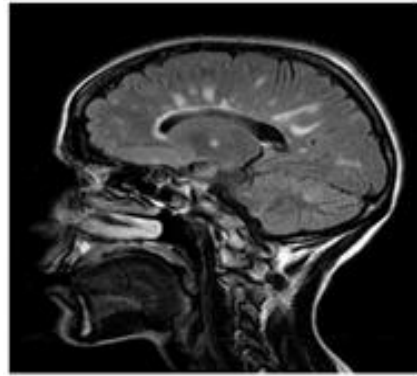


Figure I.12 : Image IRM.

9. Conclusion

Dans ce chapitre nous avons présenté d'une manière générale quelques notions essentielles de l'image, les différents types d'image, les caractéristiques d'une image numérique ainsi que l'architecture d'un système de traitement d'image et des notions de base sur l'imagerie médicale.

1. Introduction

La segmentation des images constitue le cœur de tout système de vision et une étape importante dans le processus d'analyse des images. Elle a pour objectif de fournir une description des objets contenus dans l'image par l'extraction de différents indices visuels tels que les contours des objets, les régions homogènes et les objets 3D. Ils seront exploités ensuite pour une description symbolique de la scène permettant une interprétation et éventuellement une prise de décision.

La segmentation d'images dépend fortement aussi bien du type de traitement fixé par l'utilisateur sur les objets présents dans l'image, que de la nature de l'image (présence de bruit, présence de zones texturées, contours flous....), ainsi que des primitives que l'on cherche à extraire de l'image et qui dépendent des opérations situées en aval de la segmentation (localisation, calcul 3D, reconnaissance de formes, interprétation).

2. Définition de la segmentation [6]

La segmentation d'image est une opération de traitement d'image qui a pour but de séparer des objets du fond d'une image numérique et entre eux, ou d'une manière plus générale, de rassembler les différents pixels ayant des propriétés communes. Si la segmentation sépare deux classes de pixels ; elle est appelée binarisation.

L'homme sait naturellement séparer des objets dans une image. Il se base pour cela sur des connaissances de haut niveau, qui lui permettent de détecter ce qui l'intéresse dans l'image.

En traitement d'image, la segmentation est le partitionnement d'une image I en N régions R_i disjointes et homogènes selon un prédicat, ou critère P déterminé (couleur, intensité de texture....).

Formellement, l'ensemble des régions $\{R_1, R_2, \dots, R_N\}$ est une segmentation de l'image I si :

$$\checkmark \quad \bigcup_{i=1}^N (R_i) = I, \quad i = 1, \dots, N \quad (\text{II.1})$$

$$\checkmark \quad R_i \cap R_j = \emptyset \quad \forall i \neq j \quad (\text{II.2})$$

$$\checkmark \quad R_i \text{ est connexe } \forall i \in \{1, \dots, N\} \quad (\text{II.3})$$

$$\checkmark \quad P(R_i) = \text{vrai} \quad \forall i \in \{1, \dots, N\} \quad (\text{II.4})$$

$$\checkmark \quad P(R_i \cup R_j) = \text{faux} \quad \forall i \neq j \text{ et } R_i \text{ adjacente à } R_j \quad (\text{II.5})$$

Plusieurs algorithmes de segmentation existent selon les diversités des propriétés recherchées, les problèmes rencontrés et la manière d'opérer le regroupement des pixels.

3. Méthodes de segmentation

3.1. Segmentation par approches contour

Pour une image en niveau de gris, nous définissons un contour comme une frontière entre deux régions de niveaux de gris différents et relativement homogènes.

L'approche contour consiste donc à identifier les transitions entre régions.

Un contour est alors déterminé par une variation rapide des niveaux de gris des pixels consécutifs.

3.2. Segmentation par approches région

Une région est un ensemble connexe de pixels voisins de l'image ayant des propriétés communes (intensité, texture,..) qui différencient des pixels des régions voisines.

Les principales techniques utilisées dans la segmentation des images par approche région sont:

- la segmentation par division de régions;
- la segmentation par croissance de régions.
- la segmentation par division-fusion de régions.

3.2.1. Segmentation par division de région

L'approche segmentation par division de régions consiste à diviser l'image originale en régions homogènes au sens d'un critère donné.

Ce processus est récursif et considère que la région initiale (R) correspond à l'image à segmenter (I). Si une région ne respecte pas un prédicat d'homogénéité, elle est divisée en quatre sous-régions de tailles égales (structure quadrée). Chaque sous-région est ensuite analysée. L'algorithme récursif s'arrête lorsque toutes les régions respectent le prédicat d'homogénéité.

A cause de l'utilisation de la structure quadrée, cette méthode est plutôt adaptée aux images carrées ayant un nombre de lignes et de colonnes égales à une puissance de deux, et dans lesquelles les régions sont de forme rectangulaire.

3.2.2. Segmentation par croissance de région

Ce type de segmentation consiste à faire croître des régions en y ajoutant successivement les pixels adjacents qui satisfont un critère d'homogénéité. La croissance s'arrête lorsque tous les pixels ont été traités.

L'avantage de la croissance de région est de préserver la forme de chaque région de l'image. Cependant, une mauvaise sélection des régions entraîne des phénomènes de sous-segmentation ou de sur-segmentation. Un exemple de segmentation par croissance de région est donné par la figure 2.1.

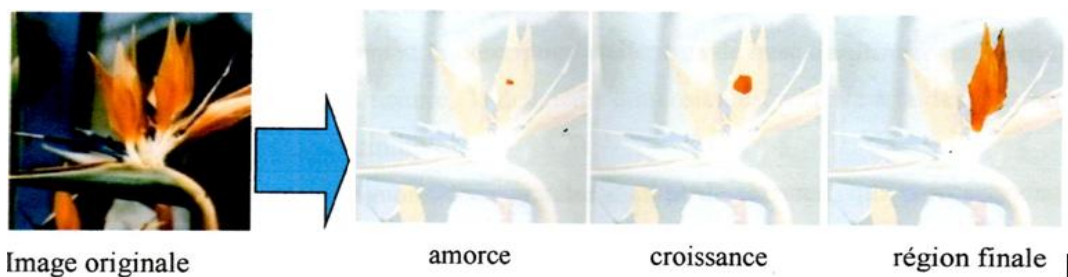


Figure II.1 : Exemple de segmentation par croissance de régions.

3.2.3. Segmentation par division-fusion de régions

La procédure de découpage décrite dans la méthode de division de région aboutit à un nombre de régions trop élevé. La cause fondamentale de cette sur-segmentation est que l'algorithme découpe les régions de manière arbitraire. Il se peut qu'il coupe de cette façon une zone homogène en deux ou en quatre parties. Alors, la phase de regroupement consiste à connecter les régions adjacentes et homogènes.

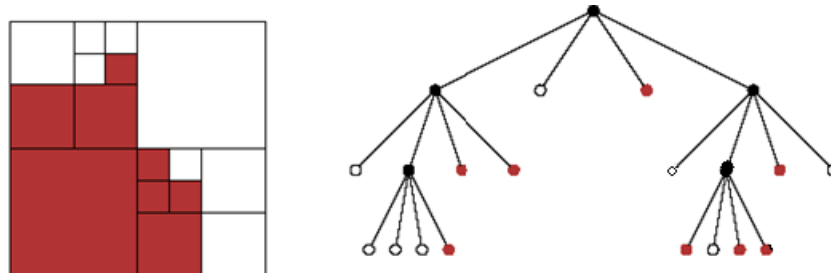


Figure II.2: Structure quatrée.

L'algorithme de fusion commence par une image pré-segmentée, pour ensuite :

1. Fusionner tout couple de régions adjacentes qui vérifie un critère d'homogénéité.
2. Définir un prédicat Fusionne (R_i, R_j) où R_i et R_j sont deux régions adjacentes.

La labellisation est caractérisée pour sa connexité à savoir 4 ou 8 :



Figure II.3 : Connexité

3.3. Segmentation par classification

Dans cette approche, la segmentation sera vue comme étant un problème de classification de données. Les techniques de classification identifient les classes des pixels homogènes dans l'espace des caractéristiques utilisées (niveau de gris, paramètre de texture,...) sous la forme de nuages de points opaques. Ensuite chaque classe est étiquetée donnant ainsi l'image des étiquettes.

Les techniques de classification automatiques seront détaillées dans le chapitre suivant.

4. Définition de la texture [7]

La texture est une des propriétés de base de l'image. Elle représente l'ordonnancement visuel du changement des niveaux de gris dans une surface de l'espace ou une structure d'un objet. **Maitre** a proposé la définition suivante : une texture est un champ de l'image qui apparaît comme un domaine cohérent et homogène. La figure II.4 montre quelques images de la texture.

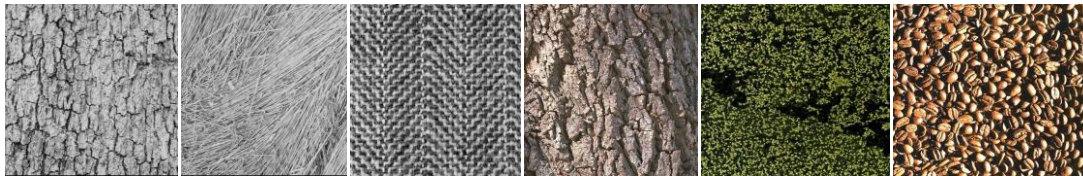


Figure II.4 : Exemples de textures.

En pratique, nous distinguons deux grandes classes de textures, qui correspondent à deux niveaux de perception : les macro-structures et les micro-structures.

- **Les macro-structures** : présentent un aspect régulier, sous forme de motifs répétitifs, spatialement placés, selon une règle précise (ex : peau de lézard, mur de brique) elles nécessitent une approche structurale déterministe pour leur analyse.
- **Les micro-structures** présentent des primitives « microscopiques » distribuées de manière aléatoire (ex : sable, laine tissée, herbe).

4.1. Analyse de texture [10]

L'analyse de texture regroupe un ensemble de techniques mathématiques permettant de quantifier les différents niveaux de gris présents dans une image en termes d'intensité et de distribution. Le but de l'analyse de texture est de formaliser les descriptifs de la texture par des paramètres mathématiques qui serviraient à l'identifier.

Comme il existe deux grandes classes de texture, deux approches seront donc nécessaires : les méthodes statistiques et les méthodes structurales.

4.1.1. Les méthodes statistiques

Elles étudient les relations entre un pixel et ses voisins. Elles sont utilisées pour caractériser des structures fines, sans régularité apparente. L'étude se fait sur des structures tout à fait aléatoires et plus souvent non homogènes, c'est pourquoi ce type de méthodes sera préférentiellement utilisé. Dans ce cas, la texture est décrite par des statistiques de la distribution de ces niveaux de gris (ou intensité).

4.1.2. Les méthodes structurales

Ces méthodes tiennent compte de l'information structurale et contextuelle d'une forme et sont particulièrement bien adaptées aux structures macroscopiques. Elles permettent de décrire la texture en définissant les primitives et les règles d'arrangement qui les relient. En effet, les textures ordonnées possèdent des primitives qui se répètent dans les images en des positions selon une certaine loi.

4.2. Les attributs statistiques de texture [11]

La distinction aisément faite par l'œil humain entre plusieurs textures est une tâche difficile à réaliser en vision par ordinateur, dans la mesure où il existe un nombre infini de textures et où chacune possède ses propres caractéristiques. Il n'existe pas de méthodes capables de classifier toutes les textures aussi bien que le ferait un observateur humain. Cependant, certaines propriétés définissant les impressions visuelles, peuvent être extraites et permettent par conséquent de reproduire au mieux la classification de texture, effectuée par cet observateur humain.

Ces propriétés sont à la base de nombreuses méthodes d'analyse de texture. Elles ont permis de définir un grand nombre d'attributs caractérisant les textures présentes dans les images. Ces attributs peuvent être géométriques, basés sur la modélisation spatiale des textures, spatiaux-fréquentiels et statistiques.

Les attributs statistiques permettent de caractériser tous les types de textures, même les textures fines et sans régularité apparente. Ces attributs peuvent être divisés en plusieurs catégories selon leur « ordre ».

L'ordre des attributs dépend du type d'interaction spatiale entre les pixels considérés. Il est donné par le nombre de pixels mis en jeu dans le calcul des paramètres. Par exemple, pour les histogrammes d'images, on ne s'intéresse qu'au pixel lui-même, ce paramètre appartient donc à la catégorie des attributs statistiques de premier ordre.

Quant au calcul des matrices de cooccurrence, se sont les couples de pixels qui sont considérés ; ce sont des attributs d'ordre deux.

4.2.1. Attributs statistiques du premier ordre [12]

De nombreux attributs peuvent directement être extraits de l'image afin de caractériser la texture qu'il contient. Les attributs statistiques du premier ordre donnent des informations générales caractérisant une image en fonction de la variation de l'intensité des pixels qui la composent.

Les plus courants sont :

- La moyenne
- La variance
- Le moment d'ordre 3
- Le moment d'ordre 4

4.2.1.1. La moyenne

La valeur moyenne (ou intensité moyenne) des niveaux de gris de tous les pixels est donnée par :

$$MOY = \frac{1}{N} \sum_{x,y} I(x,y) \quad (II.6)$$

avec $I(x,y)$ représente la valeur du niveau de gris du pixel (x,y) et N correspond au nombre total des pixels.

La figure II.5 illustre un exemple de la même image texturée mais ayant des moyennes différentes.

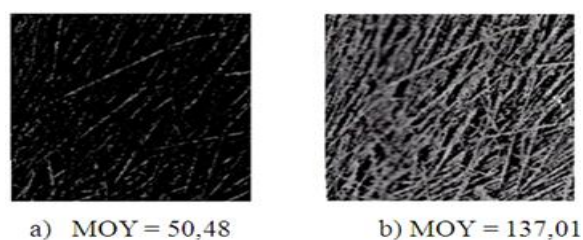


Figure II.5 : Illustration de la moyenne sur une même image de Brodatz.

4.2.1.2. La variance

Elle correspond au moment d'ordre 2. Elle mesure la répartition des niveaux de gris autour de la valeur moyenne. Sa valeur est donnée par :

$$VAR = \frac{1}{N} \sum_{x,y} (I(x,y) - MOY)^2 \quad (II.7)$$

La figure II.6 fournit deux exemples d'images ayant des variances différentes.

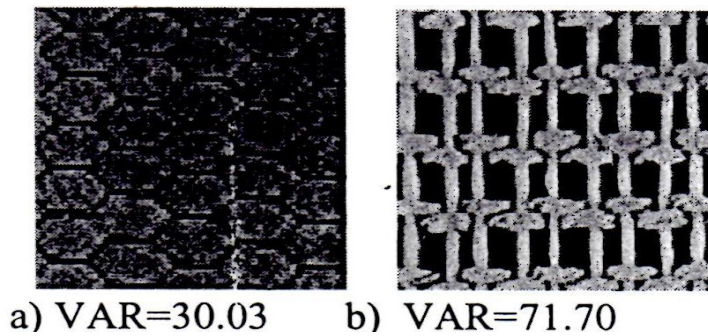


Figure II.6 : Illustration de différentes valeurs de la variance.

4.2.1.3. Le moment d'ordre 3

Le moment d'ordre 3 (en anglais skewness) centré autour de la moyenne, mesure la déviation de la distribution des niveaux de gris par rapport à une distribution symétrique. Il est donné par :

$$M_3 = \frac{1}{N} \sum_{x,y} (I(x,y) - MOY)^3 \quad (II.8)$$

Il mesure le degré (d'asymétrie) des niveaux par rapport à leur moyenne. Pour une déviation vers les valeurs élevées, le moment d'ordre 3 est positif, alors que pour une déviation vers les basses valeurs, il est négatif.

La figure 2.7 fournit deux exemples d'images et leurs histogrammes ayant des skewness différents.

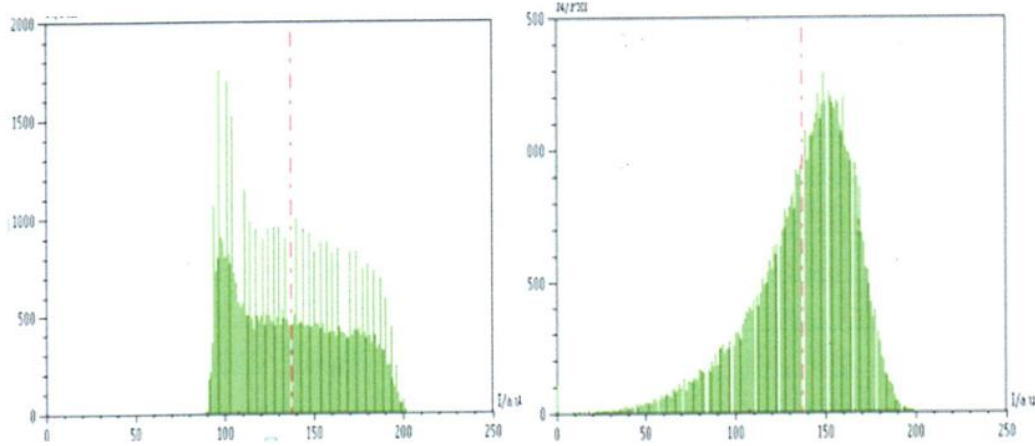
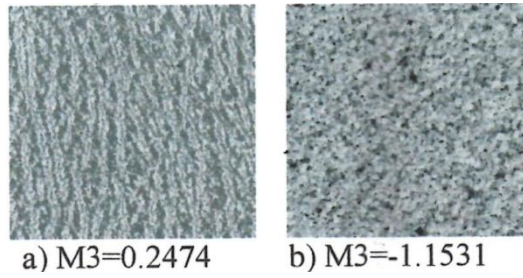


Figure II.7 : Illustration du moment d'ordre 3.

Ces deux images possèdent la même moyenne. En revanche, leurs skewness sont différents. La texture de la figure (a) qui possède un skewness positif possède un histogramme décalé vers la droite.

4.2.1.4. Le moment d'ordre 4

Le moment d'ordre 4 centré autour de la moyenne (en anglais kurtosis) est donné par :

$$KURT = \frac{1}{N} \sum_{x,y} (I(x,y) - MOY)^4 \quad (II.9)$$

Ce paramètre estime le degré de concavité ou convexité des niveaux par rapport à la moyenne. La figure II.8 illustre deux exemples d'images et leurs histogrammes ayant des kurtosis différents.

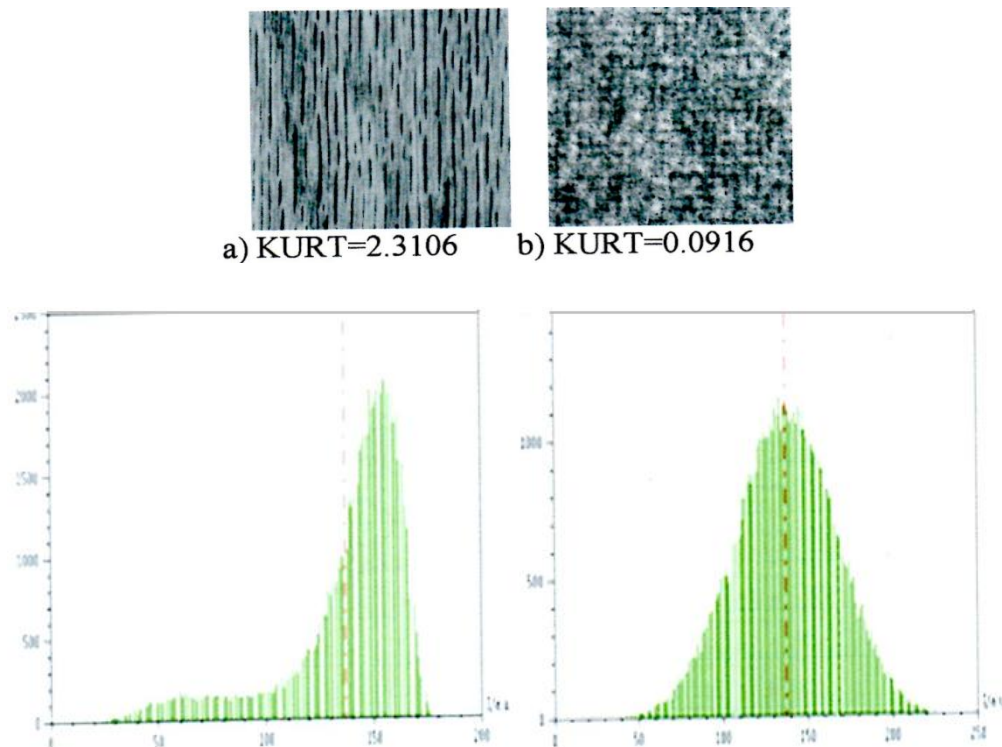


Figure II.8 : Illustration sur les différentes valeurs du moment d'ordre 4.

On peut constater que la texture qui possède un moment d'ordre 4 plus élevé (figure a) possède une distribution décentrée vers les valeurs élevées de niveaux de gris.

4.2.2 Attributs statistiques du second ordre [12]

Dans les attributs de premier ordre, il n'y a pas d'informations sur la localisation du pixel. Il est donc nécessaire d'utiliser des méthodes d'ordre supérieur pour une analyse plus précise.

En particulier, on retiendra la méthode de la matrice de cooccurrence, dite aussi méthode de dépendance spatiale des niveaux de gris.

Elle permet de déterminer la fréquence d'apparition d'un motif formé de deux pixels séparés par une certaine distance « d » dans une direction particulière θ . Afin de limiter le nombre de calculs, on prend généralement comme valeurs de la direction θ : 0° , 45° , 90° , 135° , 180° et 1 pour la valeur de « d ».

Ces indices bien que corrélés, réduisent l'information contenue dans la matrice de cooccurrence et permettent une meilleure discrimination des textures.

Le choix du vecteur de déplacement et de la taille de la fenêtre du voisinage sur laquelle s'effectue la mesure, sont les paramètres sur lesquels reposent la réussite de la méthode. La difficulté à surmonter lors de l'application de cette technique réside justement dans le choix de ces paramètres car ils varient selon le type d'images et de textures.

A chaque direction θ et pour chaque valeur de d correspond une matrice de cooccurrence $\varphi(d, \theta)$. La figure II.9 illustre un exemple de construction de la matrice de cooccurrence pour $\theta = 0^\circ$ et $d=1$

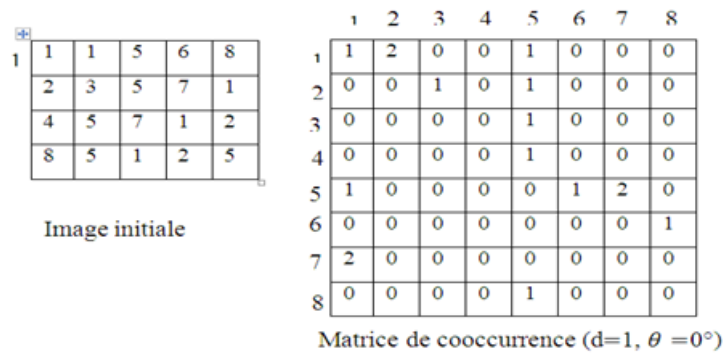


Figure II.9 : Exemple de construction d'une matrice de cooccurrence.

On définit généralement les matrices symétriques de cooccurrence. Elles sont construites à partir des constatations suivantes :

$$\varphi(d, 0) = \varphi^t(d, 180^\circ)$$

$$\varphi(d, 45^\circ) = \varphi^t(d, 225^\circ)$$

$$\varphi(d, 90^\circ) = \varphi^t(d, 270^\circ)$$

$$\varphi(d, 135^\circ) = \varphi^t(d, 315^\circ)$$

Par exemple, la matrice symétrique associée à la direction 0° sera de la forme :

$$S_0(d) = \frac{1}{2} [\varphi(d, 0^\circ) + \varphi(d, 180^\circ)] \quad (\text{II.10})$$

Une fois la matrice symétrique est réalisée, il est possible d'en extraire une quinzaine de paramètres qui contiennent des informations sur la finesse et la granularité de la texture.

Pour une texture grossière, les valeurs de la matrice sont concentrées sur la diagonale principale. Au contraire, pour une texture fine, les valeurs de la matrice seront dispersées : en effet, pour une telle texture il existe beaucoup de transitions de niveaux de gris.

A partir de cette matrice de cooccurrence, on peut extraire la moyenne, la variance, l'énergie, le contraste, la corrélation et l'homogénéité.

La moyenne est donnée par la relation :

$$MOY = \frac{1}{Ng \times Ng} \sum_{i=1}^{Ng} \sum_{j=1}^{Ng} p(i, j)^2 \quad (II.11)$$

$P(i, j)$: correspond aux éléments de la matrice de cooccurrence, c'est-à-dire à la probabilité de passer d'un pixel de niveau de gris i à un pixel de niveau de gris j .

i, j : Deux niveaux de gris.

Ng : correspond au maximum des niveaux de gris de l'image.

La variance est définie par :

$$VAR = \frac{1}{Ng \times Ng} \sum_{i=1}^{Ng} \sum_{j=1}^{Ng} (i - MOY)^2 p(i, j) \quad (II.12)$$

Elle caractérise la distribution des niveaux de gris autour de la valeur moyenne MOY calculée précédemment.

L'énergie ou moment angulaire d'ordre deux mesure l'homogénéité de l'image. L'énergie a une valeur d'autant plus faible qu'il y a peu de zones homogènes : dans ce cas, il existe beaucoup de transitions de niveaux de gris.

$$E = \sum_{i=1}^{Ng} \sum_{j=1}^{Ng} p(i, j)^2 \quad (II.13)$$

Le contraste ou (l'inertie) mesure les variations locales des niveaux de gris. Si elles sont importantes (c'est-à-dire s'il existe peu de régions homogènes), alors le contraste sera élevé.

Ce paramètre permet aussi de caractériser la dispersion des valeurs de la matrice par rapport à sa diagonale principale.

$$CONT = \sum_{i=1}^{Ng} \sum_{j=1}^{Ng} (i - j)^2 p(i, j) \quad (II.14)$$

La corrélation est donnée par: $\mu_i \mu_j$

$$COR = \frac{\sum_{i=1}^{Ng} \sum_{j=1}^{Ng} i j p(i, j) - \mu_i \mu_j}{\sigma_i \sigma_j} \quad (II.15)$$

Où μ_i et μ_j représentent les moyennes respectivement des lignes et des colonnes de la matrice et σ_i et σ_j représentent les écarts types respectivement des lignes et des colonnes de la matrice.

L'homogénéité a un comportement inverse du contraste. Plus la texture possède de régions homogènes, plus le paramètre est élevé. Elle est donnée par :

$$HOM = \sum_{i=1}^{Ng} \sum_{j=1}^{Ng} \frac{1}{1 + |i - j|} p(i, j) \quad (II.16)$$

4.2.3 Attributs statistiques d'ordre supérieur [13]

Les attributs statistiques d'ordre supérieur sont extraits en étudiant les interactions entre plusieurs pixels. La méthode des longueurs de plages de niveaux de gris est l'une des plus connues. Elle consiste à compter le nombre de plages d'une certaine longueur j , de niveau de gris i dans une direction θ donnée (généralement θ vaut $(0^\circ, 45^\circ, 90^\circ$ ou $135^\circ)$. Une plage de niveau de gris

correspond donc à l'ensemble des pixels d'une image ayant la même valeur de niveau de gris. La longueur de la plage correspond au nombre de pixels appartenant à la plage. A chaque direction θ correspondra donc une matrice $R(\theta)$ donnée par :

$$R(\theta) = [r_{\theta}(i, j)] \quad (II.17)$$

Avec $r(i, j)$: élément de la matrice $R(\theta)$, représente le nombre de plages de niveau de gris i , de longueur j dans la direction θ . La figure suivante donne un exemple de construction de la matrice de longueur de plages de niveau de gris.

0	1	2	3
0	2	3	3
2	1	1	1
3	0	3	0

Image initiale

i \ i	1	2	3	4
0	4	0	0	0
1	4	2	1	0
2	3	0	0	0
3	5	1	0	0

Matrice $R(0^\circ)$ de longueur de plages associées

Figure 2.10 : Exemple de construction de la matrice de longueur de plages de niveaux de gris.

Les principaux paramètres issus de la matrice de longueur de plages sont le « poids » des plages courtes, le poids des plages longues et la distribution des niveaux de gris.

Le poids des plages courtes caractérise la finesse de la texture. Plus la valeur de cet attribut est grande, plus les zones ayant le même niveau sont petites. Il est défini par la relation :

$$SRE = \frac{1}{nr} \sum_{i=0}^{M-1} \sum_{j=1}^N \frac{r(i, j)}{j^2} \quad (II.18)$$

Avec $nr = \sum_{i=0}^{M-1} \sum_{j=1}^N r(i, j)$: représente le nombre total de plages de l'image.

M correspond au nombre de niveaux de gris dans l'image et N à la longueur de la plage maximale.

Le « poids » des plages longues caractérise les textures grossières. Plus la valeur de cet attribut est grande plus il y'a des zones étendues ayant le même niveau. Il est donné par :

$$LRE = \frac{1}{nr} \sum_{i=0}^{M-1} \sum_{j=1}^N j^2 r(i, j) \quad (II.19)$$

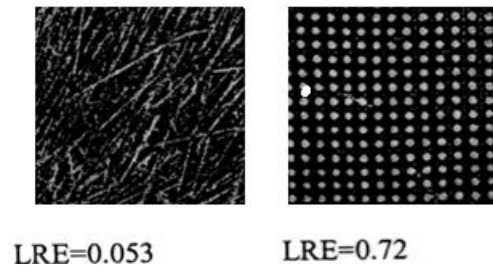


Figure II.9 : Illustration du paramètre LRE

Nous constatons bien que l'image qui possède un LRE plus grand possède plus de plages, c'est-à-dire plus de zones étendues ayant le même niveau de gris.

La distribution des niveaux de gris mesure l'uniformité de la distribution des plages. Cet attribut est minimal lorsque les plages sont également distribuées entre les différents niveaux.

Elle est définie par :

$$RLDIST = \frac{1}{nr} \sum_{i=0}^{M-1} [\sum_{j=1}^N r(i, j)]^2 \quad (II.20)$$

Ce paramètre augmente si le nombre de plages de même augmente

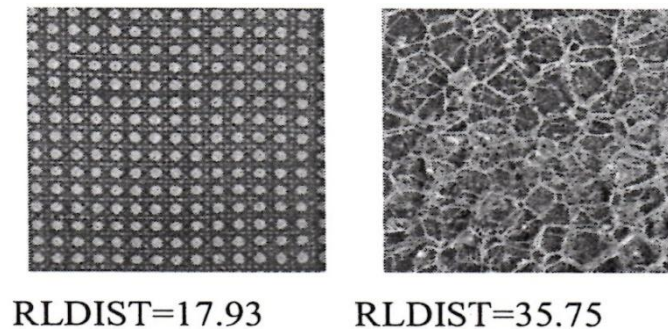


Figure II.10: Illustration du RLDIST

5. Conclusion :

Dans ce chapitre, nous avons présenté d'une manière générale quelques notions sur la segmentation, les notions et les attributs de texture qui décrivent les régions, puis les principales méthodes d'analyse de texture. Nous avons défini aussi les différentes techniques de segmentation.

L'une des stratégies possibles pour la segmentation est celle de la classification des pixels. Nous utiliserons dans notre travail une méthode de classification non supervisée k-means pour classer les pixels d'une image médicale. Nous allons présenter dans le chapitre suivant notre méthode ainsi que d'autres méthodes de classification.

1. Introduction

En reconnaissance de formes aussi bien qu'en taxinomie, la classification d'objets, d'individus ou d'observations issues d'une population, constitue une grande préoccupation.

Les objets sont généralement représentés sur un vecteur de mesure de dimension N correspondant aux réponses d'individus à un ensemble de paramètres ou variables

caractéristiques définies à priori. Ces observations sont fréquemment représentées comme des points projetés dans un espace à N dimensions.

2. Définition d'une classe

Une classe est composée d'un certain nombre d'objets similaires qu'on peut regrouper ensemble. On a:

- Une classe est un ensemble d'entités qui sont semblables, alors que les entiers provenant de classes différentes ne sont pas semblables.
- Une classe est un agrégat de points dans l'espace de représentation de données tel que la distance entre le centre et un point de cette classe est moins importante que celle entre le centre de cette classe et n'importe quel point d'une autre classe.
- Les classes peuvent être décrites comme des régions connexes de l'espace contenant relativement une grande densité de points.

3. Définition de la classification [14]

La classification consiste à organiser un ensemble de données multidimensionnelles en un ensemble fini de classes selon un ou plusieurs critères de classification.

La figure III.1 illustre le principe de la classification. Les données sont représentées par des points (vecteurs) dans un espace à N dimensions (dans notre exemple $N=2$). En sortie de la classification, nous obtenons m classes de points selon les critères donnés (dans notre exemple $m=3$).

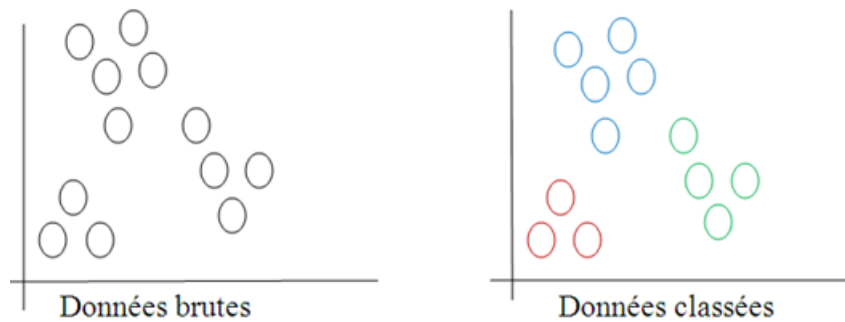


Figure III.1 : Le principe de la classification.

En classification, les objets ou observations sont définis par un vecteur de mesure. Ce vecteur de dimension P correspond à un ensemble de P attributs ou variables (caractéristiques) définies à priori. Ces observations sont généralement représentées comme des points projetés dans un espace à P dimensions.

Les observations ayant des attributs semblables, appartiennent à la même classe et forment un nuage de points très dense dans l'espace à P dimensions. Ainsi, la classification en général est définie comme étant l'action de construire une collection d'objets similaires au sein d'un même groupe, dissimilaires quand ils appartiennent à des groupes différents.

Elle est communément connue sous le nom anglophone « clustering ».

Mathématiquement, la classification est définie de la manière suivante :

Soit $X = \{X_1, X_2, \dots, X_N\}$ l'ensemble de N observations à classer.

Chaque observation X_i est caractérisée par P paramètres.

$X_i = \{x_{i1}, x_{i2}, \dots, x_{iP}\}$ tel que $X_i \in R^P$.

Soit $C = \{C_1, C_2, \dots, C_k\}$ l'ensemble des k classes.

La classification consiste à répartir l'ensemble de N observations en k classes tel que :

$$\begin{array}{ll} \blacksquare & C_i \neq \emptyset \text{ pour } i = 1, 2, \dots, k \end{array} \quad \text{(III.1)}$$

$$\blacksquare \quad C_i \cap C_j = \emptyset \quad \forall i \neq j \quad \text{(III.2)}$$

$$\blacksquare \quad \bigcup_{i=1}^n C_i = X \quad \text{(III.3)}$$

De plus, on exige de la classification de vérifier deux propriétés :

- Compacité : les points représentant une classe donnée sont plus proches entre eux que des points de toutes les autres classes.
- Séparabilité : les classes sont bornées et il n'y a pas de recouvrement.

4. Les techniques de classification

Pour partitionner l'ensemble des observations en un certain nombre de classes distinctes, on peut faire appel aux différentes méthodes de classification automatique qui peuvent être distinguées par leur caractère supervisé ou non supervisé.

4.1. Classification supervisée (Semi automatique) [4]

La classification est dite supervisée si l'on dispose à priori de l'information telle le nombre de classes possibles et l'appartenance de chaque individu à l'une de ces classes. La plupart des algorithmes d'apprentissage supervisés tentent donc de trouver un modèle qui explique le lien entre les données d'entrée et les classes de sortie.

Deux méthodes existent :

4.1.1. Méthodes métriques

Elle consiste à représenter les classes par des domaines délimités par des surfaces. Les surfaces les plus utilisées sont les surfaces linéaires hyper-planes. A partir des équations de ces dernières, on construit un ensemble de règles permettant d'affecter les observations aux différentes classes correspondantes. Ces règles sont généralement similaires à :

$$\left\{ \begin{array}{ll} (X) \in C_k & \text{Si } d_k(X) > d_k \\ & \text{Où } d_k(X) \text{ est appelée la fonction discriminante} \end{array} \right. \quad (III.4)$$

Où $d_k(X)$ est appelée la fonction discriminante

$$X \in D_k \equiv X \in C_k$$

$d_k(X) = d_k^*(X)$: correspond à l'équation de l'hyperplan.

4.1.2. Méthodes statiques

Dans le cadre des méthodes statistiques, on fait appel aux caractéristiques statistiques de la distribution des observations. La théorie de la décision constitue une approche fondamentale des problèmes de classification qui se trouve posée en terme probabiliste. Elle permet d'effectuer un classement optimal à partir de la règle de Bayes qui elle-même basée sur la connaissance des probabilités conditionnelles associées à chaque classe.

Soit $X = (x_1, x_2, \dots, x_N)$ une observation caractérisée par N paramètres à classer parmi un nombre K de classes. Si on désigne par $P(X/C_K)$ la fonction de densité de probabilité sous jacente à la distribution des observations provenant de la classe k et si $P(C_K)$ la probabilité à priori de cette classe, alors la probabilité $P(C_K/X)$ que l'observation x appartienne à la classe k est selon la règle de Bayes:

$$P\left(\frac{X}{C_k}\right) = \frac{P\left(\frac{X}{C_k}\right)}{\sum_{k=1}^K P\left(\frac{X}{C_k}\right)P(C_k)} \quad (\text{III.5})$$

Ainsi, une observation X est affectée à la classe C_{k1} si elle vérifie la règle suivante :

$$P(X/C_{k1}) P(C_{k1}) > P(X/C_{k2}) P(C_{k2}) \quad \forall k_2 \in [1, K] \text{ et } k_1 \neq k_2 \quad (\text{III.6})$$

La décision Bayésienne est théoriquement optimale car elle permet de minimiser le risque d'erreurs.

4.2. Classification non supervisée (automatique) [15]

L'expression (non supervisée) fait référence au fait qu'aucun superviseur ou label est utilisé pour préciser à quelle classe appartient un individu. En conséquence, le nombre de classes existant dans un ensemble d'individus est à priori inconnu. De ce fait, l'un des problèmes les plus délicats à propos des méthodes de classification non supervisées concerne le choix du nombre de classes à retenir. Pour palier cet écueil, il existe des artifices permettant d'approcher le bon nombre de classes. Une fois le nombre de classes est choisi, le prochain pas dans le processus de classification consiste à évoluer la qualité de la partition obtenue.

La classification non supervisée (classification automatique, regroupement ou clustering en anglais) a pour but de regrouper des individus en classes homogènes en fonction de l'analyse des caractéristiques qui décrivent les individus. Par classes homogènes, on entend regrouper des individus qui se ressemblent et séparer ceux qui sont dissemblables.

Ces méthodes se subdivisent en deux catégories :

4.2.1. Méthodes statiques

L'hypothèse faite dans toute approche probabiliste de la classification non supervisée est de considérer que les données forment un échantillon aléatoire $x_1, \dots, x_N$, issu d'une population, et de s'appuyer sur l'analyse de la distribution de probabilité de cette population pour définir une classification. Différentes approches probabilistes de la classification ont été envisagées parmi lesquelles on peut distinguer les approches paramétriques et les approches non paramétriques.

4.2.1.1. Méthodes non paramétriques [16]

La classification s'apparente dans ce cas à la recherche des modes de la fonction de densité de probabilités sous-jacente à l'ensemble des observations. Comme la fdp est inconnue, on fait appel à l'estimateur du noyau de Parzen ou l'estimateur des k plus proches voisins. Étant donné que les noyaux des classes correspondent à des régions de l'espace caractérisées par une densité élevée d'observations, leurs localisations sont basées sur la détection des modes (pics) de la fdp. Afin de détecter les pics, plusieurs méthodes ont été adaptées:

- **Recherche des maxima locaux:** basée sur la notion du gradient, elle consiste à identifier les pics de la fdp (les classes) et d'assigner les observations au mode le plus proche.
- **Analyse de la convexité:** dans ce cas les modes de la fdp sont caractérisés par des domaines de l'espace des données où la fdp est concave.
- **Extraction des contours des modes:** Après filtrage de la fdp, des opérateurs différentiels multidimensionnels identiques à ceux utilisés en traitement d'images (opérateurs Sobel, Prewitt) permettent d'extraire des modes.
- **Morphologie mathématique:** Postaire et Al ont proposé d'étendre la morphologie mathématique utilisée en analyse d'images, pour la détection des modes de la fdp binaire. Une série d'opérations telle que l'érosion et la dilatation permettent de rehausser les modes de la fdp et de creuser les vallées.

4.2.1.2. Méthodes paramétriques [16]

Dans cette approche la fdp est connue a priori. On peut distinguer plusieurs types de modèles paramétriques de la classification; les plus importants étant les modèles de mélange; le problème de la classification est de déterminer les paramètres d'un mélange de fonctions de

densité de probabilité représentant les distributions des observations provenant de chacune des classes en présence dans l'échantillon analysé.

L'estimation des paramètres se fait par des techniques d'apprentissage bayésien ou par des procédures d'estimation par le maximum de vraisemblance.

Parmi ces méthodes on peut citer:

- L'analyse de mélange Gaussien dont l'objectif consiste à déterminer les paramètres de mélange à partir d'un ensemble d'observations disponibles par maximisation de la vraisemblance.
- Estimation par maximisation dont le rôle est d'estimer la moyenne et la déviation standard pour chaque classe afin de maximiser la vraisemblance des données observées. En plus l'algorithme tend à approcher les distributions observées dans les différentes classes. Les principaux inconvénients rencontrés de l'algorithme tend EM est l'initialisation des paramètres et la nécessité de connaître le nombre exacte des classes.

4.2.2. Méthodes métriques

Basées sur des notions mathématiques non probabilistes déterministes. Ces méthodes ont pour but de montrer le degré de similarité entre les objets en utilisant la notion de distance. Dans ce cas la classification porte le nom de groupement de données.

Considérons un ensemble d'objets à classifier $X = \{X_1, X_2, \dots, X_Q\}$, chaque $X_i \in R^N$ est un vecteur attributs de N mesures réelles décrivant l'objet.

Les objets sont à classifier en K groupes $\pi = \{C_1, C_2, \dots, C_k\}$, π est une partition, K est le nombre de classes.

Pour classifier les objets, on doit définir une mesure pour laquelle les objets se ressemblent. Ainsi une petite distance entre deux objets indique une grande similarité.

Plusieurs mesures de distances ont été employées, nous citerons entre autres la distance Manhattan, la distance de Mahalanobis (forme normalisée de la distance euclidienne), cependant la plus fréquente est la distance euclidienne définie pour deux objets comme suit:

Dans la classification métrique on distingue deux approches principales. La première procède par partitionnement de l'espace des observations. La seconde établit une hiérarchie de classes.

4.2.2.1. Classification hiérarchique [4]

L'approche hiérarchique consiste à construire les familles par agglomération ascendante de classes. On construit les classes par agglomération successive de deux groupes de données dont les centres de gravité sont plus proches qu'une distance d_{max} donnée.

L'algorithme de classification hiérarchique (ou agglomération) est le suivant :

1. Regrouper deux par deux les points les plus proches.
2. Tant que les groupes ne sont pas stables :
 - a. Calculer les centres de gravité des groupes ainsi créés.
 - b. Regrouper deux par deux les groupes dont les centres ont une distance inférieure à d_{max} .

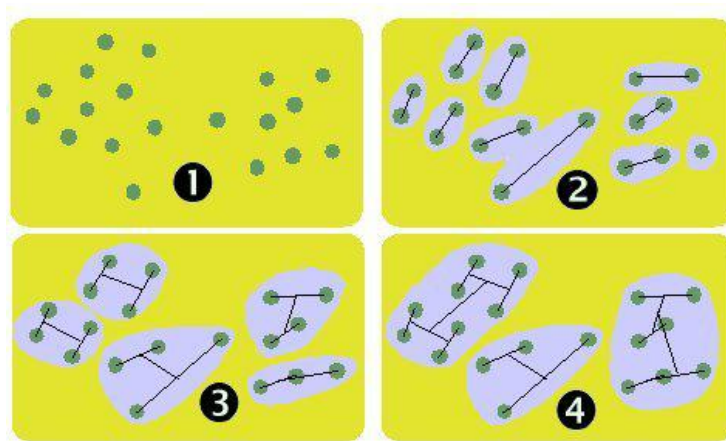


Figure III.2 : Classification hiérarchique.

La figure III.2 présente un ensemble d'agglomération de classes. L'ensemble des points (1) est traité pour regrouper les points les plus proches par classes de deux (2). Les centres des classes sont calculés et les classes regroupées par deux (3). Enfin, quand les groupes obtenus sont trop éloignés pour être fusionnés, l'algorithme s'arrête.

Cette approche nécessite beaucoup de calculs et converge vite. Par contre, elle ne sait traiter que des ensembles de données de cardinalité proche du second.

4.2.2.2. Classification non hiérarchique [17]

A l'opposé des techniques hiérarchiques qui proposent une famille de groupement dont la taille et le nombre de classes varie. Les techniques de partitionnement produisent une classification unique qui optimise un critère prédéfini ou une fonction objective. Le nombre de classes est spécifié a priori et doit être respecté le long de processus itératif.

5. Quelques méthodes de classification

Nous proposons de décrire quelques méthodes de classification supervisées et non supervisées .

5.1. Méthode des k plus proches voisins [9]

Elle est connue en anglais sous le nom K-Nearest Neighbor (KNN) (ou encore Memory Based Reasoning). C'est une méthode de classification supervisée sauf qu'elle diffère des traditionnelles méthodes d'apprentissage par le fait qu'aucun modèle n'est induit à partir des exemples. Les données restent telles qu'elles : elles sont simplement stockées en mémoire.

Pour prédire la classe d'un nouvel individu, l'algorithme cherche les k plus proches voisins de ce nouvel individu et prédit la réponse la plus fréquente de ces k plus proches voisins. La méthode utilise donc deux paramètres :

- ✓ Le nombre k.
- ✓ La fonction de similarité pour comparer le nouvel individu aux individus déjà classés.

Pour déterminer les k plus proches voisins d'un nouvel individu, nous calculons souvent la distance de cet individu aux différents individus connus.

$$d(x_i, x_j) = \sqrt{\sum_{r=1}^n (a_r(x_i) - a_r(x_j))^2} \quad (\text{III.7})$$

Avec :

$a_r(x_i)$: Attribut du nouvel individu.

$a_r(x_j)$: Attribut du $j^{\text{ème}}$ individu classé.

Avec $j = 1, 2, \dots, N$ et N : nombre d'individus classés.

r : Indice de l'attribut.

n : Nombre d'attributs.

$d(x_i, x_j)$: Distance entre le nouvel individu i et le $j^{\text{ème}}$ individu classé.

Dans KNN de base, on choisit la classe majoritairement représentée par les k plus proches voisins. Une autre solution est de pondérer la contribution de chaque k plus proche voisin en fonction de sa distance avec le nouvel individu à classer.

Noter qu'avec cette méthode de pondération, on pourrait utiliser les N exemples au lieu des k plus proches voisins : en effet, plus un exemple sera éloigné du nouvel individu à classer et moins sa classe contribuera au résultat final.

Les expériences menées avec les KNN montrent qu'ils résistent bien aux données bruitées. Par contre, ils requièrent de nombreux exemples classés.

Noter aussi que, si le temps d'apprentissage est inexistant puisque les données sont stockées telles qu'elles. La classification d'un nouvel individu est par contre coûteuse puisqu'il faut comparer ce cas à tous les exemples déjà classés.

La figure III.3 illustre un exemple de classification par KPV. On doit classer le nouvel individu X_q . Si on choisit $k = 1$, X_q sera classé +. Si $k = 5$, le même X_q sera classé -. On voit donc que le choix de k est très important dans le résultat final.

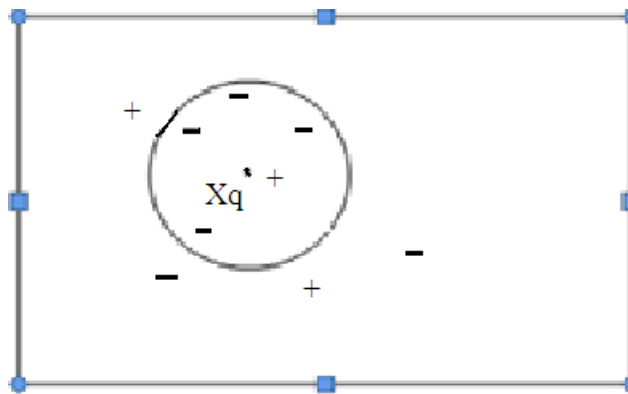


Figure III.3 : Classification par les k plus proches voisins

5.2. L'approche densité des informations [2]

Les méthodes basées sur la densité sont des méthodes de classification non supervisées, dans lesquelles on cherche à regrouper l'ensemble des points situés dans un voisinage d'un point non classé selon une certaine distance. Pour atteindre ce but, on définit une distance maximale entre deux points (voisinage) ainsi qu'un nombre de voisins minimal (densité). Ensuite, on procède au classement des données selon l'algorithme suivant :

1. Choisir au hasard un point non classé.
2. Former un groupe contenant ce point et les points de son voisinage dense.
3. Chercher le voisinage dense de chaque point du groupe et l'ajouter au groupe.
4. Tant que le groupe n'est pas stable, aller à 3.
5. S'il y a des points non classés, aller à 1 sinon finir.

La figure 3.4 illustre un ensemble de classification par l'approche densité des informations.

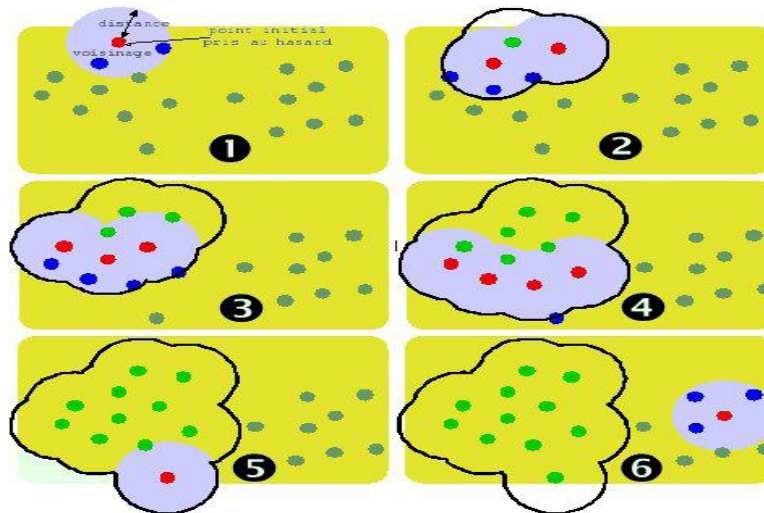


Figure III.4 : Classification à l'aide des voisinages denses.

Dans cette figure (1), on choisit au hasard un point non classé. Ce point et son voisinage sont classés dans le même groupe. En (2), on prend le voisinage de chacun des points du groupe (2) et on les classe dans le même groupe ainsi de suite (3), (4) et (5) jusqu'à la possibilité du groupe obtenu en (6). On recommence ensuite avec un autre point de départ non classé (6) et réitère le processus jusqu'à ce que tous les points soient classés.

Cette méthode est précise pour trouver le nombre de classes présentes dans les données sans indication préalable. Cependant, elle souffre de certains inconvénients : elle nécessite de nombreux calculs ; elle est très sensible au choix de la distance et du nombre minimum de voisins.

5.3. Méthode des k-means [4]

La méthode de classification non hiérarchique ou de partitionnement a pour but de fournir une partition d'éléments à classer. Le nombre de classes dans la partition à générer doit être fixé au départ. Le principe des méthodes de partitionnement est le suivant :

- ✓ A partir d'un ensemble E de N individus, on cherche à construire une partition $P = (C_1 \dots C_k)$ contenant k classes homogènes et distinctes, tout en respectant un critère basé sur une distance entre les individus.
- ✓ Doivent se retrouver dans une même classe les individus qui se ressemblent et en classes différentes, les individus divergents en termes de critère adopté.
- ✓ Une partition optimale peut être obtenue à partir de l'énumération exhaustive de toutes les partitions, ce qui devient cependant prohibitif en termes de temps de calcul.

Comme solution alternative à ce problème, des méthodes de partitionnement basées sur l'optimisation itérative du critère à respecter permettent d'obtenir des groupes bien distincts

en un temps de calcul raisonnable. Ces méthodes d'optimisation utilisent une réaffectation afin de redistribuer itérativement les individus dans les k classes.

Parmi les nombreuses méthodes de classification non hiérarchiques, la méthode des k-means est largement utilisée car elle est inspirée de la stratégie de regroupement de la vision humaine. Cette méthode offre une certaine similitude avec la méthode humaine par tâtonnement : on classe d'abord les individus en paquets grossiers, puis on cherche à affiner le résultat du classement et tendent vers une solution idéale.

L'algorithme k-means est le suivant :

1. Choisir aléatoirement k centres de classes.
2. Affecter les individus aux classes dont la distance est minimale.
3. Calculer les nouveaux centres de gravité des classes.
4. Affecter chaque individu à la nouvelle classe.
5. Tant que le critère d'arrêt n'est pas vérifié, aller à 3.

Le centre de gravité d'une classe est calculé comme suit :

$$g_i = \frac{1}{N} \sum_{j=1}^n X_{ij} \quad (\text{III.8})$$

j : indice d'attribut.

X_{ij} : attribut j de l'individu x_i .

On choisit souvent comme critère d'arrêt, la stabilité des classes. Ce critère peut être déterminé par plusieurs techniques. Parmi ces techniques, nous avons :

- ✓ Lorsque les centres des classes ne changent pas durant deux itérations successives ; c'est-à-dire lorsque $g(n, p) = g(n, p+1)$

$$n = 1, 2 \dots k$$

Avec $g(n, p)$ est le centre de gravité de la classe n à la $p^{\text{ème}}$ itération.

La figure III.5 illustre un exemple de classification par la méthode des k-means.

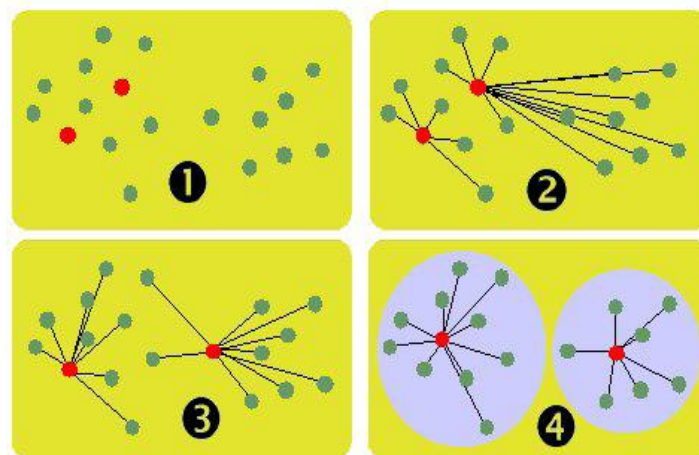


Figure III.5 : Classification par *k-means*.

L'idée est de placer au hasard un nombre de centres de classes correspondant au nombre de classes souhaitées. Pour que l'algorithme donne de bons résultats, il est préférable d'avoir une connaissance à priori du nombre de classes présentes dans les données à regrouper. La classification est automatiquement obtenue par le déplacement itératif des centres de classes.

Cette technique n'implique pas de longs calculs et converge rapidement. Par contre, elle impose de donner le nombre de classes désirées au début de l'algorithme et elle est sensible aux conditions initiales.

5.4.Méthode Fuzzy C-Means (FCM) [18]

Fuzzy C-Means (FCM) est un algorithme de classification non supervisée floue. Issu de l'algorithme des *C-moyennes* (*C-means*), il introduit la notion d'ensemble flou dans la définition des classes : chaque point dans l'ensemble des données appartient à chaque cluster avec un certain degré, et tous les clusters sont caractérisés par leur centre de gravité. Comme les autres algorithmes de classification non supervisée, il utilise un critère de minimisation des distances intra-classe et de maximisation des distances inter-classe, mais en donnant un certain degré d'appartenance à chaque classe pour chaque pixel. Cet algorithme nécessite la connaissance préalable du nombre de clusters et génère les classes par un processus itératif en minimisant une fonction objective. Ainsi, il permet d'obtenir une partition floue de l'image en donnant à chaque pixel un degré d'appartenance (compris entre 0 et 1) à une classe donnée. Le cluster auquel est associé un pixel est celui dont le degré d'appartenance sera le plus élevé.

Les principales étapes de l'algorithme Fuzzy C-means sont :

1. La fixation arbitraire d'une matrice d'appartenance.

2. Le calcul des centroïdes des classes.
3. Le réajustement de la matrice d'appartenance suivant la position des centroïdes.
4. Calcul du critère de minimisation et retour à l'étape 2 s'il y a non convergence de critère.

L'inconvénient de cette méthode réside principalement dans l'initialisation des centres, et sa sensibilité au bruit. De plus, le nombre de classes doit être connu à priori. Pour ces raisons, d'autres variantes que nous décrirons dans ce chapitre ont été proposées afin d'améliorer ce dernier.

6. Conclusion

Différentes méthodes de classification ont été présentées dans ce chapitre. Nous avons insisté particulièrement sur l'approche classification non supervisée, non probabiliste et non hiérarchique. Ces méthodes métriques représentent l'avantage d'être rapides et simples, mais leur inconvénient principal est que la solution obtenue dépend des centres initiaux et aboutissent souvent à un optimum local.

1. Introduction:

Le but de ce travail est de réaliser une segmentation d'image médicale par la méthode de classification non supervisée k-means. Il existe plusieurs méthodes de classification qui donnent de bons résultats mais on a choisi la méthode de k-means vu sa fiabilité. Nous avons appliqué cette méthode sur une image médicale IRM en changeant à chaque fois le nombre de classes.

2. Image originale:



Figure IV.1: Image originale.

3. Les différents résultats obtenus par k-means:

- Résultats pour deux classes:

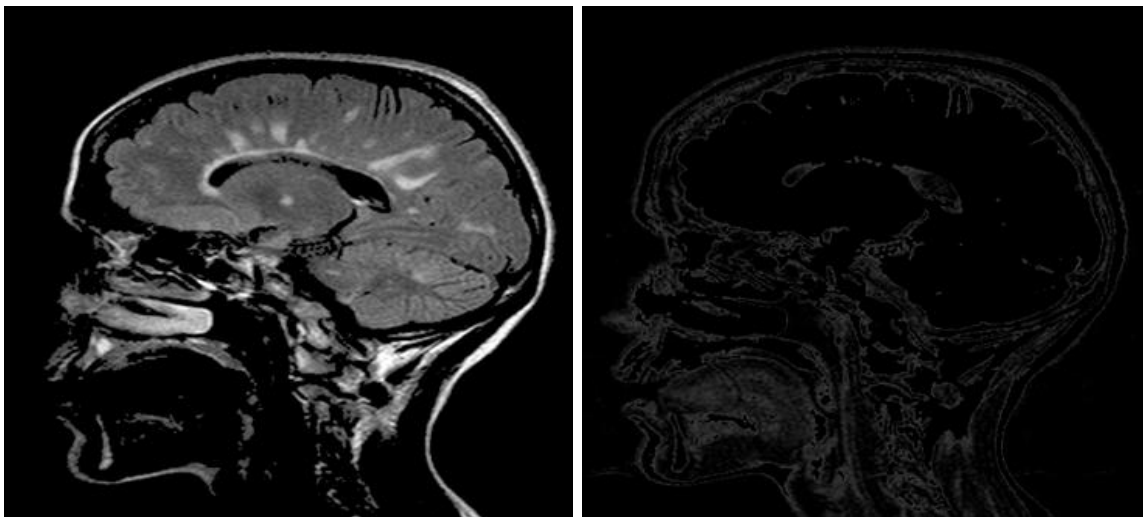
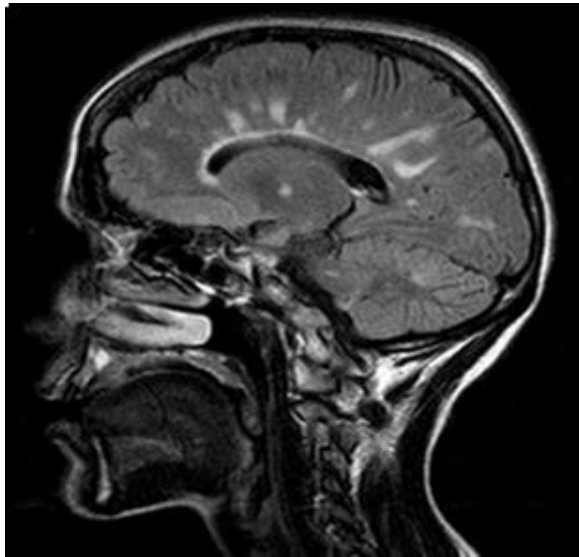


Figure IV.2: Labellisation des deux régions.



a) Image originale



b) Image segmentée par deux classes

Figure IV.3:Segmentation avec deux classes.

Cette image représente une image segmentée en considérant le nombre de classes est égal à deux. Nous remarquons que le nombre de classe choisi n'est pas assez suffisant pour extraire les différentes régions constituant l'image.

➤ Résultats pour trois classes:

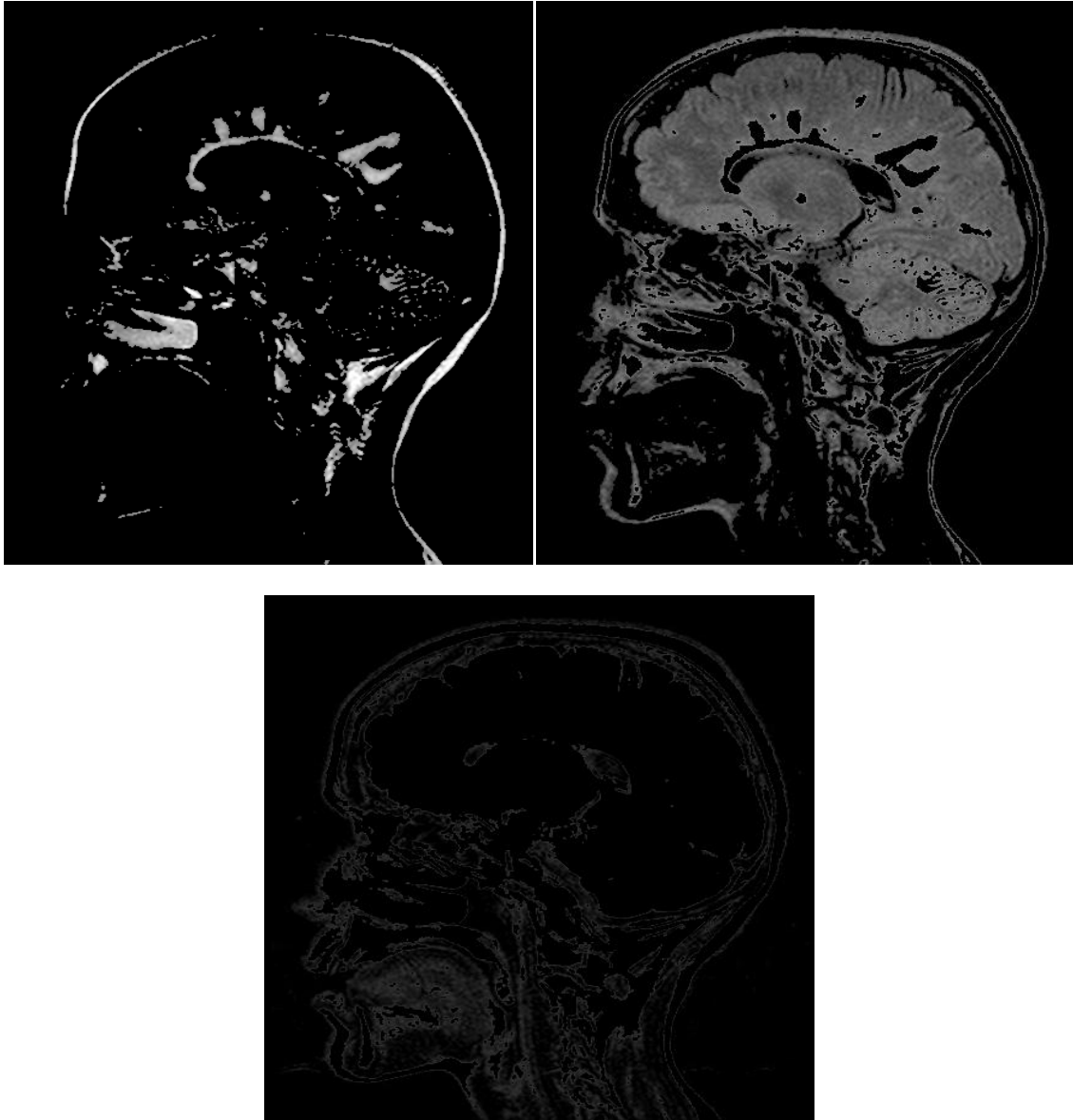
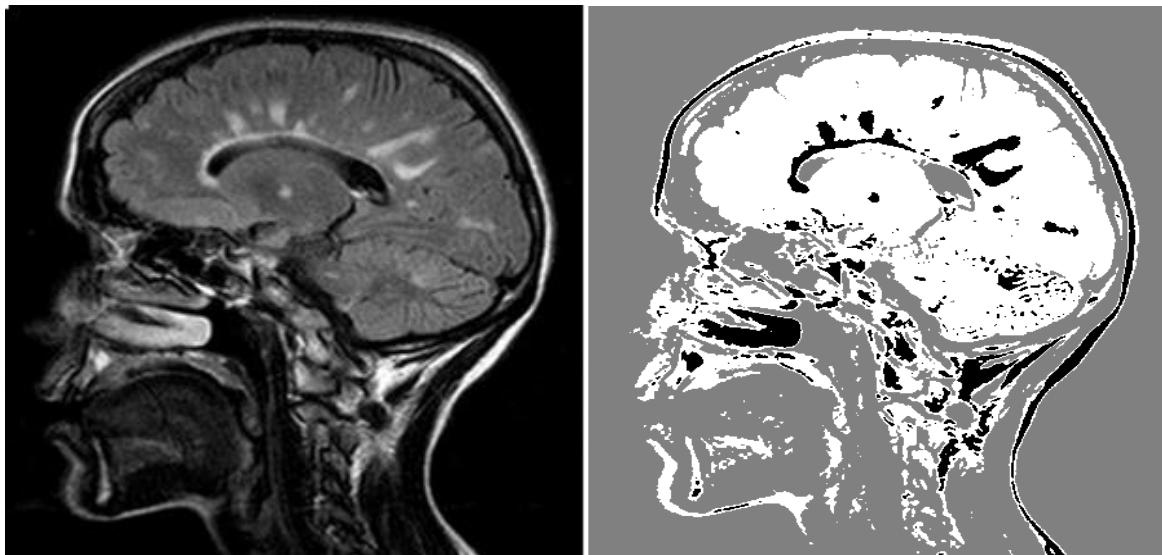


Figure IV.4: Labellisation des trois régions.



a) Image originale

b) Image segmentée par deux classes

Figure IV.5: Image segmentée avec trois classes.

➤ *Résultats pour quatre classes:*

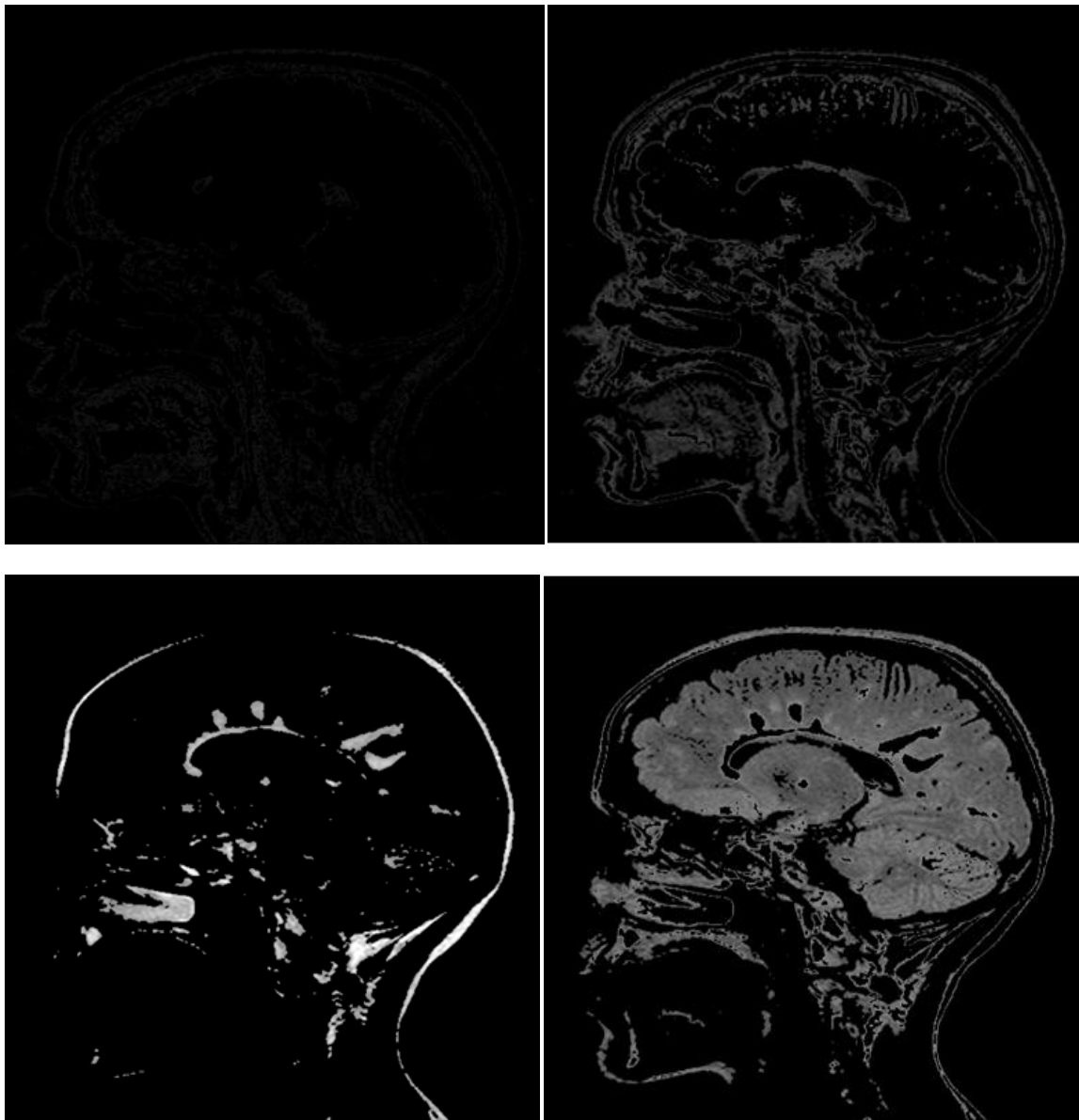
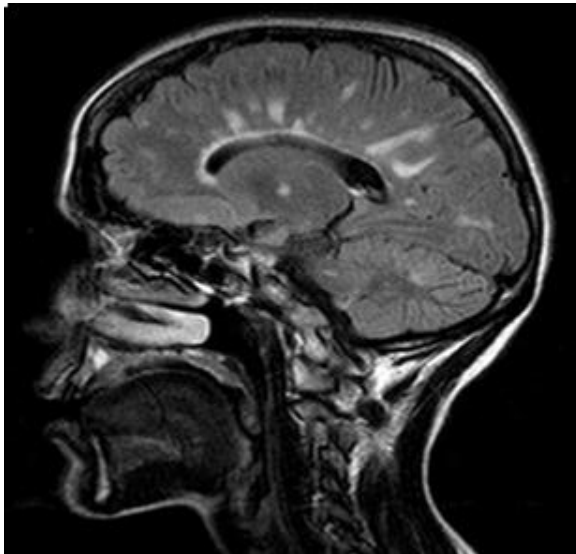


Figure IV.6: *Labellisation des régions avec quatre régions.*



a) *Image originale*



b) *Image segmentée par quatre régions*

Figure IV.7: *Image segmentée avec quatre classes.*

➤ Résultats pour 7 classes:

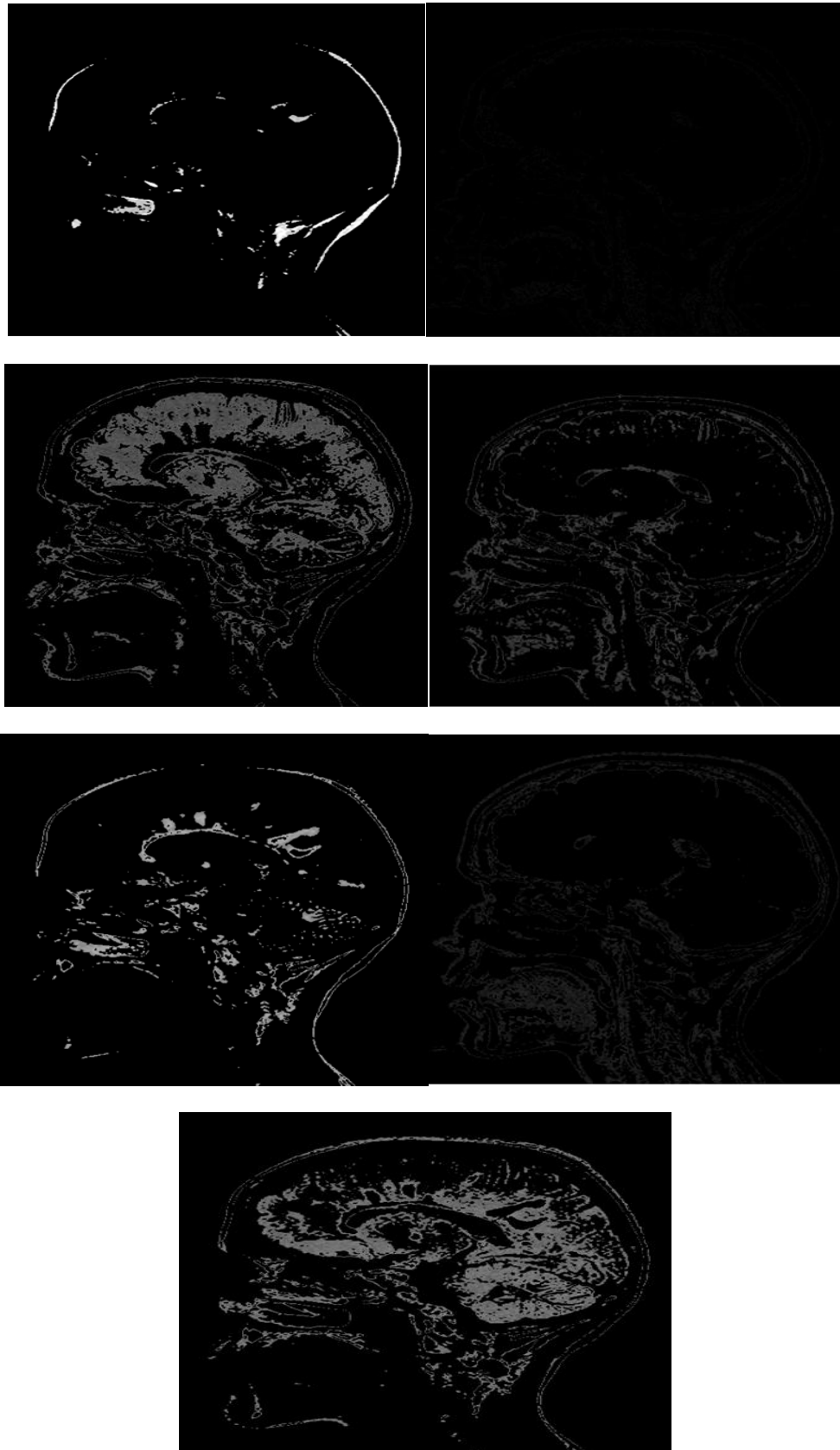
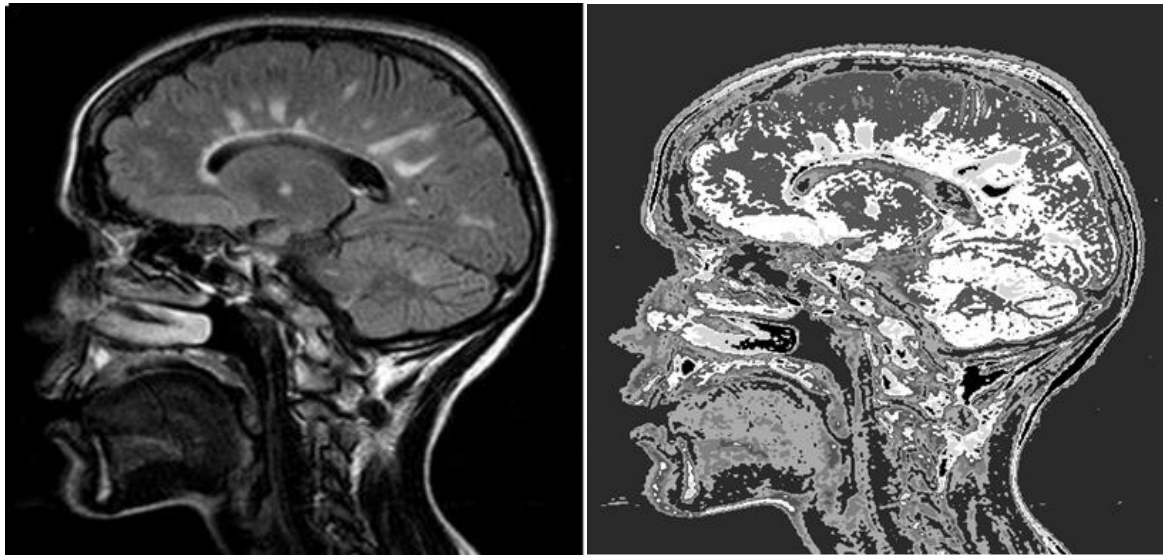


Figure IV.8: Labellisation des régions des sept régions.



a) Image originale

b) Image segmentée par sept régions

Figure IV.9: Image segmentée avec sept classes.

4. Discussion des résultats

Pour évaluer les performances de la méthode de segmentation de classification non supervisée choisie, nous l'avons appliqué sur une image test médicale en changeant à chaque fois le nombre de classes.

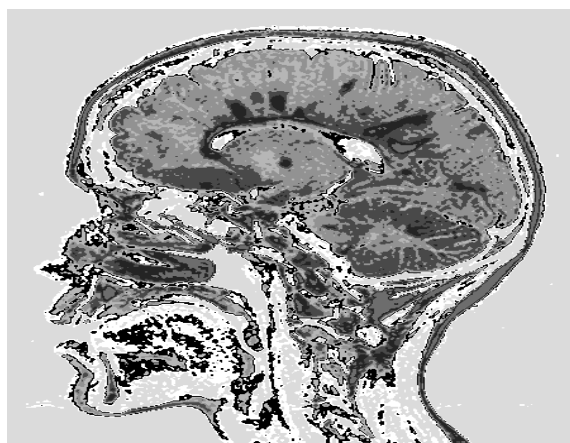
Les figures IV.2, IV.4, IV.6 et IV.8: représentent la labellisation des différentes régions constituant l'image pour les nombres de classes deux, trois, quatre et sept successivement.

Les figures IV.3, IV.5, IV.7 et IV.9: représentent les résultats de la segmentation de l'image médicale obtenue par IRM pour les nombres de classes 2, 3, 4 et 7 successivement.

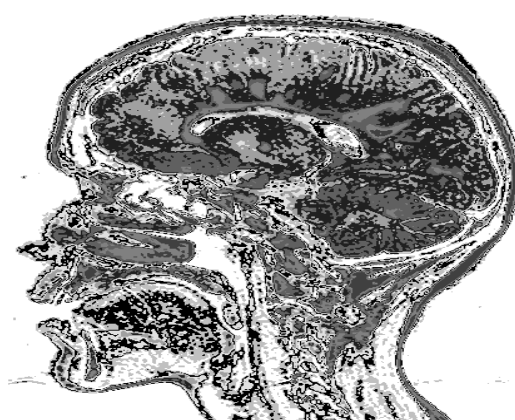
Nous remarquons que la segmentation se diffère d'une figure à l'autre en fonction du nombre de classes.

A chaque fois qu'on augmente le nombre de classes, il y'a de nouvelles régions qui apparaissent représentant de nouveaux tissus composant l'organe.

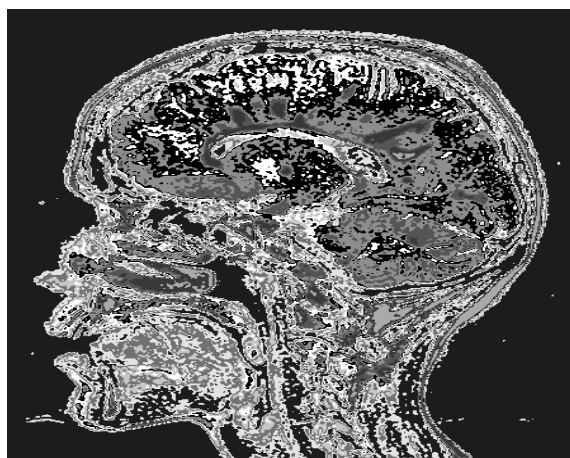
5. Quelques résultats pour d'autres nombres de classes



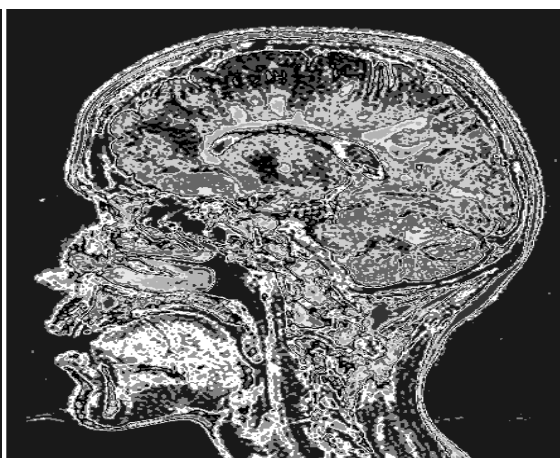
a) Image segmentée avec huit classes.



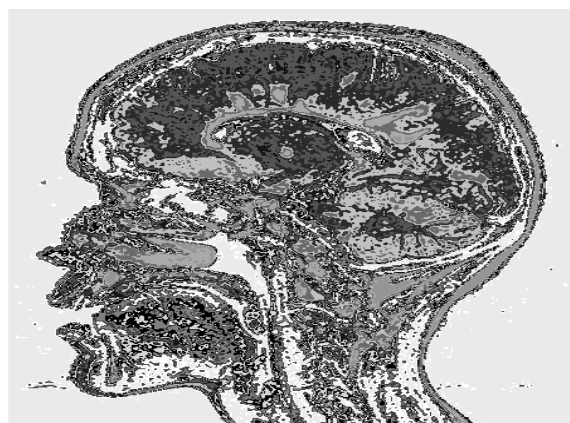
b) Image segmentée avec neuf classes.



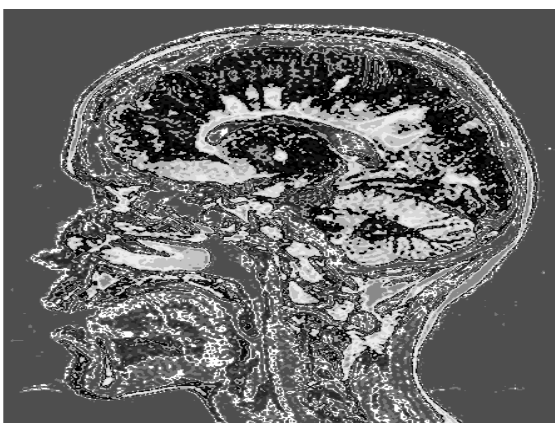
c) Image segmentée avec dix classes



d) Image segmentée avec onze classes



e) Image segmentée avec treize classes



f) Image segmentée avec quatorze classes

Figure IV.10: exemples d'image segmentée avec différents nombres de classes

6. conclusion

En observant les résultats affichés sur les figures précédentes, nous constatons qu'à chaque fois que nous augmentons le nombre de classes, il y'a de nouveaux tissus qui apparaissent; ce qui montre l'importance et l'influence du choix du nombre de classes sur la méthode de classification choisie. Le choix du nombre de classes dépend fortement du besoin de l'utilisateur et des régions qu'il veut extraire et traiter.

Vu les résultats obtenus, nous déduisons que la performance de la segmentation dépend de nombre de classes ainsi que la texture et les attributs pertinents.

Conclusion générale

Le travail effectué tout au long de ce mémoire nous a permis de se familiariser avec les techniques de traitement d'images. Nous nous sommes intéressés à la segmentation qui est une étape indispensable dans tout processus d'analyse et de traitement d'images.

Nous avons d'abord abordé les méthodes de classification supervisées et non supervisées. Ensuite, nous avons testé quelques images par la méthode de K.means, qui nous semble très utilisée vu sa fiabilité. Les tests effectués nous ont permis de constater que le choix du nombre de classe joue un rôle très important dans la segmentation.

Nous concluons que les images médicales utilisées rendent la segmentation plus compliquée, vu la complicité de leurs texture très variée, à comparer aux images de Brodatz. Nous espérons que ce travail apportera un plus aux travaux déjà fait

ENVIRONNEMENT MATLAB

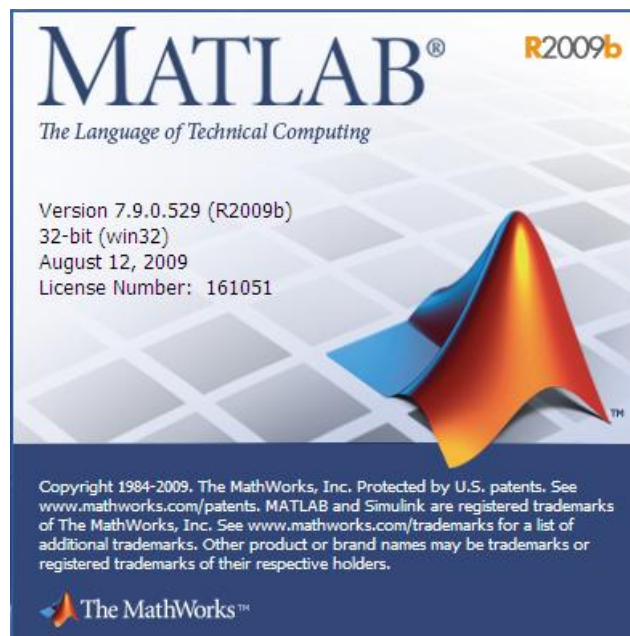
1. MATLAB

MATLAB (MATrix LABoratory) est un logiciel de calcul scientifique dédié plus particulièrement aux applications numériques. A l'origine, il a été conçu pour manipuler des données matricielles, ce qui a fait de lui un outil majeur d'analyse de données, du traitement de signal, du traitement d'images, de simulation numérique,...etc. Il possède son propre langage de programmation avec de nombreuses fonctions fournies.

2. Prise en main du logiciel MATLAB

2.1. Lancement

Pour démarrer MATLAB, il suffit de cliquer dans l'icône « Matlab » si vous êtes sous windows, ou de taper la commande Matlab si vous êtes sous Unix. Ou bien, on clique sur l'icône Matlab dans le menu démarrer.



L'interface utilisateur affiche les fenêtres suivantes :

a. Fenêtre de commande :

Dans cette fenêtre, l'utilisateur donne les instructions et Matlab retourne les résultats.

L'espace de travail de Matlab se présente alors sous la forme d'une fenêtre affichant un prompt ">>" à la suite duquel vous pouvez taper une commande qui sera exécutée après avoir tapé sur la touche "**return**".

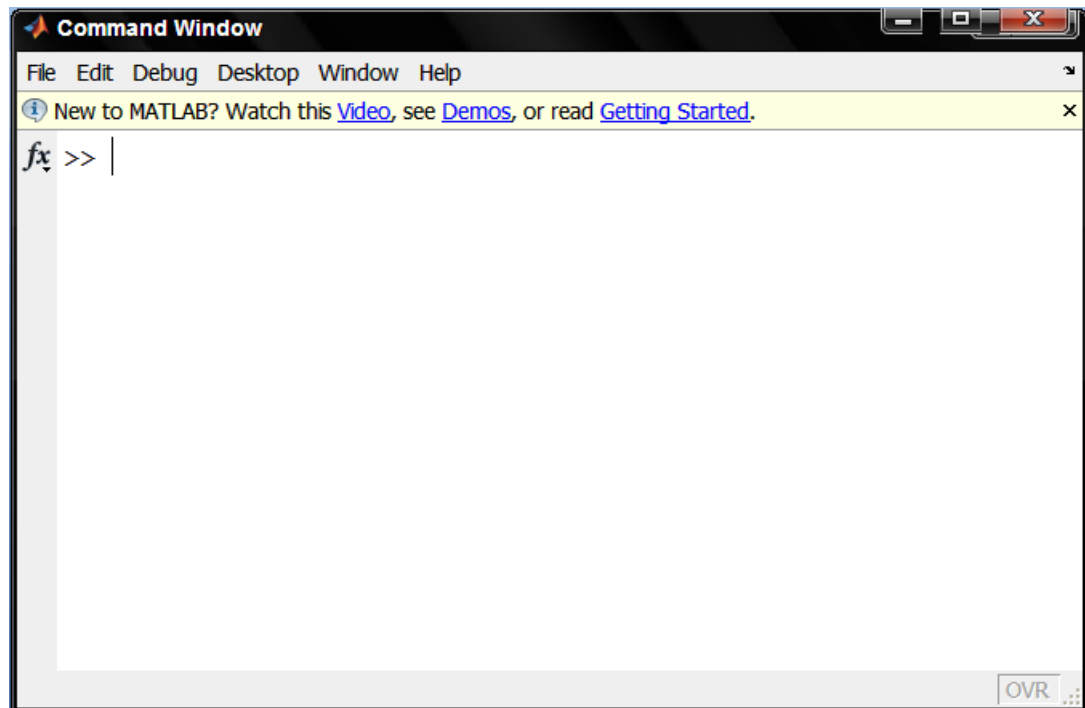
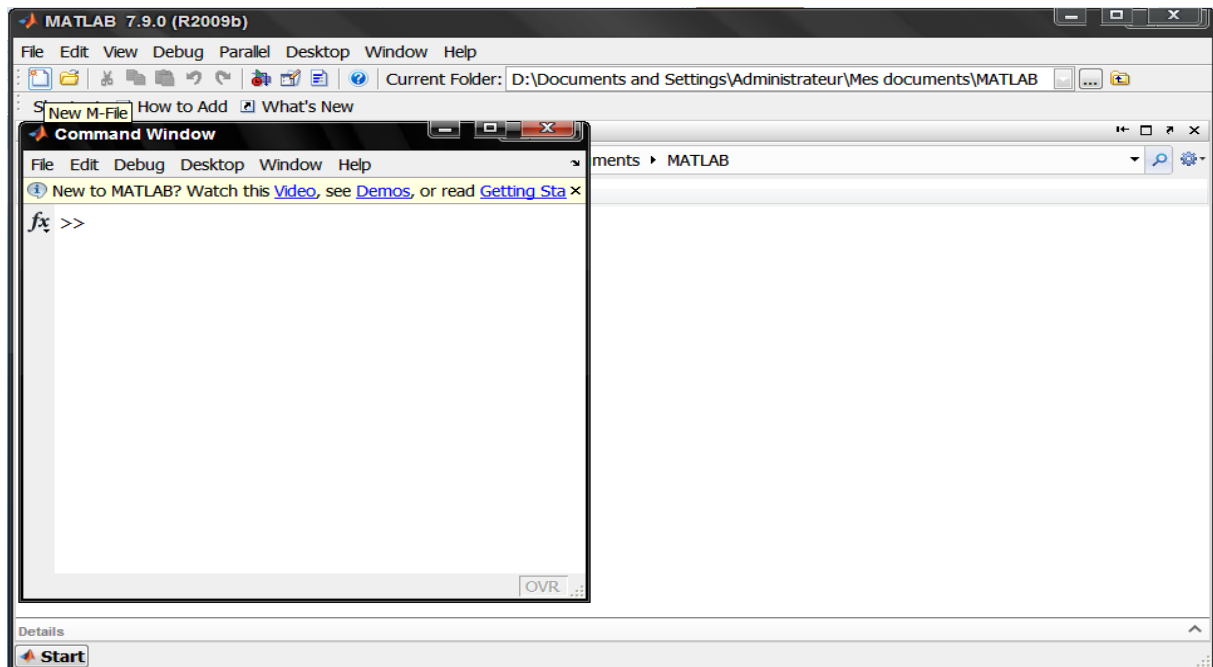


Figure1 : Fenêtre de commande.

En haut de cette fenêtre se trouve une barre de menu qui vous permet d'ouvrir un fichier texte, de définir certaines variables de travail et surtout d'accéder à l'ensemble des fichiers d'aides.

b. Editor (fichier m)

Ce sont les programmes écrits par l'utilisateur en langage Matlab. Matlab exécute ligne par ligne un fichier 'm' (programme en langage Matlab).



On clique sur « **New M-file** », ceci fait apparaître une nouvelle fenêtre (« **Editor** ») dans laquelle on éditera le texte du programme.

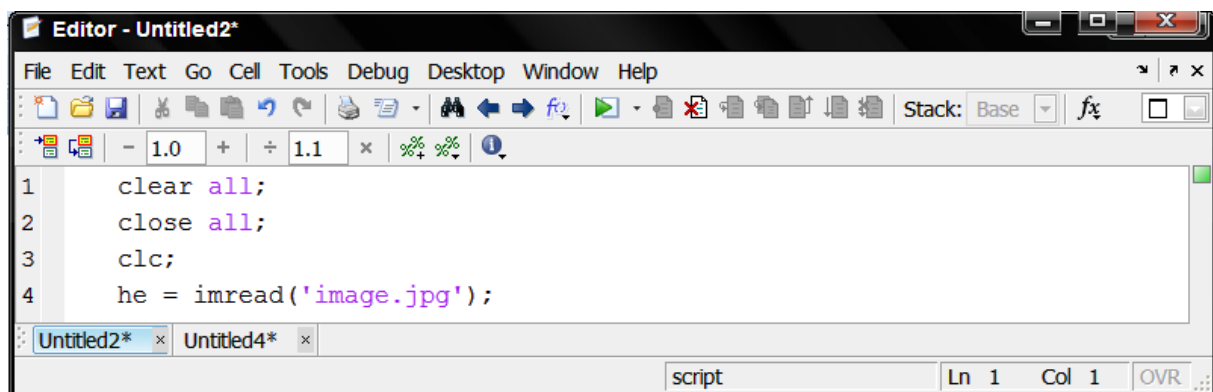


Figure2 : Fenêtre Editor

Il est important de s'habituer à la gestion des fichiers : il faut savoir dans quel dossier on range les fichiers pour pouvoir les retrouver, les modifier éventuellement et les exécuter par la suite.

Une fois qu'on a écrit un programme, il faut le sauvegarder. On choisit par exemple :

« **save as** » et le nom du fichier qui doit nécessairement avoir l'extension « **.m** » :

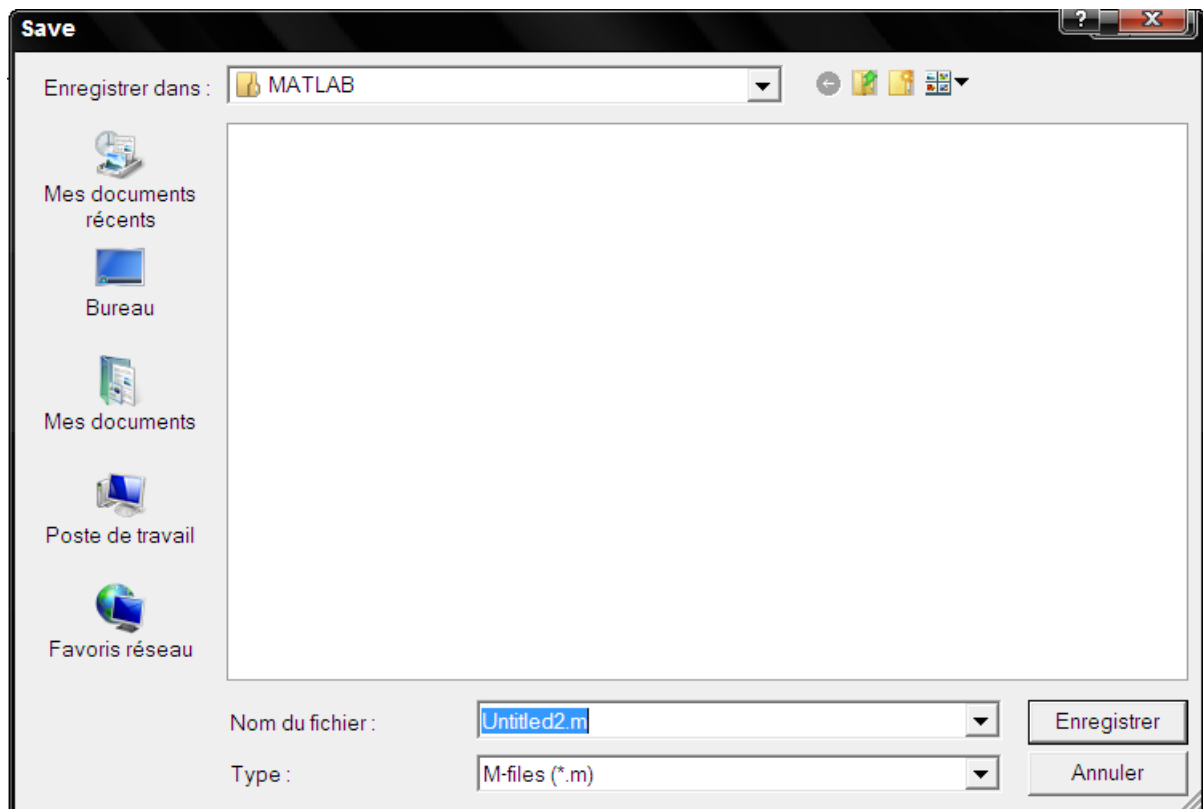


Figure3 : Fenêtre de sauvegarde d'un fichier dans le répertoire Matlab

c. Help

Permet d'obtenir de l'aide sur une commande.

- **help** : permet d'obtenir de l'aide et donne une liste thématique ;
- **help nom de fonction** : donne la définition de la fonction désignée et des exemples d'utilisation ;
- **help elfun**
- **lookfor sujet** : donne une liste de rubriques de l'aide en relation avec le sujet indiqué

Soit en allant dans le menu **help** → **matlab help**, une fenêtre s'ouvre partagée en deux. A gauche, on peut en cliquant sur signet supérieur, activer

- **Contents**
- **Search**
- **Index**

➤ Demos

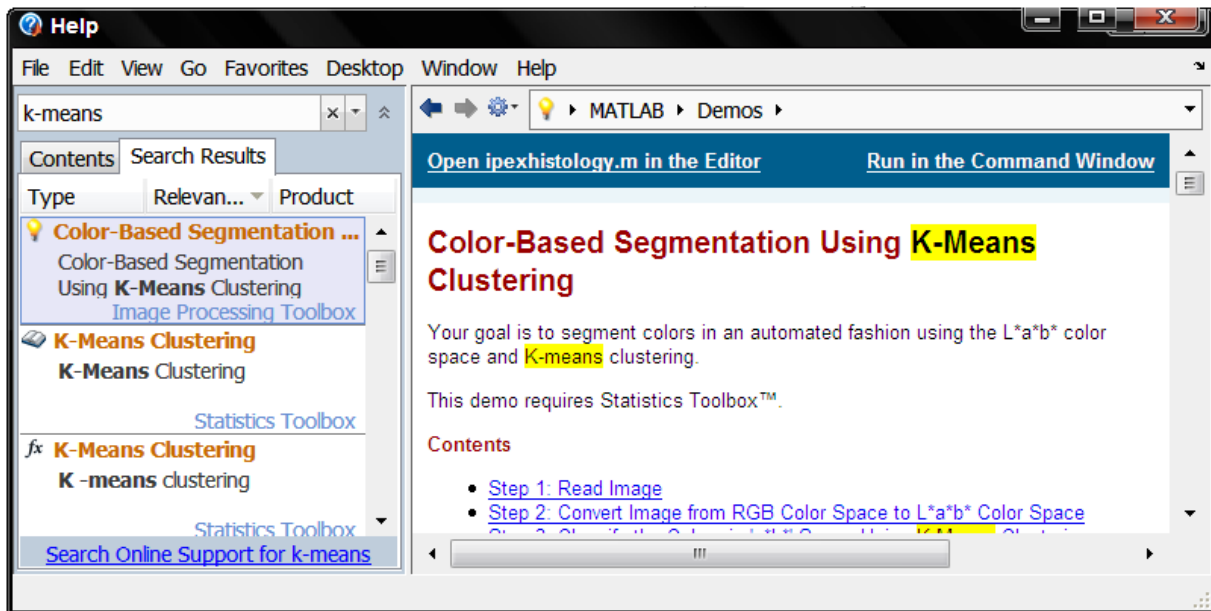


Figure4 : fenêtre help

3. Toolboxes

Cet environnement est complété par une bibliothèque de fonctions prédéfinies, très souvent sous forme de fichiers m-file, et regroupés en paquets ou **toolboxes**. A côté des toolboxes requises **local** et **matlab**, il est possible d'ajouter des toolboxes spécifiques à tel ou tel problème mathématique, *Optimization Toolbox*, *Signal Processing Toolbox* par exemple ou encore des toolboxes créés par l'utilisateur lui-même. Un système de chemin d'accès ou path permet de préciser la liste des répertoires dans lesquels MATLAB trouvera les différents fichiers m-files.

Bibliographie:

- [1] M^{me} Z. Ameur, « Codage des images en vue d'une segmentation de haut niveau, application aux images satellitaires», Thèse de doctorat, UMMTO.
- [2] H.MEDROUK et S. TOUCHERIFT, «Segmentation par classification d'une partie de l'image», Mémoire d'ingénieur en automatique, UMMTO, 2010.
- [3] K. DJOUDREZ et H. LARFI, «Segmentation d'image médicales basée sur la classification floue», Mémoire d'ingénieur en automatique, UMMTO, 2010.
- [4] Jérôme LANDRÉ, « Analyse multi-résolution pour la recherche l'indexation d'image par le contenu dans les bases images application à la base d'image paléontologique », Thèse de doctorat, Université de Bourgogne, 2006.
- [5] Sié. OUATTARA, «Stratégie de segmentation d'images multi-composantes par analyse d'histogrammes multidimensionnels, application à des images couleurs de coupes histologiques de pommes», Thèse de doctorat, Ecole doctorale STIM, 2009.
- [6] Jean- Loïc ROSE, « Croissance de région variationnelle et contraintes géométriques tridimensionnelles pour la segmentation d'images », Thèse de doctorat, l'Institut National des Sciences appliquées de Lyon, 2008.
- [7] N. RABIA, « Revue des méthodes de segmentation des images texturées: cas des images couleurs », Mémoire de magister, UMMTO, 2011.
- [8] M.LEHAMEL, « Segmentation d'image texturée à partir des attributs fractals », Mémoire de magister en électronique, UMMTO.
- [9] D. AMAZIT, « Segmentation d'images texturées par approche statistique», Thèse de magister, UMMTO, 2005.
- [10] BELKACEM et BOUBRIT, « Segmentation d'images par analyse de la texture», Mémoire d'ingénieur, UMMTO, 1998
- [11] J.P. COCQUEREZ et Sylvie Philippe, « Analyse d'images: Filtrage et Segmentation», Masson, 1995.

- [12] S. ATMANI et M.HALLI, « Extraxion rapide des attributs de texture basés sur les matrices de cooccurrence », Mémoire d'ingénieur en électronique, UMMTO, 2006.
- [13] M.DEHBI et F. CHERIFI, « Segmentation d'images texturées par les statistiques de rangs des niveaux de gris », UMMTO, 1999.
- [14] Sylvain PASIN et Bertrand GRANDGEORGE, « Image segmentation », Mémoire d'ingénieur, Ecole polytechnique fédérale de Lausanne, 2003.
- [15] Pierre Gançarski, « Classification non supervisée », Thèse de doctorat, Université de Strasbourg, 2004
- [16] S.BAZIZ, « Algorithme génétique pour la classification des données par la méthode des K-Means », Mémoire d'ingénieur en électronique, UMMTO, 2003.
- [17] H.LAGHEL, « Déploiement sur une plateforme de visualisation, d'un algorithme coopératif pour la segmentation d'images IRM basé sur les systèmes multi-agents », Mémoire d'ingénieur, USTHB, 2010.
- [18] Lotfi KHODJA, « Contribution à la classification floue non supervisée », Thèse de doctorat, Université de SAVOIE.
- [19] REZKI, SMAIL, HAHAD, « Application de la logique floue à la segmentation d'images couleur », Mémoire d'ingénieur en électronique, UMMTO, 2008.
- [20] Antoine SIMON, « Extraction et caractérisation du mouvement cardiaque en imagerie scanner multi-barrette ». Thèse de doctorat, Université de Rennes 1, 2005.
- [21] Benoît CHANTEPIE, « Étude et réalisation d'une électronique rapide à bas bruit pour un détecteur de rayons X à pixels hybrides destiné à l'imagerie du petit animal », Thèse de doctorat, Université de la méditerranée Aix- Marseille **II**, 2008.
- [22] BEKKAOUI Khalid et GRANDGEORGE Bertrand, « Développement d'une extension pour image; Segmentation d'image par classification de pixels », ..., Université des Sciences et Technologies de Lille, 2008
- [23] Fabrice Muhlenbach, « Méthode de clustering par graphe de voisinage », Thèse de doctorat, Université Lyon2, 2008.

[24].

[25] Weibi DOU, «Segmentation d'image multispectrales basée sur la fusion d'information: application aux images IRM», Thèse de doctorat, Université de CAEN, 2002.

[26] Jérémy Lecoeur et Christian Barillot, «Segmentation d'images cérébrales: Etat de l'art», Rapport de recherche, Institut National de Recherche en Informatique et en Automatique, 2008.

[27] ZOUAOUI Hakima, «Clustering par fusion floue de données appliquées à la segmentation d'image IRM», Mémoire de magister en électronique, Université M'hamed BOUGARA de Boumerdès, 2008.

[28] Henry Maitre – Le traitement des images, Hermes-science, Lavoisier, Paris, 2003

.

