

Table des matières

Introduction.....	1
Chapitre I : Etat de l'art.....	4
Introduction.....	4
I.1 Les méthodes empiriques.....	5
I.2 Les modèles déformables.....	6
I.3 Les méthodes statistiques.....	6
Chapitre II : Méthodologie.....	10
II.1 Architecture et représentation de la méthode	10
II.2 Le prétraitement	12
II.2.1 Modification d'histogramme.....	12
II.2.2 L'égalisation d'histogramme.....	12
II.2.3 L'égalisation d'histogramme adaptative.....	13
II.3 L'extraction des caractéristiques.....	14
II.3.1 La représentation par la transformée d'ondelettes.....	14
II.3.1.1 Introduction à la théorie d'ondelettes.....	15
II.3.1.2 L'analyse en ondelettes.....	15
II.3.1.3 Le choix d'ondelettes.....	18
II.3.1.4 Ondelettes de Harr.....	18
II.3.1.5 Représentation par ondelettes proposée.....	19
II.3.2 La représentation par les moments invariants de HU.....	21
II.3.2.1 Les moments géométriques.....	21
II.3.2.2 Théorie des moments.....	22
II.4 Classification SVM	24
II.4.1 Notions sur l'apprentissage statistique	24
II.4.1.1 Généralités	24
II.4.1.2 Formulation d'un problème de classification.....	25

II.4.1.3 Minimisation du risque structurel.....	25
II.4.2 Théorie des machines à vecteurs de support	27
II.4.2.1 Introduction.....	27
II.4.2.2 Principe des machines à vecteurs de support.....	27
II.4.2.3 Cas des données linéairement séparables.....	27
II.4.2.4 Cas des données non linéairement séparables.....	30
II.4.3 Architecture d'un classificateur SVM.....	30
II.4.3.1 La fonction noyau	30
II.4.3.2 Sélection de modèle SVM.....	31
II.4.3.3 Estimation de l'erreur de généralisation.....	32
II.5 Détection par SVM.....	33
II.5.1 Coefficient de corrélation.....	35
II.6 Estimation par RANSAC.....	35
II.6.1 L'algorithme RANSAC.....	36
II.6.2 Sélection des paramètres.....	37
Chapitre III. Résultats Expérimentaux :.....	39
III.1 Base d'image	39
III.2 Etape de classification.....	39
III.3 Etape de test et validation.....	42
III.4 Etape de détection.....	44
III.5 Raffinement des résultats de détection.....	47
III.5.1 Elimination des fausses détections.....	47
III.5.2 Recouvrement des vertèbres manquantes.....	49
Conclusion.....	52
Perspectives et travaux futur.....	53
Bibliographie	54



Introduction

1.1 Introduction :

Les techniques d'imagerie sont largement utilisées en médecine. Ils représentent en effet un outil très important dans de nombreux domaines, tels que la planification de chirurgie, la simulation, la planification de la radiothérapie, et le suivi de la progression des maladies, etc.

Beaucoup de travaux de recherche ont été réalisés en imagerie médicale assistée par ordinateur. Par exemple, dans le domaine de la neuro-intervention guidée par l'image, les images IRM sont analysées afin de prévoir les traitements des anévrismes cérébraux.

Pour l'analyse et la détection du cancer, les images à rayons-X, résonance magnétique (RM), et l'échographie, sont analysées afin de fournir une détection précoce, avec un suivi et un traitement du précis.

Dans le domaine de l'imagerie cardiaque, les images IRM et échographiques sont analysées afin d'obtenir des informations variables dans le temps pour l'évaluation de la perfusion tissulaire.

Ces dernières années, les systèmes d'aide au diagnostic basés sur le traitement de l'image et des techniques d'intelligence artificielle ont prouvé leur efficacité à fournir des suggestions constructives pour le radiologiste afin de l'aider dans la phase de prise de décision.

Dans ce contexte, nous proposons dans ce travail une méthode de détection automatique des régions de vertèbres dans les images à rayons-X, présentant une aide au radiologiste dans la phase de sélection des vertèbres.

1.2 Problématique et motivation :

Dans une analyse assistée par ordinateur, les objets d'intérêt doivent être isolés dans l'image. Ainsi l'extraction et la segmentation de ces objets est la première étape à établir, dans ce genre d'applications.

L'efficacité des techniques de segmentation des vertèbres dans les images à rayons-X de la colonne vertébrale est d'une importance primordiale dans l'évaluation des anomalies de la colonne vertébrale (présence de ruptures, ostéophytes antérieurs, le rétrécissement de l'espace disque intervertébrale, présence de la subluxation, présence de scoliose, etc.) pour la

radiologie de morphométrie osseuse. Cependant, une telle segmentation est difficile en raison de la structure articulée des vertèbres et de leur contact dense avec les côtes et d'autres organes, en particulier dans la colonne thoracique. Cela implique une segmentation manuelle avec certaines estimations initiales par des experts pour délimiter les différentes régions des vertèbres à analyser.

En général, la qualité des algorithmes de segmentation est affectée par trois facteurs importants :

1. La variabilité de la taille, la forme et l'orientation des vertèbres.
2. Les images à rayons-X sont généralement de mauvaise qualité, et les méthodes de segmentation confondent souvent les bords et le tissu et de la vertèbre.
3. La grande résolution de l'image: par exemple nous pouvons avoir des radiographies digitales ayant une taille originale pour le rachis cervical de 1462 par 1755, et de 2048 par 2488 pour les lombaires.

L'inconvénient majeur de ces méthodes est qu'elles s'appuient sur un jugement subjectif qui limite leurs taux de succès. Notre motivation dans ce travail est le besoin d'une méthode de détection automatique des régions de vertèbres dans les images rayons-X. Notre but final est de préparer le terrain par notre approche de détection de régions de vertèbres à une segmentation automatique, ou avec une intervention humaine minimale.

Notre travail consiste à l'étude de l'application d'une méthode d'apprentissage statistique pour la classification des régions sélectionnées à partir des images à rayons-X.

1.3. Organisation du mémoire :

Ce mémoire est composé de deux parties : une partie théorique regroupant deux chapitres présentant les principaux concepts de traitement d'images et de classification utilisés dans notre travail, et une autre partie pratique dont laquelle nous expliquons les différentes expérimentations que nous avons mis en œuvre.

Notre travail est organisé de la manière suivante :

Dans le premier chapitre nous citerons les différents travaux existants dans la littérature et qui proposent de résoudre le problème d'automatisation du processus de la segmentation des vertèbres, ainsi que les inconvénients trouvés dans chacune de ces solutions.

Dans le deuxième chapitre, nous décrivons la solution que nous proposons au problème de la détection des vertèbres, avec la base théorique des méthodes utilisées y compris la méthode de transformée en ondelette et les moments géométriques et les machine à vecteurs de support pour la classification et la détection.

Le troisième chapitre est consacré aux différents tests effectués ainsi qu'aux résultats obtenus par notre méthode de classification SVM, avec la visualisation de quelques images lors de la phase de détection et de raffinement.

Chapitre 1 : Etat de l'art

Introduction :

De nos jours ; les méthodes de segmentation des vertèbres sont en forte demande en raison de leur impact très important dans de nombreuses applications orthopédiques, neurologiques et oncologiques.

Pour cela plusieurs méthodes ont été développées dans la littérature afin de répondre aux besoins des radiologistes dans ce domaine.

En pratique dans les centres hospitaliers, la segmentation manuelle des vertèbres est effectuée par des radiologues experts, en se basant sur des informations visuelles et de leur expériences, afin de pouvoir identifier les différentes parties vertébrales : cervicales, thoraciques et lombaires, et assigner des étiquettes appropriées à chaque forme de ces vertèbres.

Pour une telle segmentation, deux méthodes sont utilisés dans la pratique clinique : la méthode de segmentation semi quantitative et la morphométrie quantitative.

Dans la segmentation semi quantitative, le radiologue peut identifier la pose et la forme général de la colonne vertébrale et quelques informations visuelles relatives à chaque vertèbre comme la forme des coins des plateaux de vertèbres qui peuvent donner des indices de présence de fracture [15], ou de séparation entre les vertèbres, etc. Cette méthode demeure subjective et dépendante de l'expérience des experts.

Pour la morphométrie quantitative, des points de marquage sont placés manuellement sur l'image radiographique pour plus ou moins délimiter les contours des vertèbres [38]. Ces points peuvent servir comme base pour les méthodes de segmentation semi-automatiques. Cette méthode peut prendre beaucoup de temps, et on peut perdre des informations sur la forme lorsque le nombre de points est minimisé.

Une solution efficace à ce problème serait l'utilisation d'un processus de segmentation d'images automatique et semi-automatique.

En effet avec l'apparition des techniques de vision par ordinateur ; différents algorithmes ont été développés dans la littérature pour répondre au besoin d'automatisation du processus de segmentation.

Ces méthodes peuvent être classées en trois grandes catégories principales :

Les méthodes empiriques, les méthodes basées sur les modèles déformables et les méthodes statistiques.

I.1. Les méthodes empiriques :

Ces méthodes se basent sur une connaissance « empirique » à priori concernant l'image et ces caractéristiques de forme.

Ces méthodes sont utilisées pour estimer la pose vertébrale afin de faire correspondre des modèles simples et également pour étiqueter les vertèbres.

Dans la littérature, beaucoup de travaux mettent en pratique ce principe, nous citerons comme exemple :

- Long et Thoma [19] ont essayé de localiser la vertèbre cervicale C2 par rapport à la position de l'arrière de crâne. Cette méthode est très utile pour réduire l'incertitude de l'emplacement des vertèbres cervicales. Plus tard, dans [18], les auteurs ont développé une méthode automatique permettant d'obtenir une approximation de premier ordre de l'endroit de la colonne vertébrale.
- G.Zamora et al [11] ont développé un algorithme pour estimer l'orientation et la position des vertèbres cervicales. Pour examiner l'efficacité de cette estimation, des points morpho-métriques ont été placés par des radiologues experts sur un ensemble de 40 films radiographiques. Cette approche a donné des résultats acceptables ce qui peut mener à un placement précis d'un modèle de forme moyen.

L'inconvénient majeur des méthodes de segmentation empiriques est qu'elles sont fortement dépendantes de la qualité de l'image radiographique d'une part, des facteurs comme le bruit et les bords faibles peuvent influencer les résultats.

D'une autre part, ces méthodes ne présentent aucune information a priori sur les caractéristiques géométrique et la forme d'objet d'intérêt.

I.2. Les modèles déformables :

Cette deuxième classe de méthodes peut être considérée comme une plate-forme pour les méthodes semi-automatiques et automatiques.

Le principe de ces approches est d'utiliser un modèle qui se déforme vers les bords de l'objet d'intérêt en se basant sur un processus de minimisation d'énergie par exemple.

De nombreux travaux ont utilisés cette classe de méthodes de segmentation:

■ Tezmol et al [1] ont utilisé la Transformée de Hough généralisé (TGH) afin de trouver la pose (position, orientation et échelle) de deux vertèbres cervicales dans les images radiographiques. Ce travail est constitué de deux étapes :

1. Un modèle de forme moyen est d'abord choisi parmi un ensemble d'apprentissage.
2. Une technique de vote est implémentée pour trouver la meilleure correspondance dans le domaine de Hough.

Le taux de succès a atteint 90% pour la détection des vertèbres cervicales.

■ Yalin Zheng et al [51] ont utilisé un algorithme génétique pour la recherche dans l'espace de Hough, et la forme de la vertèbre est caractérisée par les descripteurs de Fourier.

L'inconvénient de ce type de méthodes est l'initialisation « manuelle » qui doit être proche de l'objet d'intérêt. La THG peut apporter de bons résultats mais reste incapable de s'adapter avec les déformations des formes à cause d'usage de modèles rigides.

Une solution a été proposée par Mahmoudi et al. [41] intégrant la méthode de contours actifs « snakes » [27] pour combler les imperfections de l'approche de segmentation basée sur les modèles actifs de formes – ASM (marquage semi-automatique). Les résultats obtenus ont été satisfaisants.

I.3. Les méthodes statistiques :

La dernière classe de ces des méthodes de segmentation et qui est largement utilisé est la classe des méthodes statistiques. On peut distinguer les modèles statistiques de forme et d'apparence.

Dans les modèles statistiques de forme, la forme d'objet est représentée par un ensemble de points. Des annotations manuelles faites sur un ensemble d'images vont former l'ensemble d'apprentissage.

Le but est de construire un modèle représentant la variation de forme dans cet ensemble, qui va permettre d'analyser à la fois de nouvelles formes, et de synthétiser des formes analogues à celles de l'ensemble d'apprentissage.

Parmi les applications largement utilisées pour ce type de méthodes nous pouvons citer le modèle de forme actif (en anglais Active Shape Model ou ASM) introduites par Cootes et al [6], les méthodes basées sur les contours actifs « snakes » ainsi que les contours actifs par courbes de niveaux « Levels sets ».

Le principe des méthodes à base de contours actifs est d'essayer de déplacer un contour initial qui est créé comme modèle moyen de forme de la vertèbre segmenté manuellement vers les bords de l'objet.

Plusieurs applications remarquables dans la littérature peuvent être citées :

- ❖ Une version améliorée de l'ASM a été proposée dans [12] puis [13] pour la segmentation des vertèbres dans les images radiographiques numérisées.
 - Cette étude a mis en œuvre le principe de la morphométrie quantitative pour générer des points de repère délimitant la vertèbre : d'abord un module de la transformée de Hough généralisé est utilisé pour une estimation de la position, de l'orientation et de l'échelle de la vertèbre d'une part et remédier au problème d'initialisation d'autre part.
 - Ensuite, la variabilité de la forme des vertèbres est utilisée pour construire un modèle en utilisant un module d'ASM amélioré.
 - Finalement ce modèle va être déformé en fonction de la variabilité de formes observée dans l'ensemble d'apprentissage afin de bien capturer les détails fins des déformations locales au niveau de la vertèbre.
- ❖ Une autre étude rapportée par Klinder et al [46] pour la détection, l'identification et la segmentation automatique des vertèbres dans l'image tomographiques (en anglais CT images). La solution proposée dans ce travail est constitué des étapes suivantes :
 - L'extraction de la courbe de la colonne vertébrale, la détection de la position des vertèbres et l'identification (nomination) de chaque

vertèbre. Ce processus est finalisé par une étape de segmentation en utilisant un modèle de forme spécifique pour l'ensemble des vertèbres. Le taux de succès d'identification a atteint plus de 70% pour chaque vertèbre.

- ❖ Récemment, une seconde solution a été proposée dans [43] combinant deux méthodes afin de développer un algorithme automatique de la segmentation de colonne vertébrale y compris la partie cervicale, thoracique et lombaire dans les images CT :
 - La première méthode utilise une segmentation initiale à base de modèles déformables.
 - La deuxième méthode repose sur une segmentation fine à base de modèles de formes et d'intensités statistiques. Les résultats ont été comparés aux résultats de la segmentation manuelle. La performance de cet algorithme a été jugé équitable.

Une autre variante de la méthode d'ASM est le Modèle Actif d'Apparence (en anglais Active Appearance Model AAM) qui est aussi largement décrite dans la littérature [16].

L'utilisation de la forme et des paramètres d'apparence du modèle ont été discutés dans [26] pour la localisation automatique des vertèbres lombaires. Cependant les méthodes ASM restent plus rapides avec des résultats meilleurs.

Une autre structure hiérarchique a été proposée dans [2] afin de prouver l'efficacité des AAM pour la segmentation des vertèbres cervicales (avec un taux de succès de 65%) et lombaires (68%) dans les images à rayons-X.

Par ailleurs les méthodes statistiques, d'une façon générale, restent très sensibles par rapport à la précision de la phase d'initialisation.

Nous pouvons également citer les solutions proposée par M. Benjelloun et Saïd Mahmoudi [22][23][24] pour l'extraction des contours des vertèbres dans les images à rayons-X, en utilisant différents types de méthodes comme la signature polaire, la comparaison de modèles « Template matching », et les contours actifs.

Une étude étendue proposée par Szu-Hao Huang pour la classification des vertèbres a été présentée dans [45], afin de surmonter le problème de la sélection ou le marquage manuelle.

Cette étude représente une méthode d'apprentissage statistique pour la détection et la segmentation automatique des vertèbres dans les images à résonance magnétique(IRM).

Le system est composée de trois étapes principales :

- 1- Etape de détection basée sur l'algorithme AdaBoost (Adaptif Boosting)[40].
- 2- Etape de raffinement avec l'algorithme Ransac (RandomSAmple Consensus) [25].
- 3- Etape de segmentation à base de minimisation d'énergie et de l'algorithme coupe de graphes « Graph cut ».

Les contributions majeurs de cette méthode consiste à utiliser les apprenants bayésiens faibles pour préserver la distribution des données d'apprentissage, l'estimation d'entropie par l'arbre de décision ID3, et enfin le développement du pouvoir de représentation des caractéristiques par une architecture de couplement de caractéristiques.

Le gradient et l'intensité représente le couple de caractéristiques choisie, qui sont soumis à une procédure d'extraction de coefficients par la transformée en ondelettes de Harr.

Dans ce travail, nous proposons une méthode statistique à base des machines à vecteurs de supports (SVM) pour la détection des régions vertèbres dans les images à rayons-X.

Notre contribution consiste à la combinaison d'une telle méthode avec un moyen très puissant de la description d'image qui est la transformé en ondelettes de Haar puis un autre essai avec une description globale de forme en utilisant les moments géométriques HU.

Les résultats sont présentés afin d'évaluer la performance de détection de classifier SVM.

Le chapitre suivant va détailler les démarches suivis afin de sélectionner le modèle optimal de notre méthode de classification.

Chapitre 2:

Méthodologie



II.1 Architecture et représentation de la méthode :

La détection des vertèbres dans les images radiographiques est une tâche difficile qui est due à leurs structures complexes, et aussi au manque de contraste entre les structures osseuses et les tissus musculaires. Dans ce chapitre nous allons décrire l'architecture de la méthode proposée pour résoudre ce problème. Notre processus de détection est divisé en trois principales parties :

La première étape implique une procédure d'amélioration de la qualité des images par ajustement d'intensité, afin de clarifier les régions à extraire.

La deuxième étape qui est primordiale dans notre système, représente l'extraction de caractéristiques des régions des vertèbres.

La dernière étape met en œuvre un modèle de classification basée sur une méthode à base des machines à vecteur support –SVM – pour effectuer une bonne classification et renforcer la détection. Les résultats de cette étape sont ensuite raffinés par une méthode d'estimation RANSAC.

Un aperçu de notre système est montré dans le schéma ci-dessous (figure 1).

Dans la phase d'entraînement, le système prend comme entrée :

- 1) Un ensemble d'images contenant les vertèbres qui ont été aligné pour avoir la même taille.
- 2) Pour chaque modèle de notre base, une représentation intermédiaire, encapsulant les informations importantes concernant la classe des vertèbres, est calculée. Et cela pour donner un ensemble de vecteurs de caractéristiques positives et négatives. Ces vecteurs sont utilisés pour entraîner le classificateur afin de différencier les deux modèles de classes.

Cette représentation a été calculée par l'utilisation de deux descripteurs différents, qui sont les ondelettes, et les moments HU, afin de pouvoir distinguer l'information la plus pertinente pour une bonne caractérisation.

Afin d'améliorer de la performance de notre classificateur, la méthode de validation croisée a été utilisée dans le but d'aboutir à une bonne minimisation d'erreur de généralisation de nos données d'apprentissage.

Dans la phase de test, nous nous intéressons à la détection des vertèbres dans une nouvelle image. Le système fait glisser une fenêtre de taille fixée dans toute l'image pour décider si le modèle peut être assimilé à une région d'intérêt.

A chaque position de la fenêtre, on extrait les mêmes caractéristiques prédéfinies dans la phase d'entraînement, ensuite nous les injectons dans le classificateur. La sortie va déterminer si la région correspond ou pas à une région de vertèbre. Une méthode d'estimation est ensuite utilisée afin de raffiner les résultats de détections (fausse ou manque de détections).

Nous notons que pour une détection multi-échelle, il suffira de redimensionner l'image en entrée et effectuer le même traitement, en utilisant toujours la même fenêtre de balayage.

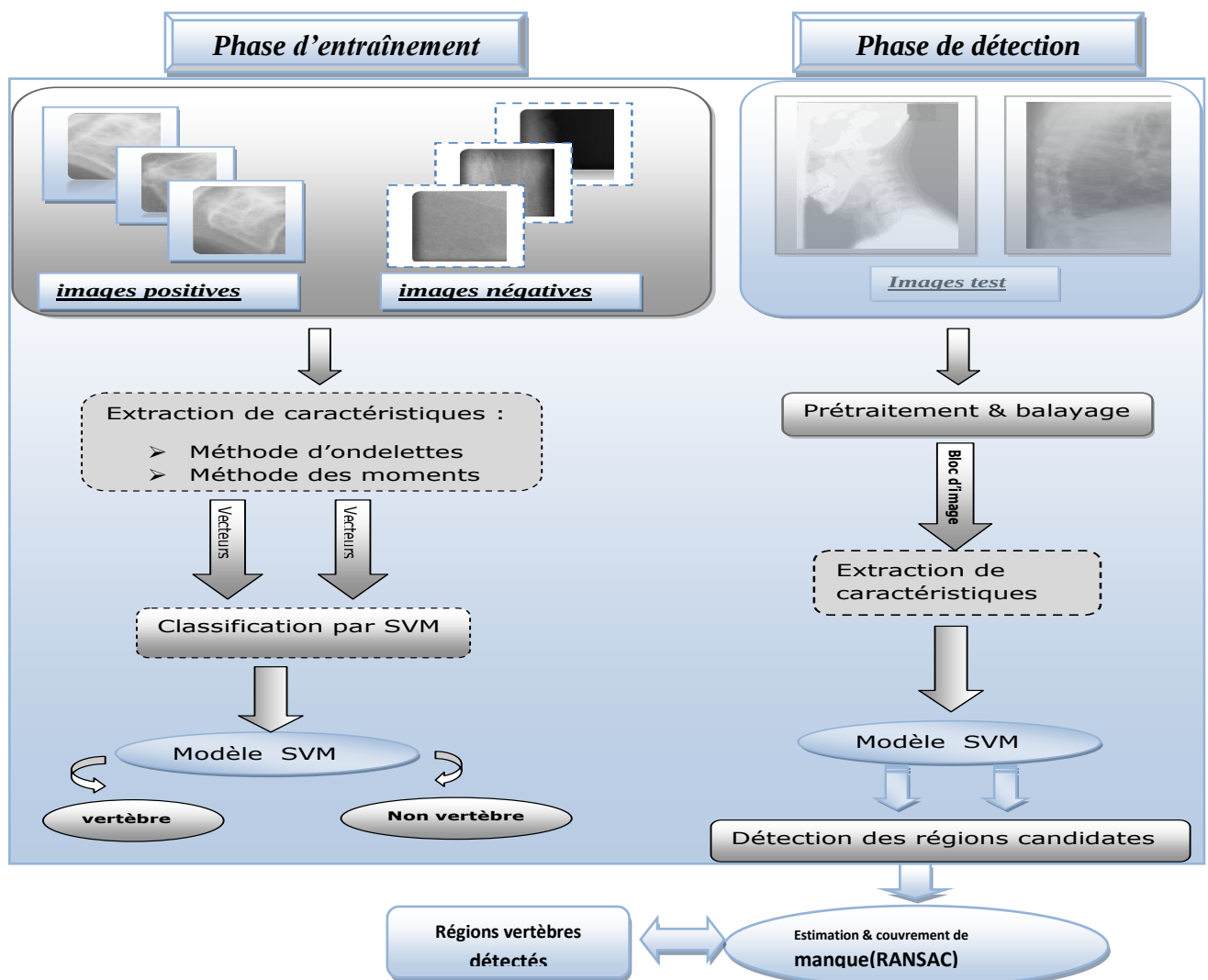


Figure1 : Schéma général de la méthode de détection des vertèbres par SVM.

II.2 Le prétraitement :

Le but de cette étape est d'éliminer, le plus possible, les informations non pertinentes dues au bruit résultant de l'acquisition de l'image, et par conséquent, faciliter l'extraction des informations utiles à l'analyse. Un traitement de rehaussement de contraste est appliqué aux images rayon X de notre base. Ce traitement consiste à accentuer les caractéristiques d'une image afin de rendre son affichage plus convenable à l'analyse.

L'objectif est d'augmenter le contraste d'images afin d'accroître la séparabilité des régions.

La méthode que nous avons retenue est une égalisation adaptative d'histogramme :

II.2.1 Modification d'histogramme :

L'histogramme d'une image est une fonction H définie sur l'ensemble des entiers naturels :

$$H(x) = \text{card} \{p: I(p) = x\}$$

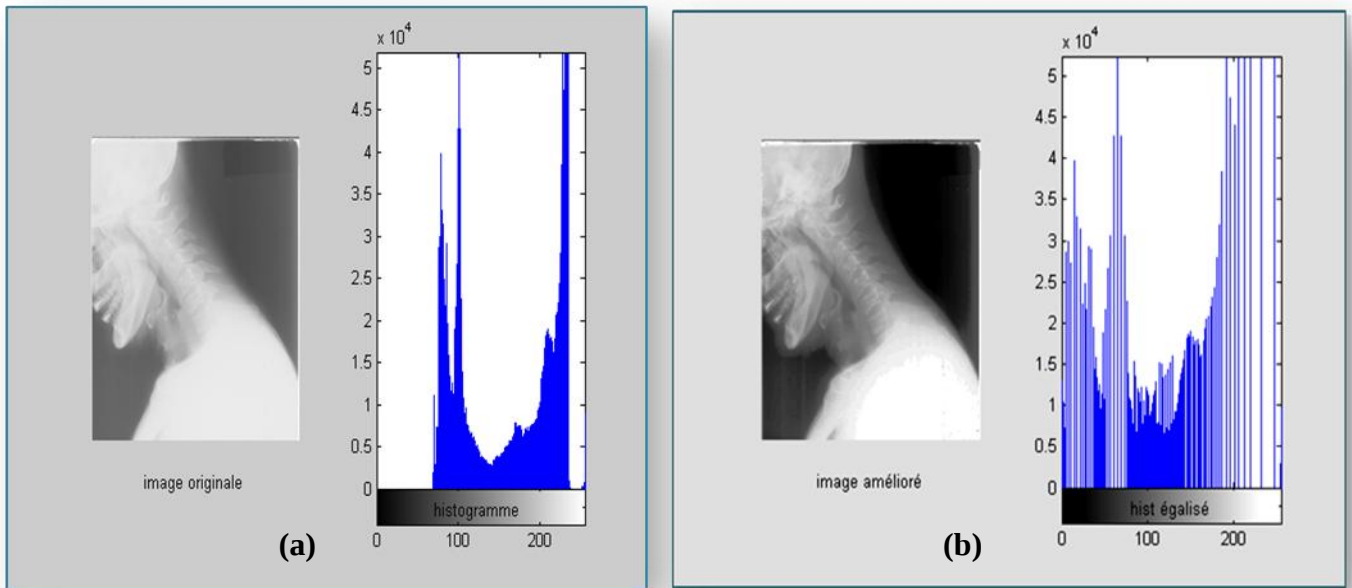
$H(x)$ correspond au nombre d'occurrences de niveaux de gris présent dans l'image. Autrement dit, l'histogramme est une représentation graphique de la distribution des valeurs des niveaux de gris.

Le principe de la modification d'histogramme est d'appliquer une linéarisation, afin de répartir uniformément les valeurs des pixels sur l'ensemble de l'histogramme.

II.2.2 L'égalisation d'histogramme :

Cette opération consiste à calculer à partir de l'histogramme $H(x)$ de l'image I une fonction de rehaussement des niveaux de gris f telle que l'image rehaussée J , définie par : $J(p) = f(I(p))$ puisse avoir son histogramme H_j se rapprochant le plus possible d'une fonction plate.

La figure 2 suivante présente un exemple d'une image de la base avant et après l'égalisation d'histogramme.



**Figure 2 : Exemple d'image et son histogramme (a) : avant
égalisation (b) après égalisation**

II.2.3 L'égalisation adaptative d'histogramme :

Son principe consiste à appliqué sur chaque pixel ainsi que « sa région contextuel » une égalisation d'histogramme.

Cette région représente en effet les pixels voisions entourant le pixel traité.

Toutes les images de notre base ont été améliorées pour mettre en valeur les caractéristiques visuelles des différentes régions des vertèbres (Figure 3) afin de faciliter la tâche d'extraction des caractéristiques.

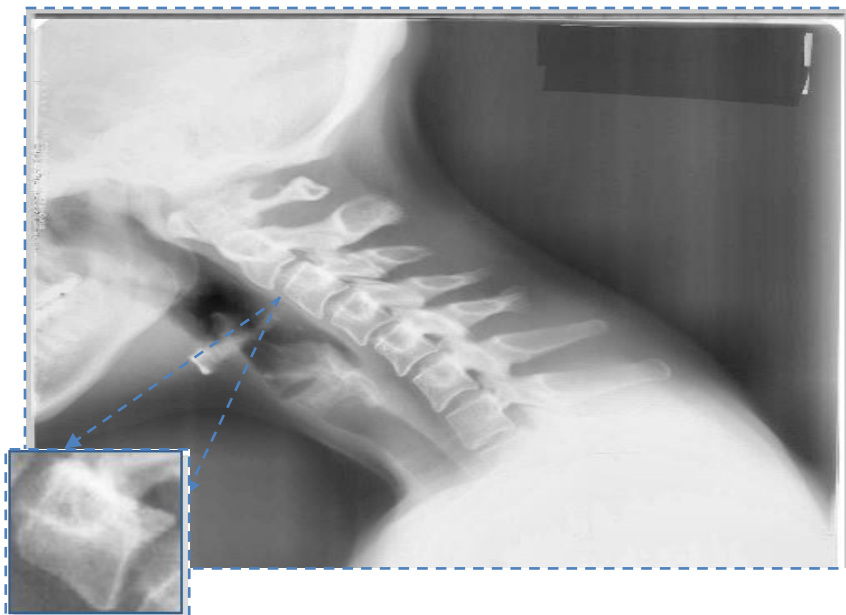


Figure 3 : Exemple d'image amélioré par égalisation d'histogramme adaptative

II.3. L'extraction des caractéristiques :

L'analyse des images se focalise généralement autour des attributs de bas niveaux tels que la texture, la forme, et la couleur.

Il y'a principalement deux approches de caractérisation :

La première est la construction de descripteurs **globaux** à toute l'image c.à.d. fournir des observations sur la totalité de l'image.

L'avantage de ces descripteurs est la simplicité de leur mise en œuvre, ainsi que le nombre réduit d'observations que l'on obtient. Cependant leur inconvénient majeur est la perte de l'information de localisation des éléments de l'image.

La seconde approche est **locale** et consiste à calculer des attributs sur des portions restreintes de l'image.

L'avantage des descripteurs locaux est de conserver une information localisée dans l'image, évitant ainsi que certains détails ne soit noyés par le reste de l'image. L'inconvénient majeur de ces techniques est que la quantité d'information produite est très grande.

Le choix des caractéristiques extraites est souvent guidé par la volonté d'invariance ou robustesse par rapport aux transformations de l'image.

L'approche utilisée dans notre solution est une approche à la fois locale car basée sur le choix de la texture et la forme par la méthode d'ondelettes et globale par l'utilisation des moments de HU.

II.3.1. La représentation par la transformée en ondelettes :

La texture peut être vue comme un ensemble de pixels (niveaux de gris) spatialement organisés selon un certain nombre de relations spatiales, ainsi créant une région homogène. De ce fait, la modélisation des textures est portée sur la caractérisation de ces relations spatiales.

Parmi les modèles les plus connus nous pouvons citer : les méthodes statistiques (Matrice de cooccurrence, différence de niveaux de gris, etc.), les méthodes

fréquentielles (Transformée de Fourier, Filtre de Gabor, et les ondelettes, qui sont utilisées dans notre travail).

II.3.1.1. Introduction a la théorie d'ondelette :

L'imagerie médicale a révolutionné les pratiques médicales. Néanmoins, de nombreux problèmes liés au traitement d'image sont encore ouverts et leur résolution (même partielle) peut aboutir à une amélioration des diagnostics et des actes chirurgicaux.

Nous pouvons citer par exemple :

- Le problème de la réduction des radiations administrées lors d'un examen scanner (problème de la tomographie locale).
- La chirurgie assistée par ordinateur (incluant des problèmes de segmentation automatique, de recalage de données et de reconstruction temps-réel 3D).
- La détection et l'analyse de structures malignes dans des données d'échographie, mammographie, ou spectroscopie RMN (incluant par exemple l'analyse d'images texturées), etc.

L'utilisation des bases d'ondelettes en traitement d'images c'est généralisée durant les vingt dernières années [33]. Leur intérêt pour la compression et le dé-bruitage a été démontré puisqu'elles ont intégré le dernier standard de compression des images numériques JPEG 2000.

Leur application à l'imagerie médicale date de 1992 et c'est largement répandu depuis [14,32]. Dans ce contexte les ondelettes sont utilisées pour la compression et le dé-bruitage, mais aussi pour l'analyse fonctionnelle des données médicales (en vue d'établir un diagnostic), la tomographie locale, la segmentation et le rehaussement d'images, ou encore la description de textures.

Nous nous intéressons ici à la contribution des ondelettes à l'analyse et à la caractérisation des régions des vertèbres dans les images radiologiques.

II.3.1.2. L'analyse en ondelettes :

Depuis les travaux de Grossman et Morlet [14], la transformation en Ondelettes est apparue comme un outil performant permettant de résoudre des problèmes relevant de

différents domaines d'application [45]. Très tôt, un intérêt soutenu c'est manifesté à l'égard de la Transformation en ondelettes en traitement d'images [20,32].

La notion d' "Ondelettes" ou "Wavelets" a été utilisée pour la première fois au début des années '80' par le géophysicien français J.Morlet [32] pour désigner des fonctions mathématiques utilisées dans la représentation des données sismiques. Les ondelettes sont des fonctions de base de variation multi-échelles, ou multi-résolution, utilisées dans le but de l'approximation et/ou de la compression des données.

La transformée en ondelettes décompose le signal d'entrée, équation (1), en une série de fonctions d'ondelettes $\psi_{a,b}(t)$ qui dérivent d'une fonction mère $\psi(t)$ donnée par des opérations de dilatation et de translation, équation (2).

$$Ca, b = \int_{-\infty}^{+\infty} x(t) \psi_{a,b}(t) dt \quad (1)$$

$$\psi_{a,b}(t) = \frac{1}{\sqrt{a}} \psi\left(\frac{t-b}{a}\right) \text{ avec } a > 0 \quad (2)$$

Où les paramètres :

- **a** est le facteur d'échelle.
- **b** est le paramètre de translation.

L'analyse par ondelettes est un outil mathématique capable de transformer un signal d'énergie finie dans le domaine spatial en un autre signal d'énergie finie dans le domaine spatio-fréquentiel.

Les composantes de ce nouveau signal sont appelées *les coefficients d'ondelettes*. Ces coefficients renseignent sur la variation locale des niveaux de gris autour d'un pixel donné de l'image. Ils sont d'autant plus grands que cette variation est importante.

En 1989, Mallat a proposé un algorithme de décomposition multi-résolution basé sur la transformation en ondelettes. L'algorithme décompose une image en entrée en un ensemble d'images de détails et une image d'approximation. À chaque niveau de décomposition la taille des images transformées est réduite par un facteur de deux.

Mallat montra que la transformée en ondelette discrète (TOD) peut être implémentée grâce à un banc de filtres comprenant un filtre passe bas (PB) et un filtre passe-haut (PH).

Dans la figure 4 qui illustre ces bancs de filtre, le signal d'entrée subit un filtre passe-haut et un filtre passe-bas. Après une opération de sous-échantillonnage, le filtrage reprend sur chaque sous-bande.

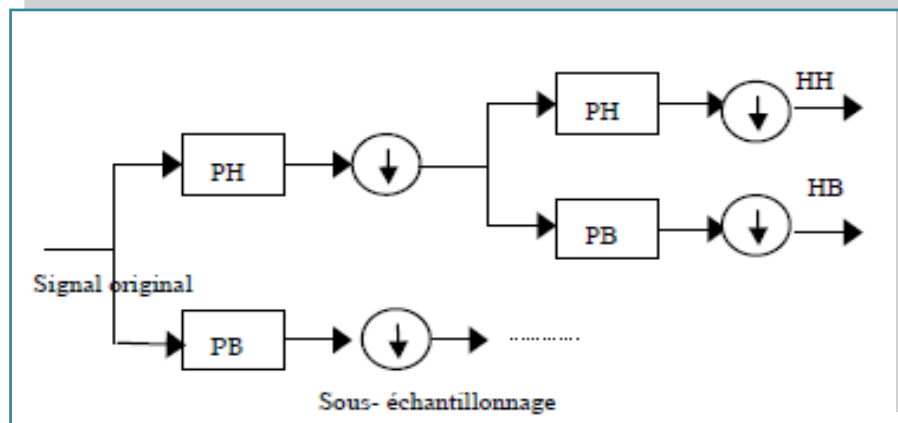


Figure 4 : Décomposition d'un signal en approximation et détails

Pour le cas d'un signal 2D, la TOD est appliquée d'abord ligne par ligne, puis colonne par colonne. Quatre images sont alors générées à chaque niveau comme le montre la figure 5 :

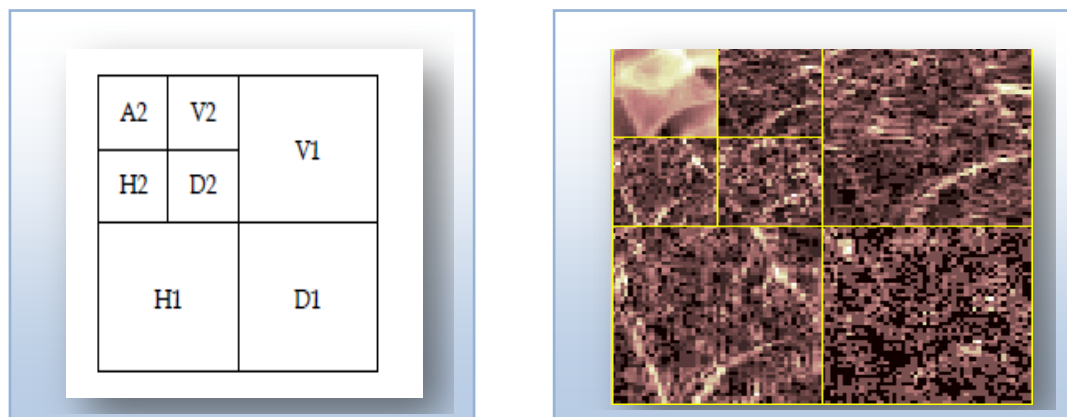


Figure 5 : Un exemple de décomposition de l'image sur deux niveaux.

II.3.1.3 Choix de l'ondelette :

Dans la littérature, nous pouvons trouver une multitude de fonctions de base d'ondelettes. Le choix de l'ondelette dépend essentiellement du type de l'application (compression, segmentation, etc.). Prenons l'exemple de la classification de textures, les conclusions des auteurs ayant utilisé différents types d'ondelettes sont souvent contradictoires. En effet, de bons résultats ont été trouvés par exemple pour des ondelettes à support étendu [20], tandis que pour ce même type d'ondelettes, les résultats obtenus n'apportaient aucune amélioration significative [47].

Dans notre travail, nous avons utilisé les ondelettes de Harr. Nous ne sommes pas fixés sur la recherche de l'ondelette optimale mais notre objectif est plutôt de mettre l'accent sur les phases d'extraction et d'analyse après une transformation en ondelettes sur des images gradients.

La figure 6 présente quelques exemples d'ondelettes.

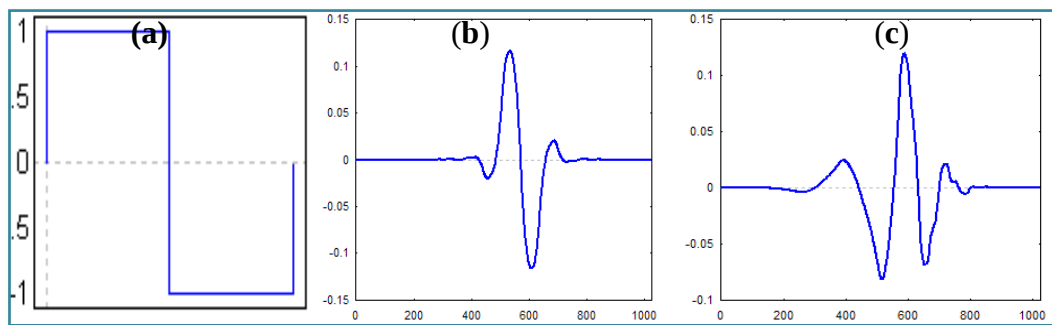


Figure 6 : Exemples classiques d'ondelettes : (a) ondelettes de Harr, (b) ondelette bi-orthogonale, (c) ondelettes de Daubechies

II.3.1.4 Ondelettes de Haar :

Le mérite revient à Alfred Haar d'avoir construit en 1909 des bases considérées aujourd'hui comme le fondement de la théorie des ondelettes (Meyer, 1994). Haar a défini une fonction $h(x)$ telle que présentée dans l'équation 3 :

$$h(x) = \begin{cases} 1 & \text{pour } 0 \leq x \leq 1/2 \\ -1 & \text{pour } 1/2 \leq x \leq 1 \\ 0 & \text{ailleurs} \end{cases} \quad (3)$$

Pour $n \geq 1$, il a construit une base orthonormée de $L^2 [0 ; 1]$ avec des fonctions définies par l'équation 4 :

$$\begin{cases} h_n(x) = 2^{j/2} \cdot h(2^j x - k) \text{ Avec} \\ n = 2^j x + k, j \geq 0, 0 \leq k \leq 2^j \end{cases} \quad (4)$$

Dans l'équation (4), $L^2 [0 ; 1]$ est l'espace des fonctions de carré intégrable sur l'intervalle $[0 ; 1]$. Le support de $h_n(x)$ est l'intervalle dyadique défini par l'équation 5 :

$$I_n = [k2^{-j} ; (k+1)2^{-j}] \subset [0; 1[\quad (5)$$

$h_n(x)$ peut aussi s'écrire selon l'équation (6) :

$$h_n(x) = 2^{j/2} \cdot h(2^j x - k) = \frac{1}{\sqrt{2^{-j}}} h\left(\frac{x - k2^{-j}}{2^{-j}}\right) \quad (6)$$

Par analogie avec la formule d'ondelettes définie précédemment, le facteur d'échelle de h_n est $a = 2^{-j}$ et le paramètre de dilatation est $b = k2^{-j}$.

II.3.1.5 Représentation par ondelettes proposée:

Notre représentation va s'intéresser aux petites régions locales contenant les vertèbres, en observant les différences d'intensités présentes.

La transformation par ondelettes de Harr est appliquée à chaque image, et le résultat est un ensemble de coefficients à plusieurs échelles (Levels), qui représente les différentes orientations de différence d'intensité : verticale, horizontale, et diagonale. Par conséquent nous allons avoir trois ensembles de coefficients représentant les orientations d'ondelette. (Voir figure 7).

La figure 8 montre différentes décompositions et détails d'ondelettes pour un exemple de notre base d'apprentissage.

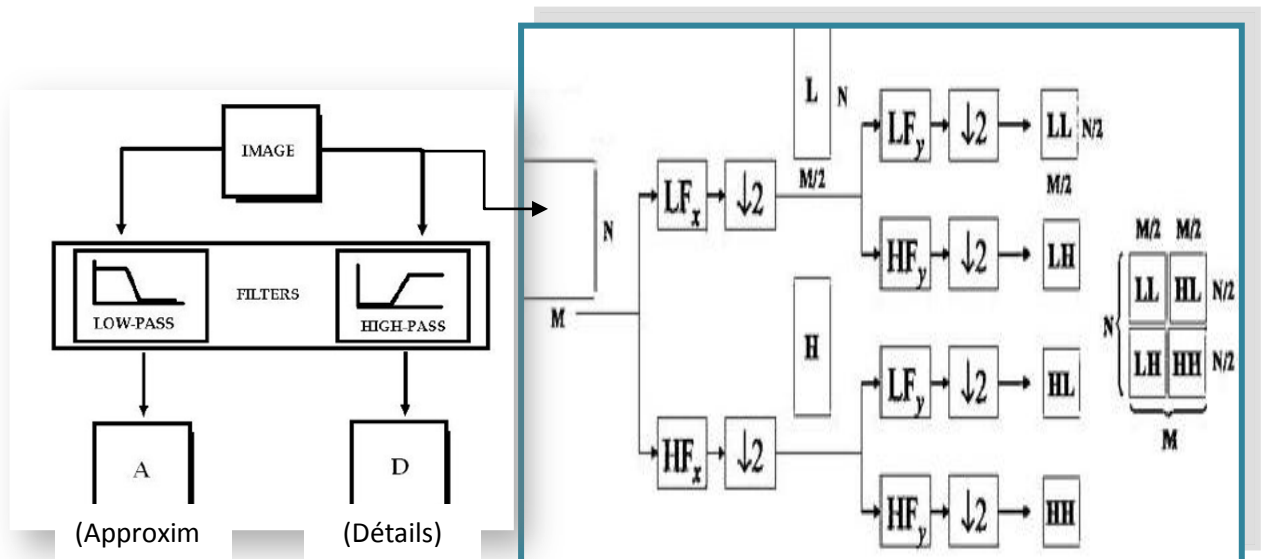


Figure 7 : Décomposition de l'image par transformée en ondelette

- LF : représente le filtre passe bas de haar [0.71, 0.71].
- HF : représente le filtre passe haut de haar [-0.71, 0.71].

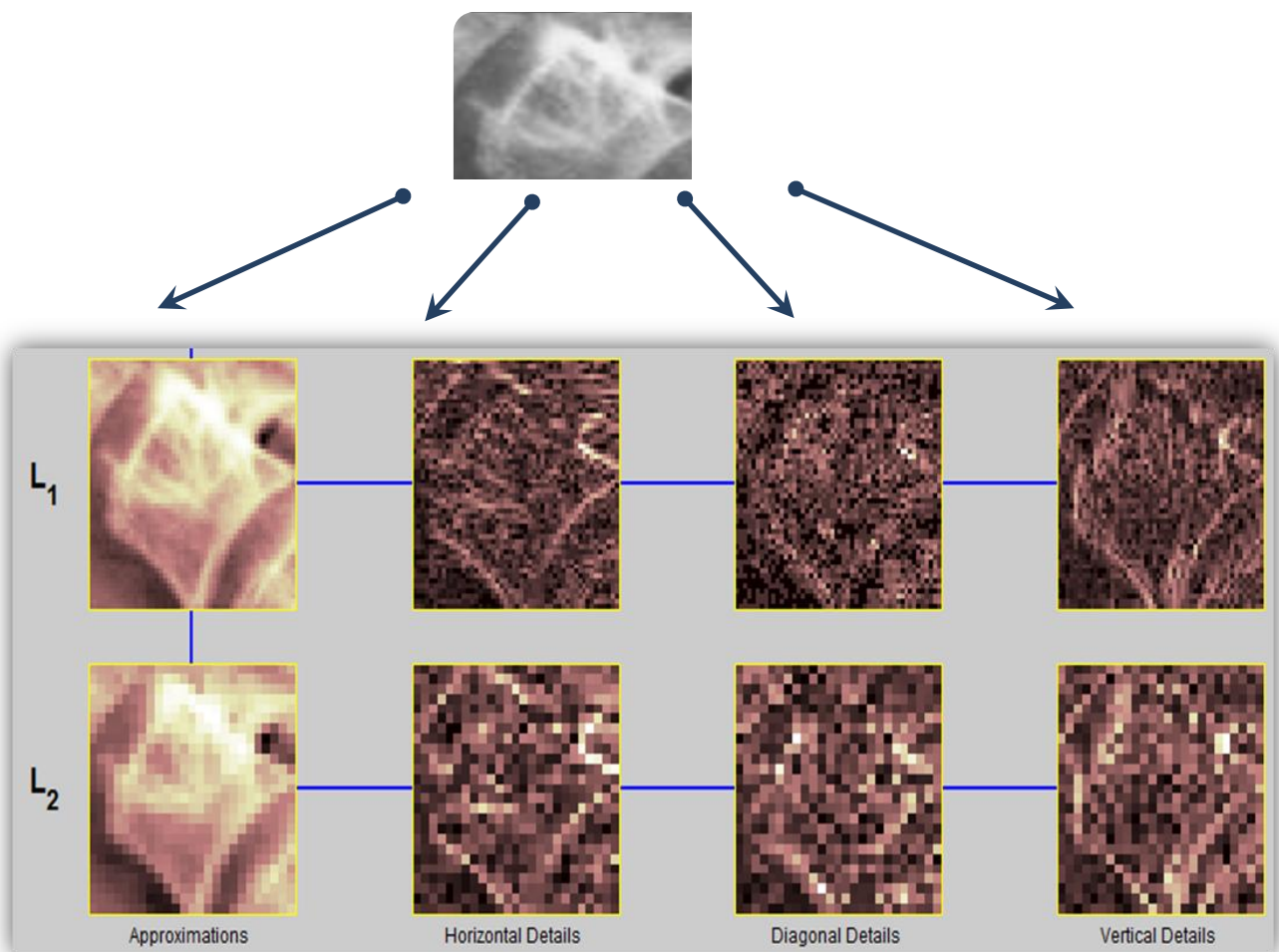


Figure 8 : Exemple d'analyse en ondelette à différents niveaux de détails aux 1^{ers} et 2^{èmes} niveaux

Au même titre que la texture, la forme est généralement une description très riche d'un objet.

De nombreuses solutions ont été proposées dans la littérature pour représenter une forme. Nous distinguons deux catégories de descripteurs de formes : les descripteurs basés sur les frontières (ou contours), et les descripteurs basés sur les régions.

Les premiers font référence aux descripteurs de Fourier et portent sur une caractérisation des contours de la forme. La seconde approche fait référence, entre autres, aux moments invariants qui sont utilisés pour caractériser l'intégralité de la forme région. Ces attributs sont robustes aux transformations géométriques comme la rotation et le changement d'échelle. La méthode que nous utilisons dans cette étude concerne une description de forme à base de région par le moyen de moments géométriques.

II.3.2 La représentation par les moments invariants de HU :

II.3.2.1 les moments géométriques:

Les moments géométriques [29] permettent de décrire une forme à l'aide de propriétés statistiques. Ils représentent les propriétés spatiales de la distribution des pixels dans l'image. Ils sont facilement calculés et implémentés, par contre cette approche est très sensible au bruit et aux déformations.

Les moments sont utilisés depuis longtemps pour calculer la position et le centre d'une distribution mais aussi sa variance. En vision par ordinateur, les moments permettent de calculer la position et l'orientation d'un objet.

A partir des moments géométriques, HU [28] a proposé un ensemble de sept moments invariants aux translations, rotations et changement d'échelle.

Une première utilisation des moments pour la reconnaissance de motifs géométriques a été proposée par HU [35].

Les moments offrent un cadre théorique puissant pour résoudre des problèmes rencontrés dans plusieurs applications d'imagerie y compris l'imagerie médicale, surtout lorsqu'il s'agit d'observer un objet rigide (comme les vertèbres) à différents

niveaux d'échelle, rotations et translations dans une image 2D (tomographie, projection, etc.).

Récemment, une application remarquable des moments a été présentée par [21] en combinaison avec la transformée en ondelettes pour la catégorisation des images médicales, en utilisant les cartes d'auto-organisation avec un taux de réussite de 81.8%.

II.3.2.2. Théorie des moments :

La notion de moment en mathématiques (notamment en calcul des probabilités) a pour origine la notion de moment en physique. Le moment m_{pq} d'ordre $n = p + q$ d'une fonction de distribution $f(x, y)$ est définie par l'équation (7) suivante :

$$m_{pq} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x^p y^q f(x, y) dx dy \quad (7)$$

Pour une image digitale $g(x, y)$ de taille $M \times N$, la formule ci-dessus devient :

$$m_{pq} = \sum_{y=0}^{N-1} \sum_{x=0}^{M-1} x^p y^q g(x, y) \quad (8)$$

Le moment centrale est défini par :

$$\mu_{pq} = \sum_{y=0}^{N-1} \sum_{x=0}^{M-1} (x - \bar{x})^p (y - \bar{y})^q g(x, y) \quad (9)$$

Les moments apportent différentes informations statistiques sur la forme :

- ✚ ordre 0, surface de la forme : m_{00} .
- ✚ ordre 1, centre de gravité de la forme, et est calculé par la formule (10) :

$$\bar{x} = \frac{m_{10}}{m_{00}}, \bar{y} = \frac{m_{01}}{m_{00}} \quad (10)$$

Pour caractériser la forme des images, nous utilisons la famille de moments HU [35] qui sont invariants aux différents changements d'image :

Les moments de Hu sont 7 moments issus de produits et quotients des moments centrés normés d'ordre 3, définis comme suit :

Moments d'ordres 2 :

$$\Phi_1 = m_{20} + m_{02}$$

$$\Phi_2 = (m_{20} + m_{02})^2 + 4m_{11}^2$$

Moments d'ordres 3 :

$$\Phi_3 = (m_{30} - 3m_{12})^2 + (3m_{21} - m_{03})^2$$

$$\Phi_4 = (m_{30} + m_{12})^2 + (m_{21} + m_{03})^2$$

$$\Phi_5 = (m_{30} - 3m_{12})(m_{30} + m_{12})[(m_{30} + m_{12})^2 + (m_{21} + m_{03})^2] - 3(m_{21} - m_{03})(m_{30} - m_{12})(m_{21} + m_{03})$$

$$\Phi_6 = (m_{20} - m_{02})[(m_{30} + m_{12})^2 - (m_{21} + m_{03})^2] - 4m_{11}(m_{30} - m_{12})(m_{21} + m_{03})$$

$$\Phi_7 = (3m_{21} - m_{03})(m_{30} + m_{12})[(m_{30} + m_{12})^2 - 3(m_{21} + m_{03})^2] - 3(m_{21} - m_{03})(m_{30} - m_{12})[(m_{30} + m_{12})^2 - (m_{21} + m_{03})^2]$$

Les six premiers moments sont invariants aux translations, aux changements d'échelles, et aux rotations ainsi qu'aux réflexions. Or, l'invariance aux réflexions peut être problématique quand il s'agit de reconnaître des images « miroirs ». C'est pourquoi M. K. Hu exploite aussi dans son système le 7^{ème} moment qui n'est pas invariant aux

réflexions. Ce dernier change de signe lorsqu'une telle transformation est appliquée à l'image et permet donc de détecter celle-ci.

II.4.Classification SVM :

L'enjeu essentiel de l'apprentissage artificiel est l'aptitude à généraliser des résultats obtenue à partir d'un échantillon limitée. Dans ce travail, nous choisissons la méthode à base des machines à vecteurs de support comme un moyen opérationnel pour ce problème.

Nous allons présenter dans cette section le fondement théorique de cette méthode ainsi que sa mise en œuvre dans notre problème de détection des régions vertèbres dans les images à rayons-X.

II.4.1 Notions sur l'apprentissage statistique :

2.4.1.1 Généralités :

D'après Mari & Napoli (1996) [1], effectuer une classification, c'est mettre en évidence, d'une part, des relations entre des objets et, d'autre part les relations entre ces objets et leurs paramètres. Il s'agit de construire une partition de l'ensemble des objets en un ensemble de classes qui soient les plus homogènes possible.

La classification, a donc deux objectifs à atteindre :

- Trouver un modèle capable de représenter la répartition des données (catégorisation).
- Définir de manière formelle l'appartenance à l'une ou l'autre des classes de toute nouvelle donnée (généralisation).

En effet la classification a pour but de réduire l'espace de recherche dans une base de données lors du processus d'identification.

En pratique, on peut rencontrer deux catégories de problèmes de classification, la classification supervisée et non-supervisée :

- **Apprentissage supervisé** : dans ce type d'apprentissage, on cherche à :
 - Estimer une fonction $f'(x)$ qui est la relation entre les objets et leurs classes.
 - Les objets utilisés comme données d'apprentissages sont accompagnés par la classe à laquelle ils appartiennent.



- **Apprentissage non supervisé** : on ne cherche pas cette fois à estimer une fonction mais on cherche à regrouper les objets ayant des caractéristique commune, les objets utilisés comme données d'apprentissage sont présentés sans leur classes.

II.4.1.2 Formulation d'un problème de classification (supervisé) :

Le classificateur doit estimer une fonction $f(x)$ qui est l'estimation de la fonction qui représente la relation entre l'objet et sa catégorie. Cette fonction est appelée fonction de décision :

$$f : X \rightarrow Y$$

X: L'ensemble des objets à classifier (appelé espace d'entrée).

Y: L'ensemble des catégories (appelé espace d'arrivée).

2.4.1.3 Minimisation du risque structurel :

Deux types de données sont utilisés pour un problème d'apprentissage : les données d'entraînement (données d'origine pour calculer le modèle) et les données de test (pour évaluer la performance de généralisation du modèle).

La qualité de ce modèle est alors jugée par rapport à sa capacité à réduire l'erreur de test ou de « généralisation ». Cependant, comme le modèle n'est pas construit en utilisant l'ensemble de test, l'erreur de généralisation ne peut pas être évaluée exactement car elle dépend de la distribution de probabilité des données.

Suivant la théorie de Vapnik [2], nous supposons que les données sont générées selon une distribution de probabilité inconnue $P(x,y)$. De plus, nous supposons que les données sont indépendantes et identiquement distribuées (iid).

L'erreur moyenne commise sur toute la distribution $P(x,y)$ par la fonction $f(x)$ est donnée par:

$$R[f] = \int \frac{1}{2} Q(x) dP(x, y) \quad (11)$$

Où:

- Q : est la fonction d'erreur (erreur absolue dans le cas des SVM).
- x : est le vecteur d'entrée.
- y : est l'ensemble des classes.

Ainsi la fonction f devra être optimale : la fonction f^{opt} devra être calculée de sorte que l'erreur moyenne sur toute la distribution soit minimale.

$$f^{opt} = \arg \min_f (R[f]) \quad (12)$$

Le critère formulé dans (12) est malheureusement inutilisable en pratique. En effet, pour calculer le risque, nous devrions disposer d'une estimation de la distribution $P(x,y)$, ce qui n'est pas le cas. La seule information dont nous disposons comme évaluation de l'erreur est l'erreur d'entraînement appelée Risque empirique :

$$R_{emp} = \frac{1}{l} \sum_{i=1}^l \frac{1}{2} Q(x_i) \quad (13)$$

Où :

- l est le nombre d'objets d'entraînements.

Ceci n'est pas suffisant. La raison en est que l'on peut facilement trouver un modèle minimisant l'erreur d'entraînement mais pour lequel l'erreur de généralisation sera très grande. Donc cette dernière est liée à la famille de fonction utilisée comme modèle. Cette dépendance est nommée « risque structurel ».

Dans la théorie de l'apprentissage statistique, Vapnik et Chervonenkis ont prouvé qu'il est possible de définir une majoration du risque structurel en fonction de la famille de fonctions utilisée pour le modèle [2]. L'une de ces majorations peut être calculée en utilisant la dimension de Vapnik–Chervonenkis (dimension VC) qui représente le plus grand nombre de points pouvant être séparés de toutes les façons possibles par un membre de l'ensemble de fonctions de F . Cela veut dire qu'il doit exister une configuration de h ($=VC(F)$) points, telle que les fonctions $f \in F$ peuvent leur assigner les 2^h combinaisons des labels (classes) possibles.

Dans la section suivante, nous allons présenter une technique de classification automatique supervisée basée sur la minimisation de risque structurel : SVM.

II.4.2 Théorie des machines à vecteurs de support (SVM) :

2.4.2.1. Introduction :

Les Machines à Vecteurs de Support -SVM - [49,50] ont été largement utilisées et appliquées aux problèmes de classification [4, 8, 16, 17,44] et la régression non linéaire [36, 37, 48].

Nous nous intéressons ici à l'application des SVM dans le domaine bio-informatique et médical.

Les SVM ont été utilisées comme méthode de prédiction de cancer du sang (leucémie, lymphome, etc.) dans le contexte de la technologie des puces d'ADN (ou micro-array) [39, 42].

Liang et Lin ont prouvé l'efficacité des SVMs pour la classification des données ECG pour la détection d'une vidange gastrique tardive dans [50].

Une autre application en diagnostic médicale est présentée pour la détection des micro-calcifications dans les mammographies dans [3, 10].

II.4.2.2 Principe des machines à vecteurs de support :

Les SVM constituent une classe d'algorithmes basée sur le principe de minimisation de « Risque structurel » décrit par la théorie de l'apprentissage statistique qui utilise la séparation linéaire. Cela consiste à séparer par un hyperplan des individus représentés dans un espace de dimension égale au nombre de caractéristiques, les individus étant alors séparés en deux classes. Cela est possible quand les données à classer sont linéairement séparables. Dans le cas contraire, les données seront projetées sur un espace de plus grande dimension afin qu'elles deviennent linéairement séparables.

II.4.2.3 Cas de données linéairement séparables :

Considérons « l » points $\{(x_1, y_1), (x_2, y_2), \dots, (x_l, y_l)\}$, $x_i \in R^N$

Avec $i=1 \dots L$ et $y_i \in \{\pm 1\}$

Ces points sont classés en utilisant une famille de fonctions linéaires définies par :

$$\langle w, x \rangle + b = 0 \quad (14)$$

Avec $w \in \mathbb{R}^N$ et $b \in \mathbb{R}$ de telle sorte que la fonction de décision concernant l'appartenance d'un point à l'une des deux classes soit donnée par :

$$f(x) = \text{sgn}(\langle w, x \rangle + b) \quad (15)$$

La fonction (14) représente l'équation de l'hyperplan H. La fonction de décision (15) va donc observer de quel côté de H se trouve l'élément de x.

On appelle la marge d'un élément la distance euclidienne prise perpendiculairement entre H et x. Si on prend un point quelconque t sur H, cette marge peut s'exprimer en :

$$M_x = \frac{w}{\|w\|} (x - t) \quad (16)$$

La marge de toutes les données est définie comme étant :

$$M = \min_{x \in E} M_x \quad (17)$$

L'approche de classification par SVM tend à maximiser cette marge pour séparer le plus clairement possible deux classes. Intuitivement, avoir une marge la plus large possible sécurise mieux le processus d'affectation d'un nouvel élément à l'une des classes.

Un SVM fait donc partie des classificateurs à marge maximale.

Un classificateur à marge maximale est un classificateur dont l'hyperplan optimal séparant deux classes est une solution du problème d'optimisation mathématique suivant (forme primale) :

$$\text{minimiser } \frac{1}{2} \|w\|^2 (\langle w, x \rangle + b) \geq 1 \quad x \in E \quad (18)$$

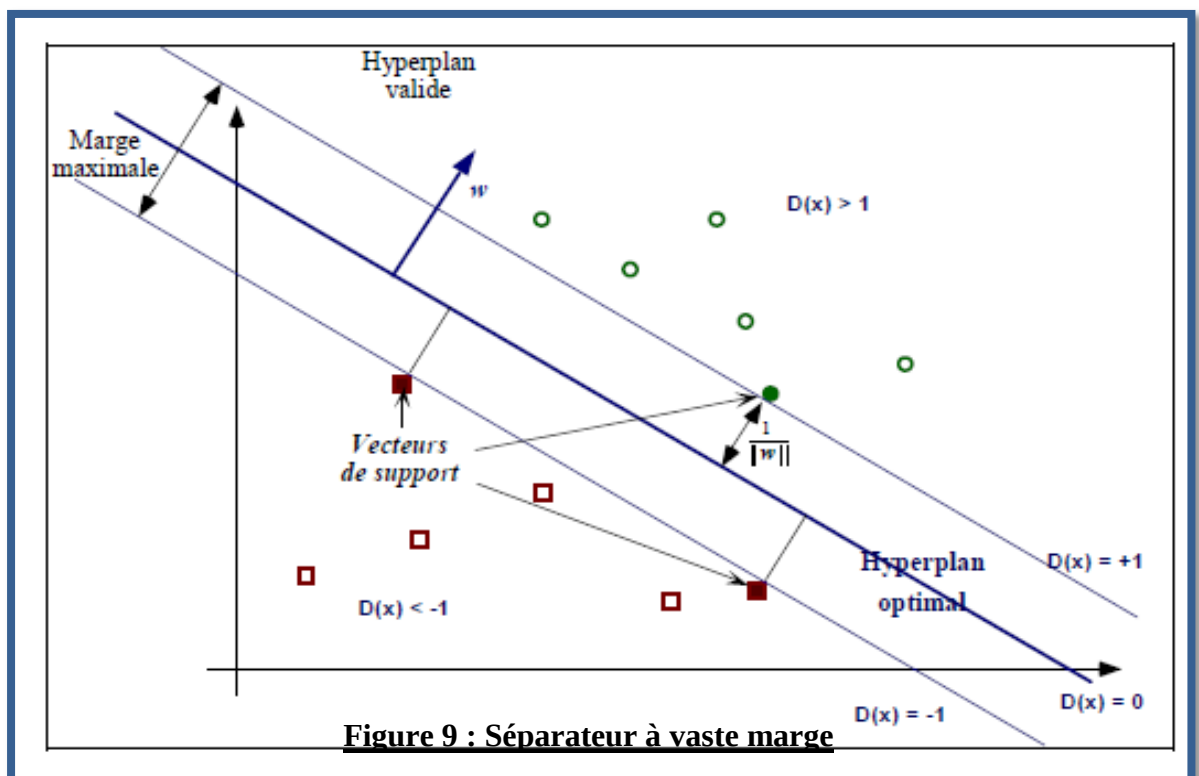
La fonction objective de ce problème est le carré de l'inverse de la double marge qu'on veut maximiser. La contrainte unique correspond au fait que les éléments x doivent être bien placés.

La résolution de ce problème nécessite de fixer les paramètres w et b qui constituent les variables α_i de la machine d'apprentissage.

Les classificateurs à marge maximale donnent de bons résultats lorsque les données sont linéairement séparables.

La tâche de discrimination est de trouver un hyperplan qui sépare deux (ou plus) ensembles de vecteurs (Figure 9).

Pour la détection et la reconnaissance des vertèbres, ces deux classes peuvent être région de vertèbre ou non.



Forme duale :

La formulation primale peut être transformée en formulation duale en utilisant les multiplicateurs de Lagrange. L'équation (18) s'écrit alors sous la forme suivante :

$$L(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^l \alpha_i (y_i (\langle w, x_i \rangle + b) - 1) \quad (19)$$

II.4.2.4 Cas des données non-linéairement séparables :

En pratique, il est assez rare d'avoir des données linéairement séparables. Afin de traiter également des données bruitées ou non linéairement séparable, les SVM ont été généralisées grâce à deux outils : la marge souple (soft margin) et les fonctions noyau (kernel functions).

Le principe de la marge souple est d'autoriser des erreurs de classification. Le nouveau problème de séparation optimale est reformulé comme suit :

$$\left\{ \begin{array}{l} \text{MIN}_{w,b,\epsilon} \frac{1}{2} w^T W + C \sum_{i=1}^l \epsilon_i, C \geq 0 \text{ sous contraintes} \\ y_i (\langle w, x \rangle + b) \geq +1 - \epsilon_i \\ \epsilon_i \geq 0 \text{ pour } i = 0, \dots, l \end{array} \right. \quad (20)$$

Un terme de pénalité est introduit dans la formule (20), Le paramètre C est défini par l'utilisateur. Il peut être interprété comme une tolérance au bruit de classificateur.

Remarque :

L'idée de base pour les données non linéairement séparable, est de projeter l'espace d'entrée (espace des données) dans un espace de plus grande dimension appelé espace de caractéristiques (feature space) afin d'obtenir une configuration linéairement séparable (à l'approximation de la marge souple près) de nos données, et d'appliquer alors l'algorithme SVM.

II.4.3 Architecture d'un classificateur SVM :

II.4.3.1 la fonction noyau :

Afin de résoudre le problème de données non linéaires, la fonction noyau joue le rôle central de liaison des vecteurs d'entrées à l'espace de caractéristiques de grande dimension. (Voir figure 10).

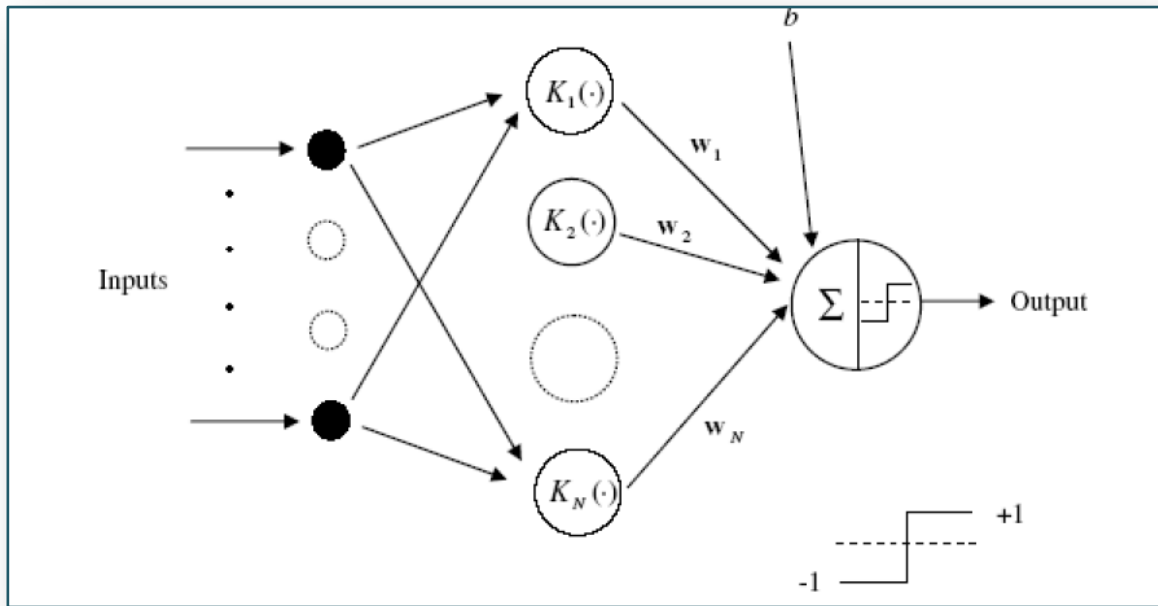


Figure 10 : Architecture d'une machine à vecteurs de support (d'un nombre N)

Les choix typiques pour la fonction noyau sont :

Noyau gaussien à base radiale (ou RBF : Radial Basis Function) :

$$K(x, x_i) = \exp[\gamma \|x - x_i\|] \quad (21)$$

Noyau polynômiale :

$$K(x_i - x_j) = (x_i^T x_j + 1)^p \quad (22)$$

Avec p : une constante qui spécifie le degré du polynôme.

Remarque :

Si $p=1$: la formule (22) devient la fonction d'un noyau linéaire.

Dans cette étude nous avons utilisé le noyau Gaussien, linéaire, et polynômial de degré 4.

II.4.3.2 Sélection de modèle SVM :

Une machine à vecteur de support binaire sépare les exemples positifs des exemples négatifs dans la phase d'apprentissage.

Les multiplicateurs de Lagrange α_i pour chaque machine binaire sont déterminés par

la minimisation de la fonction coût donnée par la formule (23) suivante :

$$P(w) = \frac{1}{2} \|w\|^2 \quad (23)$$

Ce problème est résolu à partir de la forme duale qui est exprimée par :

$$L_D = \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i x_j \quad (24)$$

Cela revient à chercher : $\sum_i \alpha_i y_i = 0$ et $0 \leq \alpha \leq C$

Si la valeur du paramètre de régularisation « C » qui contrôle la tolérance aux erreurs de classification dans la phase d'apprentissage est élevée, donc plus de pénalité sera donnée à l'erreur.

Le vecteur d'apprentissage x_i qui a une valeur de α_i non nul est appelé « **vecteur de support** ».

II.4.3.3 Estimation de l'erreur de généralisation :

La technique la plus populaire pour l'estimation de l'erreur de généralisation est la validation croisée (en anglais cross validation ou Leave One Out (LOO)) qui est utilisée indépendamment de la nature de la machine d'apprentissage utilisée. Le principe de cette méthode consiste à séparer en N sous-ensembles les échantillons de la base d'apprentissage, à apprendre sur N - 1 sous-ensembles, à valider sur le sous-ensemble restant, puis à faire tourner les sous-ensembles de façon à ce que chacun ait pu être testé. Au final, on note le nombre d'erreurs de classification de la procédure LOO par $l(X_1, Y_1, \dots, X_n, Y_n)$.

Il a été démontré que dans [7] que cette procédure donne une estimation presque non biaisée de l'espérance de l'erreur de généralisation.

En effet, l'espérance $E(\cdot)$ de la probabilité p_{err}^{n-1} de l'erreur de test par une machine entraînée à partir de $n-1$ exemples est donnée par la formule (25):

$$E(p_{err}^{n-1}) = \frac{1}{n} E(l(X_1, y_1, \dots, X_n, y_n)) \quad (25)$$

Cependant la procédure est coûteuse en calcul, car nécessitant « n » apprentissages.

Une procédure simple d'estimation d'erreur de généralisation est la validation croisée dite K -fold. Elle consiste à diviser l'ensemble des données en k sous-ensembles mutuellement exclusifs de taille approximativement égale.

L'apprentissage de la machine est effectué en utilisant $k-1$ sous ensemble et le test est effectué sur le sous-ensemble restant. Cette procédure est utilisée une fois pour le test. La moyenne des k taux d'erreur obtenus estime l'erreur de généralisation.

Nous avons montré dans notre travail que la procédure de la validation croisée donne de bonnes estimations de l'erreur de généralisation. Le nombre d'ensembles utilisé dans notre travail est égal à 5 (5fold).

II.5 Détection par SVM:

Notre approche de détection est basée sur l'extraction automatique des fenêtres rectangulaires. Le contenu de ces fenêtres est transformé en vecteurs de caractéristiques en appliquant les méthodes d'extraction citées précédemment. Ces vecteurs sont ensuite injectés dans notre classificateur SVM afin de déterminer si une fenêtre correspond ou pas à une région de vertèbres.

Les fenêtres extraites sont de taille fixe, avec un chevauchement pour parcourir la totalité de l'image, respectant l'algorithme suivant (voir figure 11):

1. Mettre une fenêtre de taille fixe $W_{sub} \times H_{sub}$ dans le coin supérieur gauche de l'image.
2. Extraire une région de l'image dans cette fenêtre.
3. Décaler la fenêtre à droite avec un pas de S_x ($S_x < W_{sub}$).

4. Extraire une nouvelle région de l'image dans la fenêtre. Refaire l'étape 3 et 4 jusqu'à ce que la fenêtre soit au bord droit de l'image.
5. Décaler la fenêtre au bord gauche de l'image.
6. Décaler la fenêtre en bas avec un pas de S_y ($S_y < H_{sub}$).
7. Extraire une nouvelle région de l'image dans la fenêtre.
8. Refaire l'étape 3 à 8 jusqu'à ce que la fenêtre soit au bord droit de l'image.

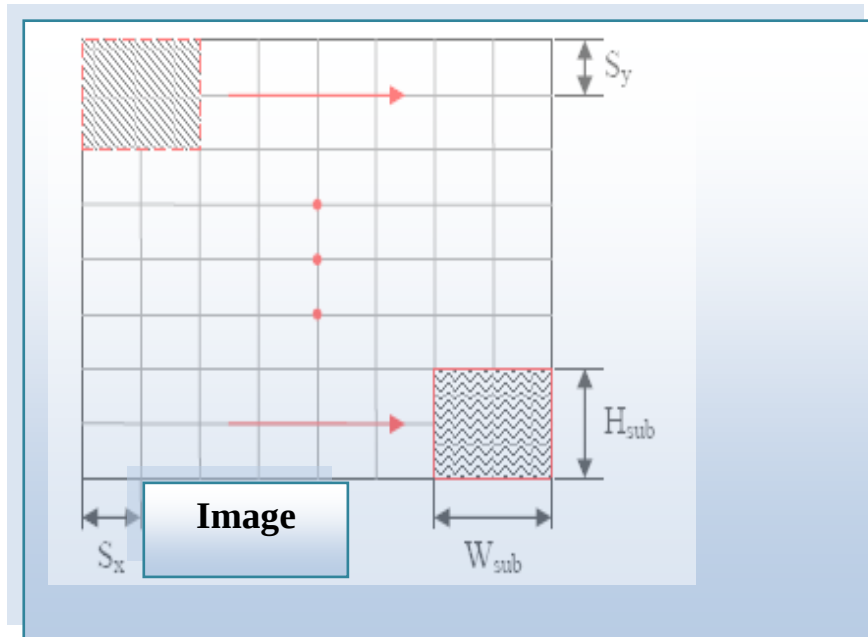


Figure 11 : Processus d'extraction de fenêtres à la phase de détection.

La figure (12) présente ce processus d'extraction de différentes fenêtres et de balayage appliqué à un exemple de notre base :

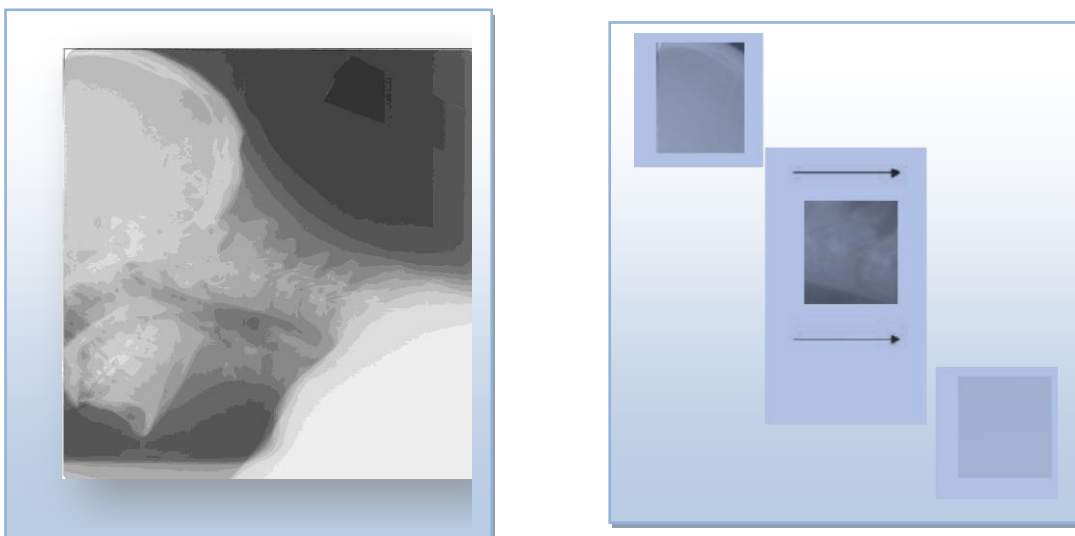


Figure 12 : Exemple de processus d'extraction de fenêtres à la phase de détection.

Post traitement : Afin d'éliminer les fausses réponses du classifieur SVM une méthode d'analyse à base de coefficients de corrélation est utilisée.

II.5.1 Le coefficient de corrélation : est un indice de mesure de l'intensité d'un lien qui peut exister entre deux variables. Le coefficient de corrélation peut prendre une valeur comprise entre -1 et +1. S'il est égal à 0, cela signifie qu'il n'existe aucun lien entre ces 2 variables. Il est très généralement utilisé dans le cadre de l'analyse de variables quantitatives.

Dans notre cas les deux variables représentent la fenêtre de détection qui a été mal classée et un modèle d'une région positive contenant la vertèbre.

Notons A la matrice représentant la fenêtre de balayage et B la matrice représentant notre modèle moyen.

La formule utilisée pour calculer ce coefficient 'r' est :

$$r = \frac{\sum_m \sum_n (A_{mn} - \bar{A})(B_{mn} - \bar{B})}{\sqrt{\left(\sum_m \sum_n (A_{mn} - \bar{A})^2\right) \left(\sum_m \sum_n (B_{mn} - \bar{B})^2\right)}} \quad (26)$$

La valeur du Coefficient de Corrélation est toujours comprise entre -1 et +1

Plus il est proche de 1 ou de -1, plus les points sont proches d'une droite. S'il est égal à ± 1 , les points sont strictement alignés.

II.6 Estimation par RANSAC :

Dans le but du raffinement des résultats de la phase précédente, une méthode d'estimation « robuste » est utilisée afin d'éliminer les fausses détections du classifieur et encore retrouver les vertèbres manquantes.

Un algorithme d'estimation est dit « robuste » s'il garde ses propriétés malgré les incertitudes sur le modèle, les erreurs de mesure et les changements de l'environnement.

La méthode RANSAC est utilisée afin de répondre aux besoins de raffinement.

Nous l'utilisons dans notre travail pour estimer une courbe qui passe par les vertèbres détectées.

II.6.1 L'algorithme RANSAC :

La méthode RANSAC [25] (en anglais Random Sample Consensus) est une méthode de vote probabiliste. Le principe de cette dernière est l'utilisation d'un minimum de données nécessaires pour l'estimation.

Chaque estimation, avec un jeu de données particulier, correspond à un "vote" pour les paramètres obtenus. Le jeu de paramètres élu, i.e. le plus "voté", est retenu comme résultat de l'estimation.

La formulation basique de l'algorithme RANSAC, permet d'estimer les paramètres d'un modèle mathématique à partir d'un ensemble de données observées qui contient des valeurs aberrantes (en anglais « outliers ») illustrés dans la figure 13.

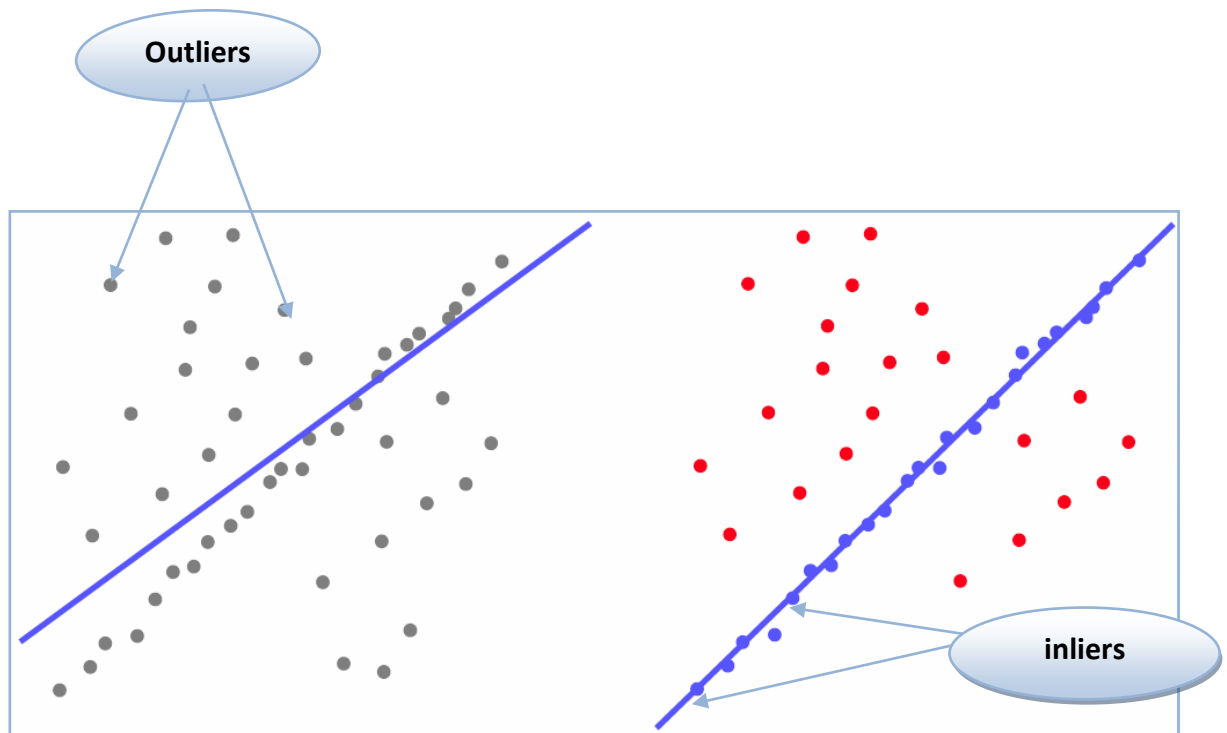


Figure 13 : Exemple d'estimation d'une droite avec RANSAC

Son principe de fonctionnement est le suivant :

L'algorithme effectue d'une façon itérative, un tirage aléatoire d'un échantillon dont la taille est suffisante pour estimer les paramètres d'un modèle mathématique.

Après l'estimation du modèle, le nombre de points « proches » du modèle, est comptabilisé. On parlera dans ce cas des points valables (« en anglais « inliers »).

Le modèles ayant le plus grand nombre de points est choisi et considéré comme étant le meilleur modèle présent dans le jeu de données. La taille de l'échantillon correspond au nombre de points minimum requis pour initialiser le modèle.

Ensuite, l'ensemble des points valables est déterminé via le calcul de la distance qui sépare chaque point du model estimé. Tout point se trouvant à une distance inférieur à une tolérance donnée du modèle initialisé est considéré comme un point valable.

L'algorithme RANSAC de base peut se résumer par les étapes suivantes :

Données : ensemble S d'éléments.

- **Répéter N fois :**
 1. Sélectionner au hasard un sous ensemble de S (s_k)
 2. Instancier le modèle à partir de s_k .
 3. Déterminer l'ensemble S_k contenu dans S qui sont plus près du modèle instancié qu'une distance t (c'est-à-dire pour chaque donnée de S , calculer sa distance du modèle, et si elle est plus petite que t , on le sélectionne pour S_k).
 4. S_k est un ensemble consensus et définit les données régulières de S
 5. Si $|S_k| > T$, on sort de la boucle
- **Retenir $\max|S_k|$ pour tout k , et ré-estimer le modèle à partir de tous les éléments du S_k retenu.**

II.6.2 Sélection des paramètres :

Après « N » itérations de l'algorithme RANSAC, les paramètres retenus sont ceux qui minimisent le nombre d'erreurs, c'est à dire le nombre de points distants du modèle estimé.

Le nombre d'itérations, N , est choisie suffisamment élevé pour faire en sorte que la probabilité p (normalement fixé à 0.99) fait qu'au moins un des ensembles d'échantillons aléatoires S_k ne comprend pas une valeur aberrante (« outliers »).

Supposons ω , la probabilité qu'une donnée soit régulière. Alors $\varepsilon = 1 - \omega$. On tire un échantillon de s données ($|s_k| = s$).

Ainsi, ω^s est la probabilité d'avoir s données régulières, et $1 - \omega^s$ est la probabilité d'avoir au moins une donnée aberrante.

'N' échantillons sont tirés à partir de s données, alors :

$(1 - \omega^s)^N$: Avoir N fois au moins une données aberrante.

$1 - (1 - \omega^s)^N$: Avoir au moins 1 échantillon avec aucune données aberrante.

Ainsi :

Le nombre N peut être calculé par la formule (27) suivante :

$$N = \frac{\log(1 - p)}{\log(1 - \omega^s)} \quad (27)$$

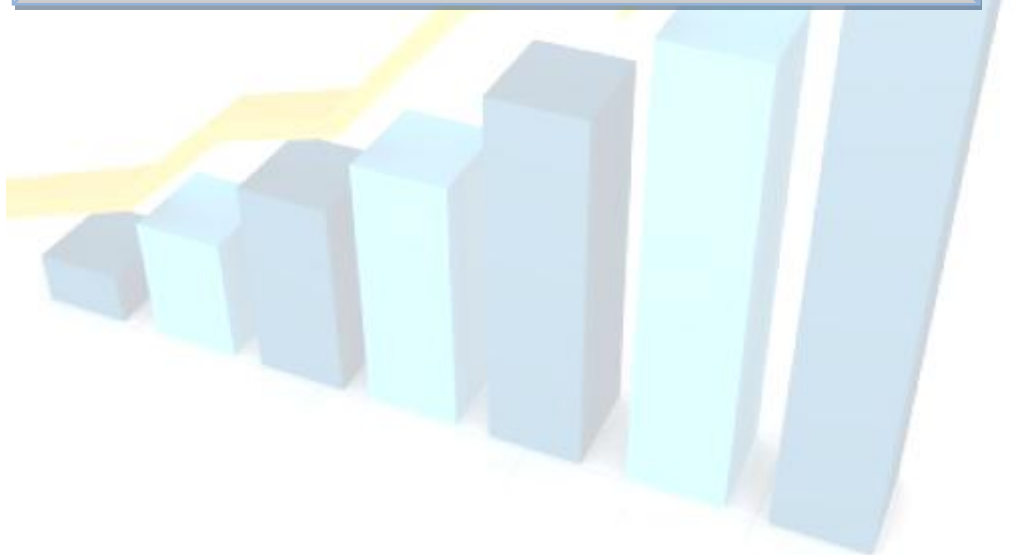
Le paramètre T est posé :

$T = \text{nombre présumé de données régulières dans } S = \omega * |s|$ (règle heuristique).

Dans le chapitre suivant nous allons présenter les résultats de classification des régions de vertèbres par SVM en utilisant la procédure de 5 - fold pour le réglage des différents paramètres de notre modèle, suivie des résultats de test de détection et de raffinements sur différents images par l'algorithme RANSAC.

Chapitre 3:

Tests et résultats



Dans le chapitre précédent nous avons décrit la méthodologie que nous proposons pour réaliser la détection des vertèbres à partir des images de radiographie. Cette méthode utilise une description à base des coefficients d'ondelettes et des moments combinée à une méthode de classification statistique : machines à vecteurs de supports (ou SVM).

Ce chapitre présente les résultats obtenus avec la méthode proposée et qui sont appliqués aux images radiographiques cervicales.

III.1 Base d'images :

La base de données utilisée dans ce travail est constituée d'un ensemble de 100 images radiographique de vertèbres cervicales choisies à partir de la base NLM (National Library of Medicine).

La taille d'une image est de 1752*1462 pixels. Ces images ont été redimensionnées à 250*250 dans la phase de test, afin de réduire le temps d'exécution.

Les ressources :

Les expérimentations et les tests ont été exécutés et réalisés sur un « Intel Core i3 » possédant une mémoire de 3Giga.

Choix de langage MATLAB (version 7.9.0):

MATLAB : est un langage de développement informatique particulièrement dédié aux applications scientifiques « traitement du signal, imagerie, etc. », d'où le choix de ce langage qui nous permettra de traiter les images avec efficacité et rapidité.

MATLAB est doté d'un environnement simple et convivial et contient de nombreuses boîtes à outils (réseaux de neurones, bio-informatique, ondelettes, statistiques, etc.).

III.2 Etape d'apprentissage :



L'ensemble d'apprentissage est formé de la façon suivante :

1. Pour la classe « région vertèbre » (désignée classe 1), un ensemble de 50 sous images est formé par l'extraction de fenêtres de taille 25*25

contenant la vertèbre à partir de chaque image radiographique. Cet ensemble représente les exemples positifs.

2. Un autre ensemble est formé par l'extraction aléatoire des fenêtres de même taille dans des zones différentes dans l'image (autre que la vertèbre), pour la classe « région non vertèbre » (désignée classe2). Cet ensemble représente les exemples négatifs.



Pour chaque exemple de cette base, deux méthodes d'extraction de caractéristiques ont été appliqués :

1. Méthode d'ondelettes : qui consiste à extraire un ensemble de coefficients d'ondelettes comprenant : les détails horizontaux, verticaux, diagonaux et les coefficients d'approximation.
2. Méthode des moments géométriques.

Pour la méthode d'ondelettes les vecteurs caractéristiques (descripteurs) utilisées pour entraîner le classifieur SVM ont les dimensions suivantes :

1. $D=13*13$ pour le premier niveau de décomposition d'ondelette.
2. $D=7*7$ pour le deuxième niveau.
3. $D=4*4$ pour le troisième niveau.

$D=1*7$ est la dimension du descripteur de la deuxième méthode pour les septes moments HU géométriques.



Chaque orientation de détails est entraînée et classifiée séparément en utilisant le classificateur SVM avec différents noyaux y compris le noyau linéaire, Gaussien, et polynomial de degré 04 (quadratique).

Formule noyau polynomial :

$$K(x_i - x_j) = (x_i^T x_j + 1)^p$$

Avec $P > 0$: constante.

(Noyau linéaire $p=1$).

Formule noyau Gaussien (RBF) :

$$K(x, x_i) = \exp[\gamma \|x - x_i\|]$$

Les figures suivantes représentent la visualisation de quelques tests de changements de la fonction noyau SVM sur des images de validation:

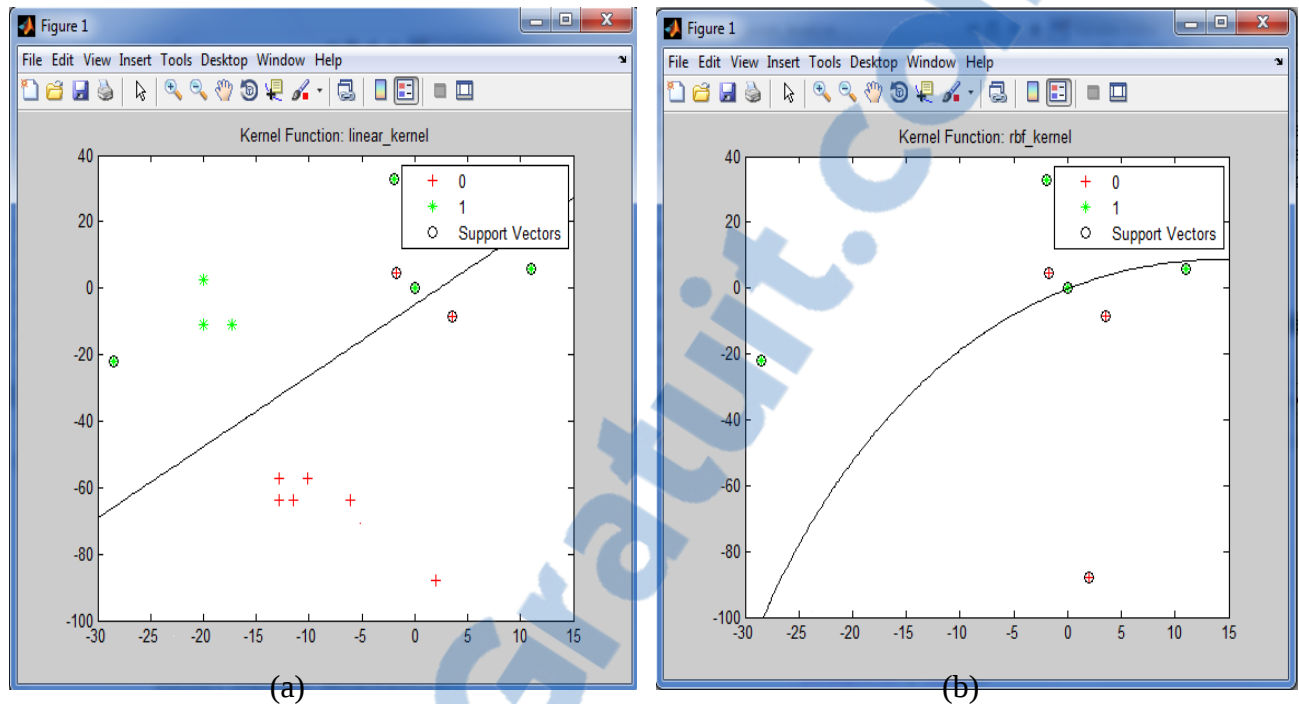


Figure14 : Application d'un noyau linéaire(a) et Gaussien(b) pour l'apprentissage.

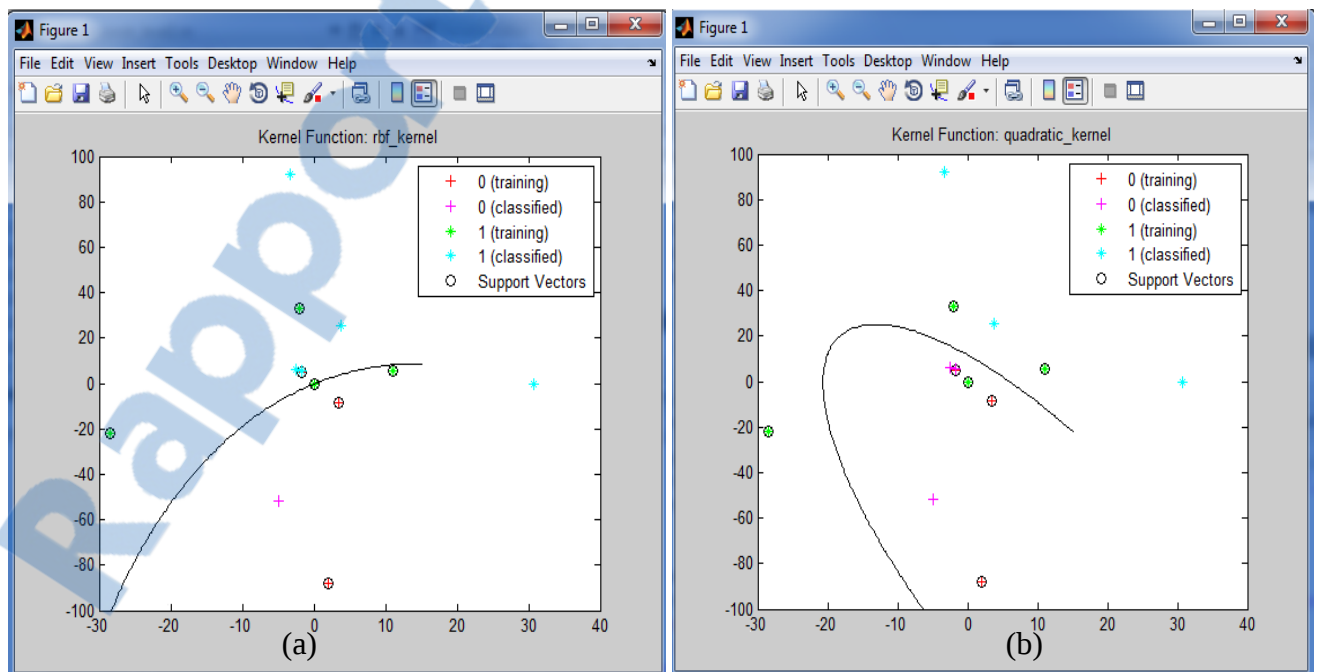


Figure15 : Application d'un noyau Gaussien(a)et quadratique (b) pour la classification.

III.3 Etape de test et validation :

La phase d'apprentissage est la plus importante, mais la détermination des bons paramètres n'est pas toujours facile.

Dans cette étude, nous nous intéressons aux paramètres du classificateur SVM afin d'aboutir à une erreur de classification minimale.

Pour fixer le modèle SVM approprié à la détection de différentes régions vertèbres, nous étudions les performances de trois types de modèles SVM : linéaire, à fonction noyau gaussien (RBF) et à fonction de noyau polynomiale (ordre 4).

Chaque modèle est évalué en utilisant les quatre sous-ensembles de paramètres (y compris les coefficients d'approximation, les détails horizontaux, verticaux, et diagonaux) issus de la méthode d'ondelettes, et un autre ensemble de moments géométriques de HU.

Les différents paramètres qui ont été étudiés durant la phase d'apprentissage représentent : la fonction noyau avec ses paramètres (exemple : la valeur sigma pour la fonction RBF), et le paramètre de pénalisation « C » par la méthode de validation croisée citée dans le chapitre précédent. Plusieurs essais ont été effectués afin de pouvoir fixer ce paramètre de régularisation. En effet ce paramètre contrôle le compromis entre l'erreur de SVM sur les données d'apprentissage et la maximisation de la marge. La figure suivante, Figure 16, illustre cette influence du changement de ce paramètre :

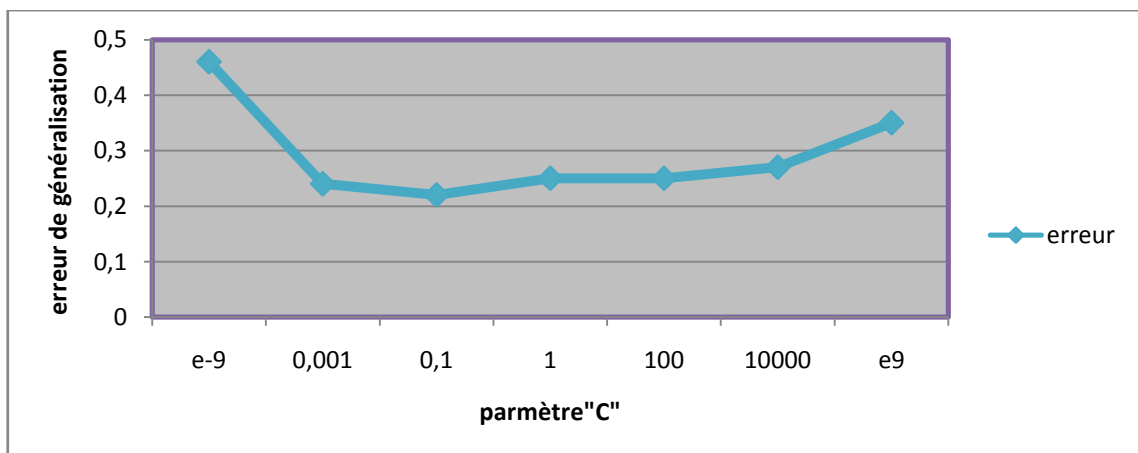


Figure 16 : Influence de paramètre « C » sur l'erreur de généralisation.

Des coefficients de corrélation sont ensuite calculés entre la fenêtre à classer et un modèle moyen (Template) de notre base d'apprentissage, afin d'identifier le degré de ressemblance entre le résultat et la région de vertèbre désirée.

La performance de notre classificateur est déterminée par le calcul de la sensibilité (rappel) et la spécificité (précision) de la classification de la manière suivante :

La sensibilité: est la capacité à détecter les "**vrais positifs**(VP)".

$$\text{Sensibilité} = VP / (VP + FN)$$

La spécificité (précision): est la capacité à éliminer les "**faux positifs**(FN)".

$$\text{Spécificité} = VN / (VN + FP)$$

Le taux total de classification :

$$\text{Taux total} = (VP + VN) / (VP + FP + FN + VN)$$

Dans notre cas :

- ⊕ **VP** ou **vrai positif** représente les exemples classés « région vertèbre » et qui le sont vraiment.
- ⊕ **FP** ou **faux positif** représente les exemples classés « région vertèbre » et qui sont en fait « non vertèbre ».
- ⊕ **VN** ou **vrai négatif** représente les exemples classés « non vertèbre » et qui le sont vraiment.
- ⊕ **FN** ou **faux négatif** représente les exemples classés « non vertèbre » et qui sont en fait « région vertèbre ».

Les résultats obtenus dans le tableau, table 01, ont été calculés en utilisant les différentes représentations possibles par les deux méthodes de caractérisation, avec une valeur de C=0.1 et un Sigma=3 pour le noyau Gaussien (RBF).



<u>Noyau</u>	<i>Lineaire</i>			<i>Quadratique</i>			<i>Gaussien(RBF)</i>		
<i>Taux</i>	Sensiv	Specif	Total	Sensiv	Specif	Total	Sensitiv	Specif	Total
<i>Coeff..H</i>	71.1	53.8	63.7	78.8	53.8	68.1	88.4	71.7	81.3
<i>Coeff.V</i>	51.9	89.7	68.1	84.6	51.2	70.3	80.7	87.1	83.4
<i>Coeff.D</i>	44.2	43.5	43.9	76.9	35.8	59.3	86.5	48.7	70.3
<i>Coeff.Apx</i>	51.9	84.6	65.9	75	82	78.0	73	97.4	83.5
<i>Moments</i>	57.1	98	78.5	71.4	57.4	64.2	71.4	99	85.7

Tableau 1 : Les résultats de classification.

Les différents tests nous donnent de bons résultats pour la représentation des moments pour le noyau Gaussien RBF 85,7%, qui donne aussi 83.5% pour la représentation par les coefficients d'approximation et les coefficients de détails verticaux pour le premier niveau de décomposition.

Pour raffiner plus les résultats issus de la représentation par ondelettes, nous procédons de la même manière pour le deuxième et le troisième niveau de décomposition.

Les meilleurs résultats restent toujours élevés en utilisant le noyau Gaussien avec un $\sigma=3$ et un taux de 84% et 83.2% pour la représentation par les coefficients d'approximation pour le niveau 2 et 3 respectivement.

III.4 Etape de détection :

Dans la figure (17) nous présentons des exemples d'application de notre système de détection de vertèbres dans les images radiographiques. Les images de tests ont été redimensionnées à 250*250 afin de réduire le temps de parcours de l'image.



Figure17 : Quelques résultats de détection sur des images test.(a)-(c)-(e) avant raffinement,(b)-(d)-(f) résultat de détection après raffinement par corrélation.

Les différents tests de détection appliqués sur des nouvelles images (différentes des images de notre base d'apprentissage) montrent une capacité acceptable de notre classifieur à distinguer les différentes régions des vertèbres.

Les résultats de détection ont montrés la puissance qu'a donné la représentation par transformée d'ondelettes au classifieur SVM. La seconde représentation par les moments a donné des résultats très proches des premiers.

Un autre test a été effectué sur une image radiographique de la partie thoracique est visualisé dans la figure18 suivante :

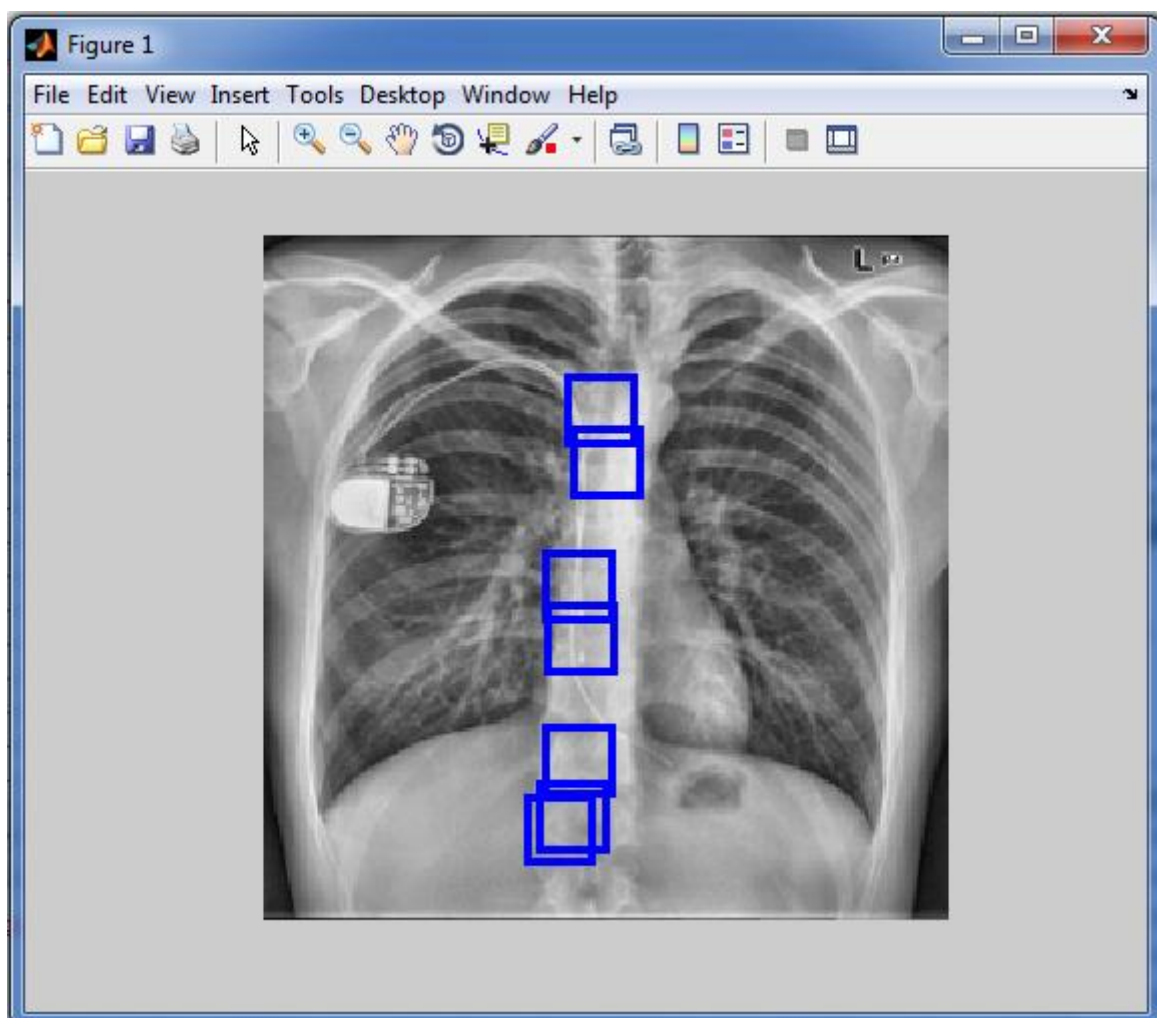


Figure 18 : Résultat de détection dans une image thoracique.

III.5 Raffinement des résultats de détection :

Les résultats de la phase de détection par SVM sont plus ou moins satisfaisants mais il reste encore des vertèbres non classés. Ce problème est réglé en utilisant la méthode d'estimation RANSAC afin d'éliminer les fausses détections du classifieur et aussi retrouver les vertèbres manquantes.

1. Elimination des fausses détections :

Une estimation du modèle est la première étape de l'algorithme pour les points valables (inliers).

Dans ce travail l'ensemble des points S représentent les centroïdes des régions détectées par le classificateur SVM.

Le modèle choisie dans notre cas est le modèle linéaire qui nécessite un minimum de deux points. Cela revient à estimer les coefficients de la droite passant par ces deux points.

Ces paramètres sont conservés comme « modèle élu » et l'algorithme est réitéré jusqu'à trouver le modèle qui passe par le maximum de points valables.

Ces points sont ensuite comparés par rapport à leur distance au modèle et les points les plus proches sont retenus.

La figure 19 suivante présente un test de l'algorithme RANSAC pour l'élimination des points aberrants (c.à.d. qui n'appartiennent pas à la courbe ou supérieur à une valeur seuil).

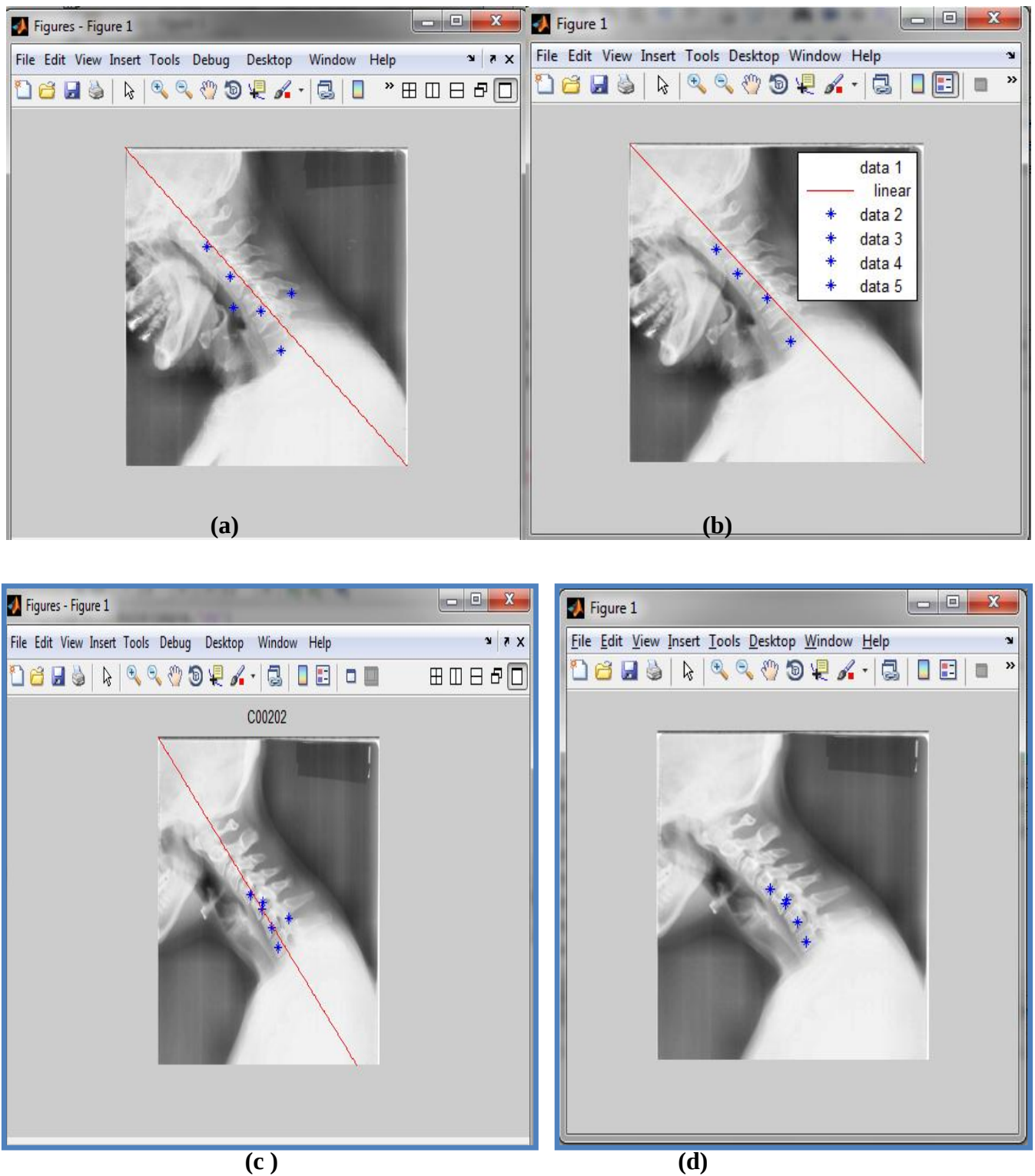


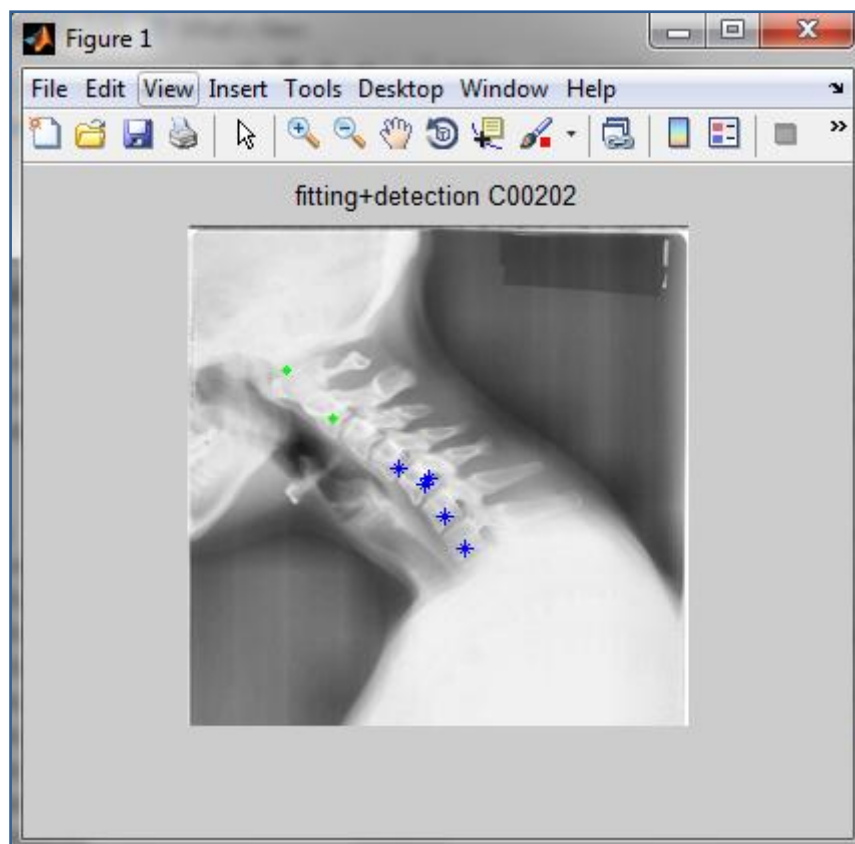
Figure 19 : Résultat de raffinement des résultats détection par RANSAC.(a), (c) Avant élimination. (b), (d) après élimination des outliers.

2. Recouvrement des vertèbres manquantes :

Cette étape présente en effet un deuxième balayage de notre détecteur SVM sur l'image test, mais cette fois ci le parcours doit se faire juste dans la partie de la courbe estimé par l'algorithme RANSAC. Ensuite poursuivre les mêmes étapes appliquées lors de la première passe.

Les figures 20 et 21 suivantes présentent un test de recouvrement par le détecteur SVM en utilisant la courbe estimée par RANSAC sur des images radiographiques.

Les points verts dans l'exemple représentent les nouvelles régions vertèbres détectés.



**Figure 20 : Résultat de raffinement des résultats détection par RANSAC
de l'image test « C00202 ».**

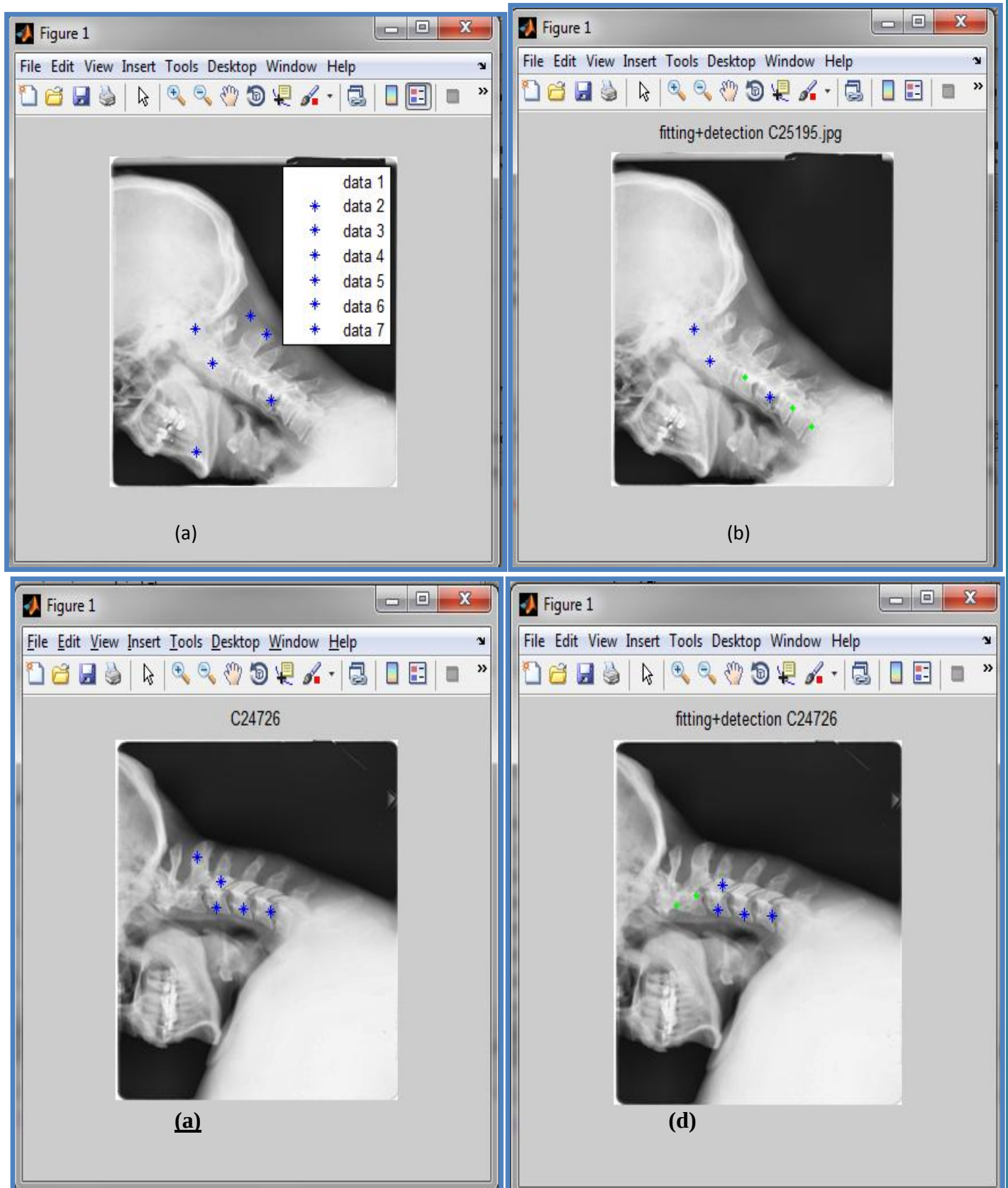


Figure 21 : Résultat de raffinement des résultats détection par RANSAC

**(a) , (c) Avant élimination, (b), (d) après élimination des outliers et couvrent
des vertèbres manqués.**

Les résultats de raffinement par l'algorithme d'estimation RANSAC présenté dans la figures 20 et la figure 21 (b) et (d) montrent un certain renforcement du classificateur SVM pour la détection des régions vertèbres manquées, et cela est dû à la robustesse d'estimation de l'algorithme RANSAC d'une part et à l'efficacité de la classification parla méthode SVM d'une autre part.



Conclusion

Conclusion :

Dans cette étude nous avons présenté une méthode d'analyse qui peut être incluse dans un système d'aide au diagnostic médical. Notre objectif était d'examiner l'application de la méthode des machines à vecteurs de supports pour la détection des régions de vertèbres dans les images radiographiques.

Parmi les raisons du choix de l'approche de classification par SVM est que cette dernière est mathématiquement bien établie, avec la mise en œuvre du principe de risque structurel qui a pour objectif la minimisation d'une limite supérieure pour l'erreur de généralisation. Une autre raison est la simplicité de sélection des paramètres qui représentent la fonction noyau et la marge « C ».

Les résultats de la représentation par ondelettes prouvent que cette dernière peut fournir des informations fréquentielles significatives concernant la classe des régions de vertèbres à détecter.

Un taux de 83.5% a été atteint en utilisant un modèle de classificateur SVM avec un noyau Gaussien entraîné par les coefficients d'approximations de la transformée en ondelette.

Ce même modèle adonné un taux de 85.7% pour la méthode basée sur les moments de HU qui a prouvé aussi son efficacité grâce à la richesse de leurs informations sémantiques.

Les résultats de la phase de détection sont renforcés par l'utilisation de l'algorithme RASAC qui a prouvé son efficacité en combinaison avec la méthode SVM.

Les résultats expérimentaux obtenus sont prometteurs et permettent d'automatiser complètement le processus de segmentation.

Nous estimons avoir atteint l'essentiel de notre objectif, et nous souhaitons que notre réalisation servira comme un outil d'automatisation de la procédure de sélection des vertèbres qui a été citée comme faiblesse des travaux de la segmentation des images rayons-X.

Perspectives et travaux futurs :

Ce travail est amené à être prolongé dans plusieurs directions :

- L'utilisation d'une grande base d'image dans l'apprentissage pour augmenter les performances de la classification.
- L'ajout ou la combinaison d'autres méthodes d'extraction de caractéristiques pour mieux présenter les propriétés de la vertèbre, afin d'améliorer le taux de classification.
- Introduction de la notion de Boosting issue de la combinaison de classificateurs pour l'optimisation de performances.
- Une future exploration des données va inclure une description étendue des classes analysées par des paramètres concernant la forme de vertèbre (par exemple les descripteurs de Fourier)
- Une étude comparative d'autres méthodes d'apprentissage statistique peut amener à sélectionner le meilleur modèle approprié à la détection.

[1] A. Tezmol, H. Sari-Sarraf, S. Mitra, R. Long. "Customized Hough Transform for Robust Segmentation of the Cervical Vertebrae from X-Ray Images". *SSIAI2002*, Proc. IEEE 1537, pp.224-228, 2002.

[2] Benjamin M. HOWE, B.S., B.S.E.E., "Segmentation of cervical and lumbar vertebrae in X-ray images using active appearance models and extensions », Thèse de Master, Université Texas tech, 2003.

[3] Campanini, R et al, "A novel approach to mass detection in digital mammograms". In Proc. of the 6th Int. Workshop on Digital Mammography ,2002

[4] Chappelle.O, Haffner.P, Vapnik VN, Support vector machines for histogram-based image classification. *IEEE Transaction on Neural Networks* **10:1055-1064, 1999.**

[5] Cootes TF, Edwards GJ, Taylor CJ. "Active appearance models ». *IEEE Transactions on Pattern Analysis and Machine Intelligence.*; vol.23(6):p.p681–685,2001.

[6] Cootes TF, Taylor CJ, Cooper DH, Graham J. "Active shape models-their training and application". *Computer Vision and Image Understanding*. **vol.61 (1): p.p38–59. 1995.**

[7] Cortes C, Vapnik VN; support vector networks Machine Learning, 1995

[8] Drucker N, Donghui W, Vapnik VN Support vector machines for spam categorization. *IEEE Transaction on Neural Networks* 10:1048-1054,**1999.**

[9] E. McCloskey, T. Spector, K. Eyres et al. "The assessment of vertebral deformity: a method for use in population studies and clinical trials". *Osteoporosis Int* 3, pp. 138-147, 1993.

[10] El-Naqa, I., Yongyi Yang, Wernick, M.N., Galatsanos, N.P., Nishikawa, R.: Support vector machine learning for detection of microcalcifications in mammograms. In: Biomedical Imaging, Proceedings. 2002 IEEE International Symposium on, pp. 201--204, 2002.

[11] G. Zamora, H. Sari-Sarraf, S. Mitra, R. Long, "Estimation of Orientation and Position of Cervical Vertebrae for Segmentation with Active Shape Models," *Medical Imaging 2001: Image Processing Proc. SPIE* 4322, 378-**387**, 2000.

[12] G. Zamora-Camarena. "Automatic segmentation of vertebrae from digitized x ray Images". Thèse de doctorat de type *Ph.D*, Université Texas Tech, 2002.

[13] Gilberto Zamora, Hamed Sari-Sarraf and L. Rodney Long, "Hierarchical segmentation of vertebrae from x-ray images", Proc. SPIE Digital Library, vol. 5032(631), 2003.

[14] Grossman .A, J. Morlet, "Decomposition of Hardy Functions into Square Integrable Wavelets of Constant Shape". *SIAM J. of Math. Anal.* vol. 15, no. 4, pp.

723-736, July 1984.

[15] H. Genant, C. Wu, C. van Kuijk et al. « Vertebral fracture assessment using a semi-quantitative technique ». *J Bone Miner Res* 8, pp. 1137-1148, 1993.

[16] Hua.S, Sum.Z, A novel method of protein secondary structure prediction with high segment overlap measure: support vector machine approach. *Journal of molecular Biology* 308:397-407,2001

[17] Kumar.R, Kulkarni.A, Jayaraman VK, Kukarni Bd Symbolization assisted SVM classifier for noisy data. *Pattern recognition Letters* 25:495-504, 2004.

[18] L. R. Long, G. Thoma, "Identification and Classification of Spine Vertebrae by Automated Methods;" *Medical Imaging 2001: Image, Proc. SPIE* 4322, 1478-1479, 2001.

[19] L. R. Long, G. Thoma, "Segmentation and Image Navigation in Digitized Spine Xrays," *Medical Imaging 2000: Image Processing Proc. SPIE* 3979, 169-175, 2000.

[20] Laine. A and Fan, Texture classification by wavelet packet signatures. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. vol. 15, n° 11, p. 1186-1191, 1993.

[21] Leandro A. Silva, Emilio Del-Moral-Hernandez, Ramon A. Moreno and Sérgio S. Furuie, "Combining wavelets transform and Hu moments with self-organizing maps for medical image categorization", *J. Electron. Imaging* 20, 043002 ,doi:10.1117/1.3645598, 2011

[22] M. Benjelloun and S. Mahmoudi, "Semi-automatic vertebra segmentation," *Handbook of Research on Developments in E-Health and Telemedicine: Technological and Social Perspectives*, pp. 110–124, 2010.

[23] M. Benjelloun and S. Mahmoudi , Spine localization in x-ray images using interest point detection, *Journal of Digital Imaging*, vol. 22, pp. 309–318, 2009.

[24] M. Benjelloun, H. Tellez, and S. Mahmoudi, "Vertebrae edge detection using polar signature," in *International Conference on Pattern Recognition*, pp. 476–479, 2006,

[25] M.A. Fischler and R.C. Bolles, Random sample consensus: A paradigm for model fitting with applications to image analysis and automated cartography. *CACM*, Vol. 24, No 6, pp.381–395, 1981.

[26] M.G.roberts,T.Coots and J.EAdams."Automatic segmentation of lumbar vertebrae on digitized radiographs using linked active appearance models". *Proc.medical Image understanding and analysis*, vol.1, p.p120-124, 2006.

[27] M.Kass,A.P Witkin and D. Terzopoulos." Snakes: Active contour models". In *1CCV87*, p.p259-268, 1987.

[28] M.Kueui .HU.Visual Pattern Recognition by moments invariants .*IRE Transaction on Information Theory* , **Volume 8,n°2:179-187,1962.**

[29] M.Sonka ,V.Hlavac, R.Boyle."Image **processing, Analysis** and Machine Vision ".PWS Publishing , **seconde edition ,1999.**

- [30] Mallat .S , Wavelet Tour of Signal Processing,– Academic Press, London, 1998.
- [31] Mallat .S, « Une exploration des signaux en ondelettes », Les Editions de l'Ecole Polytechnique, 2000.
- [32] Mallat S, "Theory of multiresolution signal decomposition: the wavelet representation" , IEEE Trans. Pattern Analysis and Machine Intelligence, vol. 11, no. 7, pp.674-693, July 1989.
- [33] Meyer Y, « Les ondelettes : Algorithmes et applications », Armand Colin .Paris, p. 18-23, 1994.
- [34] Meyer Y, Ondelettes et Opérateurs. Hermann, Paris, 1990.
- [35] Ming-Kuei Hu. Visual pattern Recognition by Moment Invariants. IRE Transactions on **Information Theory**, **IT-8:179-187**, 1962.
- [36] Mukkamala.S Sung AH, Abraham.A Intrusion detection using an ensemble of intelligent paradigms. Journal of Network and Computer Applications, In Press, **2004**.
- [37] Norinder.U Support vector machine models in drug design: application to drug transport processes and QSAR using simplex optimisation and variable selection. Neuro computing **55:337-346,2003**.
- [38] R. Eastell, S. Cedel, H. Wahner et al. "Classification of vertebral fractures". *J Bone Miner Res* 6(3), pp. 207-215, 1991.
- [39] Russo.G, Zegar.C, Giordano advantages and limitations of microarray technology in human cancer.oncogene.22:6497-6507, 2003.
- [40] S.-H. Huang, Q.-J. Wu, and S.-H.Lai, Improved AdaBoost-based image retrieval with relevance feedback via paired feature learning, Multimedia Systems, Vol. 12, pp. **14-26, 2006**.
- [41] Said mahmoudi, Medjeded merati." Une nouvelle approche de segmentation d'images médicales par combinaison des deux Modèles Déformables : ASM et Snakes. Thèse de magistère, Université D'IBN KHALDOUN – **Tiaret, 2009**
- [42] Slonim.Dk from pattern to path ways: gene expression data analysis comes of age. Nature Genetic suppl. **32:502-508,2002**.
- [43] Shouhei Hanaoka,et al , "Whole vertebral bone segmentation method with a statistical intensity -shape model based approach", Proc. SPIE 7962, 796242 (2011).
- [44] Strauss DJ, Steidl G, Hybrid wavelet-support vector classification of wave forms. JCompt and Appl 148:375-400. ,2002.

[45] Szu-Hao Huang, Yi-Hong Chu, Shan-Hong lai: A statistical learning Approach to vertebra detection and segmentation from spinal MRI, *IEEE Transact Med Imaging*, May 2007.

[46] Tobias Kinder, Jorn Ostermann et al,: **automated model based vertebrae detection, identification, and segmentation**, *Medical Image analysis* 13:471-482 ,2009.

[47] Unser. M, A. Aldroubi and A. Laine eds, Special issue on "Wavelets in Medical Imaging", *IEEE Transactions on Medical Imaging*, vol.22, no. 3, mars 2003.

[48] Van GT ,Suykens Jak , Baestaens De,Lambrech.A,Lanckriet G , vandaele.B , Vandewalle J Financial time series prediction using least squares SVM within the evidence framework.*IEEE Transactions on neural networks* 12:809-821,2001.

[49] Vapnik VN the nature of statistical learning theory springer – Verlag New York 1995.

[50] VapnikVN statistical learning theory. Wiley New York ,1998.

[51] Yalin Zheng, Mark S Nixon, Robert Allen, 'Automatic Lumbar Vertebrae Segmentation in Fluoroscopic Images via Optimized Concurrent Hough Transform", *23rd Annual International Conference of the IEEE Engineering in Medicine and Biology Society. Vol.3*, p2653-2656, 2002.

Abstract:

In this work, a support vector machine classification method is introduced for vertebrae detection of X-rays images. This classifier is a part of a computer aided diagnosis system that assists radiologists in vertebrae part selection and for the analysis of the vertebral mobility analysis.

The proposed decision making process was performed in two stages. The first consists of feature extraction by computing the wavelet coefficients and moments shape descriptors. The second stage presents a classification method using a support vector machine (SVM) classifier trained with the extracted features.

Supervised learning used in this work allows to classify vertebra regions by using SVM method, which is a learning machine approach based on statistical learning theory.

Preliminary tests on enhanced X-ray images showed a rate of 83.5% classification accuracy for wavelet representation and over than 85 % for moments representation by using the SVM with Radial Basis Function (RBF) kernel.

Also, a cross validation method and a sensitivity/specificity analysis with different kernel types were used to show the classification performance of SVM classification process used for vertebrae region detection.

Key words: Support vector machine, Wavelet transform, medical images analysis, classification.

Résumé:

Dans ce mémoire, nous proposons une méthode de détection automatique de vertèbres dans les images à rayons-X par application de la méthode machines à vecteurs de support (SVM). Le classificateur proposé peut être intégré à un système d'aide au diagnostic permettant aux radiologistes la sélection des régions vertèbres.

Le processus de traitement passe par deux étapes consistant à l'extraction de caractéristiques et la classification.

La première étape consiste à caractériser les régions des vertèbres par l'utilisation des coefficients d'ondelettes caractérisant les fréquences d'image et les moments comme descripteur de forme. Ces descripteurs ont été introduits à un classifieur SVM dans la deuxième étape de classification.

La méthode d'apprentissage utilisée, SVM, est basée sur la théorie d'apprentissage statistique supervisé et utilisée dans ce travail pour la classification et la détection des régions vertèbres.

Les tests préliminaires montrent un taux de classification de 83.5% pour la représentation par ondelettes et plus de 85% pour la représentation des moments en utilisant un classificateur SVM avec un noyau RBF.

Les performances de notre classifieur sont examinées par une méthode de validation croisée et une analyse de sensibilité/spécificité pour la détection des régions vertèbres.

Mot clés : Machines à vecteurs de support, transformée en ondelettes, imagerie médicale, classification.

ملخص:

في هذا العمل، نقوم بتطبيق طريقة "جهاز الدعم الناقل" للكشف عن الفقرات في صور الأشعة السينية. هذا المصنف هو جزء من نظام التشخيص بالحاسوب، الذي يساعد أطباء الأشعة في تعيين واختيار الفقرات لتحليل حركة العمود الفقري.

تتم هذه العملية بمرحلتين: الأولى تتمثل في استخراج الميزات من خلال معاملات الموجات ومميزات الشكل اللحظية، أما المرحلة الثانية فهي التصنيف باستخدام آلات النقل الداعم التي تعتمد على نظرية التعلم الإحصائية.

أظهرت النتائج الأولية نسبة تصنيف بمعدل 83.5% بالنسبة للموجات وأكثر من 85% بالنسبة للحظات، باستخدام نواة ذات الأساس الشعاعي لجهاز الدعم. كما استخدم أسلوب التحقق من الصحة عبر تحليل الحساسية / الخصوصية مع أنواع نواة مختلفة لإظهار أداء التصنيف للكشف عن منطقة الفقرات.

الكلمات المفتاحية: آلات الدعم الناقل، تحويل الموجات، التصوير الطبي، التصنيف.