



Les Forêts Aléatoires en Apprentissage Semi-Supervisé (Co-forest) pour la segmentation des images rétinienne

Settouti Nesma, Bechar Mohammed El Amine

► To cite this version:

Settouti Nesma, Bechar Mohammed El Amine. Les Forêts Aléatoires en Apprentissage Semi-Supervisé (Co-forest) pour la segmentation des images rétinienne. [Rapport de recherche] Biomedical Engineering Laboratory, Tlemcen University Algeria. 2015. hal-01294064

HAL Id: hal-01294064

<https://inria.hal.science/hal-01294064>

Submitted on 26 Mar 2016

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Les Forêts Aléatoires en Apprentissage Semi-Supervisé (Co-forest) pour la segmentation des images rétinienne

Nesma Settouti^{1,2} et Mohammed Amine Bechar¹

¹ Biomedical Engineering Laboratory, Tlemcen University Algeria

² LIMOS, CNRS, UMR 6158, 63173, Aubiere, France

Table des matières

1	Objectifs	2
2	État de l'art du domaine	2
2.1	Segmentation des images rétinienne	2
2.2	Classification pixellaire semi-supervisée des images médicales	6
3	L'approche proposée	8
3.1	Méthodes d'extraction des caractéristiques	10
3.1.1	Les espaces couleurs	10
3.2	Calcul des points d'intérêt	11
3.3	Méthodes de classification	12
3.3.1	Arbre de décision	12
3.3.2	Méthode d'auto-apprentissage "SETRED"	12
3.3.3	Méthode d'ensemble : Forêts Aléatoires	14
3.3.4	Méthode des Forêts Aléatoires en apprentissage semi-supervisé "co-Forest" . .	14
4	Résultats et expérimentations	17
5	Conclusion	21

1 Objectifs

Le glaucome est une maladie due à la neuro-dégénérescence du nerf optique qui conduit à la cécité. Elle peut être évaluée en surveillant la pression intra-oculaire (PIO), par le champ visuel et l'aspect de la papille optique (rapport Cup sur Disque). Le glaucome augmente le taux de Cup/Disque (CDR), ce qui affecte la perte de vision périphérique.

La segmentation des images rétiniennes, en particulier les fonds d'œil, est une étape importante dans le suivi médical du glaucome. Nous proposons dans ce rapport de recherche une méthode de segmentation et de reconnaissance automatique des régions Disques et Cups pour la mesure du rapport CDR dans un contexte semi-supervisé.

L'intervention d'un expert ophtalmologue est importante dans l'identification des régions Cups et Disques en image rétinienne. Dans ce travail, l'expert réalisera un fenêtrage sur 5% de la base (3 images rétiniennes). Cette approche contribue à automatiser la segmentation des régions du glaucome par les techniques intelligentes.

De ce fait, une étude comparative de plusieurs techniques est proposée. Le principe repose sur une croissance de région en classifiant les pixels voisins à partir des pixels d'intérêt de l'image par apprentissage semi-supervisé. Les points d'intérêt sont détectés par l'algorithme Fuzzy C-means (FCM).

Quatre classifieurs de principe différents sont appliqués dans ce travail : Arbre de décision (mono-classifieur en supervisé), Forêt aléatoire (méthode d'ensemble en supervisé), SETRED (méthode d'auto-apprentissage en semi-supervisé) et l'algorithme *co-Forest* (méthode d'ensemble en semi-supervisé). Cette étude permettra d'adapter la meilleure approche pour la segmentation des images rétiniennes avec une intervention minimale de l'expert.

Ce rapport est organisé de la manière suivante : dans la section 2, une revue de quelques méthodes de segmentation et classification d'images en semi-supervisé est effectuée. Nous exposons ensuite dans la section 3, le processus général de notre approche proposée et ses différentes étapes (caractérisation, pré-traitement et méthodes de classification SSL). Nous validons notre approche et les choix que nous avons réalisés par une phase d'expérimentation. Nous montrons la capacité de notre méthode à la segmentation automatique par l'application de plusieurs méthodes. Finalement, nous terminons par une conclusion qui présente une synthèse des contributions apportées.

2 État de l'art du domaine

2.1 Segmentation des images rétiniennes

Le glaucome est une maladie de l'œil liée à l'augmentation anormale de la pression du liquide oculaire. Cette pression anormale conduit de façon progressive et le plus souvent sans douleur à une atteinte irréversible de la vision. Elle peut être évaluée en surveillant la pression intra-oculaire (PIO), par le champ visuel et l'aspect de la papille optique (rapport Cup sur Disque).

La valeur CDR augmente avec l'augmentation de la dégénérescence neuro-rétinienne et la vision est perdue complètement à la valeur de CDR= 0,8. Plusieurs procédés d'extraction de caractéristiques à partir d'images du fond d'œil sont rapportés dans la littérature [1, 2, 3, 4, 5, 6]. Les techniques décrites dans la littérature pour la localisation du Disque optique sont généralement destinées à identifier soit le centre approximatif du Disque optique ou de placer le Disque dans une région spécifique, comme un cercle ou un carré.

Lalonde et al., utilisent le détecteur de Canny [3], Ghafar et al. [4], la transformée de Hough pour détecter le Disque optique (DO). Bock et al. [2], font appel au concept de l'analyse en composantes principales (ACP), bitsplines et l'analyse de Fourier pour l'extraction de caractéristiques et au séparateur à vaste marges (SVM) comme classifieur pour la prédiction du glaucome. Cette méthode atteint une précision de 88%.

Acharya et al. [7] ont développé un système de détection du glaucome automatisé en combinant la texture et les spectres d'ordre supérieur (HOS). Naive bayes, SVM, Random forest et l'optimisation minime séquentielle sont utilisés pour effectuer la classification supervisée. Après la normalisation z-score et la sélection des paramètres, les résultats révèlent que les caractéristiques de texture et HOS combinées avec le classifieur Random forest, effectuent de meilleures performances que les autres classifieurs. Cette méthode diagnostique correctement les images de glaucome avec une précision de 91%.

L'analyse automatique d'image rétinienne est en train de devenir un outil de dépistage important pour la détection précoce des maladies oculaires. L'examen manuel du Disque Optique (OD) est une procédure standard utilisé pour détecter le glaucome.

Le meilleur moyen de contrôler la maladie du glaucome est à l'aide de l'appareil rétinographe numérique. Ces images sont stockées en format RGB, qui est divisé en trois canaux rouge, vert et bleu. D'autres travaux se sont intéressés aux techniques de traitement d'image pour diagnostiquer le glaucome en se basant sur l'évaluation CDR d'images couleurs du rétinographe.

Madhusudhan et al. [8] ont développé un système pour le traitement et la classification automatique des images basé sur la pratique habituelle dans la routine clinique. De ce fait, trois différentes techniques de traitement d'image à savoir multi-seuillage, méthodes de segmentation basée sur la région de contours actifs sont proposées pour la détection du glaucome.

Mohammad el al. [9] présentent une approche qui comprend deux étapes principales. En premier lieu, un procédé de classification de pixels pour identifier les pixels qui peuvent appartenir à la limite de la papille optique. En second lieu, une mise en correspondance du gabarit circulaire pour estimer le rapprochement de la limite circulaire du Disque optique. Les caractéristiques du pixel utilisé sont basées sur la texture, calculée à partir des différences d'intensité des image locales. Le C-Moyens floues (FCM) et Naive Bayes sont utilisées pour regrouper et classer les pixels de l'image.

Chandrika et al. [10] ont adopté une technique d'identification automatique du Disque optique des images rétiniennes par le calcul du ratio CDR. En premier lieu, un seuillage est appliqué puis la segmentation d'images est effectuée en utilisant k-means et la transformée de Gabor en onde-

lettes. En second lieu, le lissage du contour du Disque et du Cup optique est effectué en utilisant différentes caractéristiques morphologiques.

La détection précoce du glaucome est essentielle pour réduire au minimum le risque de perte de vue. Il a été démontré qu'un bon prédicteur du glaucome est le rapport tasse à Disque de la tête du nerf optique. [11] présente une méthode automatisée pour la segmentation du Disque optique. Notre approche utilise la sélection de fonction de pixel pour former un ensemble de fonctionnalités afin de reconnaître la région du Disque. La classification pixellaire est utilisée pour générer un tableau de probabilité du Disque. Une nouvelle fonction de coût est élaborée afin de maximiser la probabilité de la région à l'intérieur du Disque.

La segmentation de l'image est effectuée en utilisant un nouvel algorithme de recherche par graphe capable de détecter la frontière et de maximiser la probabilité du Disque. La combinaison de la recherche graphique et la classification de pixels permet d'incorporer des ensembles de fonctionnalités dans la conception de la fonction de coût, ce qui est essentiel pour la segmentation du Disque optique. Les résultats sont validés sur 82 ensembles de données de référence et comparés aux segmentations manuelles de trois cas glaucomateux.

Dans [12], une méthode de segmentation du Disque optique basée sur la classification pixellaire est proposée. Dans la classification, l'histogramme du contraste et les statistiques du centre à partir de cartes de différence sont proposés comme caractéristiques pour déterminer si chaque superpixel sera considéré comme région Disque ou non. Dans l'étape d'apprentissage, le bootstrapping est adopté pour traiter la question du déséquilibre des clusters dû à la présence d'une atrophie péri papillaire.

Un score de fiabilité d'auto-évaluation est calculé pour évaluer la qualité de la segmentation du Disque optique automatisé. La méthode proposée a été testée sur une base de données de 650 images avec le contour du Disque optique marqué manuellement par des professionnels. Les résultats expérimentaux montrent un chevauchement avec une erreur moyenne de 9,5%, améliorant ainsi les procédés antérieurs.

D'une manière plus classique, Hatanaka et al. [13] proposent une méthode pour mesurer le rapport Cup/ Disque avec un profil vertical sur le Disque optique. Le bord du Disque optique est ensuite détecté par l'utilisation d'un filtre de détection de bord de Canny. Le profil obtenu est réalisé autour du centre du Disque optique dans la direction verticale. Par la suite, le bord de la zone du Cup sur le profil vertical est déterminé par une technique de seuillage.

Joshi et al. [14] mettent en œuvre une technique automatique pour le paramétrage du Disque optique (OD) en fonction des régions Cup et Disque segmentées obtenues à partir des images rétinienne monoculaires. Une nouvelle méthode de segmentation du OD est proposée, elle intègre les informations de l'image locale autour de chaque point d'intérêt dans l'espace de fonction multidimensionnelle. Cela dans le but de fournir de la robustesse contre les variations trouvées dans et autour de la région de OD. Il est proposé également une méthode de segmentation du Cup, fondée sur des preuves anatomiques telles-que les vessel bends à la frontière du Cup, jugés pertinents par les experts en glaucome. Une stratégie en multi-étapes est utilisée pour obtenir un sous-ensemble fiable de vessel bends appelé r-bends suivi par un raccord pour dériver la limite du Cup souhaitée.

La méthode dans [14] a été évaluée sur 138 images comprenant 33 cas normaux et 105 images de glaucome désignées par trois experts de glaucome. Les résultats de segmentation montrent de la cohérence dans le traitement de diverses variations géométriques et photométriques trouvées dans l'ensemble de données. L'erreur d'estimation de la méthode du rapport du diamètre vertical du Cup/Disque est de 0,09 / 0,08 (moyenne / écart-type), tandis que pour le ratio de la zone cup sur Disque, il est de 0,12 / 0,10. Dans l'ensemble, les résultats qualitatifs et quantitatifs obtenus montrent de l'efficacité à la fois dans la segmentation et au paramétrage du OD pour l'évaluation de glaucome.

Burana-Anusorn et al. [15] ont mis au point une approche automatique pour le calcul du ratio CDR à partir d'images du fond d'œil. L'idée est d'extraire le Disque optique en utilisant une approche de détection de contour et l'approche level-set par variation individuelle. Le Cup optique est ensuite segmenté en utilisant une méthode d'analyse de composante de couleur et de la méthode de seuillage level-set. Après avoir obtenu les contours, une étape d'ajustement par ellipse est introduite pour lisser les résultats obtenus. La performance de cette approche est évaluée en comparant le CDR calculé de manière automatique avec celui calculé manuellement. Les résultats indiquent que l'approche de Burana-Anusorn et al. atteint une précision de 89% pour l'analyse du glaucome. En conséquence, cette étude a un bon potentiel dans les systèmes automatisés de dépistage pour la détection précoce du glaucome.

Récemment, Abdul Khalid et al. [16] proposaient le déploiement de la dilatation et l'érosion avec c-Moyens flous (FCM) comme une technique efficace pour la segmentation du Cup et Disque optique des images couleurs du fond d'œil. Des travaux antérieurs ont identifié le canal vert comme le plus adéquat en raison de son contraste. De ce fait, en premier lieu le canal vert extrait est segmenté avec la FCM. Dans un autre essai, l'ensemble des images sont pré-traitées avec dilatation et l'érosion afin de supprimer le réseau vasculaire. La segmentation est évaluée sur la base des étiquettes décrites par les ophtalmologistes. Les mesures de CDR sont calculées à partir du rapport de diamètre du Cup et du Disque segmenté. L'évaluation montre que l'omission de la zone vernaculaire améliore la sensibilité, la spécificité et la précision du résultat segmenté.

Dans un registre plus global, Sivaswamy et al. [17] s'intéressent au problème de la segmentation de la tête du nerf optique (ONH) qui est d'un intérêt crucial pour l'évaluation du glaucome automatisé. Le problème de la segmentation implique la segmentation du Disque optique et du Cup de la région ONH. Les auteurs mettent en évidence la difficulté d'évaluer et de comparer les performances des méthodes existantes en raison d'absence d'un ensemble de données de référence. De ce fait, un ensemble complet de données d'images rétinienne qui comprennent yeux normaux et glaucomateux par segmentations manuelles de plusieurs experts a été mis en œuvre. Les deux mesures d'évaluation à base de superficie et de contour sont présentées afin d'évaluer une méthode sur divers aspects du problème de l'évaluation du glaucome.

Plusieurs approches s'exercent à la détermination du Cup mais ne sont pas toujours couronnées de succès, en particulier pour les yeux normaux. Muramatsu et al. [18] ont développé une méthode informatisée pour mesurer de manière plus précise le CDR sur des photographies du fond d'œil stéréo. Dans cette étude, une nouvelle méthode pour quantifier la probabilité de Disques glaucomateux basée sur la similarité des scores pour les modèles de glaucome et non glaucome.

2.2 Classification pixellaire semi-supervisée des images médicales

Les approches automatiques de segmentation par classification pixellaire peuvent être regroupées en trois grandes familles :

Méthode de classification par histogrammes est une technique couramment utilisée dans la segmentation d'images couleur car elle présente l'avantage de ne pas requérir de connaissance a priori sur l'image. Les méthodes d'analyse d'histogrammes se différencient par l'espace couleur choisi ou la composante couleur la plus significative [19, 20].

Méthodes de classification supervisée et non supervisée elles présentent l'espace couleur en sous-espaces homogènes selon un critère de ressemblance des couleurs de pixels. Nous citons des algorithmes de classification de pixels non-supervisés comme les k-moyennes proposé par Macqueen [21], C-moyens flous [22, 23], Fisher [24, 25] ou bien des algorithmes de classification de pixels supervisée comme celui de Bayes, les Réseaux de Neurones Multi-Couches, les Machines à Supports de Vecteurs (SVM), les k plus proches voisins (k-PPV) et les arbres de décision.

Méthodes de classification par apprentissage semi-supervisé en apprentissage supervisé, les algorithmes infèrent un modèle de prédiction à partir de données préalablement étiquetées. Cependant, l'étiquetage est un processus long et coûteux qui nécessite souvent l'intervention d'un expert. Cette phase contraste avec une acquisition automatique des données. Ce n'est alors pas rare de se retrouver avec un volume important de données dont seulement une petite partie a pu être étiquetée. Par exemple, en recherche d'images par contenu, l'utilisateur souhaite étiqueter le minimum d'images pour fouiller une base aussi grande que possible. Dans ce contexte, l'apprentissage semi-supervisé (SSL) intègre les données non-étiquetées dans la mise en place du modèle de prédiction. En ce sens, l'apprentissage semi-supervisé est à mi-chemin entre les apprentissages supervisé et non-supervisé : il cherche à exploiter les données non-étiquetées pour apprendre la relation entre les exemples et leur étiquette.

Les méthodes de classification pixellaire ont été fréquemment utilisées dans la segmentation d'images avec deux approches supervisée et non supervisée. Les méthodes de segmentation supervisées conduisent à une grande précision, mais elles ont besoin d'une grande quantité de données étiquetées, ce qui est difficile, coûteux et lent à obtenir. En outre, elles ne peuvent pas exploiter les données non étiquetées.

D'autre part, les méthodes de segmentation non supervisées n'ont pas de connaissance préalable et conduisent à un faible niveau de performance.

Cependant, l'apprentissage semi-supervisé qui utilise un ensemble de données étiquetées avec une grande quantité de données étiquetées réalise une plus grande précision. Plusieurs travaux se sont intéressés à cette approche dans la segmentation des images médicales, nous citerons ci-dessous quelques-uns de ces derniers.

Niaf et al. [26] ont mis en œuvre un système d'aide au diagnostic automatique pour le dépistage du cancer de la prostate sur les images multi-paramétriques de la résonance magnétique (mp-MR).

Basé sur une base de données d'apprentissage d'images mp-MR annotées de 30 patients. Ils utilisent le classifieur (SVM)-inspired qui permet simultanément d'avoir une discrimination linéaire optimale et un sous-ensemble de variables prédictives (ou caractéristiques) qui sont les plus pertinentes pour la tâche de classification, tout en favorisant la fluidité spatiale des cartes de prévision de malignité du cancer de prostate.

Suite à leur travail préliminaire [27], ils ont reformulé le problème d'optimisation de la sortie de SVM en introduisant à la fois une norme ℓ_1 dans le terme de régularisation, ainsi que la modification de la fluidité spatiale par un terme de coût supplémentaire qui code le quartier spatial des voxels, afin d'éviter la génération de cartes de prévisions bruyantes. L'apprentissage semi-supervisé appliqué permet même l'utilisation de l'information structurelle détenu par les voxels non marquées (ou images), qui dans une approche supervisée régulière, sont tout simplement mises à l'écart.

Les expérimentations démontrent une étude visuelle et numérique claire sur l'ensemble de données cliniques des images mp-MR du cancer de la prostate.

La segmentation des images médicales écho-graphiques est une problématique difficile en raison de la présence de bruit en forme de tacheture et l'atténuation du signal. Dans [28] Xbresson et al. présentent une méthode de segmentation intelligente basée sur le concept de la semi-supervision pour identifier les anomalies en imagerie écho-graphique. Les auteurs ont proposé une initialisation des étiquettes assistée par un expert du domaine suivie d'une utilisation du graphe de patches d'image qui agit en tant que probabilités correctives d'intensité pour une bonne représentation de l'image ultra-sonore. Ensuite, ils ont formulé la segmentation comme un problème de réduction minimale continue et cela est résolu avec un algorithme d'optimisation efficace de « Continuous graph cut model ». Ce dernier est proposé comme un algorithme de segmentation qui vise à minimiser l'opérateur de coupe (ou norme à base de graphes) appliqué sur les objets d'intérêt. Ce travail est validé dans le cadre de segmentation des imageries cliniques écho-graphiques (prostate, fœtus, tumeurs du foie et des yeux). Les résultats obtenus ont une grande similarité avec la vérité de terrain fournie par les délimitations d'experts médicaux dans toutes les applications.

Dans [29], Khashman et Al-Zgoul. proposent l'utilisation de l'analyse morphologique des images microscopiques de cellules sanguines leucémiques à des fins d'identification. Les auteurs présentent une première phase d'un système d'identification automatisé de forme de leucémie, qui est la segmentation d'images de cellules infectées. Le processus de segmentation fournit deux images améliorées pour chaque cellule du sang ; contenant le cytoplasme et les noyaux régions. Les caractéristiques uniques pour chaque forme de leucémie peuvent ensuite être extraites des deux images et utilisées pour l'identification.

Dans le domaine médical, Azmi et al. [30] ont adapté les approches semi-supervisées pour une segmentation par classification pixellaire dans le but d'aide au diagnostic. Plus précisément, ils ont proposé une nouvelle approche améliorée de classification semi-supervisée interactive de type auto-apprentissage pour la segmentation des lésions suspectes en IRM du sein. L'idée de ce travail est de faire intervenir en premier lieu un radiologue expert pour une segmentation semi-complète de la région d'intérêt (sélection 20% de la région pour quelques images). En second lieu,

une caractérisation de texture pour chaque pixel de l'image est faite utilisant trois catégories : matrice de co-occurrence, variable statistique d'histogramme et lenteur d'exécution matricielle. Avant d'aborder la partie apprentissage semi-supervisé, ils ont proposé un seuillage des images d'apprentissage pour faire générer les pixels de potentiel.

Un auto-apprentissage est appliqué par la suite par un classifieur bayésien entraîné sur les pixels expertisés, renforcé par une mesure de confiance probabiliste pour améliorer la robustesse du classifieur, permettant par la suite de segmenter les lésions avec précision. En terme de performances, les auteurs ont démontré que la segmentation semi-supervisée est plus fiable que la segmentation supervisée (KNN, SVM et réseau bayésien) et non supervisée (FCM). Par contre, ce que nous reprochons à ce papier, c'est qu'il traite d'une manière non exhaustive la partie du calcul de confiance, et qu'aucun détail n'est présenté sur la fonction de mesure de ce dernier, nous rappelons que l'auto-apprentissage est basé essentiellement sur le ré-apprentissage du classifieur qui utilise des données nouvellement classées avec une confiance. Ces dernières sont sélectionnées après des mesures mathématiques adaptées telle que la probabilité d'appartenance qui est utilisée en basic self training [31] ou celle de SETRED [32] qui utilise l'information du graphe de voisinage comme mesure de confiance.

Azmi et al. [33] reprennent l'idée de la segmentation semi-supervisée mais cette fois sur les images (IRM) du tissu cérébral. Ils font appel à une méthode ensembliste en semi-supervisé pour la segmentation d'images (IRM), par plusieurs classifieurs semi-supervisés simultanément. Ainsi, dans cet article, les auteurs utilisent deux algorithmes semi-supervisés : L'algorithme de filtrage espérance-maximisation et MCo_Training qui est une version améliorée de la méthode semi-supervisée Co_Training afin d'augmenter la précision de la segmentation. Les résultats expérimentaux montrent que la performance de la segmentation dans cette approche est plus élevée que les deux méthodes supervisées et les classifieurs semi-supervisés individuels.

3 L'approche proposée

L'objectif est de reconnaître automatiquement les régions Disques et Cups (figure 1) qui sont primordiales pour la mesure de la progression du glaucome. Pour ce faire, nous proposons une approche basée essentiellement sur une classification pixellaire en apprentissage semi-supervisé.

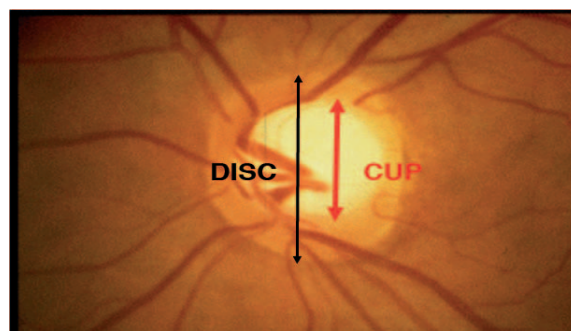


FIGURE 1 – Excavation papillaire glaucomateuse (augmentation du rapport Cup/Disque)

L'intervention d'un expert ophtalmologue est importante dans l'identification du Disque et du Cup en image rétinienne. Nous nous sommes inspiré du même principe proposé par Reza et al. [30], en faisons appel à l'expert pour un fenêtrage de région d'intérêt sur un minimum d'images, dans notre application nous sélectionnons 3 images aléatoirement soit 5% de la base d'images. La sélection est réalisée sur trois fenêtres (Disque, Cup et fond). Par la suite, une phase de caractérisation prend part, où chaque pixel de l'image est représenté par un vecteur de variables type couleur et spatiale.

Notre contribution est l'application des techniques de classification semi-supervisée sur les pixels expertisés par le médecin afin de former une hypothèse robuste et fiable qui permet la classification pixellaire. L'idée est de réaliser en premier lieu un pré-traitement qui nous permet de marquer quelques pixels des régions Cups et Disques en utilisant la méthode des c-moyens flous. Pour un meilleur apprentissage, une classification sera appliquée sur le voisinage de chaque région d'intérêt qui a été marquée dans la phase de pré-traitement. Cette phase sera réalisée par quatre différentes approches de classification (Arbre de décision, la forêt aléatoire, l'algorithme d'auto-apprentissage SETRED, et la forêt en apprentissage semi-supervisé *co-Forest*).

L'algorithme proposé est illustré dans la figure 2.

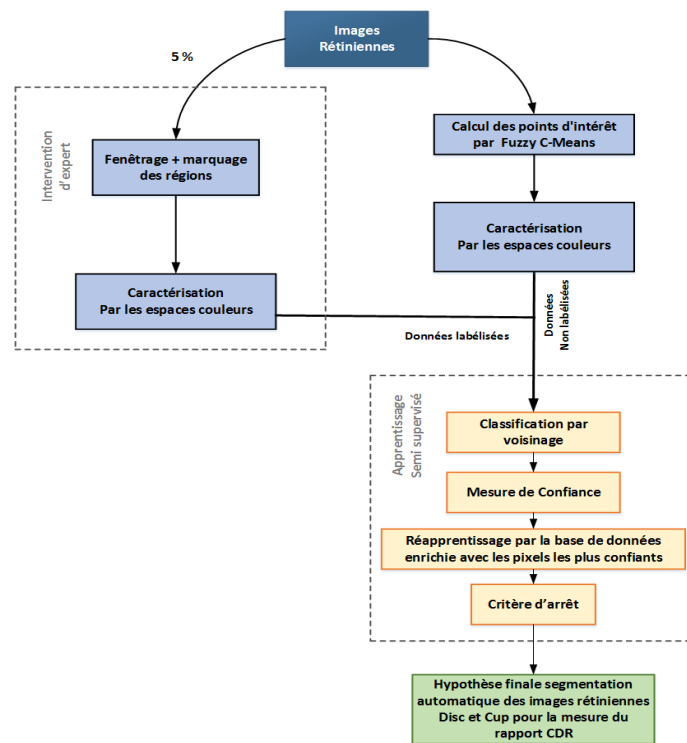


FIGURE 2 – L'approche proposée pour la segmentation automatique par Classification pixellaire en apprentissage semi-supervisé

3.1 Méthodes d'extraction des caractéristiques

3.1.1 Les espaces couleurs

Différents espaces couleurs ont été utilisés dans la segmentation par classification pixellaire, mais beaucoup d'entre eux partagent des caractéristiques similaires. Par conséquent, dans ce travail, nous nous intéressons à cinq espaces couleurs plus représentatifs qui sont couramment utilisés dans le domaine de traitement d'image [34].

Types de variables	Forme
<i>RGB</i>	$R(i, j)$ $G(i, j)$ $B(i, j)$
<i>LUV</i>	$L = 116(\frac{Y}{Y_n})^{1/3} - 16 \quad \text{Si } \frac{Y}{Y_n} > 0.008856$ $= 903.3(\frac{Y}{Y_n}) \quad \text{Si } \frac{Y}{Y_n} \leq 0.008856$ $U = 13L(U' - U'_n)$ $V = 13L(V' - V'_n)$
<i>HSV</i>	$H = \frac{G-B}{(Max-Min)} \quad \text{Si } R = \text{Max}$ $= \frac{B-R}{(Max-Min)} + 2 \quad \text{Si } G = \text{Max}$ $= \frac{R-G}{(Max-Min)} + 4 \quad \text{Si } B = \text{Max}$ $S = \frac{Max(R, G, B) - Min(R, G, B)}{Max(R, G, B)}$ $V = Max(R, G, B)$
<i>YUV</i>	$Y = 0.2989R + 0.5866G + 0.1145B$ $U = 0.5647(B - Y) = -0.1687R - 0.3312G + 0.5B$ $V = 0.7132(R - Y) = 0.5R - 0.4183G - 0.0817B$

TABLE 1 – Tableau des paramètres de caractérisation

L'espace RGB est un système additif qui décompose les couleurs en trois quantités des trois couleurs primaires : le rouge, le vert et le bleu (Table 1). C'est le système le plus utilisé dans les images couleurs et les moniteurs. Le modèle RVB utilise le système de coordonnées cartésiennes. Le point (1, 1, 1) représente le blanc, le point (0, 0, 0) représente le noir et la diagonale représente les niveaux de gris.

L'espace HSV est un modèle de représentation dit "naturel", c'est-à-dire proche de la perception physiologique de la couleur par l'œil humain. Il consiste à décomposer la couleur selon des critères physiologiques (Table 1).

L'espace HSL (Teinte, Saturation, Luminosité / luminance), est assez similaire à l'espace HSV, également connu sous le TSL (Teinte, Saturation, Luminosité) (Table 1). La différence est que la luminosité d'une couleur pure est égale à la luminosité du blanc, tandis que la luminosité d'une couleur pure est égale à la légèreté d'un gris moyen.

L'intention de l'espace couleur Luv est pour produire un espace de couleur plus linéaire que les autres espaces couleurs existants (Table 1). Linéaire perceptuelle signifie qu'un changement de la même quantité dans une valeur de couleur doit produire une variation dont la même importance visuelle.

L'espace YUV est destiné principalement à la vidéo analogique, ce modèle de représentation est utilisé dans les standards vidéo PAL et SECAM. La luminance est représentée par Y, tandis que les chrominances U et V sont issues de la transformation de l'espace RGB (Table 1).

3.2 Calcul des points d'intérêt

La prise en compte du temps de calcul surtout dans un context semi-supervisé, nous a conduit à proposer une phase de traitement dite détection des points d'intérêt. L'objectif de cette phase est de minimiser le temps de calcul et d'autre part de commencer l'apprentissage semi-supervisé par des pixels de potentiel qui appartiennent à notre cible. Nous proposons l'application des centres de cluster par la méthode des C moyens flous "FCM".

La méthode Fuzzy C-Means (FCM) est un algorithme de classification non-supervisée floue. Issu de l'algorithme des C-Moyens (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 point. Cet algorithme nécessite la connaissance préalable du nombre de clusters et génère les classes par un processus itératif en minimisant une fonction objective.

La classification floue des objets est décrite par une matrice floue à n lignes et c colonnes dans laquelle n est le nombre d'objets de données et c est le nombre de grappes. μ_{ij} , l'élément de la i -ème ligne et la colonne j . La procédure de classification est représentée dans le schéma suivant : dans μ , représente le degré de fonction d'appartenance de l' i -ème objet avec le pôle j .

Les étapes de l'algorithme FCM (Dunn, 1974[35] ; Bezdek, 1981[36]) sont énumérées ci-dessous :

Algorithm 1 Pseudo code de Fuzzy C-Means

- 1: **for** $t = 1, 2, \dots$ **do**
 - 2: **Step1** Calcul des centres de clusters $c_i^{(t)} = \frac{\sum_{j=1}^N (\mu_{ij}^{(t-1)})^m x_j}{\sum_{j=1}^N (\mu_{ij}^{(t-1)})^m}$
 - 3: **Step2** Calcul des distances D_{ijA}^2 avec :

$$D_{ijA}^2 = (x_j - c_i)^T A (x_j - c_i), \quad 1 \leq i \leq n_c, 1 \leq j \leq N.$$
 - 4: **Step3** Mise à jour de la matrice des partitions floues : $\mu_{ij}^{(t-1)} = \frac{1}{\sum_{k=1}^{n_c} (D_{ijA} / D_{kjA})^{2/(m-1)}}$
 - 5: **end for**
 - 6: **return** $\|U^{(t)} - U^{(t-1)}\| < \epsilon$
-

3.3 Méthodes de classification

Avec la disponibilité des données non étiquetées et la difficulté d'obtenir des étiquettes, des méthodes d'apprentissage semi-supervisé ont acquis une grande importance. A la différence de l'apprentissage supervisé, l'apprentissage semi-supervisé vise les problèmes avec relativement peu de données étiquetées et une grande quantité de données non étiquetées. La question qui se pose est alors de savoir si la seule connaissance des points avec labels est suffisante pour construire une fonction de décision capable de prédire correctement les étiquettes des points non étiquetés. Différentes approches proposent de déduire des points non étiquetés, des informations supplémentaires et de les inclure dans le problème d'apprentissage.

Dans cette partie classification, nous nous focalisons sur l'amélioration des performances de la classification supervisée en employant les données non étiquetées (SSL). Nous mettons en place d'abord la classification semi-supervisée dans le cadre des problèmes de classification, limitée à l'utilisation des méthodes d'ensembles en classification semi-supervisée.

En effet, la combinaison de différentes méthodes d'apprentissage/classification permet de tirer avantage de leurs forces tout en contournant leurs faiblesses. Aujourd'hui, force est de constater que ce qu'on appelle maintenant systèmes multi-classifieurs (MCS) constituent une des voies les plus prometteuses de l'apprentissage automatique [37].

De ce fait, nous proposons dans ce travail d'appliquer les systèmes multi-classifieurs de type méthode d'ensemble en comparaison aux méthodes mono-classifieurs en apprentissage supervisé et semi-supervisé : Arbre de décision, la forêt aléatoire, l'algorithme d'auto-apprentissage SETRED, et la forêt en apprentissage semi-supervisé *co-Forest*.

3.3.1 Arbre de décision

Les arbres de décision représentent une méthode très efficace d'apprentissage supervisé. Il s'agit de partitionner un ensemble de données en des groupes les plus homogènes possibles du point de vue de la variable à prédire. On prend en entrée un ensemble de données classées, et on fournit en sortie un arbre qui ressemble beaucoup à un diagramme d'orientation. Un arbre de décision est composé d'une racine qui est le point de départ de l'arbre, des nœuds, des branches qui relient : la racine avec les nœuds, les nœuds entre eux et les nœuds avec les feuilles.

Il existe plusieurs algorithmes présents dans la littérature, comme : CART [38], ID3 [?] et C4.5 [39]. Dans ce travail nous nous limitons à l'application de l'algorithme CART (Classification and Regression Tree).

3.3.2 Méthode d'auto-apprentissage "SETRED"

SETRED (self-training with data editing) est un algorithme d'auto-apprentissage proposé par Ming Li et al. [32] Dans ce travail, les auteurs ont étudié le potentiel des techniques de filtrage des données comme une mesure de confiance, cela permet de déminuer le risque d'ajouter des données bruitées à la base d'apprentissage. cet algorithme utilise les techniques d'édition des données pour identifier et supprimer les données mal classées. En détail, à chaque itération dans le proces-

sus d'auto-apprentissage, l'algorithme fait appel au principe de la règle du plus proche voisin et coupure des arêtes pondérées pour mesurer la confiance et faire enrichir la base d'apprentissage.

Règle du plus proche voisin La règle du plus proche voisin (*ppv*) a été proposée par Fix et Hodges [40], c'est une méthode non paramétrique, où la règle de classification est obtenue en posant que la classe d'une donnée non étiquetée est celle de la plus proche parmi les étiquettes de ses voisins dans les échantillons d'apprentissage. La détermination de leur similarité est basée sur des mesures de distance.

Par la suite, cette règle a été développée en *k-ppv*, elle est basée sur la définition d'un nombre *k* de voisinage et l'étiquette d'une donnée non-classée est celle qui est majoritaire parmi les étiquettes de ses *k* plus proches voisins.

Graphe de voisinage relatif Le graphe de voisinage est un outil issu de la géométrie computationnelle qui a été exploité dans nombreuses applications d'apprentissage automatique. Par définition un graphe de voisinage $G = (S, E)$ [41] associé à un ensemble de données dont les sommets « *S* » composent l'ensemble des arêtes *E*.

Chaque donnée dans un graphe de voisinage est représentée par un sommet, il existe des arêtes entre les sommets x_i et x_j si la distance entre les deux satisfait l'équation 1

$$(x_i, x_j) \in E \Leftrightarrow dist(x_i, x_j) \leq \max (dist(x_i, x_k), dist(x_j, x_k)), \forall x_k \in TR, k \neq i, j \quad (1)$$

Avec : $dist(x_i, x_j)$: la distance entre x_i et x_j .

Cette définition signifie qu'il n'y a pas de sommet à l'intérieur de l'intersection de deux cercles de centre x_i et x_j ainsi de rayon $dist(x_i, x_j)$.

D'après la définition ci-dessus, on peut construire un graphe de voisinage relatif dans lequel le voisinage de chaque donnée est un ensemble d'échantillons connectés avec des arêtes. Muhlenbach et al. [42] ont exploité l'information des arêtes pour calculer un poids statistique afin de couper les arêtes des données connexes de différentes classes, dans le cas des algorithmes de filtrage des données bruitées [42] et les mesures de confiance suivant le principe de SETRED [32].

En premier lieu une hypothèse est apprise sur les données étiquetées. Suivie d'une application du processus d'auto-apprentissage qui fait appel à l'algorithme de cut edge weight statistic [42], ce dernier utilise les équations (2, 3, 4,5) afin de calculer le rapport R_i . Pour juger si la donnée est bien classée, le rapport R_i doit être supérieur à un seuil qui est fixé par l'utilisateur. Pour une compréhension de l'algorithme consulter les travaux de [32, 42].

$$R_i = \frac{J_i}{I_i} \quad (2)$$

Avec :

$$I_i = \sum_{(j \in Neighborhood(x_i))} w_{ij} \quad (3)$$

$$J_i = \sum_{(j \in \text{Neighborhood}(x_i), y_j \neq y_i)} w_{ij} \quad (4)$$

$$w_{ij} = \frac{1}{((1 + \text{dist}(x_i, x_j)))} \quad (5)$$

3.3.3 Méthode d'ensemble : Forêts Aléatoires

Une Forêt Aléatoire est un prédicteur constitué d'un ensemble de classifieurs élémentaires de type arbres de décision. Dans les cas spécifiques des modèles CART (arbres binaires), Breiman [43] propose une amélioration du bagging avec un algorithme d'induction de forêts aléatoires (Forest-Random Input — pour Random Forest - Random Input —) qui utilise le principe de randomisation "Random Feature Selection" proposé par Amit et Geman [44]. L'induction des arbres se fait sans élagage et selon l'algorithme CART [38], toutefois, au niveau de chaque nœud, la sélection de la meilleure partition, basée sur l'index de Gini, s'effectue uniquement sur un sous-ensemble d'attributs de taille préfixée (généralement égale à la racine carrée du nombre total d'attributs) sélectionné aléatoirement depuis l'espace originel des caractéristiques [45]. La prédiction globale de la forêt aléatoire est calculée en prenant la majorité des votes de chacun de ses arbres. Cet algorithme appartient à la famille la plus large des forêts aléatoires défini comme suit (Algorithme 2) par Breiman [43].

Algorithm 2 Pseudo code de l'algorithme des forêts aléatoires

Entrées : L'ensemble d'apprentissage L , Nombre d'arbres N .

Sortie : Ensemble d'arbres E

Processus :

for $i = 1 \rightarrow N$ **do**

$T^i \leftarrow \text{BootstrapSample}(T)$

$C^i \leftarrow \text{ConstructTree}(T^i)$ où à chaque nœud :

- Sélection aléatoire de $K = \sqrt{M}$ Variables à partir de l'ensemble d'attributs M
- Sélection de la variable la plus informative K en utilisant l'index de Gini
- Création d'un nœud fils en utilisant cette variable

$E \leftarrow E \cup \{C^i\}$

end for

Retourner E

3.3.4 Méthode des Forêts Aléatoires en apprentissage semi-supervisé "co-Forest"

L'algorithme *co-Forest* repose sur le paradigme du *co-Training* [46], où deux classifieurs sont d'abord formés à partir de L , puis chacun d'eux choisit les exemples les plus confiants en U de son point de vue. Il met par la suite à jour les autres classifieurs avec ces exemples nouvellement étiquetés. Un des aspects les plus importants dans *co-Training* est d'estimer la confiance d'un exemple donné non étiqueté.

Dans *co-Training* standard, l'estimation de la confiance profite directement à partir de deux sous-ensembles d'attributs suffisants et redondants, où la confiance d'étiquetage d'un classifieur

pourrait être considérée comme sa confiance pour un exemple non étiqueté. Lorsque la condition des deux sous-attributs suffisants et redondants n'est pas présente, la validation croisée est appliquée à chaque itération d'apprentissage afin d'estimer la confiance pour les données non étiquetées [47]. L'inefficace estimation de la confiance réduit considérablement l'applicabilité de l'étendue de l'algorithme *co-Training* dans des applications telles que le diagnostic assisté par ordinateur.

Cependant, si un ensemble de classifieurs N , qui est désigné par H^* , est utilisé dans le *co-Training* au lieu de deux classifieurs, la confiance peut être estimée de manière efficace. Lors de la détermination des exemples les plus confiants étiquetés pour un classifieur de l'ensemble de $H_i (i = 1, \dots, N)$, tous les classifieurs sauf h_i sont utilisés. Ces classifieurs forment un nouvel ensemble, qui est appelé l'ensemble de concomitance de H^* , noté par H_i . Notez que diffère H_i de H^* seulement par l'absence de h_i . Maintenant la confiance pour un exemple sans étiquette peut être simplement estimée par le degré d'accords sur l'étiquetage, c'est à dire le nombre de classifieurs qui sont d'accord sur l'étiquette assignée par H_i . En utilisant cette méthode *Co-forest* [48], l'ensemble aborde un ensemble de classifieurs sur L , puis affine chaque classifieur avec des exemples non étiquetés choisis par son ensemble concomitant.

Le fonctionnement de *Co-forest* peut se résumer dans les étapes suivantes :

Étape 1 *Co-forest* lance l'apprentissage des H^* sur des bootstrap¹ de L Figure 3.

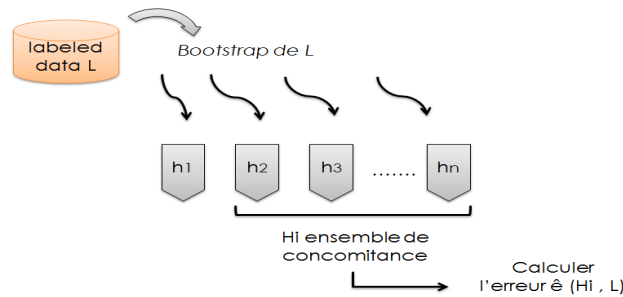


FIGURE 3 – Apprentissage des arbres sur les données labellisées L

Étape 2 l'ensemble de concomitance examine chaque exemple de U

Si le nombre de votant d'accord sur l'étiquette de $xu > \theta$ **Alors** xu est labellisé et copié dans un nouveau ensemble L'

Remarque : On pourrait être confronté à une situation où $L' \geq U$ cela affecte les performances de h_i

1. Un échantillon bootstrap L est, par exemple, obtenu en tirant aléatoirement n observations avec remise dans l'échantillon d'apprentissage L_n , chaque observation ayant une probabilité $1/n$ d'être tirée.

Solution : introduire un poids par la prédiction de confiance de l'ensemble de concomitance Figure 4.

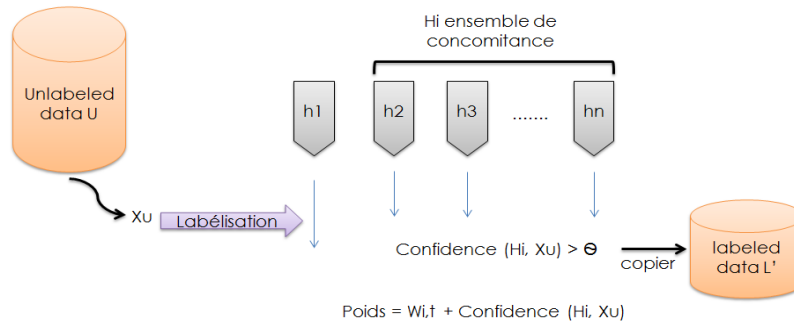


FIGURE 4 – Labellisation des données non labellisées U par l'ensemble de concomitance

Étape 3 Chaque arbre aléatoire est raffiné avec des exemples nouvellement marqués LL' sélectionnés par son ensemble de concomitance sous la condition suivante :

$$e_{i,t} \cdot W_{i,t} < e_{i,t-1} \cdot W_{i,t-1}$$

Où : $W = \sum w_{ij}$ et w_{ij} : la confiance prédictive de H_i sur x_i dans L' Figure 5.

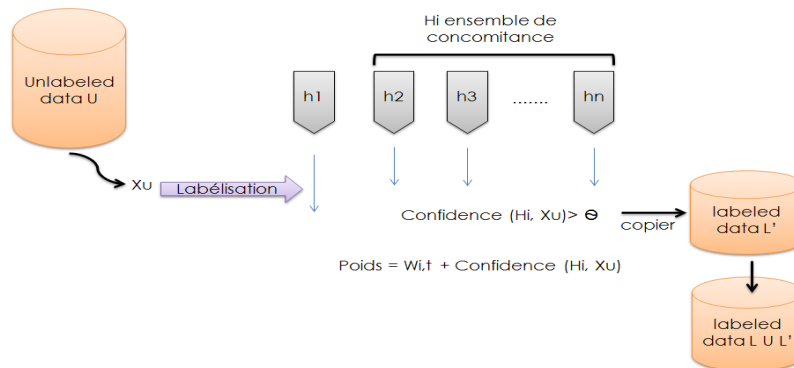


FIGURE 5 – Ré-apprentissage par les exemples nouvellement marqués $L L'$

Pour que le succès de cette méthode d'ensemble soit présent, il faut que deux conditions soient satisfaites :

- Chaque prédicteur individuel doit être relativement bon.
- Les prédicteurs individuels doivent être différents les uns des autres.

En plus simple, il faut que les prédicteurs individuels soient de bons classificateurs. Et là où un prédicteur se trompe, les autres doivent prendre le relais sans se tromper.

Dans *Co-forest* afin de maintenir la diversité l'idée étant l'application des forêts aléatoires. Elle permettent d'injecter l'aléatoire dans son principe d'apprentissage, pour maintenir cette condition, les auteurs de *Co-forest* ont fixé un seuil pour la labélisation des U où juste les U dont le total de poids $< \frac{e_{i,t-1} \cdot W_{i,t-1}}{e_{i,t}}$ seront sélectionnés (voir Figure 6).

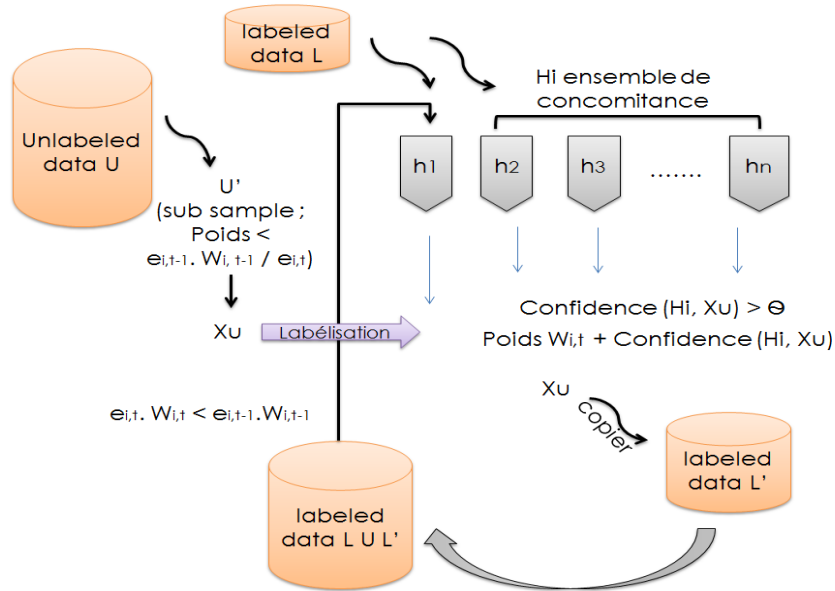


FIGURE 6 – Schéma de principe de l'algorithme *Co-forest*

4 Résultats et expérimentations

La base de données d'images rétinienne a été construite à partir d'images réelles acquises au sein de la clinique ophtalmologique locale (Clinique LAZOUNI Tlemcen). Fond d'oeil RTVue XR 100 Avanti Edition de la société Optovue permet d'obtenir des images couleur RGB de taille 1609x1054 pixels.

Cent trois images rétinienne du fond d'œil ont été utilisées pour tester l'algorithme de segmentation proposé. Nous avons construit une base d'apprentissage, où l'expert sélectionne 3 régions : Cup, Disque (régions cibles) et fond (critère d'arrêt) (Figure 7).

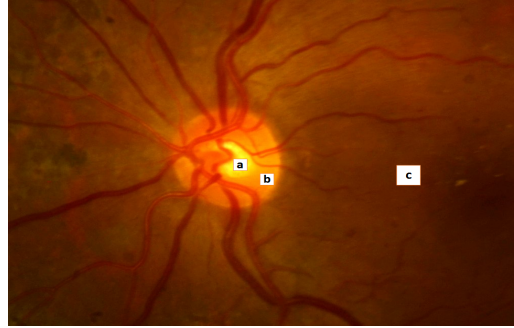


FIGURE 7 – (a) Cup, (b) Disque et (c) fond.

Le tableau (Table 2) résume les propriétés des jeux de données construites par l’approche Fuzzy C-means. Il présente pour chaque région le nombre de pixels, d’un total de 15667 pixels étiquetés.

Région	Pixels par Fuzzy C-means
<i>Cup</i>	2956
<i>Disque</i>	3995
<i>Fond</i>	8716

TABLE 2 – Description de la base d’apprentissage par les deux approches de traitement

Dans nos expériences, nous avons sélectionné 5% de la base de données (3 images) pour réaliser l’apprentissage. L’expert ophtalmologue interviendra dans l’étiquetage de ces trois images par fenêtrage de taille [576 - 50466] permettant ainsi une meilleure perception des zones d’intérêt. Dans la partie d’apprentissage semi-supervisé, une classification sera appliquée sur le voisinage d’un degré égal à 50. L’application des méthodes d’ensemble (forêt aléatoire et *co-Forest*) avec un nombre d’arbre égal à 100 étant choisi. L’évaluation est réalisée avec une validation croisée égale à 5. Les détails des paramètres d’expérimentation sont résumés dans le tableau (Table 3).

Base étiquetée	3 images
Base d’apprentissage	53 images
Base de test	50 images
Nombre de clusters	3
Degré de Voisinage	50
Degré de confiance	75%
Nombre d’arbres	100
Validation croisée	5

TABLE 3 – Paramètres de classification

Le Tableau (Table 4) résume les performances de classification par les quatre approches supervisées et semi-supervisées sur 50 images de test pour la reconnaissance des régions : cup et Disque. Les résultats réalisés par les approches d’ensemble *Forêt Aléatoire* et *co-Forest*, garantissent une plus grande précision de segmentation des deux régions ciblées. Par contre, des résultats médiocres ont

été obtenus par l'arbre de décision *CART* mais le mode semi-supervisé a permis d'améliorer les performances pour la méthode *SETRED-CART*. Notre approche par l'algorithme *co-Forest* réalise les meilleures performances de segmentation pour les deux régions et plus particulièrement la zone cup qui est la plus difficile à l'extraction. Ces résultats nous permettent de consolider notre proposition et affirmer sa rigueur et sa robustesse pour la tâche de croissance de région par classification pixellaire en apprentissage semi-supervisé.

TECHNIQUES	MOYENNE DES PERFORMANCES SUR LES 50 IMAGES	
	Cup	Disque
ARBRE CART	0,5350	0,7297
FORÊT ALÉATOIRE (100 ARBRES)	0,75609	0,9228
SETRED - CART	0,7556	0,60746
<i>co-Forest</i> (100 ARBRES)	0,8900	0,93503

TABLE 4 – Performances de classification par les techniques supervisées et semi-supervisées avec 5% d'images étiquetées

Discussion

Pour plus d'exhaustivité, et afin d'établir une évaluation visuelle des performances de notre approche, nous sélectionnons aléatoirement six images de la base de test (Figure 8) pour discuter les performances et la qualité de segmentation par les quatre approches mono et multi-classifieur en mode supervisé et semi-supervisé à savoir respectivement l'Arbre de décision et *Forêt Aléatoire* ainsi que la méthode *SETRED* et *co-Forest*.

De ce fait, une étude comparative de plusieurs techniques est proposée. Le principe repose sur une croissance de région en classifiant les pixels voisins à partir des pixels d'intérêt de l'image par apprentissage semi-supervisé. Les points d'intérêt sont détectés par l'algorithme Fuzzy C-means (FCM). Dans notre processus, nous avons utilisé les points d'intérêt afin de réaliser la segmentation des régions ciblées. Les points d'intérêt sont détectés par l'algorithme Fuzzy C-means (FCM) en images 2,3,4,6 (Figure 8) ont induit à une mauvaise segmentation par l'utilisation des mono-classifieurs (CART et SETRED) ; par contre les multi-classifieurs (Forêt Aléatoire et *co-Forest*) ont bien pu séparer le cup avec succès.

Un autre avantage d'utilisation des méthodes d'ensemble est la puissance de séparation du réseau vasculaire ainsi que la région fond comme il est démontré en images 3,4,6 (Figure 8). Par contre, comme on peut le constater aussi dans l'image 4 (Figure 8), l'apport des pixels non étiquetés à la méthode d'ensemble *co-Forest* a permis une bonne identification du Disque contrairement à la Forêt Aléatoire. Nous remarquons également dans l'image 1 (Figure 8), une mauvaise classification des pixels du cup par la globalité des approches à l'exception de l'algorithme *co-Forest*.



FIGURE 8 – Exemples d’images de segmentation automatique par les différentes techniques

Dans sa globalité, ce travail, nous a permis de voir une multitude de pistes de recherche qui s'offrent pour la segmentation automatique des images. L'idée d'extrapoler la segmentation de région par classification pixellaire au contexte d'apprentissage semi-supervisé par l'approche *co-Forest*, nous a permis d'exploiter les données non-étiquetées dans la mise en place du modèle de prédiction ensembliste. En ce sens, les données non-étiquetées ont renforcé la reconnaissance des régions d'intérêts pour un rapport cup/Disque calculé approximant celui de l'ophtalmologiste.

Une comparaison des rapports Cup/Disque réalisés par un ophtalmologiste et les méthodes proposées sur quinze images aléatoires est représentée dans la Figure 9. Les différences de mesures des ratios CDR de l'ophtalmologiste en comparaison avec les méthodes supervisées proposées (CART et Forêt Aléatoire) sont d'un grand écart dans les cas de glaucome et cas normaux. Bien que les ratios CDR calculés par l'approche d'ensemble sont plus approchant de l'optimal (mesure d'expert).

Cependant, nos résultats par les algorithmes semi-supervisé *co-Forest* et *SETRED* étaient légèrement différents de la valeur de l'ophtalmologiste, mais notre méthode *co-Forest* tendait à montrer les plus petites variations dans les cas normaux et les cas glaucome, ayant une mesure presque semblable à celle de l'expert.

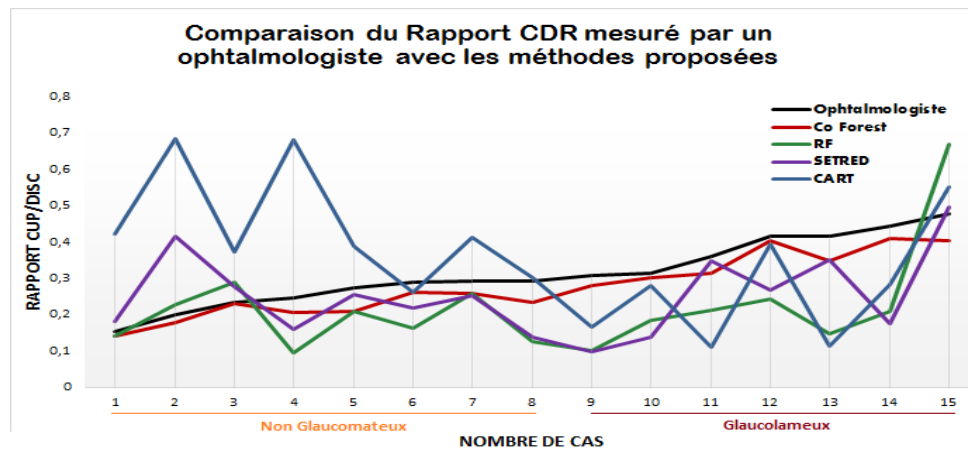


FIGURE 9 – Comparaison du rapport Cup/Disque mesuré par un ophtalmologiste et les méthodes proposées. Les cas 1-8 sont des cas normaux, et 9-15 cas sont des cas de glaucome.

5 Conclusion

Dans ce rapport de recherche, nous avons proposé une méthode de segmentation automatique des régions Cup et Disque des images rétinienne par classification pixellaire en apprentissage semi-supervisé. L'objectif visé est d'impliquer l'expert à l'apprentissage de notre modèle pour une meilleure discrimination des régions d'intérêt, l'évaluation et la segmentation des images de tests sont réalisées à l'aide de l'algorithme Fuzzy C-Means. Une croissance de région est développée en classifiant les pixels voisins par l'application de quatre classifieurs : Arbre de décision et Forêt

Aléatoire ainsi que la méthode SETRED et *co-Forest*.

Les résultats obtenus sont très convaincants et encourageants, indiquant une grande capacité de reconnaissance et de segmentation des régions ciblées, cela étant plus clair par l'application de la forêt aléatoire en apprentissage semi-supervisé *co-Forest*, l'heuristique de ces méthodes d'ensemble permet en utilisant plusieurs classifieurs, d'explorer grandement l'espace des solutions, et qu'en agrégeant toutes les prédictions, on récupère un classifieur qui prend en considération toute cette exploration. L'apport des données non-étiquetées dans la mise en place du modèle de prédiction permet de renforcer l'apprentissage et la reconnaissance des relations entre les pixels et leur région.

Références

- [1] A.A.-H. Abdel-Razik Youssif, A.Z. Ghalwash, and A.A.S. Abdel-Rahman Ghoneim, "Optic disc detection from normalized digital fundus images by means of a vessels' direction matched filter," *Medical Imaging, IEEE Transactions on*, vol. 27, no. 1, pp. 11–18, Jan 2008.
- [2] Rudiger Bock, Jorg Meier, Laszla G. Nyul, Joachim Hornegger, and Georg Michelson, "Glaucoma risk index :automated glaucoma detection from color fundus images," *Medical Image Analysis*, vol. 14, no. 3, pp. 471 – 481, 2010.
- [3] M. Lalonde, M. Beaulieu, and L. Gagnon, "Fast and robust optic disc detection using pyramidal decomposition and hausdorff-based template matching," *Medical Imaging, IEEE Transactions on*, vol. 20, no. 11, pp. 1193–1200, Nov 2001.
- [4] R.A.A. Ghafar, T. Morris, T. Ritchings, and I. Wood, "Detection and characterization of the optic disc in glaucoma and diabetic retinopathy," in *Medical Image Understand Annual Conference, London, UK*, pp. 23-24, September 2004.
- [5] A. Hoover, V. Kouznetsova, and M. Goldbaum, "Locating blood vessels in retinal images by piecewise threshold probing of a matched filter response," *Medical Imaging, IEEE Transactions on*, vol. 19, no. 3, pp. 203–210, March 2000.
- [6] A.M. Mendonca and A. Campilho, "Segmentation of retinal blood vessels by combining the detection of centerlines and morphological reconstruction," *Medical Imaging, IEEE Transactions on*, vol. 25, no. 9, pp. 1200–1213, Sept 2006.
- [7] U. Rajendra Acharya, Sumeet Dua, Xian Du, Subbhuraam Vinitha Sree, and Chua Kuang Chua, "Automated diagnosis of glaucoma using texture and higher order spectra features," *IEEE Transactions on Information Technology in Biomedicine*, vol. 15, no. 3, pp. 449–455, 2011.
- [8] Mishra Madhusudhan, Nath Malay, S.R. Nirmala, and Dandapat Samerendra, "Image processing techniques for glaucoma detection," in *Advances in Computing and Communications*, Ajith Abraham, JaimeLloret Mauri, JohnF. Buford, Junichi Suzuki, and SabuM. Thampi, Eds., vol. 192 of *Communications in Computer and Information Science*, pp. 365–373. Springer Berlin Heidelberg, 2011.
- [9] S. Mohammad, D.T. Morris, and N. Thacker, "Texture analysis for the segmentation of optic disc in retinal images," in *Systems, Man, and Cybernetics (SMC), 2013 IEEE International Conference on*, Oct 2013, pp. 4265–4270.

- [10] S. Chandrika and K. Nirmala, "Analysis of cdr detection for glaucoma diagnosis," *International Journal of Engineering Research and Applications (IJERA)*, vol. NCACCT-19, pp. ISSN : 2248-9622, March 2013.
- [11] Michael B. Merickel, Jr., Michael D. Abr?moff, Milan Sonka, and Xiaodong Wu, "Segmentation of the optic nerve head combining pixel classification and graph search," *Proc. SPIE*, vol. 6512, pp. 651215-651215-10, 2007.
- [12] Jun Cheng, Jiang Liu, Yanwu Xu, Fengshou Yin, Damon Wing Kee Wong, Ngan-Meng Tan, Ching-Yu Cheng, Yih Chung Tham, and Tien Yin Wong, "Superpixel classification based optic disc segmentation," in *Proceedings of the 11th Asian Conference on Computer Vision - Volume Part II*, Berlin, Heidelberg, 2013, ACCV'12, pp. 293-304, Springer-Verlag.
- [13] Yuji Hatanaka, Atsushi Noudo, Chisako Muramatsu, Akira Sawada, Takeshi Hara, Tetsuya Yamamoto, and Hiroshi Fujita, "Automatic measurement of vertical cup-to-disc ratio on retinal fundus images," in *Proceedings of the Second International Conference on Medical Biometrics*, Berlin, Heidelberg, 2010, ICMB'10, pp. 64-72, Springer-Verlag.
- [14] Gopal Datt Joshi, Jayanthi Sivaswamy, and S. R. Krishnadas, "Optic disk and cup segmentation from monocular color retinal images for glaucoma assessment," *IEEE Trans. Med. Imaging*, vol. 30, no. 6, pp. 1192-1205, 2011.
- [15] Chalinee Burana-Anusorn, Waree Kongprawechnon, Toshiaki Kondo nad Sunisa Sintuwong, and KanokvateTungpimolrut, "Image processing techniques for glaucoma detection using the cup-to-disc ratio," *Thammasat International Journal of Science and Technology*, vol. Vol. 18, No. 1, March 2013.
- [16] Noor Elaiza Abdul Khalid, Noorhayati Mohamed Noor, and Norharyati Md. Ariff, "Fuzzy c-means (fcm) for optic cup and disc segmentation with morphological operation," *Procedia Computer Science*, vol. 42, no. 0, pp. 255 - 262, 2014, Medical and Rehabilitation Robotics and Instrumentation (MRR2013).
- [17] Jayanthi Sivaswamy, S. R. Krishnadas, Gopal Datt Joshi, Madhulika Jain, and A. Ujjwaft Syed Tabish, "Drishti-gs : Retinal image dataset for optic nerve head(ONH) segmentation," in *IEEE 11th International Symposium on Biomedical Imaging, ISBI 2014, April 29 - May 2, 2014, Beijing, Chin, Beijing, China*, 2014, pp. 53-56.
- [18] Chisako Muramatsu, Yuji Hatanaka, Kyoko Ishida, Akira Sawada, Tetsuya Yamamoto, and Hiroshi Fujita, "Preliminary study on differentiation between glaucomatous and non-glaucomatous eyes on stereo fundus images using cup gradient models," *Proc. SPIE*, vol. 9035, pp. 903533-903533-6, 2014.
- [19] Raimondo Schettini, "A segmentation algorithm for color images," *Pattern Recognition Letters*, vol. 14, no. 6, pp. 499 - 506, 1993.
- [20] Patrick Lambert and L Macaire, "Filtering and segmentation : the specificity of color images," *International Conference on Color in Graphics and Image Processing*, vol. 1, pp. 57-71, 2000.
- [21] J. B. MacQueen, "Some methods for classification and analysis of multivariate observations," in *Proc. of the fifth Berkeley Symposium on Mathematical Statistics and Probability*, L. M. Le Cam and J. Neyman, Eds. 1967, vol. 1, pp. 281-297, University of California Press.
- [22] Young Won Lim and Sang Uk Lee, "On the color image segmentation algorithm based on the thresholding and the fuzzy c-means techniques," *Pattern Recognition*, vol. 23, no. 9, pp. 935-952, 1990.

- [23] Jean-Pierre Cocquerez and Sylvie Philipp, *Analyse d'images : filtrage et segmentation*, Masson, 1997.
- [24] Walter D. Fisher, "On grouping for maximum homogeneity," *Journal of the American Statistical Association*, vol. 53, no. 284, 1958.
- [25] R. A. Fisher, "The use of multiple measurements in taxonomic problems," *Annals of Eugenics*, vol. 7, no. 7, pp. 179–188, 1936.
- [26] E. Niaf, R. Flamary, A. Rakotomamonjy, O. Rouvière, and C. Lartizien, "Svm with feature selection and smooth prediction in images : application to cad of prostate cancer," in *IEEE International Conference on Image Processing (ICIP)*, 2014.
- [27] R. Flamary and A. Rakotomamonjy, "Support vector machine with spatial regularization for pixel classification," in *International Workshop on Advances in Regularization, Optimization, Kernel Methods and Support Vector Machines : theory and applications (ROKS)*, 2013.
- [28] A. Ciurte, X. Bresson, O. Cuisenaire, N. Houhou, S. Nedevschi, J.P. Thiran, and M. Bach Cuadra, "Semi-supervised segmentation of ultrasound images based on patch representation and continuous min cut," *PLoS ONE*, vol. 9(7), 2014.
- [29] Adnan Khashman and Esam Al-Zgoul, "Image segmentation of blood cells in leukemia patients," in *Proceedings of the 4th WSEAS International Conference on Computer Engineering and Applications*, Stevens Point, Wisconsin, USA, 2009, CEA'10, pp. 104–109, World Scientific and Engineering Academy and Society (WSEAS).
- [30] Reza Azmi, Narges Norozi, Robab Anbiaee, Leila Salehi, and Azardokht Amirzadi, "Impst : A new interactive self-training approach to segmentation suspicious lesions in breast mri," *Journal of Medical Signals and Sensors*, vol. 1 :2, pp. 138–148, 2011.
- [31] David Yarowsky, "Unsupervised word sense disambiguation rivaling supervised methods," in *Proceedings of the 33rd Annual Meeting on Association for Computational Linguistics*, Stroudsburg, PA, USA, 1995, ACL '95, pp. 189–196, Association for Computational Linguistics.
- [32] Ming Li and Zhi-Hua Zhou, "Setred : Self-training with editing.," in *PAKDD*, Tu Bao Ho, David Wai-Lok Cheung, and Huan Liu, Eds. 2005, vol. 3518 of *Lecture Notes in Computer Science*, pp. 611–621, Springer.
- [33] Reza Azmi, Boshra Pishgoo, Narges Norozi, and Samira Yeganeh, "Ensemble semi-supervised frame-work for brain magnetic resonance imaging tissue segmentation," *Journal of Medical Signals and Sensors*, vol. 3(2), pp. 94–106, 2013.
- [34] James D. Foley, Andries van Dam, Steven K. Feiner, and John F. Hughes, *Computer Graphics : Principles and Practice (2Nd Ed.)*, Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA, 1990.
- [35] J. Dunn, "A fuzzy relative of the isodata process and its use in detecting compact, well-separated clusters," *Journal of Cybernetics*, vol. vol. 3,, pp. 32–57, 1974.
- [36] James C. Bezdek, *Pattern Recognition with Fuzzy Objective Function Algorithms*, Kluwer Academic Publishers, Norwell, MA, USA, 1981.
- [37] Antoine Cornuéjols and Laurent Miclet, *Apprentissage artificiel : Concepts et algorithmes*, Eyrolles, June 2010.
- [38] L. Breiman, J. H. Friedman, R. A. Olshen, and C. J. Stone, *Classification And Regression Trees*, Chapman and Hall, New York, 1984.

- [39] J. Ross Quinlan, *C4.5 : Programs for Machine Learning*, Morgan Kaufmann, 1993.
- [40] Evelyn Fix and Jr, "Discriminatory analysis : Nonparametric discrimination : Consistency properties," Tech. Rep. Project 21-49-004, Report Number 4, USAF School of Aviation Medicine, Randolph Field, Texas, 1951.
- [41] L. Devroye, L. Györfi, and G. Lugosi, *A Probabilistic Theory of Pattern Recognition*, Springer, 1996.
- [42] Fabrice Muhlenbach, Stéphane Lallich, and Djamel A. Zighed, "Identifying and handling mislabelled instances.," *J. Intell. Inf. Syst.*, vol. 22, no. 1, pp. 89–109, 2004.
- [43] L. Breiman, "Random forests," *Machine Learning*, vol. 45, pp. 5–32, 2001.
- [44] Yali Amit and Donald Geman, "Shape quantization and recognition with randomized trees," *Neural Computation*, vol. 9, no. 7, pp. 1545–1588, 1997.
- [45] N. Sirikulviriyaya and S. Sinthupinyo, "Integration of rules from a random forest.," in *International Conference on Information and Electronics Engineering IPCSIT vol.6 (2011) IACSIT Press, Singapore*, 2011.
- [46] Avrim Blum and Tom Mitchell, "Combining labeled and unlabeled data with co-training," in *Proceedings of the eleventh annual conference on Computational learning theory*, New York, NY, USA, 1998, COLT' 98, pp. 92–100.
- [47] Sally A. Goldman and Yan Zhou, "Enhancing supervised learning with unlabeled data," in *Proceedings of the Seventeenth International Conference on Machine Learning*, San Francisco, CA, USA, 2000, ICML '00, pp. 327–334, Morgan Kaufmann Publishers Inc.
- [48] Ming Li and Zhi-Hua Zhou, "Improve computer-aided diagnosis with machine learning techniques using undiagnosed samples," *Trans. Sys. Man Cyber. Part A*, vol. 37, no. 6, pp. 1088–1098, Nov. 2007.