



Nouvelle approche d'extraction de l'axe médian dans les traits manuscrits : application à la constitution du code book

Hani Daher, Véronique Eglin, Stéphane Bres, N. Vincent

► To cite this version:

Hani Daher, Véronique Eglin, Stéphane Bres, N. Vincent. Nouvelle approche d'extraction de l'axe médian dans les traits manuscrits : application à la constitution du code book. Colloque International Francophone sur l'Écrit et le Document, Mar 2010, Sousse, Tunisie. pp.425-432. hal-00592017

HAL Id: hal-00592017

<https://hal.science/hal-00592017>

Submitted on 11 May 2011

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.

Nouvelle approche d'extraction de l'axe médian dans les traits manuscrits :

Application à la constitution du code book

H. Daher², V. Eglin², S. Bres², N. Vincent¹

1 Université René Descartes – CRIP5 - Systèmes Intelligents de Perception - 75270 Paris Cedex 06

*2 Université de Lyon, CNRS, INSA-Lyon, UMR5205, F-69621 France
{hani.daher, veronique.eglin, stephane.bres} @insa-lyon.fr
nicole.vincent@math-info.univ-paris5.fr*

RÉSUMÉ. *Pour obtenir une bonne décomposition des traits d'écriture dans les images de manuscrits, l'extraction de l'axe médian représente une étape importante. Dans cet article, nous proposons une approche de repérage de l'axe médian de l'écriture, directement appliquée aux images en niveaux gris. Cette approche est basée sur l'estimation de l'orientation et de l'épaisseur du trait en différents points du tracé et converge à chaque itération en des points positionnés sur la ligne médiane du tracé. Elle est robuste aux dégradations des traits, au bruit de l'arrière plan et aux irrégularités d'imprégnation des encres dans le support papier des documents. Nous procédons ensuite à la classification des graphèmes de l'écriture, obtenus à partir de l'analyse des points de l'axe. Ceux-ci sont exploités à des fins d'identification de style d'écriture. Ce travail s'inscrit dans le cadre du projet ANR Graphem.*

ABSTRACT. *To obtain a good decomposition of handwritten strokes in manuscript images, the extraction of the medial axis represents a very important step. In this paper, we propose an automatic approach for the extraction of the median axis directly from the gray level images. This approach is based on the orientation and the stroke thickness analysis at different points of the handwritten shapes. This method is robust to degradations of the strokes, background noise and ink irregularities in the support paper. Then we propose to consider all glyphs and to classify them from the analysis of median axis points. We use this decomposition to produce an identification scheme from various handwriting styles. This work is supported by the ANR Graphem project.*

MOTS-CLÉS Images de manuscrits, repérage de l'axe médian des traits, matrice hessienne, code book de l'écriture.

KEYWORDS. *Handwriting images, median axis location, Hessian matrix, handwriting codebook*

1. Introduction

Dans le domaine de l'analyse des images de traits, l'extraction de l'axe médian des traits représente une étape très importante. Elle conditionne la qualité de la localisation des éléments de contenus ainsi que leur caractérisation. Par exemple, en bioinformatique, l'axe médian est utilisé pour réaliser un suivi des vaisseaux sanguins ou des neurones. Pour l'analyse des empreintes digitales, l'axe médian est extrait pour l'identification de la personne. Pour l'analyse des images satellites, on retrouve l'axe médian pour l'extraction des routes. En ce qui concerne les images de manuscrits, l'extraction de l'axe médian constitue une étape importante de caractérisation permettant ainsi de calculer aisément l'épaisseur d'un trait, de construire un graphe représentant une forme ou encore de décomposer les traits qui forment une lettre. Dans cet article, nous nous intéressons aux images de traits manuscrits et plus précisément aux manuscrits anciens latins du Moyen Age. Dans ce type de manuscrits, on est confronté à des problèmes très courants, comme le vieillissement des supports et des encres. Dans ce papier, nous proposons d'extraire l'axe médian des traits présents directement sur les images des manuscrits en niveaux de gris, alors que la littérature regorge d'approches procédant par une première étape de binarisation rendant la localisation des traits plus facile. Nous nous distinguons des méthodes utilisées de nos jours en évitant ces approches conventionnelles procédant par seuillage et squelettisation de l'image. Dans ce domaine, il existe de nombreuses approches permettant d'améliorer la qualité des algorithmes des binarisation et de squelettisation. Ces méthodes peuvent donner de bons résultats sur des textes peu dégradés ou réguliers. Mais face aux problèmes d'imprégnation irrégulière des encres et à la qualité souvent médiocre des supports, ce type d'approches n'est pas réellement adapté. Nous proposons une alternative aux approches plus classiques qui permet de procéder à l'extraction de l'axe médian sur des images de manuscrits de qualité, de résolution et de contenus très variables.

Le papier est organisé de la façon suivante. La section 1 introduit les concepts de notre étude et présente les principales contributions antérieures à la notre en relation avec les aspects de squelettisation et de recherche de l'axe médian des traits d'écriture. La section 2 décrit nos contributions principales en trois parties : nous présentons pour commencer l'algorithme original de suivi du tracé en le validant par rapport aux performances estimées des approches de squelettisation les plus fréquemment employées. Nous détaillons ensuite notre approche de décomposition des traits exclusivement basée sur l'épaisseur du trait estimé en chaque point de l'axe médian. Cette section se poursuit avec la construction du code book à partir des traits segmentés en utilisant comme caractéristiques l'écart types des orientations des points de l'axe médian et leur direction principale. La section 3 montre enfin comment la constitution d'un code book pour chaque écriture peut être appliquée à la comparaison des écritures en analysant les variabilités inter-scripteurs. Pour conclure ce travail, la section 4 dresse un bilan de l'outil de décomposition et d'analyse des écritures que nous avons choisi de faire porté sur l'axe médian et annonce les perspectives à court terme de ces contributions.

2 Travaux antérieurs

Nous présentons dans cette section trois méthodes différentes classiquement utilisées pour effectuer le suivi de traits : la squelettisation, les méthodes heuristiques et les méthodes qui utilisent le contour pour la navigation.

2.1 Les méthodes de squelettisation

Ces méthodes appliquent dans un premier temps une binarisation pour que seules les structures importantes apparaissent. D'une manière générale, à travers la binarisation des images en niveaux de gris, des informations importantes peuvent être perdues [1]. D'après Maio et al [2], la binarisation et la squelettisation sont très coûteuses en temps, et les techniques de binarisation qu'ils ont expérimentées se sont avérées être peu satisfaisantes. La principale raison de ces mauvais résultats est liée à la mauvaise qualité des documents utilisés. Les régions étroites disparaissent ce qui conduit à la production de caractères cassés ou de caractères qui se touchent [1]. La binarisation échoue souvent lors de la séparation des caractères avec l'arrière plan et produit ainsi des bords fortement irréguliers. Ces pixels des bords ajoutent quelques faux tracés et trous avec les algorithmes de squelettisation conventionnels et provoquent des erreurs significatives dans le processus d'appariement de traits [3]. Il est de plus en plus communément admis que la squelettisation des traits ne doit plus systématiquement passer par une segmentation, [4]. On peut citer dans ce contexte, les méthodes basées sur les champs potentiels 2D permettant d'identifier le squelette par l'exploitation de modèle potentiel plutôt que des transformées en distance, [4]. Les pixels aux frontières des traits sont alors considérés comme des points de charge générant un champ potentiel et utilisant alors au cœur du trait des fonctions de champs électrostatiques, [14]. Ces approches nécessitent la plupart du temps des opérations de pré traitement des images par lissage des bords des traits, elles ont néanmoins apportés des solutions à l'analyse d'images bruitées, [4].

2.2 Les méthodes heuristiques

On retrouve souvent ce type d'algorithme pour l'extraction du tracé dans les empreintes digitales [2] [5]. Nous commencerons par les présenter pour ce domaine d'application puis nous passerons à l'extraction des caractéristiques des manuscrits. La première méthode qui a utilisé le principe Direct-Gray Scale pour avancer le long des traits dans les empreintes digitales est celle de Maio et Maltoni [2]. Le principe de cette méthode est de naviguer le long des arêtes dans une image en niveau de gris sans avoir besoin de faire une binarisation et d'extraire l'axe médian des arêtes qui forment l'empreinte digitale. Leurs résultats sont bien supérieurs à ceux des méthodes conventionnelles de binarisation et de squelettisation, en termes d'efficacité et de robustesse. Cette méthode nécessite le réglage de 7 paramètres ce qui constitue un problème pour une utilisation automatique et ajoute de la complexité à l'algorithme. Par ailleurs, il est difficile d'obtenir une direction d'arête précise dans les images de mauvaise qualité. Cette méthode utilise un masque gaussien de convolution avec les traits de pixels de l'arête. Puisque la convolution peut conduire à changer la distribution de l'échelle de niveaux de gris locale, les structures des lignes d'arête peuvent ainsi être déformées dans quelques régions de

l'image. Concernant les images de manuscrits, il existe deux approches majeures permettant d'extraire directement les caractéristiques des images en niveaux de gris et d'en déduire la présence des arêtes des traits [6]. La première approche utilise les caractéristiques topographiques. Si on considère l'image en niveau de gris comme une surface, alors ses caractéristiques topographiques correspondent à des caractéristiques de forme de l'image originale. Par exemple : les traits fins vont correspondre à des arêtes et un trou étroit va correspondre à un point de selle. Les caractéristiques topographiques ont démontré qu'elles sont efficaces pour la segmentation et la reconnaissance des caractères qui se touchent ou se superposent, [1], [3], [7], [8]. Cependant, il faut noter que les caractéristiques topographiques qui utilisent les gradients et les courbures en niveau de gris sur la surface de l'image sont sensibles aux bruits et aux défauts. La seconde approche de la catégorie de méthodes heuristiques utilise une corrélation entre portions d'images ayant subi une transformation affine pour aligner les images en niveaux de gris qui s'apparient aisément. Cette méthode a montré qu'elle peut résoudre le problème de distorsions géométriques locales et qu'elle est très résistante au bruit [6].

2.3 Les méthodes basées sur le contour

La méthode présentée dans [9] fait partie de cette famille d'approches. Elle est utilisée pour l'extraction de l'axe médian dans les images de neurones, à partir d'un algorithme itératif et se base sur la corrélation entre une ligne et ses deux bords (qui sont parallèles) [9]. Elle est appliquée directement sur l'image en niveau de gris, sans binarisation préalable. L'algorithme est basé sur les nouvelles méthodes de segmentation du soma, détection des points de semences c.à.d. les points de départs où il faut commencer à faire le tracé, une détection récursive des points qui vont former l'axe médian et le lissage 2D des courbes. L'algorithme est entièrement automatique sans aucune interaction humaine, et suffisamment robuste pour une utilisation sur des images de mauvaise qualité, telles que celles présentant un faible contraste. Mais si l'image de traits contient des contours déformés, comme c'est le cas dans les écritures, alors on se retrouve face à la difficulté de trouver le point suivant à partir du point courant, lors de la progression le long du trait. Des méthodes de ce type sont aussi utilisées dans le suivi des routes sur les photos satellites [10] [11]. L'extraction de la route exploite généralement les caractéristiques spécifiques des routes et des réseaux routiers. Les méthodes présentées dans cette section exploitent toutes des aspects différents des images de traits et sont dédiées à des applications spécifiques en imagerie. Dans le cas très particulier de l'analyse des écritures anciennes, il a été nécessaire de redéfinir une approche robuste permettant de compenser les limites des approches précédentes sur ces données particulières.

3. Méthode proposée pour l'extraction de l'axe médian

3.1 Prétraitement et binarisation

Cette partie consiste à expliquer les différentes étapes de prétraitement qu'on a appliqué sur les images, avant de passer à l'étape d'extraction des caractéristiques. Dans cette étape, on a utilisé la méthode de Frangi [13]. Cette méthode présente l'avantage de préserver les différentes structures de l'image même les structures dégradées. Dans leurs travaux, les auteurs ont travaillé sur des structures tubulaires qui représentent des vaisseaux sanguins. Dans nos travaux, notre contribution porte sur l'extraction de l'axe médian au sein de traits noirs qui ont des propriétés comparables à celles des vaisseaux sanguins. Cette approche va notamment nous permettre de calculer le rayon à chaque point de l'image qui sera utilisé pour décomposer les caractères en fragments distincts. Pour initialiser la méthode, on procède à une convolution de l'image avec la dérivée seconde de la gaussienne G_{xx} , G_{xy} , G_{yy} . Dans notre cas, nous avons réglé la valeur de sigma à 4, car la valeur maximale du rayon des traits qu'on a rencontrés était égale à 4.

Puis, on construit la matrice Hessienne. Les auteurs ont montré que les valeurs propres de la matrice Hessienne qui vérifient certaines propriétés (d'ordonnement : $|\lambda_1| \leq |\lambda_2|$ et $|\lambda_1| \approx 0$) indiquant la structure idéale d'un trait noir. La valeur du rayon est ensuite calculée de façon à maximiser la fonction V_0 définie de la façon suivante : La fonction V_0 consiste en un produit de fonctions exponentielles suivant différentes valeurs d'échelle S (sigma).

$$V_0(s) = \begin{cases} 0 \\ \exp\left(\frac{-R^2\beta}{2\beta^2}\right) \left(1 - \exp\left[\frac{-S^2}{2c^2}\right]\right) \end{cases}$$

Où V_0 est égale à 0 pour $\lambda_2 \succ 0$, $R_\beta = \frac{\lambda_1}{\lambda_2}$, $S = \lambda_1^2 + \lambda_2^2$ et c , la norme de la matrice (Frobenius). On choisit $\beta = 0.5$, le ratio R_β est utilisé pour faire la distinction entre les structures qui forment des lignes et les structures plates. La norme $S = \lambda_1^2 + \lambda_2^2$ (Second order structureness) est petite quand on est dans le fond (background) avec des valeurs propres individuellement petites par manque de contraste. Dans les régions à fort contraste cette norme augmente avec la présence d'au moins une des valeurs propres de plus grande amplitude.

Pour chaque valeur de sigma, on a, pour chaque pixel de l'image, une matrice qui contient les valeurs de la fonction V_0 qui sont représentées par ψ . On fait de même pour les valeurs propres maximales λ_2 , ainsi que pour les directions qui sont calculées à partir des vecteurs propres I_x, I_y où la direction en chaque point est égale à $\arctan 2(I_x, I_y)$. Après avoir construit les différentes matrices pour chaque sigma on compare les valeurs de ψ . La valeur de ψ maximale pour chaque pixel va nous permettre d'extraire la valeur maximale de sigma qui représente la valeur de l'épaisseur du trait en un point donné. Si par exemple la valeur maximale de ψ est celle de σ_1 alors le rayon sera dans ce cas égal à σ_1 .

3.2. Extraction des points de départ et suivi du tracé

Pour l'identification des points de départ, on extrait le squelette en utilisant la méthode de Zhang [16] et on localise les points de fin de trait dans le squelette. Ces points vont servir de points de départ à notre méthode d'extraction d'axe médian.

La procédure d'extraction de l'axe médian se décompose en trois étapes : initialisation, détermination du point suivant et ajustement de la position par itérations successives qui cessent dès que l'on rencontre un point de bifurcation et que ce point a déjà été visité. Pour plus de détails, voir [12].

Pour le calcul de la direction à suivre au fur et à mesure des itérations, nous avons utilisé une combinaison de la direction géométrique et de la direction d'intensité. En utilisant un seul type de direction, nous aurions pu nous trouver face à des difficultés liées à des changements brusques d'orientation ou de diamètre (épaisseur) aux différents points P calculés de l'axe médian et qui sont liés au chevauchement ou à la superposition des traits ou encore à la dégradation de l'image. La direction géométrique est définie comme la direction du point actuel au nouveau point. Elle est estimée et corrigée. Ces estimations et corrections successives, conduisent à une direction géométrique définie entre le point courant P_k et le point suivant P_{k+1} .

Cette direction est dépendante de la topologie géométrique du trait. Quand il y a, par exemple, un changement brusque de courbure entre P_k et P_{k+1} , la direction géométrique seule peut causer de grandes déviations. Pour résoudre ce problème, on la combine avec la direction d'intensité. La matrice Hessienne, déjà utilisée dans l'analyse des vaisseaux sanguins, est définie par

$$H = \begin{bmatrix} I_{xx} & I_{xy} \\ I_{yx} & I_{yy} \end{bmatrix}$$

où I_{uv} représente la convolution de l'image avec une fonction dérivée de gaussienne du second ordre. De la matrice Hessienne, on va extraire deux valeurs propres. Considérant que les traits d'écriture sont plus sombres que le fond, la direction d'intensité est définie comme le vecteur propre de la matrice Hessienne qui correspond à la plus petite valeur propre. La direction de suivi séquentiel est obtenue en multipliant le vecteur propre de la matrice Hessienne par le signe du

produit scalaire de celui-ci et la direction de suivi actuelle $\hat{\psi}_k$, de sorte que la direction d'intensité obtenue pointe toujours sur les sections dans le prolongement du trait. Soit λ_1 et λ_2 les valeurs propres de la matrice Hessienne au point P_{k+1} tel que $|\lambda_1| \leq |\lambda_2|$ et soit v_1 et v_2 les vecteurs propres. La direction d'intensité est

définie par : $\hat{h}_{k+1} = \text{sign}(\hat{\psi}_k v_1) v_1$. Du fait de la dégradation de l'image et du chevauchement (points de croisement) des traits, l'utilisation de la direction d'intensité toute seule peut causer des problèmes. C'est pour cela qu'on la combine à la direction géométrique.

Nouvelle approche pour l'extraction de l'axe médian dans les traits manuscrits

Dans l'algorithme que nous proposons, les directions de suivi sont donc mises à jour en utilisant une combinaison pondérée des indications géométriques (pondération α) ci-dessus et des directions d'intensité (pondération $(1-\alpha)$). De façon pratique, nous utilisons $\alpha=0.5$.

Le profil d'intensité est utilisé dans l'étape de correction pour obtenir un nouveau point P_{k+1} appartenant à l'axe médian. Un point sur le profil d'intensité $\tilde{g}_{k+1}$ dont le résultat de filtrage est maximal est identifié comme le nouveau point de l'axe médian. Pour cela, la taille de la fenêtre joue un rôle très important. Une fenêtre de taille dynamique est utilisée. L'idée principale de l'utilisation d'une fenêtre de taille dynamique est de pouvoir faire face aux différentes configurations géométriques qu'on peut rencontrer durant le suivi d'un trait.

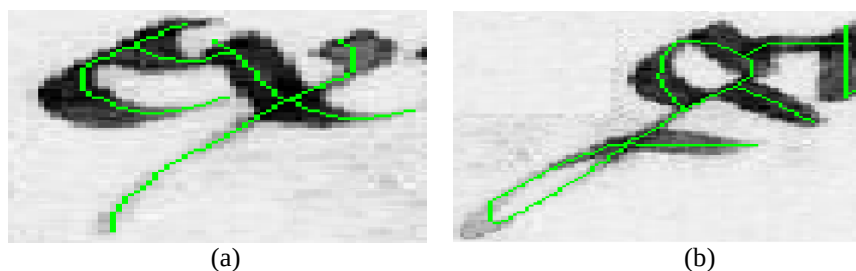


Figure 1- Exemple de suivi du tracé (Croisement et dégradation de l'encre)

On peut voir dans la figure ci dessus (Fig. 1) deux exemples zoomés de suivi de tracé. Dans ce cas précis, nous avons pris en compte le cas de croisement qu'on rencontre dans les manuscrits et le cas de dégradation de l'encre. On peut voir que l'on a pu extraire l'axe médian là où l'encre est la plus dégradée.

Nous avons testé l'algorithme sur des manuscrits médiévaux au sein desquels les traits sont dégradés et difficiles à extraire avec un algorithme de squelettisation standard. Sur la figure 2, nous présentons les résultats de l'extraction de l'axe médian selon notre approche. La première colonne contient les résultats de notre méthode. On peut voir qu'on a pu extraire l'axe médian même dans les situations où l'encre était dégradée. La seconde colonne représente le résultat de la binarisation selon la méthode de Sauvola que nous avons choisie ici car elle a donné les meilleurs résultats sur les documents dégradés anciens, [15]. La troisième colonne illustre la squelettisation par la méthode de Zhang [16] qui a donnée les meilleurs résultats parmi les 8 méthodes de squelettisation testées. Ces deux dernières approches conduisent à de moins bons résultats que notre méthode, en particulier pour des fins de traits quasi confondus avec l'arrière plan.



Figure 2-Tableau de comparaison des trois méthodes, notre Méthode (Colonne 1), Binarization « Sauvola » (colonne 2), squelettisation « Zhang » (colonne 3)

4. Méthode proposée pour l'extraction des graphèmes

En paléographie, les règles d'exécution des écritures sont très strictes : certaines lettres et combinaison de lettres ne peuvent être produites que selon une unique dynamique d'exécution. Il a donc été nécessaire dans cette étude de tenir compte de ces particularités d'exécution afin de produire une décomposition des formes cohérente, évitant notamment certains gestes de rebroussement (retour en arrière du mouvement de la plume). D'un point de vue méthodologique, la segmentation des traits est réalisée de la façon suivante : entre chaque point de départ et d'arrêt, tous les points impliqués dans la formation d'un trait seront sauvegardés dans une liste avec leurs directions et l'épaisseur. Les points d'épaisseur minimale (minimum local) sont ensuite marqués et proposés comme point de découpage, comme cela est effectivement le cas dans la formation d'un trait, voir figure 3. Sur cette figure, chaque segment du tracé a une coloration différente. On identifie par cette approche les zones de croisement, les points de levers et de posers de plume (voir le zoom de la figure 3). La décomposition de la figure 4 montre que les lettres sont constituées de fragments adjacents rattachés en points d'épaisseur minimale supposées correspondre à des points de poser et de lever de plume. Cette décomposition a été soumise à la validation des experts paléographes et qu'elle a obtenu leur approbation.

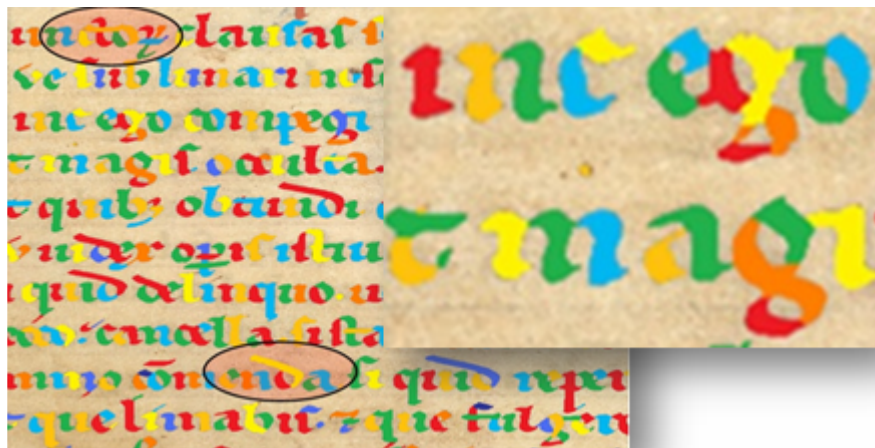


Figure 3- Décomposition des traits en graphèmes

5. Construction du code Book et utilisation

A partir de l'ensemble des graphèmes extraits à l'étape précédente, on va procéder à la construction de notre Code book. On produit ainsi une description vectorielle de chacun des graphèmes, qui sera ensuite traitée pour définir les critères de similarité. Nous avons choisi de nous démarquer des approches de caractérisation par texture usuellement employées sur des blocs de texte homogènes, [3][6], ou des approches de caractérisation différentielles [1]. Nous avons privilégié une analyse individualisée portant sur les graphèmes. Le choix des caractéristiques est un point déterminant pour la qualité finale de l'algorithme. Étant donné que les graphèmes peuvent avoir des tailles variables, l'étape suivante consiste à décrire ces images par un vecteur de caractéristiques de dimension fixe, ici 10. Les estimations de ces descripteurs sont traduites par des valeurs réelles et chacune d'elles décrit une propriété géométrique (taille, courbure, surface, orientation ...) ou statistique (moyennes d'épaisseurs, écart types d'orientations ...) de la forme.

Ces descripteurs définissent un espace de caractéristiques dans lequel chaque graphème est représenté par un point. Pour la classification des graphèmes, nous avons choisi l'algorithme des K-means. La position finale des centres des classes données par le K-Means signe les caractéristiques d'une écriture donnée. Nous avons utilisé 25 classes pour la classification des documents présentés dans cet article. Chaque document sera donc représenté par un code book de 25 clusters. On peut alors estimer les similarités entre documents décrits par l'intermédiaire de ces signatures en estimant une distance entre ces signatures. La distance $d(D1, D2)$ entre le document D1 et le document D2 est ici calculée par la distance de chacun des centres des 25 clusters de D1 au plus proche voisin des centres de D2. La distance $d(D2, D1)$ entre D2 et D1 sera le plus souvent différente. On prendra comme distance symétrique entre D1 et D2, la plus grande des deux distances $d(D1, D2)$ et $d(D2, D1)$.

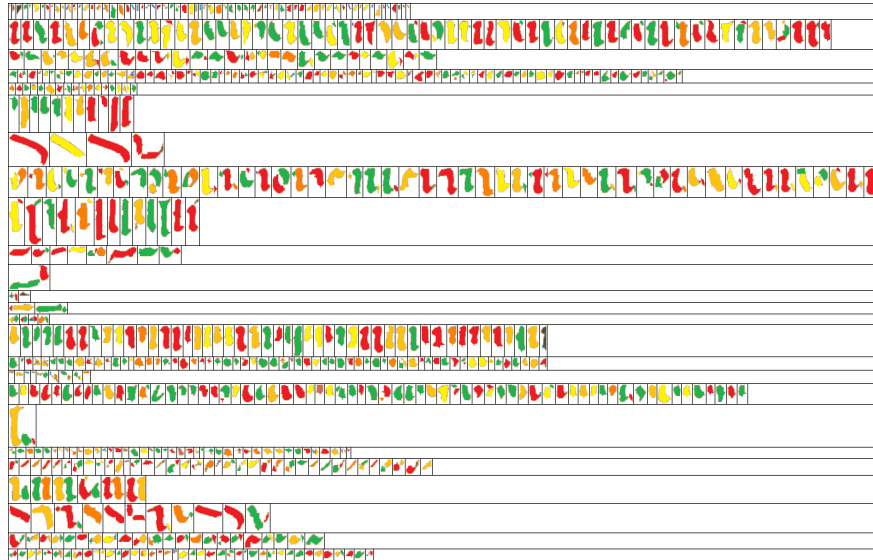


Figure 4: classification des graphèmes en 25 clusters

6. Exemple de comparaisons

Pour illustrer cette distance et l'estimation de la similarité de document ou d'écriture, nous présentons dans le tableau 1 suivant les distances estimées entre les documents de 4 classes. Les documents des classes 1, 3 et 4 appartiennent presque à la même époque mais sont écrits par différents scripteurs. En revanche, la classe 2 appartient à une époque différente, ce qui va conduire à une grande distance entre les classes 1, 3, 4 et la classe 2. La figure 5 si dessous nous montre un exemplaire de chacune de ces 4 classes.

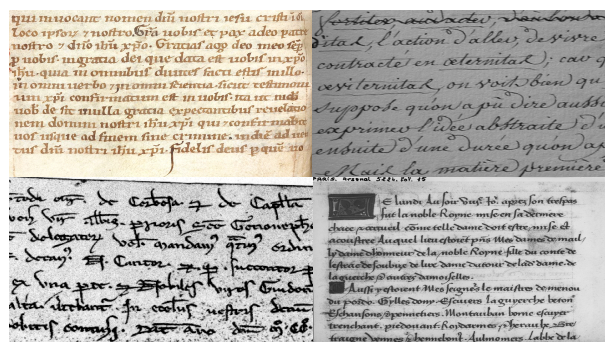


Figure 5- exemplaires des différentes classes : de gauche à droite ligne du haut, classes 1 et 2, ligne du bas classe 3 et 4.

	c11	c12	c13	c14	c21	c22	c23	c24	c31	c32	c33	c34	c41	c42	c43	c44
c11	0.0	2.2	1.5	2.2	15.1	12.6	13.3	13.5	7.0	7.7	6.3	7.9	8.2	8.6	8.8	8.3
c12	2.2	0.0	1.8	2.3	12.9	13.9	13.1	11.1	5.9	7.9	5.7	7.6	7.6	7.8	8.1	7.5
c13	1.5	1.8	0.0	2.3	15.4	13.4	11.3	13.9	8.1	6.7	5.8	7.6	5.9	6.3	7.2	6.0
c14	2.2	2.3	2.3	0.0	13.7	14.1	12.6	13.9	6.2	7.5	6.5	7.4	7.3	6.0	7.4	7.2
c21	15.1	12.9	15.4	13.7	0.0	2.1	2.4	2.3	12.6	13.8	12.2	11.6	15.4	16.3	15.5	17.1
c22	12.6	13.9	13.4	14.1	2.1	0.0	2.2	1.3	15.4	13.4	11.4	12.7	17.9	15.9	15.7	14.3
c23	13.3	13.1	11.3	12.6	2.4	2.2	0.0	1.4	13.9	15.1	12.3	11.2	16.0	17.9	17.2	15.7
c24	13.5	11.1	13.9	13.9	2.3	1.3	1.4	0.0	13.9	17.6	11.3	14.8	14.4	17.1	15.9	15.9
c31	7.0	5.9	8.1	6.2	12.6	15.4	13.9	13.9	0.0	2.3	1.5	2.6	7.7	8.6	8.3	8.5
c32	7.7	7.9	6.7	7.5	13.8	13.4	15.1	17.6	2.3	0.0	1.9	3.2	7.6	6.8	6.2	6.4
c33	6.3	5.7	5.8	6.5	12.2	11.4	12.3	11.3	1.5	1.9	0.0	2.3	7.7	7.5	7.9	6.8
c34	7.9	7.6	7.6	7.4	11.6	12.7	11.2	14.8	2.6	3.2	2.3	0.0	6.6	7.9	8.1	7.6
c41	8.2	7.6	5.9	7.3	15.4	17.9	16.0	14.4	7.7	7.6	7.7	6.6	0.0	1.8	2.5	2.1
c42	8.6	7.8	6.3	6.0	16.3	15.9	17.9	17.1	8.6	6.8	7.5	7.9	1.8	0.0	2.4	1.7
c43	8.8	8.1	7.2	7.4	15.5	15.7	17.2	15.9	8.3	6.2	7.9	8.1	2.5	2.4	0.0	1.6
c44	8.3	7.5	6.0	7.2	17.1	14.3	15.7	15.9	8.5	6.4	6.8	7.6	2.1	1.7	1.6	0.0

Tableau 1- comparaison des différentes classes

Dans le tableau 1, nous avons reporté la distance entre les différents documents. Cxy désigne le document numéro y de la classe x. On constate sur ces valeurs de distance que les documents issus d'une même classe présentent une similarité plus importante (distance plus faible) que les documents issus de classes différentes. Par ailleurs, les documents de la classe C2 présentent une similarité plus faible avec les autres classes, ce qui vérifie notre hypothèse que la classe C2 appartient à une époque et à un style d'écriture très différent.

7. Conclusion

D'après les résultats, on a prouvé que notre méthode a donné de meilleurs résultats de suivi du tracé que les méthodes conventionnelles de squelettisation. Les traits dégradés ont été pris en considération dans notre méthode et on a pu extraire l'axe médian dans les zones où cela est difficile à faire avec les autres méthodes. De même, notre méthode nous a permis de travailler directement sur les images en niveaux de gris sans passer par les étapes de prétraitement, ce qui nous a encore permis de gagner du temps dans l'analyse. Grâce aux traits continus que l'on extrait, nous construisons un ensemble de graphèmes que nous regroupons dans un CodeBook par classe de similarité. Ces classes sont ensuite utilisées comme signature du type du document. Nous utilisons principalement cette méthode pour la classification des documents anciens, selon leur origine et leur type catégorie d'écriture. Cependant, nous testons aussi une extension de cette analyse sur des documents manuscrits modernes, pour lesquels la décomposition en graphèmes pose des problèmes spécifiques. L'objectif sera d'identifier le scripteur du document.

Bibliographie

- [1] Dong-June Lee, Seong-Whan Lee, "A new methodology for gray-scale character segmentation and recognition," ICDAR'95 vol. 1, pp.524,1995.
- [2] Dario Maio, Davide Maltoni, "Direct Gray-Scale Minutiae Detection In Fingerprints," IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 19, no. 1, pp. 27-40, Jan. 1997, doi:10.1109/34.566808
- [3] Suh, J.W., Kim, J.H., Downton, A.C., Impedovo, S., "Stroke extraction from gray-scale character image" Progress in Handwriting Recognition 593-598, 1997
- [4] F Lebourgeois, H. Emptoz. "Skeletonization by gradient regularization and diffusion", ICDAR'07, 23-26 Sept., Curitiba, Brazil. pp. 1118-1122. 2007.
- [5] Yaxuan Qi: Fingerprint Ridge Line Reconstruction. Intelligent Information Processing 2004: 211-220
- [6] T. Wakahara, Y. Kimura, A. Tomono, Affine-invariant recognition of gray-scale characters using global affine transformation correlation, IEEE Trans. Pattern Recognition Mach. Intell. 23 (2001) 384–395.
- [7] L. Wang, T. Pavlidis, "Direct Gray-Scale Extraction of Features for Character Recognition," IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 15, no. 10, pp. 1053-1067, Oct. 1993.
- [8] Seong-Whan Lee, Young Joon Kim, "Direct Extraction of Topographic Features for Gray Scale Character Recognition," IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 17, no. 7, pp. 724-729, July 1995.
- [9] Y. Zhang, X. Zhou, A. Degterev, M. Lipinski, D. Adjero and J. Yuan et al., A novel tracing algorithm for high throughput imaging Screening of neuron-based assays, J Neurosci Methods 160 (2007), pp. 149–162.
- [10] Dal Poz, A.P., Zanin, R.B., do Vale, G.M., 2006, Automated extraction of road network from medium- and high-resolution images, Pattern Recognition and Image Analysis, v 16, n 2, April-June 2006, p 239-48.
- [11] Peteri, R., Celle, J., Ranchin, T., 2003, Detection and extraction of road networks from high resolution satellite images, Proceedings 2003 International Conference on Image Processing (Cat. No.03CH37429), 2003, p I-301-4 vol.1.
- [12] Y. Xu, H. Zhang, H. Li, G. Hu, An improved algorithm for vessel centerline tracking in coronary angiograms, Computer Methods and Programs in Biomedicine, Volume 88, Issue 2, Pages 131-143.
- [13] A.F. Frangi, W.J. Niessen, K.L. Vincken, and M.A. Viergever, "Multiscale Vessel Enhancement Filtering," Medical Image Computing and Computer-Assisted Intervention (MICCAI '98), pp. 130-137, 1998.
- [14] T. Grigorishin, G. Abdel-Hamid and Y.-H. Yang, "Skeletonisation: An Electrostatic Field Based Approach", Pattern Analysis & Applications, vol 1, pp. 163-177, 1998.
- [15] Stathis, P. Kavallieratou, E. Papamarkos, N, "An evaluation survey of binarization algorithms on historical documents" Pattern Recognition, 2008. ICPR 2008. 19th International Conference on Publicationm 2008.
- [16] T.Y. ZHANG et C.Y. SUEN : A fast parallel algorithm for thinning digital patterns. Communications of the ACM, 27(3):236–240, mars 1984