



Approche complexe de l'analyse de documents anciens

Mickaël Coustaty, Nguyen Giap, Vincent Courboulay, Jean-Marc Ogier

► To cite this version:

Mickaël Coustaty, Nguyen Giap, Vincent Courboulay, Jean-Marc Ogier. Approche complexe de l'analyse de documents anciens. Extraction et gestion des connaissances, Jan 2010, Hammamet, Tunisie. pp.597-608. hal-00454559

HAL Id: hal-00454559

<https://hal.science/hal-00454559>

Submitted on 8 Feb 2010

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.

Approche complexe de l'analyse de documents anciens

Mickael Coustaty*, Giap NGuyen*, Vincent Courboulay*, Jean-Marc Ogier*

*Laboratoire L3i - Université de La Rochelle

Pôle Science et Technologie, Avenue Michel Crépeau, 17042 La Rochelle

{mcoustat,vcourbou,giap.nguyen,jmogier}@univ-lr.fr

Résumé. Cet article présente une méthode complexe pour la caractérisation et l'indexation d'images graphiques de documents anciens. A partir d'un bref état de l'art, une méthode pour décrire ces images en tenant compte de leur complexité est proposée. Trois étapes principales de ce traitement sont détaillées dont une méthode novatrice d'analyse, de segmentation et de description des traits. Les résultats sont issus de travaux en cours et sont encourageants.

1 Introduction

L'indexation de documents issus du patrimoine représente un enjeu actuel pour la sauvegarde de nos mémoires. Dans cet article, nous nous intéressons précisément aux images contenues dans ces documents pour leur complexité et leurs caractéristiques particulières. Nous présentons tout d'abord quelques méthodes existantes pour l'analyse de ces images. Puis nous nous intéressons plus particulièrement dans les parties 3 et 4 à la complexité de ces documents et à la manière de décrire leurs éléments constitutifs. Enfin, la dernière partie de cet article présente notre représentation de la complexité dans les documents. Pour cela nous nous appuyons sur une description des images en "*unités et échelles fonctionnelles*". La dernière unité fonctionnelle proposée est novatrice puisqu'elle propose de caractériser les images non plus en utilisant les pixels mais des primitives de plus haut niveau : des traits. Les travaux présentés dans cet article sont réalisés dans le cadre du projet ANR NaviDoMass (2010) et les premiers résultats sont encourageants.

2 Les documents anciens

2.1 Un enjeu actuel

De plus en plus de bibliothèques nationales, de projets nationaux, européens ou mondiaux (NaviDoMass (2010); Europeana; Passe-Partout; Google Livres; Open Content Alliance,...) cherchent à préserver leur patrimoine documentaire. Dans cette optique, de grandes campagnes de numérisation sont actuellement menées par ces différents projets qui cherchent à sauvegarder en masse, au format image, des copies de ces documents. Ces campagnes numérisent de grandes masses de documents et nécessitent de plus en plus souvent des services de navigation pour permettre aux usagers de retrouver les documents. Ces services impliquent la nécessité

Approche complexe de l'analyse de documents anciens

d'organiser les bases d'images, en les indexant sur la base de leurs contenus visuels et doivent répondre aux nécessités de :

1. Conditionner l'image ou le signal souvent fortement dégradé pour l'amener à une représentation exploitable pour la suite des traitements : débruitage, filtrage, ...
2. Caractériser le contenu graphique d'une image pour représenter son contenu visuel - avec tous les outils existants de CBIR, de points d'intérêts, les approches globales, spatiales, vectorielles (statistiques, structurelles), ...
3. Définir des métriques pour mesurer la ressemblance entre les images, sur la base de la technique de caractérisation de contenus retenue (isomorphisme exact ou inexact dans le cas de graphes, distance dans le cas de vecteurs ou d'histogrammes, ...)
4. Structurer les espaces de caractéristiques retenus (graphes, signatures vectorielles, ...) afin de proposer une réponse en temps acceptable à l'utilisateur et afin d'éviter une recherche exhaustive en cas de navigation dans la base (clustering statistiques, clustering de graphes, approches à base de noyaux de graphe, ...)
5. Concevoir des Interfaces Homme-Machine en interaction avec l'utilisateur dans un but d'adaptation du système par exploitation des techniques de bouclages de pertinences

2.2 Des documents de la renaissance

Dans le cadre de cet article, nous nous intéressons plus particulièrement à la recherche par le contenu d'images graphiques de documents du XV^{ème} et XVI^{ème} siècle. Ces documents nous sont fournis par le Centre d'Études Supérieures de la Renaissance (CES) et le sont sur support papier. Le support et la période de création de ces documents font apparaître deux grands types de particularités :

Le support : Le papier présente des traces de vieillissement des documents telles que le jaunissement du papier, la fragilisation, des déchirements ou arrachement ou l'amincissement des pages.

La période : Les documents de cette période étaient imprimés à l'aide de tampons en bois taillés à la main et encrés pour être pressés sur le papier. Ces tampons permettaient de créer des documents en noirs et blancs mais n'autorisaient pas de nuances de gris. Afin de créer des nuances et des ombres, les imprimeurs ont remplacés les nuances de gris par des traits parallèles, typique de la période.

2.3 Les lettrines

Dans cet article, nous nous intéressons particulièrement aux lettrines. Elles correspondent à des images très utilisées dans les ouvrages et très réutilisées au cours du temps. Une lettrine est une lettre ornementale qui débute un chapitre ou un paragraphe et peut-être vue comme une image binaire composée de traits. Quelques travaux (Journet et al. (2008); Pareti et al. (2006); Uttama et al. (2006); Bigun et al. (1996)) ont été menés pour caractériser ce type d'image graphique en particulier.

3 Comment analyser les documents anciens ?

3.1 L'approche des historiens

Les historiens décomposent les lettrines en quatre éléments (Jimenes (2008)). Elles peuvent être vues comme une superposition de trois couches (fond, motif et lettre) insérées dans un cadre (voir Figure 1).

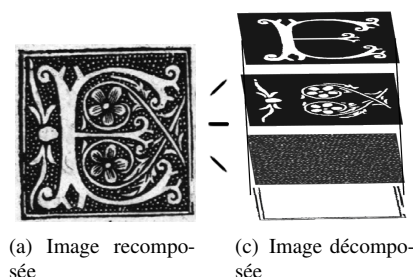


FIG. 1 – Illustration de l'assemblage des éléments constitutifs des lettrines

Le fond : correspond à l'arrière plan et peut-être plein (noire ou blanc), hachuré ou criblé.

Le motif : correspond aux ornements. Il peut être décoratif ou figuratif.

La lettre : est un élément clef qui peut-être noir ou blanc et de différentes polices.

Le cadre : composé d'aucun, un ou deux filets.

Les historiens cherchent à retrouver les images dans de grandes masses de données pour retracer l'histoire et la vie de l'époque. Ils extraient cette information en identifiant les changements qui apparaissent dans les images.

Ces dégradations sont intéressantes pour les historiens puisqu'elles permettent de déterminer un ordre chronologique entre les documents. Tous ces changements sont identifiés manuellement par les historiens à l'aide d'un thésaurus. Cette tâche est fastidieuse et non envisageable dans un contexte d'analyse automatique de grandes masses de données. Dans ce but, des outils informatiques de détection d'éléments particuliers ont été développés, nous allons les présenter ci-dessous.

3.2 L'approche informatique

Au cours des trois dernières années, plusieurs travaux ont cherché à caractériser et indexer les images de documents anciens par leur contenu. Ceux-ci se sont tout d'abord attardés sur l'extraction de la structure des documents (Journet et al. (2008)) à l'aide d'une fonction d'auto-corrélation. Une deuxième catégorie de travaux, appliqués cette fois sur les images graphiques peut également être distinguée. Ainsi, Bigun et al. (1996) extraient dix coefficients de Fourier sur six radiogrammes à différentes orientations et calcule une distance à partir de ces coefficients ; Chen et al. (2003) extraient des points d'intérêt d'images du XVII^{ème} siècle et décrivent leurs alentours à l'aide de moments de Zernike, et Baudrier et al. (2007) ont proposé une méthode de comparaison d'images à partir de cartes de dissimilarités locales. Enfin,

des travaux appliqués spécifiquement aux lettrines (Pareti et Vincent (2006); Chouaib et al. (2009)) utilisent une loi de Zipf à partir de motifs. Ces motifs sont soit utilisés pour construire la courbe de Zipf, soit pour extraire des caractéristiques qui sont ensuite utilisées pour indexer les images. Toutes ces méthodes sont résumées dans le tableau 1. Elles reposent sur une approche cartésienne du problème d'indexation, avec une simplification de l'image sans prise en compte de sa complexité.

Réf.	Type de données	Description
Journet et al. (2008)	Pages entières	Séparation des éléments d'une page à l'aide de descripteurs textures
Baudrier et al. (2007)	Images graphiques	Comparaison d'images à partir de leur dissimilarités
Pareti et Vincent (2006)	Images graphiques	Identification du style des lettrines à partir d'une loi puissance
Chouaib et al. (2009)	Lettrines	Identification du motif à partir d'une loi puissance
Chen et al. (2003)	Images du XV^{me}	Recherche de points d'intérêts dans des images non composées de traits
Bigun et al. (1996)	Images d'ornements	Identification des structures linéaires parallèles

TAB. 1 – Tableau récapitulatif des méthodes de la littérature

4 Une approche complexe

4.1 La complexité

Pour reprendre un extrait d'Edgar Morin (1996), sociologue, philosophe et grand penseur français de la complexité, *le principe de simplicité impose de disjoindre, le principe de complexité enjoint de relier, tout en distinguant*. Toutes les approches scientifiques jusqu'à la moitié du XX^{me} siècle cherchaient à simplifier les problèmes compliqués pour les résoudre. Il faut dissocier les problèmes complexes des problèmes compliqués. La complexité par définition signifie *ce qui est tissé ensemble, ce qui est relié*, ce qui n'impose en rien d'être composé de problèmes compliqués. Au contraire, là où un problème compliqué va nécessiter une simplification sans chercher à respecter son environnement, un problème complexe nécessitera des allers-retours entre description certaine par simplification et description incertaine de son contexte. Il paraîtrait donc essentiel d'intégrer, dans un schéma d'ensemble, séparabilité et logique, avec une séparabilité de l'information pour gagner en certitude et un respect de la logique globale.

4.2 La complexité dans les documents anciens

La complexité dans les documents s'articule autour d'une décomposition. Mais la complexité impose également d'ordonner et de remettre cette certitude dans son contexte, afin d'obtenir une description cohérente de l'image dans celui-ci. La description d'une lettrine d'un point de vue texture uniquement, sans respect de la lettre ne paraît pas cohérente. Elle correspond à une vision de traicteur d'images qui consiste à ne s'appuyer que sur des descripteurs bas niveau (de couleurs, de formes, de textures, ...). L'objectif global est d'essayer de revenir à des informations utilisées dans la production des lettrines pour obtenir une esthétique visuelle. Dans ce cadre, l'idée est de regrouper par "catégories" les différentes couches visuelles (dans notre cas : textures, formes, contours, ...).

Une méthode développée par Utama et al. (2006), propose de décomposer une image pour identifier le contenu des lettrines d'une manière complexe. Le processus de segmentation repose sur une décomposition en différentes couches d'informations (formes, textures, contours, ...) inspiré des principes de perception visuelle. A partir de cette étape, chaque couche va contenir les zones propres à un type d'information. Deux méthodes ont été implémentées pour décrire les images à partir de ces couches, tout en conservant une description globale de l'image. Une première repose sur le calcul de la longueur de l'arbre couvrant minimum (MST, cf. Hero et Michel (1998) tandis que la seconde méthode calcule un histogramme des fréquences d'apparitions des relations entre composantes connexes de l'image.

Cette dernière approche extrait des zones d'intérêt mais seules les relations spatiales entre ces zones sont décrites sans prendre en compte leur contenu. Partant de ce constat, nous avons cherché à étudier chaque élément de l'image et à décrire son contenu, tout en caractérisant les relations qui unissent les différentes zones d'intérêt.

5 L'approche complexe pour l'analyse de documents anciens

Le cadre général de ce travail cherche à décrire les images graphiques de documents anciens pour les indexer en tenant compte de leur complexité. La complexité implique un entremêlement d'éléments hétérogènes. Pour décrire chacun de ces éléments, nous les séparons et les étudions. Cette séparation repose sur une analyse en "*Échelles et Unités Fonctionnelles*", qui permettent de naviguer au sein de la description entre incertitude et imprécision. Cette analyse permet de mettre en avant des éléments particuliers de l'image auxquels sont associés des traitements particuliers et adéquats. Chacun de ces traitements sera intégré dans une unité fonctionnelle, rassemblées elles-mêmes au sein d'échelles fonctionnelles. Chaque échelle permettra d'avancer dans la description de l'image entre imprécision et incertitude sur la nature de son contenu. Le schéma de la Figure 2 présente cette démarche dans son ensemble.

5.1 Échelles et Unités fonctionnelles

Nous appelons *unité fonctionnelle* une unité capable d'extraire et de caractériser un (ou des) élément(s) aux spécificités semblables. Ces unités peuvent être généralisées et/ou adaptées.

Approche complexe de l'analyse de documents anciens

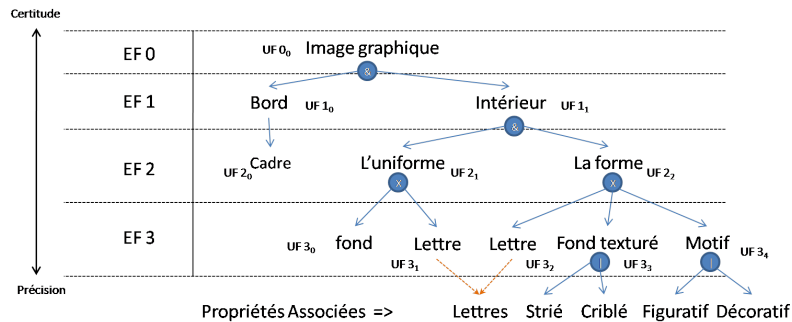


FIG. 2 – Illustration du schéma de caractérisation des éléments constitutifs des lettrines

Nous appelons *échelle fonctionnelle* un niveau d'observation et/ou d'analyse du contenu de l'image réunissant des unités fonctionnelles cohérentes.

Transitivité de l'information la plus certaine Puisque les unités fonctionnelles s'appuient sur une analyse multi-échelles, l'information la plus certaine doit pouvoir se retrouver dans les niveaux inférieurs de la pyramide d'échelle. Cette propriété implique l'utilisation de coefficients de caractérisation décimables, c'est-à-dire qu'un sous-ensemble de ces coefficients doit permettre d'obtenir une information certaine, là où l'ensemble des coefficients doit permettre une description précise.

Coefficients de caractérisation Les coefficients de caractérisation doivent permettre de suivre le principe d'incertitude/imprécision et ainsi de pondérer, dans la suite du traitement, le poids du contexte ou de la précision (caractéristiques globales ou locales / vectorielles ou structurales). Il faut noter qu'il existe un très grand nombre de descripteurs dans la littérature et que nombre d'entre eux ont été initialement définis pour les images naturelles. La possibilité de les utiliser avec des images de documents anciens (et encore plus dans le cas d'images de traits) n'est pas triviale. Par exemple, l'utilisation d'outils de description des textures usuellement utilisés dans la littérature (Fourier, Gabor, Hermitte, Corrélation, ...) est définie pour des images en niveaux de gris. Dans le cas d'images de traits, qui initialement ne comportaient que deux niveaux (encre ou pas encre), ces outils ne fonctionnent plus directement et nécessite une adaptation, comme dans Journet et al. (2008).

Dans la suite, nous présentons deux méthodes développées dans l'analyse et la caractérisation des images de lettrines.

5.2 Séparation de la forme - UF_{22}

Cette unité fonctionnelle permet de séparer les différents éléments des éléments appartenant aux formes des images. Celle-ci se compose de zones lisses, c'est-à-dire avec de faibles variations de niveaux de gris, et de zones texturées, c'est-à-dire avec des variations rapides de

niveaux de gris. Cette unité fonctionnelle repose sur une décomposition basée sur l'algorithme d'Aujol et Chambolle (Meyer (2001)). Les images de lettrines sont principalement composées de lignes, ce qui rend les approches habituelles inappropriées pour séparer les traits des zones uniformes. Nous avons donc utilisé l'approche développée par Dubois et al. (2008) pour séparer l'image en différentes couches d'informations plus facilement caractérisables.

Les couches Notre but est de séparer la composante géométrique de la composante texture, indépendamment du bruit. La décomposition de Meyer (voir Figure 3) permet d'obtenir trois composantes pour l'image :

- La couche régularisée correspond aux zones de l'image qui ont une faible variation de niveaux de gris. Cette couche permet de mettre en évidence la composante géométrique qui correspond aux formes de l'images.
- La couche oscillante qui correspond aux zones aux variations rapides de niveaux de gris. Dans notre cas, cette couche permet de mettre en évidence les textures des lettrines, c'est-à-dire les zones composées de traits parallèles.
- La couche très oscillante qui correspond au bruit dans l'image. En fait, cette couche est composée de tout ce qui n'appartient pas aux deux premières couches. Dans notre cas, on peut retrouver dans cette couche le texte du verso de la page et les problèmes dus au vieillissement du papier.

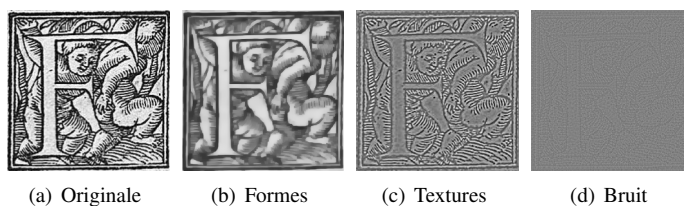


FIG. 3 – Un exemple de décomposition de lettrine

Traitement spécifique Chaque couche est alors traitée comme une image composée d'éléments uniformes (première couche composée uniquement de formes tandis que la deuxième n'est composée que de textures). Chacune de ces couches sera traitée par une méthode appropriée au type de contenu. Dans le cas des formes, nous extrayons la lettre en utilisant une loi de Zipf. La couche texture est elle analysée et décrite par l'agencement des différents traits qui la composent.

5.3 Extraction de la lettre dans la forme - UF_{3_1}

Les lettrines sont composées de nombreuses zones homogènes ce qui rend impossible l'extraction de la lettre par un algorithme d'extraction de composantes connexes. Dans notre cas, nous avons basé l'extraction de la lettre sur la couche régularisée obtenue après décomposition. Cette couche contient toutes les formes de l'image non segmentées. Une segmentation reposant sur une loi de Zipf est utilisée pour extraire la lettre et caractériser les lettrines. Cette unité

fonctionnelle, comme décrit dans Coustaty et al. (2009), repose sur un processus en quatre étapes :

1. Extraction des motifs à l'aide d'une loi de Zipf
2. Séparation de l'image en couches
3. Segmentation des formes
4. Extraction de la lettre

Des exemples de résultats d'extraction de la lettre peuvent être observés Figure 4.



FIG. 4 – Exemples de lettres extraites automatiquement à partir des lettrines

5.3.1 Expérimentations

L'évaluation d'un tel système est un point fondamental puisqu'il garantit son utilisabilité et qu'il permet d'avoir un regard objectif sur le système. Dans un objectif d'indexation de lettrines, nous avons décidé d'évaluer la qualité de l'extraction de la lettre par son taux de reconnaissance par des OCR. Pour l'évaluation, nous avons utilisé deux OCR grands publics, un commercial (FineReader) et un open-source (Tesseract). Nous avons expérimenté notre approche sur une base de 916 images avec les dictionnaires standards français de ces OCR. Les résultats obtenus (voir tableau 2) sont encore insatisfaisants mais réellement encourageants. Nous travaillons actuellement à l'amélioration de la chaîne de traitement (segmentation et sélection des formes). Cependant, il n'existe pas à l'heure actuelle de système similaire au nôtre et les historiens ont validé nos résultats.

	FineReader	Tesseract
Taux de reconnaissance	72,8%	67,9%

TAB. 2 – Taux de reconnaissance de lettrines en utilisant deux types d'OCR

5.4 Extraction et caractérisation des traits dans les lettrines - UF_{32}

Nous proposons dans cette dernière partie une nouvelle approche pour la caractérisation des lettrines reposant sur l'analyse des traits qui les composent. La stratégie appliquée consiste à repenser l'algorithme classique de l'analyse des images, qui s'appuie généralement sur des informations pixellaires, pour l'adapter aux images de traits. La méthodologie retenue consiste donc à considérer que l'information élémentaire n'est plus le pixel, mais le trait. Le postulat amène donc à repenser toute la stratégie d'analyse en intégrant cette propriété particulière de nos images. Cette méthode comprend trois étapes : l'extraction des traits, leur caractérisation et le regroupement des traits aux caractéristiques similaires, pour la construction de régions.

5.4.1 Extraction des traits

L'extraction des traits dans les images repose sur une série de traitements. Tout d'abord nous binarisons l'image à l'aide du critère d'Otsu (1979). A partir de cette image binaire, nous calculons le nombre d'Euler (Pratt (2007)) sur un découpage fin de l'image. Le signe du nombre d'Euler permet, en soustrayant le nombre de composantes connexes blanches au nombre de composantes connexes noires, d'obtenir la couleur du plus grand nombre de composantes connexes. Celle-ci est alors associée aux traits. Enfin, afin de faciliter la caractérisation des documents tout en conservant leurs propriétés, nous squelettisons les images à l'aide d'une transformée en distance (Breu et al. (1995)).

5.4.2 Caractérisation des traits

Une fois les traits extraits, nous les étiquetons à l'aide de caractéristiques pour les différencier et les caractériser. Nous avons retenu des caractéristiques basées sur l'épaisseur, l'orientation et la courbure des traits pour leur représentation similaire à celle de la vision humaine (Graham et Sutter (1998)). L'épaisseur est obtenue à partir de la transformée en distance de l'étape précédente, l'orientation principale est obtenue par une transformée de Radon (Helgason (1994)) et la courbure correspond à la valeur maximale du rapport entre l'orientation et l'orientation orthogonale du trait.

5.4.3 Classification des traits

Chaque trait est maintenant défini par un vecteur de trois caractéristiques. Nous utilisons ce vecteur associé à un classifieur hiérarchique pour définir des classes de traits. Ces classes permettent d'identifier les traits en fonction des caractéristiques extraites. Nous utilisons un classifieur hiérarchique car celui-ci ne nécessite aucune connaissance à priori sur les images et le nombre de classes souhaité.

Nous avons défini une métrique entre les vecteurs des traits qui permet d'exprimer la similarité entre deux traits. Comme l'illustre la Figure 5, l'épaisseur d'un trait a une influence sur la signification de sa longueur. C'est pourquoi, nous avons défini la longueur relative d'un trait dans la formule 1 avec l la longueur du squelette du trait (en pixels) et e la moitié de l'épaisseur du trait (voir Figure 5(c)).

$$l_{relative} = \frac{l}{2 * e} \quad (1)$$

A partir de ce critère, nous en déduisons la règle de pondération des différentes caractéristiques du vecteur (l'épaisseur a un poids constant égal à 1 ; p_o = Poids associé à l'orientation et p_c = Poids associé à la courbure) :

1. $p_o = p_c = 1$ si $l_{relative} \geq 2$,
2. $p_o = p_c = l_{relative} - 1$ si $1 \leq l_{relative} < 2$

$$d_{traits} = \sqrt{(e_1 - e_2)^2 + p_o(o_1 - o_2)^2 + p_c(c_1 - c_2)^2} \quad (2)$$

Pour classer des traits par une méthode hiérarchique, nous devons tout d'abord construire un arbre où chaque nœud représente des classes de traits. Au début de la construction de l'arbre,

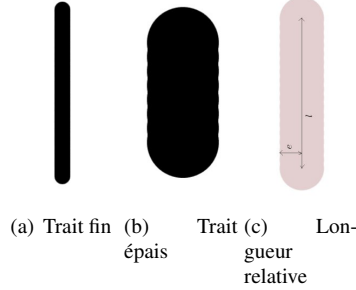


FIG. 5 – Importance de l'épaisseur des traits pour la vision humaine

chaque nœud représente un trait de l'image. A chaque itération de l'algorithme, si le nombre de nœuds est supérieur à un, les nœuds les plus proches sont fusionnés au sein d'un nouveau nœud. La notion de proximité repose sur deux critères :

- Les plus proches au sens de la métrique définie ci-dessus
- Respect de la condition d'inconsistance des nœuds

L'inconsistance d'un nœud de l'arbre peut être défini par l'équation 3 avec d_{arbres} , la distance entre ses deux sous-arbres, $\bar{d}_{arbres}$ la moyenne et σ l'écart-type des distances entre les sous-arbres de ses sous-arbres. Ce calcul revient à comparer les variances intra-classes et inter-classes des nœuds et évite la fusion d'éléments décorrélés.

$$I = \frac{d_{arbres} - \bar{d}_{arbres}}{\sigma} \quad (3)$$

5.4.4 Regroupement des traits pour la segmentation

Une fois tous les traits classés à l'aide de l'arbre hiérarchique, chaque trait se retrouve associé à une classe. La dernière étape consiste donc à rassembler au sein d'un même groupe tous les traits appartenant à la même classe proches spatialement. Dans une image I , on définit le voisinage d'un trait comme étant la partie du fond adjacente au trait. Si deux traits partagent le même voisinage et la même classe, alors ils sont groupés. Un exemple de résultat obtenu en utilisant cette approche est présenté dans la Figure 6. Une première évaluation par un expert sur 228 images composées principalement de traits a permis d'obtenir les résultats proposés dans le tableau 3. Les premiers résultats obtenus sont encourageants et des expérimentations plus poussées sont en cours.

6 Conclusion et perspectives

Cet article présente une nouvelle méthode pour l'analyse d'images graphiques complexes de documents anciens. Celle-ci repose sur une décomposition de l'image en échelles et unités fonctionnelles. Trois unités fonctionnelles sont présentées pour caractériser des éléments particuliers. En particulier, une approche basée sur les traits et non sur les pixels des images

	Sous seg- mentées	Bien seg- mentées	Sur seg- mentées
Nombre d'images	16	193	19

TAB. 3 – Nombre d'images bien segmentées, sous ou sur segmentées dans une base de 228 images

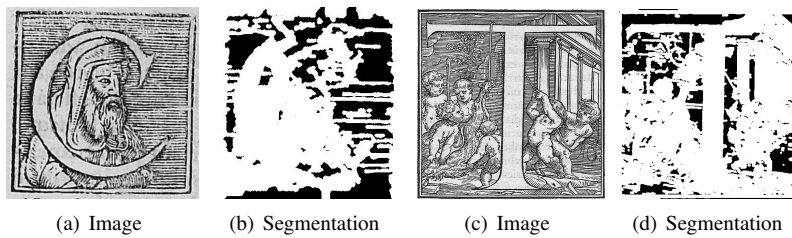


FIG. 6 – Exemple de segmentation d'une texture dans une lettrine

est utilisée. Cette approche est novatrice et les premiers résultats sont encourageants. En perspective de ces travaux, nous prévoyons de combiner les différentes descriptions obtenues dans chaque échelle fonctionnelle et d'indexer nos images à partir de ces différents niveaux de descriptions. Ainsi, à chaque même image sera associée une représentation d'ensemble et précise des images.

Références

- Les bibliothèques virtuelles humanistes - centre d'étude supérieur de la renaissance - <http://www.bvh.univ-tours.fr/>.
- Baudrier, E., N. Girard, et J.-M. Ogier (2007). A Non-symmetrical Method of Image Local-Difference Comparison for Ancient Impressions Dating. In *Seventh IAPR International Workshop on Graphics Recognition (GREC'07)*, Curitiba Brésil, pp. 257–265.
- Bigun, J., S. Bhattacharjee, et S. Michel (1996). Orientation radiograms for image retrieval : An alternative to segmentation. *Pattern Recognition, International Conference on* 3, 346.
- Breu, H., J. Gil, D. Kirkpatrick, et M. Werman (1995). Linear time euclidean distance transform algorithms. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 17, 529–533.
- Chen, V., A. Szabo, et M. Roussel (2003). Recherche d'images iconique utilisant les moments de zernike. In *Actes de CORESA'03*, Lyon.
- Chouaib, H., F. Clopet, et N. Vincent (2009). Graphical drop caps indexing. In *Eighth IAPR International Workshop on Graphics Recognition*, La Rochelle, pp. 179–185.
- Coustaty, M., J.-M. Ogier, R. Pareti, et N. Vincent (2009). Drop caps decomposition for indexing : an new letter extraction. In *International Conference on Document Analysis and Recognition*, Volume 1, Barcelona, Spain, pp. 476–480.

Approche complexe de l'analyse de documents anciens

- Dubois, S., M. Lugiez, R. Péteri, et M. Ménard (2008). Adding a noise component to a color decomposition model for improving color texture extraction. *CGIV 2008 and MCS'08 Final Program and Proceedings*, 394–398.
- Europeana. <http://www.europeana.eu/portal/>.
- Google Livres. <http://books.google.fr/>.
- Graham, N. et A. Sutter (1998). Spatial summation in simple (fourier) and complex (non-fourier) texture channels. *Vision Research* 38(2), 231–257.
- Helgason, S. (1994). *Geometric Analysis on Symmetric Spaces*. American Mathematical Society.
- Hero, A. et O. Michel (1998). Asymptotic theory of greedy approximations to minimal k-point random graphs. *IEEE Trans. on Inform. Theory* 45(6), 1921–1939.
- Jimenes, R. (2008). Les bibliothèques virtuelles humanistes et l'étude du matériel typographique. Technical report, Centre d'Etude Supérieur de la Renaissance.
- Journet, N., J.-Y. Ramel, R. Mullot, et V. Eglin (2008). Document image characterization using a multi-resolution analysis of the texture : application to old documents. *IJDAR* 11(1), 9–18.
- Meyer, Y. (2001). *Oscillating patterns in image processing and nonlinear evolution equations*. The fifteenth dean jacqueline B. Lewis Memorial Lectures.
- Morin, E. (1996). Pour une réforme de la pensée. 49(2), 10–14.
- NaviDoMass (2007-2010). Navigation, into documents masses - <http://l3iexp.univ-lr.fr/navidomass/>.
- Open Content Alliance. <http://www.opencontentalliance.org/>.
- Otsu, N. (1979). A Threshold Selection Method from Gray-level Histograms. *IEEE Transactions on Systems, Man and Cybernetics* 9(1), 62–66.
- Pareti, R. et N. Vincent (2006). Ancient initial letters indexing. In *ICPR '06*, Washington, DC, USA, pp. 756–759. IEEE Computer Society.
- Pareti, R., N. Vincent, S. Uttama, J.-M. Ogier, J.-P. Salmon, S. Tabbone, L. Wendling, et S. Adam (2006). On defining signatures for the retrieval and the classification of graphical drop caps. *International Workshop on Document Image Analysis for Libraries*, 220–231.
- Passe-Partout. Recherche d'ornement - <http://www2.unil.ch/bcutodai/app/todai.do>.
- Pratt, W. (2007). *Digital Image Processing*. Wiley.
- Uttama, S., P. Loonis, M. Delalandre, et J. Ogier (2006). Segmentation and retrieval of ancient graphic documents. In *Graphics Recognition. Ten Years Review and Future Perspectives*, pp. 88–98.

Summary

This article proposes a complex method on characterizing and indexing old documents graphical images. The proposed method permit to describe elements included in images and their complexity. Three principals steps are detaillled and one of them analyse, segment and describe images using strokes. First results are interesting.