

Recherche par le contenu dans une base d'images fixes : l'intérêt des règles d'association

Anicet Kouomou Choupo*, Annie Morin*, Laure Berti-Équille*

*IRISA Campus de Beaulieu 35042 Rennes-France
{akouomou, Annie.Morin, Laure.Berti-Equille}@irisa.fr
<http://www.irisa.fr/texmex/>

Résumé. Une base d'images fixes peut être décrite de plusieurs façons, notamment par des descripteurs visuels globaux de couleur, de texture, ou de forme (niveau pixel). Les requêtes les plus fréquentes impliquent et combinent les résultats de plusieurs types de descripteurs : par exemple, "trouver toutes les images ayant une couleur et une texture semblables à celles d'une image requête donnée". Pour retrouver plus efficacement et plus rapidement une image dans une base d'images fixes, nous exploitons des combinaisons appropriées de descripteurs. Dans ce papier, nous étudions l'intérêt des règles d'association entre clusters de descripteurs lors de la recherche dans une base d'images fixes. Après avoir décrit le système de recherche qui nous sert de cadre d'étude, nous nous focalisons essentiellement sur le temps de réponse à des requêtes puis nous confrontons notre système à l'approche classique de recherche séquentielle avec fusion des résultats.

Mots-clé : règles d'association, recherche par le contenu, image fixe, clustering.

1 Introduction

La problématique de la recherche d'informations multimédias par le contenu est rendue plus complexe lorsque la base devient conséquente. Le temps de réponse aux requêtes n'est plus négligeable et il convient de déployer des techniques d'indexation adaptées aux données multimédias en grandes quantités. Dans ce contexte, nous avons envisagé d'exploiter les résultats de deux techniques de fouille de données dans la mesure où elles nous permettent d'utiliser des informations complémentaires sous forme de règles sur les données traitées. L'extraction de connaissances à partir de données multimédias apporte cependant une complexité nouvelle liée à la représentation des données analysées. La représentation des documents multimédias est soit indépendante du contenu (format, origine des données, date, auteur, conditions de création du document, ...), soit directement liée au contenu (mots-clés, fichiers inversés, descripteurs d'images, ...). Nous nous sommes intéressés à cette dernière forme de représentation pour une base d'images fixes.

Une base d'images fixes peut être décrite de différentes façons, notamment par des descripteurs visuels globaux de couleur, de texture, ou de forme (au niveau pixel). Plusieurs descripteurs sont proposés dans la littérature (Manjunath et al. 2002, Obeid et al. 2001, Tao 1999) et certains sont standardisés comme MPEG-7 (Manjunath et al. 2002). Chacun d'eux est défini selon l'information que l'on souhaite extraire de l'image. On privilégiera par exemple un descripteur de forme si l'on souhaite retrouver toutes les images contenant un clavier d'ordinateur. Par contre, un descripteur de couleur sera indiqué si l'on recherche des images de la mer ou

du coucher du soleil. Les requêtes les plus fréquentes impliquent et combinent les résultats de plusieurs types de descripteurs : par exemple, "retrouver toutes les images ayant une couleur et une texture semblables à celles d'une image requête donnée". Le traitement de ce type de requête exige de considérer plusieurs aspects parmi lesquels le choix des bons descripteurs, la recherche sur chacun des descripteurs choisis, et la définition d'une bonne stratégie de fusion des résultats obtenus sous forme de liste ordonnée. Il est légitime de se demander si l'on peut trouver des relations entre les descripteurs existants et si ces relations permettent une recherche plus rapide et plus efficace d'une image dans une base donnée.

L'objectif de notre travail est de combiner et d'exploiter deux techniques de fouille : le clustering et la recherche de règles d'association pour améliorer la recherche dans une base de données complexes telles les images. Après avoir décrit le système de recherche que nous développons afin de valider notre étude, nous nous focalisons essentiellement sur le temps de réponse à des requêtes, puis nous confrontons notre système à l'approche classique de recherche séquentielle avec fusion des résultats.

Cet article est organisé de la manière suivante : dans la section 2 nous présentons sommairement le principe de description des images fixes ainsi qu'une synthèse de l'état de l'art des différents usages multimédias des règles d'association. Nous présentons ensuite notre méthode de recherche dans la section 3 et discutons de nos résultats. La section 4 conclut l'article et présente nos perspectives à moyen et long terme.

2 Description d'images et usage des règles d'association

2.1 Description d'images fixes

La description d'une image a pour but de rassembler tous les éléments nécessaires à la caractérisation de l'image dans un contexte d'utilisation donné. Dans le domaine médical par exemple, la détection automatique du cancer du sein à partir des images de mammographie passe par l'identification et la localisation des cellules anormales. La nature et le type de descripteurs proposés par la communauté du traitement d'images dépendent totalement et précisément du but de la recherche (par exemple un descripteur pour trouver le visage de G. Bush). Quatre niveaux de description sont cependant envisageables (Zhang et al. 2001) : le pixel, l'objet, le niveau sémantique, et la connaissance. Au niveau pixel, une image est considérée comme un ensemble de points ou de pixels dont on calcule les caractéristiques de couleur, de texture, et/ou de forme. Il peut s'agir par exemple de la détection d'un cercle dans une image. Le but au niveau de l'objet est de pouvoir identifier dans une image des objets comme un ballon. Au niveau sémantique, l'on s'appuie sur la connaissance du domaine (médical, spatial, sportif, ...) pour extraire des concepts de haut niveau à partir des objets identifiés dans une image. Il est ainsi possible de distinguer les images d'un match de football. Le niveau de la connaissance s'appuie sur le niveau sémantique auquel il est rajouté des informations complémentaires décrivant le contenu. À une image de football par exemple, on peut rajouter des mots décrivant les équipes en compétition, les actions de tir au but, et le score du match.

Notre étude est volontairement limitée à la description visuelle des images qui sont traitées comme des ensembles de pixels. Dans ce cadre, étant donnée une image i , la description de i consiste à trouver un vecteur $d = f(i)$ qui résume les caractéristiques de l'image i .

Le descripteur d de l'image i est un vecteur de réels ou d'entiers de dimension n et f la fonction de calcul du descripteur. Cette fonction doit être robuste aux variations des conditions de production de l'image et à l'orientation de celle-ci (Smeulders et al. 2000). Nous limiterons notre étude à 5 descripteurs pour nos images décrivant respectivement : la couleur (*Color Layout*, *Scalable Color*), la texture (*Homogeneous Texture*, *Edge Histogram*) et la forme (*Region Shape*) ; tous proposés dans le standard MPEG-7 (Manjunath et al. 2002).

2.2 Usage des règles d'association

Le but de la technique des règles d'association est la génération d'associations informatives à partir d'une grande masse de données. Une formalisation de la technique est introduite pour la première fois dans (Agrawal et al. 1993).

Soit $I = \{i_1, i_2, \dots, i_N\}$ un ensemble de N éléments distincts appelés items, objets ou attributs, et D un ensemble de transactions ayant comme attributs les éléments de I . Une règle d'association est une implication de la forme $A \rightarrow B$, où $A, B \subset I$, et $A \cap B = \emptyset$. A est appelé *antécédent* et B *conséquent* de la règle. Un ensemble d'objets ou d'items est appelé *itemset*. A chaque itemset, on associe une mesure statistique appelée *support*, et notée *supp*. Pour un itemset $A \subset I$, $\text{supp}(A) = s$, si la fraction de transactions de D contenant A est égale à s . Une mesure de confiance permet de caractériser la pertinence d'une règle. Pour $A \rightarrow B$, elle est définie par $\text{conf} = \text{supp}(A \cup B) / \text{supp}(A)$.

Déterminer des règles d'association consiste à extraire toutes les règles de la forme $A \rightarrow B$ ayant un support et une confiance supérieurs ou égaux à leurs seuils respectifs *minsupp* et *minconf* fixés à l'avance. Plusieurs algorithmes d'extraction des règles d'association tels que *Apriori* (Agrawal et al. 1993), *Pascal* (Bastide et al. 2002) sont proposés dans la littérature. Ces méthodes d'extraction marchent très bien avec les bases de données transactionnelles dont la structure est bien définie. Mais leur application aux images nécessite des adaptations. Une idée consiste à transformer les images et à appliquer les techniques classiques d'extraction de règles. Une image sera par exemple identifiée à une transaction et les différentes régions la constituant à des objets (Ordóñez 1999). La répétition d'objets peut dans certaines conditions être porteuse d'informations. Un grand nombre de régions de couleur bleue dans une image donne une idée de l'intensité de la couleur bleue. La définition classique de règles d'association ne tient pas compte de cette possibilité de répétition et il est nécessaire dans ce cas de développer de nouveaux formalismes. C'est ainsi que Zaïane et al. (Zaïane et al. 2000) proposent une définition de *règle d'association avec objets récurrents* comme une implication de la forme :

$$\alpha P_1 \wedge \beta P_2 \wedge \dots \wedge \gamma P_n \rightarrow \delta Q_1 \wedge \lambda Q_2 \wedge \dots \wedge \mu Q_m$$

où $P_i, i \in [1..n]$ et $Q_j, j \in [1..m]$ sont des descripteurs d'images, et $\alpha, \beta, \gamma, \delta, \lambda, \mu$ sont des entiers indiquant le nombre d'occurrence des descripteurs. L'algorithme *MaxOccur* décrit dans (Zaïane et al. 2000) est une adaptation de *Apriori* à l'extraction des règles d'association avec objets récurrents.

Nous proposons une application des règles d'association à l'indexation et à la recherche par le contenu dans une base d'images fixes. Notre approche s'inscrit dans le cadre de la transformation des données multimédias et de l'application des techniques classiques d'extraction de règles. Nous utilisons plusieurs descripteurs et les images sont regroupées en clusters pour chaque descripteur. Nous identifions ensuite une image à une transaction et un numéro de clus-

ter pour un descripteur donné à un attribut. Le problème de répétition d'objets ou d'attributs ne se pose pas puisque les clusters sont identifiés de manière unique par descripteur. Parmi les travaux sur les règles d'association dans un contexte de recherche par le contenu multimédia, Djeraba utilise une technique d'indexation voisine de la nôtre en ce sens qu'elle combine la construction des clusters et la détermination des règles d'association (Djeraba 2003). Sa base de travail est formée de plusieurs sous-groupes d'images classées thématiquement (animaux, plantes, ...) et l'idée consiste à caractériser ces groupes par des règles d'association. Lors de la recherche, une analyse de la requête à partir des règles d'association oriente le choix du sous-groupe d'images. Le travail de Djeraba a pour objectif l'amélioration de la qualité de recherche. Il ne s'intéresse pas au problème de fusion des résultats lorsque la recherche se fait suivant plusieurs descripteurs. Le temps de la recherche n'est pas évalué. De plus, l'approche est contrainte par la construction semi-automatique d'un dictionnaire nécessitant une labélisation des clusters. Dans notre approche par contre, nous restons au niveau du pixel et aucun étiquetage n'est nécessaire.

D'autres travaux exploitent les règles d'association (Bouet 2002, Aouiche et al. 2003). Le premier tente de capturer le niveau sémantique par la découverte des relations existantes entre les images d'une base de données pour réduire l'espace de recherche. L'approche d'Aouiche et al. destinée à l'auto-administration d'une base de données relationnelle, fonde le calcul de l'index sur les ensembles d'attributs les plus fréquemment utilisés dans les requêtes. Cette approche détermine des relations entre attributs (sous forme de motifs fréquents) à partir de la fréquence d'utilisation de ceux-ci. Elle fait donc une administration à partir d'un journal d'utilisation de la base (intervention indirecte de l'utilisateur) alors que nous faisons une indexation sans intervention de l'utilisateur. Nous ne supposons aucune organisation préalable de la base d'images fixes et nous ne limitons pas le nombre ainsi que le type de descripteurs utilisables.

3 Description de la méthode d'indexation et de recherche

Nous proposons une méthode utilisant la classification automatique et la recherche de règles d'association pour améliorer la recherche par le contenu dans une base d'images fixes. Notre approche se décompose en deux principales étapes fonctionnelles autour desquelles s'articule l'architecture de notre chaîne de traitement des images (voir Figure 1) : la procédure d'indexation et la procédure de recherche que nous détaillons ci-après.

3.1 L'indexation

L'indexation de la base se fait hors-ligne. Elle est composée de trois modules : le descripteur de données, le classifieur automatique, et le générateur de règles. Le descripteur de données calcule tous les descripteurs à partir d'un ensemble de fichiers d'images stockés dans un répertoire. Nous travaillons avec les descripteurs MPEG-7 de couleur (*ColorLayout*, *ScalableColor*), de texture (*HomogeneousTexture*, *EdgeHistogram*) et de forme (*RegionShape*). Notre système génère un fichier décrivant toutes les images pour chacun des descripteurs utilisés (étape I1 de la Figure 1). Le classifieur automatique organise les fichiers de descripteurs en clusters pour chaque descripteur, construit un index d'accès à chaque groupe formé et le stocke dans un fichier (étape I2). Chaque cluster est identifié par un numéro et la nature du descripteur pour lequel le cluster est calculé.

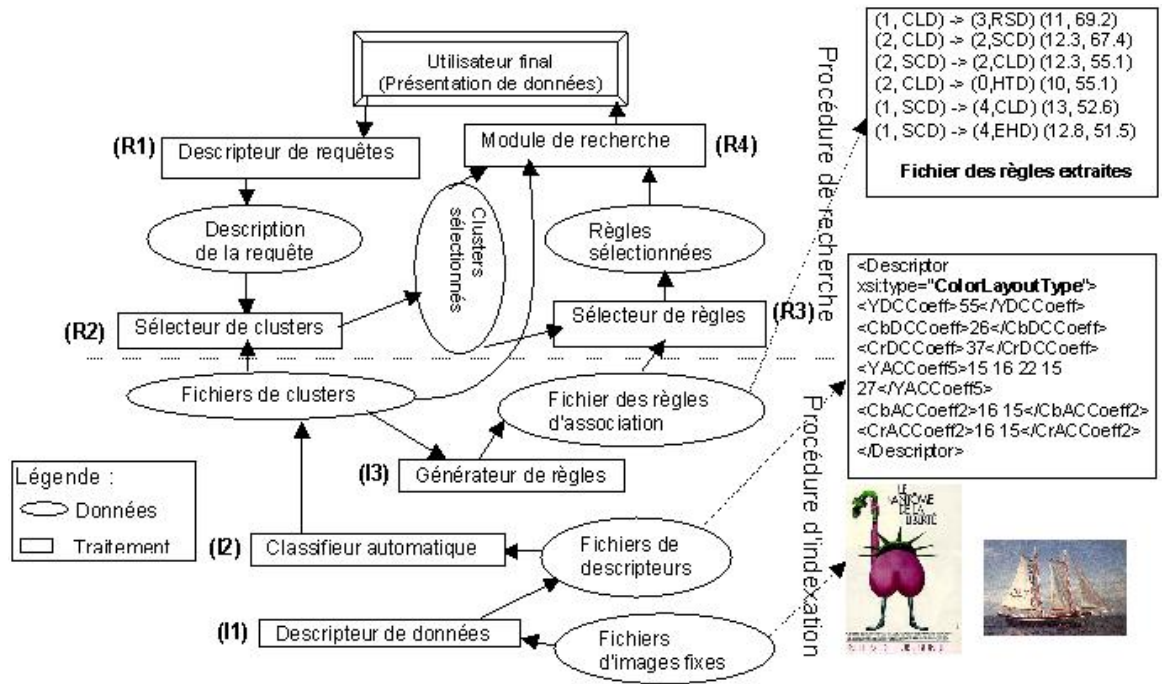


FIG. 1 – Processus d'indexation et de recherche d'images fixes.

Le générateur de règles utilise les fichiers de clusters produits par le classifieur automatique pour extraire des relations entre clusters sous forme de règles d'association (étape I3). Une règle d'association ainsi générée est une implication de la forme :

$(\langle \text{num cluster} \rangle, \langle \text{nom descripteur} \rangle) [(\langle \text{num cluster} \rangle, \langle \text{nom descripteur} \rangle) \dots] \rightarrow (\langle \text{num cluster} \rangle, \langle \text{nom descripteur} \rangle) (\langle s \rangle, \langle c \rangle)$

ayant la sémantique suivante :

$\langle \text{num cluster} \rangle$ est le numéro de cluster pour le descripteur $\langle \text{nom descripteur} \rangle$;
 $\langle s \rangle$ et $\langle c \rangle$ sont des pourcentages indiquant respectivement le support et la confiance de la règle étudiée. La partie gauche d'une règle est composée d'un ou de plusieurs couples $(\langle \text{num cluster} \rangle, \langle \text{nom descripteur} \rangle)$ identifiant les clusters. Nous limitons la partie droite à un seul couple.

3.2 La recherche

La recherche se fait en-ligne. L'utilisateur soumet au système une image requête. Son but est de retrouver toutes les images semblables à la requête spécifiée. Deux cas peuvent se présenter : soit l'utilisateur choisit les descripteurs suivant lesquels la recherche doit être faite, soit l'utilisateur n'a aucune idée des descripteurs calculés. Nous penchons pour le second cas

le plus fréquent. Le module de description de requête traite la requête soumise et produit un descripteur de chaque type dans la limite des descripteurs gérés par la procédure d'indexation (étape R1). Le sélecteur de clusters utilise les fichiers de clusters et la description de la requête afin de déduire pour chaque descripteur, le cluster dans lequel la requête se trouve (étape R2).

Nous utilisons les règles d'association pour réduire le nombre de clusters dans lesquels nous faisons de la recherche séquentielle. Le sélecteur de règles choisit parmi les règles disponibles celles qui décrivent les relations entre les clusters fournis par le sélecteur de clusters (étape R3). Une règle est choisie si tous les clusters impliqués dans la règle sont des éléments de l'ensemble des clusters sélectionnés à l'étape R2. Pour faciliter l'écriture, nous adoptons les notations abrégées *CLD*, *SCD*, *HTD*, *EHD*, et *RSD* pour les descripteurs respectifs *Color Layout*, *Scalable Color*, *Homogeneous Texture*, *Edge Histogram*, et *Region Shape*. Les clusters sont numérotés de 0 à 4 pour chaque descripteur. Supposons par exemple une requête I_q et les valeurs suivantes fournies par le sélecteur de clusters : $(2, CLD)$, $(2, SCD)$, $(1, HTD)$, $(3, EHD)$, et $(3, RSD)$. Les règles sélectionnées seront donc $(2, CLD) \rightarrow (2, SCD)$ (12.3, 67.4) et $(2, SCD) \rightarrow (2, CLD)$ (12.3, 55.1) (voir Figure 1 pour l'ensemble de toutes les règles du système).

Les clusters et les règles sélectionnés sont transmis au module de recherche (étape R4). Si aucune règle n'est sélectionnée, alors la recherche séquentielle est faite dans tous les clusters sélectionnés. Dans le cas contraire, la recherche séquentielle se fera uniquement dans les clusters n'apparaissant pas en partie droite de règle. L'exemple précédent est un cas particulier où la recherche séquentielle est faite uniquement dans les clusters $(1, HTD)$, $(3, EHD)$, et $(3, RSD)$. On aurait souhaité que la recherche se fasse aussi dans $(2, CLD)$ ou $(2, SCD)$ (mais pas dans les deux). Cette situation n'est pas encore gérée par le système et fera l'objet d'une étude plus élaborée. Les résultats de la recherche séquentielle dans les clusters sont ensuite fusionnés selon le principe suivant (Nepal 1999) :

1. pour chaque cluster C_i fourni par le sélecteur de clusters, retourner le prochain couple $(I, s_{C_i}(I))$ où I est l'image la plus proche de l'image requête I_q pour la mesure de score s_{C_i} (comprise entre 0 et 1), puis calculer $s_{C_j}(I)$ pour les clusters C_j où $i \neq j$
2. calculer $f_h = f(s_{C_i}(I_i), \dots, s_{C_m}(I_m))$ où I_i est l'image du cluster C_i à l'itération courante. f est une fonction de synthèse des scores. Nous l'avons prise égale à la fonction $\min$
3. calculer $s_Q(I) = f(s_{C_i}(I), \dots, s_{C_m}(I))$ pour tout I de l'itération courante et mettre à jour l'ensemble $Y = \{I | s_Q(I) > f_h\}$
4. répéter 1, 2, et 3 jusqu'à ce que Y contienne le nombre d'images désiré, puis retourner $\{(I, s_Q(I)) | I \in Y\}$ à l'utilisateur.

3.3 Discussion

La méthode présentée ci-dessus est en cours d'implémentation en C++ sous Linux. Nous travaillons pour l'instant avec une base de 7727 images fixes mais nous envisageons plusieurs centaines de milliers d'images à long terme. Pour chacun des 5 descripteurs MPEG-7 utilisés, nous regroupons les images en clusters avec un algorithme de type $k - means$. La complexité temporelle du regroupement en clusters d'images ne nous préoccupe pas ici car l'indexation se fait hors-ligne. Nous avons maintenu le seuil de la confiance à 50% pour le calcul des règles d'association entre clusters avec l'algorithme *Apriori* (Agrawal et al. 1993) afin de garantir

la bonne qualité des implications. La figure 2 montre les variations du nombre de règles en fonction du nombre de clusters pour les trois valeurs suivantes du seuil du support : 0.1%, 1%, et 10%. Nous avons travaillé avec le même nombre de clusters pour tous les 5 descripteurs. L'expérimentation peut être étendue en considérant des nombres de clusters différents pour les descripteurs. L'on constate que plus le nombre de clusters est élevé, moins on obtient de règles. Dans la suite des expérimentations, nous avons fixé le nombre de clusters à 5 pour chacun des descripteurs et le seuil du support à 10%. Ce sont en effet des valeurs pour lesquelles nous avons obtenu des règles de support significatif. Dans ces conditions, le système produit 6 règles dont le support varie entre 10% et 13%. Ce support relativement faible s'explique. En effet, la valeur du support est une fonction décroissante du nombre de clusters choisi par descripteur. Si l'on suppose par exemple une répartition uniforme des images dans chacun des 5 clusters pour chacun des descripteurs, alors le support des règles est majoré par 20%.

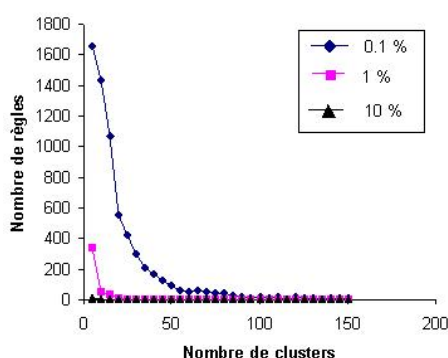


FIG. 2 – Variations du nombre de règles en fonction du nombre de clusters pour 3 valeurs du seuil du support.

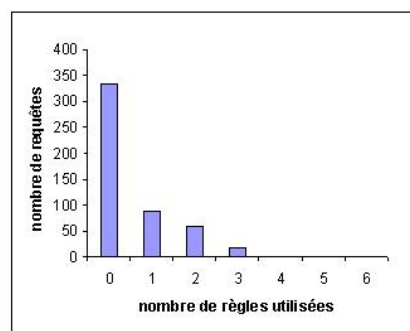


FIG. 3 – Histogramme d'utilisation des règles.

Le système conçu a été interrogé par 500 requêtes, toutes déjà présentes dans la base de 7727 images. Pour 165 d'entre elles, le système fait usage des règles d'association calculées, soit un taux d'utilisation de 33%. L'histogramme d'utilisation des règles est présenté à la figure 3. On constate que le système utilise rarement plus de trois règles dans le traitement d'une requête.

De façon générale la méthode proposée présente un temps de recherche inférieur au temps de recherche séquentielle avec fusion des résultats comme présenté dans le tableau 1. D'autre part, l'usage des règles d'association réduit le nombre de descripteurs à explorer et donc le temps recherche aussi.

Le tableau 2 décrit le comportement général de notre méthode et de celle de la recherche séquentielle sur un ensemble de 500 requêtes. L'on note par exemple que le temps moyen de recherche par notre approche reste voisin de celui de la recherche séquentielle. Nous pensons donc qu'il reste encore à faire pour améliorer ce résultat, par exemple en étudiant beaucoup plus finement les stratégies d'exploitation des règles d'association. L'introduction de la sélection des règles à partir d'une fonction de coût de calcul est envisageable. La recherche guidée par des règles d'association offre donc une perspective intéressante à explorer.

Règles d'association et recherche par le contenu

	Temps moyen de recherche séquentielle avec fusion	Temps moyen de recherche avec le système proposé
Usage des règles	46.50 secondes	44.77 secondes
Pas d'usage des règles	46.71 secondes	45.25 secondes

TAB. 1 – *Comportement de notre système par rapport à la recherche séquentielle avec fusion des résultats.*

	Recherche séquentielle avec fusion	Recherche avec le système proposé
Temps minimum (secondes)	3.96	2.12
Temps maximum (secondes)	257.82	257.90
Temps moyen (secondes)	46.64	45.09
Écart-type	23.26	23.77

TAB. 2 – *Récapitulatif du temps d'exécution des requêtes.*

Le problème du calcul automatique du nombre de clusters reste entier. Dans ce travail, nous adoptons une approche manuelle, mais la démarche pour déterminer le nombre de clusters reste introuvable. Une tentative de résolution du problème est présentée dans (Fernandez et al. 2002). Mais la complexité du regroupement en clusters est plus grande.

4 Conclusion et perspectives

Nous décrivons le principe de fonctionnement d'un système de recherche d'images par le contenu en cours de développement. Ce système sert de cadre pour l'étude de l'intérêt des règles d'association dans la recherche par le contenu. Dans notre démarche, nous utilisons 5 descripteurs MPEG-7 pour décrire à terme plusieurs milliers d'images fixes. Dans un premier temps, nous déterminons des clusters d'images pour chaque descripteur. Nous générons ensuite sous forme de règles d'association, les relations entre les différents clusters obtenus. Nous montrons que les règles d'association sont envisageables pour améliorer le temps de recherche.

- Nous nous proposons de travailler sur une base plus grande d'images dans la suite de nos expérimentations. Ainsi, le temps d'indexation ne sera plus négligeable et nous essayerons des techniques de clustering plus performantes que le *k - means*.
- D'après notre étude, les règles d'association obtenues sont d'un support faible. Il nous paraît donc intéressant d'explorer d'autres mesures afin d'en déduire les plus adaptées à l'évaluation des règles. Il est également utile d'étudier des stratégies plus fines de sélection et d'exploitation des règles pour en améliorer le taux d'utilisation par le système lors de la recherche.
- Nous nous sommes limité au cours de notre étude à l'évaluation du temps de recherche. La comparaison de la qualité des résultats du système proposé et de celle de la recherche

séquentielle avec fusion sera la prochaine étape de la validation de notre approche. Cette étape passe par une définition des mesures adaptées à l'évaluation de la qualité des résultats.

- Enfin, l'expérimentation des techniques statistiques voisines des règles d'association nous semble indispensable pour évaluer notre approche. Nous pensons par exemple à l'analyse des correspondances multiples.

Références

- Agrawal R., Imielinski T., Swami A. (1993), Mining Association Rules Between Sets of Items in Large Databases, ACM SIGMOD International Conference on Management of Data, 1993, pp 207-216.
- Aouiche K., Darmont J., Gruenwald L. (2003), Extraction de motifs fréquents pour l'auto-administration des bases de données, Revue des Sciences et Technologies de l'Information, 2003, p 547.
- Bastide Y., Taouil R., Pasquier N., Stumme G., Lakhal L. (2002), Pascal : un algorithme d'extraction des motifs fréquents, Technique et science informatique, Vol 21, No 1, 2002, pp 65-95.
- Bouet M., Khenchaf A. (2002), Traitement de l'information multimédia : Recherche du média image, RSTI - Ingénierie des systèmes d'information. Bases de données et multimédia, 7(5-6), 2002, pp 65-90.
- Djeraba C. (2003), Association and Content-Based Retrieval, IEEE Transactions on Knowledge and Data Engineering, Vol. 15, No.1, 2003, pp 118-135.
- Fernandez G., Meckaouche A., Peter P., Djeraba C. (2002), Intelligent Image Clustering, Lecture Notes in Computer Science, Vol. 2490, 2002, pp 406-419.
- Manjunath B.S., Salembier P., Sikora T. (2002), Introduction to MPEG-7, John Wiley & Sons, 2002.
- Nepal S., Ramakrishna M.V. (1999), Query Processing Issues in Image (Multimedia) Databases, Proceedings of the ICDE, 1999, pp 22-29.
- Obeid M., Jedynak B., Daoudi M. (2001), Image indexing and retrieval using intermediate features, Proceedings of the 9th ACM/ICM, 2001, pp 531-533.
- Ordonez C., Omiecinski E. (1999), Discovering Association Rules based on Image Content, Proceedings of the IEEE Advances in Digital Libraries Conference(ADL'99), 1999, pp 38-49.
- Smeulders A.W.M., Worring M., Santini S., Gupta A., Jain R. (2000), Content-Based Image Retrieval at the End of the Early Years, Vol. 22, No.12, 2000, pp 1349-1380.
- Tao Y., Grosky W. (1999), Object-Based image retrieval using point feature maps, Proceedings of the International Conference on Database Semantics(DS-8), 1999, pp 59-73.
- Zaïane O., Han J., Zhu H. (2000), Mining Recurrent Items in Multimedia with Progressive Resolution Refinement, Proceedings of the 16th IEEE International Conference on Data Engineering(ICDE'00), 2000, pp 461-476
- Zhang J., Hsu W., Lee M. L. (2001), Image Mining: Issues, Frameworks and Techniques, Proceedings of the 2nd International Workshop on Multimedia Data Mining, 2001, pp 13-20.

Summary

Fixed image database can be described in several ways by global visual descriptors of color, texture, or form (pixel level). Most frequent queries imply and combine results of several type of descriptors such as: "retrieve all images that have similar color and similar texture to the given example image". To retrieve more efficiently and more effectively an image of a large database, we exploit combinations of descriptors. In this papier, we study the effet of association rules between clusters of descriptors on fixed image content based retrieval. We first describe the system retrieval which is used as our study framework, then we set our focus on query response time. We finally compare our system with the sequential retrieval followed by result fusion.

Keywords: association rules, content based retrieval, fixed image, clustering.