

Vers une nouvelle génération d'outils d'analyse et de recherche d'information

A New Generation of Analysis and Retrieval Tools

Hacia una nueva generación de herramientas de análisis y de investigación de información

Dominic Forest

Volume 55, numéro 2, avril-juin 2009

URI : <https://id.erudit.org/iderudit/1029091ar>

DOI : <https://doi.org/10.7202/1029091ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Association pour l'avancement des sciences et des techniques de la documentation (ASTED)

ISSN

0315-2340 (imprimé)

2291-8949 (numérique)

[Découvrir la revue](#)

Citer cet article

Forest, D. (2009). Vers une nouvelle génération d'outils d'analyse et de recherche d'information. *Documentation et bibliothèques*, 55(2), 77-89. <https://doi.org/10.7202/1029091ar>

Résumé de l'article

Les récents efforts visant à favoriser la diffusion et la circulation de l'information en format numérique ont contribué au phénomène de l'infobésité (*information overload*). Il est désormais important de concevoir des outils de recherche d'information plus adaptés aux besoins des utilisateurs afin de leur permettre de récupérer les documents pertinents répondant à leurs besoins informationnels. Dans cet article, nous ferons état, dans un premier temps, de certaines observations sur les conséquences découlant des limites des outils traditionnels en recherche d'information numérique. Dans un deuxième temps, nous exposerons les concepts et les techniques de base du domaine de la fouille de textes, en insistant sur les opérations de classification et de catégorisation automatiques. Finalement, nous montrerons comment certaines techniques de fouille de textes peuvent contribuer au développement d'une nouvelle génération d'outils de recherche d'information.

Vers une nouvelle génération d'outils d'analyse et de recherche d'information

DOMINIC FOREST

Professeur adjoint
École de bibliothéconomie et des sciences de l'information
Université de Montréal
dominic.forest@umontreal.ca

RÉSUMÉ | ABSTRACTS | RESUMEN

Les récents efforts visant à favoriser la diffusion et la circulation de l'information en format numérique ont contribué au phénomène de l'infobésité (information overload). Il est désormais important de concevoir des outils de recherche d'information plus adaptés aux besoins des utilisateurs afin de leur permettre de récupérer les documents pertinents répondant à leurs besoins informationnels. Dans cet article, nous ferons état, dans un premier temps, de certaines observations sur les conséquences découlant des limites des outils traditionnels en recherche d'information numérique. Dans un deuxième temps, nous exposerons les concepts et les techniques de base du domaine de la fouille de textes, en insistant sur les opérations de classification et de catégorisation automatiques. Finalement, nous montrerons comment certaines techniques de fouille de textes peuvent contribuer au développement d'une nouvelle génération d'outils de recherche d'information.

A New Generation of Analysis and Retrieval Tools

The recent efforts to increase the dissemination and circulation of information in numeric format have led to a phenomenon known as information overload. It is now imperative to develop retrieval tools that are better adapted to the users' needs and that will enable them to retrieve relevant documents that meet their information needs. In this article, we will begin with a summary of observations of the consequences of the limits of the traditional tools used in numeric information retrieval. Following this, we will describe the concepts and techniques of textual searching, emphasizing automatic classification. Lastly, we will demonstrate how certain textual searching techniques contribute to the development of a new generation of information retrieval tools.

Hacia una nueva generación de herramientas de análisis y de investigación de información

Los recientes esfuerzos tendientes a favorecer la difusión y la circulación de la información en formato digital han contribuido al fenómeno de la infobesidad (sobrecarga de información). Es importante, de aquí en adelante, diseñar herramientas de búsqueda de información que se adapten mejor a las necesidades de los usuarios a fin de facilitarles la recuperación de documentos pertinentes que respondan a sus necesidades informacionales. En este artículo, realizaremos, en un primer momento, ciertas observaciones sobre las consecuencias que derivan de los límites de las herramientas tradicionales para la búsqueda de información digital. En un segundo momento, expondremos los conceptos y las técnicas de base del dominio del registro de textos, insistiendo sobre las operaciones de clasificación y de categorización automáticas. Finalmente, mostraremos cómo determinadas técnicas de registro de textos pueden contribuir al desarrollo de una nueva generación de herramientas de búsqueda de información.

Introduction

LES DIX DERNIÈRES ANNÉES ont été caractérisées par une hausse importante du nombre d'initiatives visant à numériser et à rendre disponible, très souvent sur Internet ou à l'intérieur d'un Intranet, le patrimoine informationnel des organisations. De nos jours, la majorité des gestionnaires de l'information ne remettent plus en question les nombreux avantages que peuvent procurer les documents en format numérique (reproductibilité quasi parfaite et illimitée, dissémination rapide de l'information, etc.). Parallèlement, les nombreux avantages associés à la documentation numérique ont motivé le développement d'infrastructures performantes dédiées à la numérisation, ainsi qu'à la diffusion de l'information. À cet égard, les manifestations du numérique sont nombreuses et des plus variées. Ainsi, Bibliothèques et Archives nationales du Québec (BAnQ) est un des principaux acteurs d'un projet visant à permettre la diffusion en format numérique de plusieurs dizaines de milliers de documents généalogiques¹. À l'échelle internationale, il importe de souligner le succès, tant en termes de la qualité de la numérisation que de la quantité de documents disponibles et consultés du projet *Européana*², de la bibliothèque numérique *Gallica*³ ou encore du projet ARTFL⁴. Dans le domaine de la recherche universitaire, l'intérêt pour la diffusion numérique des résultats de la recherche scientifique est à la base de nombreux travaux en cours cherchant à promouvoir le développement de cyberinfrastructures pour la diffusion numérique des résultats de la recherche. Dans la Francophonie, les deux principaux chantiers dans ce domaine sont les projets Synergies⁵ et Adonis⁶.

L'évolution du numérique, indissociable du Web, est d'une telle ampleur qu'il est très difficilement quantifiable⁷. Cependant, les inconvénients du numérique,

1. <http://www.banq.qc.ca/portal/dt/genealogie/genealogie.jsp>

2. <http://dev.europeana.eu/>

3. <http://gallica.bnf.fr/>

4. <http://humanities.uchicago.edu/orgs/ARTFL/>

5. <http://www.synergiescanada.org/>

6. <http://www.tge-adonis.fr/>

7. Certaines études, dont la plus citée a été menée à l'Université de Berkeley en 2003 (*How much information?*), ont tenté de chiffrer la croissance du numérique. Si toutes les études démontrent que la quantité de documents disponibles en format numérique augmente radialement, aucun consensus n'émerge des études lorsqu'il s'agit de chiffrer précisément cette évolution.

Dans cet article, nous définirons d'abord le domaine de la fouille de textes, puis exposerons les principes fondamentaux des deux processus de fouille à la base de nouvelles stratégies de recherche d'information. Finalement, nous présenterons deux exemples d'outils de recherche sur le Web qui reposent en grande partie sur l'utilisation de techniques de classification et de catégorisation automatique pour la recherche d'information en ligne.

La fouille de textes⁸

Définition du domaine

La fouille de textes (aussi nommée forage de textes, ou encore *text mining*) est un nouveau domaine de recherche interdisciplinaire dont l'objectif général est le développement de techniques, d'algorithmes, d'applications informatiques et de méthodologies applicatives afin d'extraire et de structurer automatiquement, ou semi-automatiquement, de nouvelles informations à partir de grands corpus de documents textuels non structurés ou semi-structurés.

Ce domaine récent a fait l'objet de plusieurs définitions variant en fonction de l'angle sous lequel il est abordé. Ainsi, les définitions qui ont été proposées par la communauté informatique accordent une plus grande importance aux questions traitant des algorithmes de fouille. En contrepartie, les définitions proposées par les spécialistes du traitement automatique des langues (TAL) sont davantage focalisées sur les enjeux linguistiques et terminologiques. Malgré les divergences de points de vue, il semble possible de dégager un noyau commun consensuel. Celui-ci réside dans les opérations d'identification, d'extraction et de mise en relation de nouvelles informations présentes dans de grands corpus de documents. À cet égard, la définition proposée par Hearst (2003) permet de bien cerner les principaux traits caractéristiques du domaine de la fouille de textes. Selon Hearst, « *la fouille de textes est la découverte (à l'aide d'outils informatiques) de nouvelles informations en extrayant différentes données provenant de plusieurs documents textuels. Un élément fondamental de ce processus réside dans les relations identifiées entre les informations extraites afin d'identifier de nouveaux faits ou de nouvelles hypothèses à explorer* » (2003, notre traduction).

La première étape de tout processus de fouille de données réside dans le développement ou la constitution d'un corpus de documents.

Les origines de la fouille de textes proviennent principalement d'un autre domaine, celui de la fouille de données. Le domaine de la fouille de textes peut d'ailleurs être perçu comme une variante plus complexe que celui de la fouille de données. Si ces deux domaines font appel à des algorithmes et à des processus de traitement identiques, ou à tout le moins très comparables, ils sont radicalement différents en ce qui a trait à la nature des données qu'ils traitent. En effet, les données traitées dans le domaine de la fouille de données sont normalement de nature structurée (bases de données, entrepôts de données, etc.). Par contre, les processus de fouille de textes sont appliqués sur des données textuelles (de grands corpus de documents) dont l'une des principales caractéristiques est d'être beaucoup moins structurées.

Démarche méthodologique

La démarche méthodologique caractéristique des applications de fouille de textes est également inspirée de celle que l'on retrouve au cœur de nombreux projets de fouille de données. Cette démarche est de nature itérative.

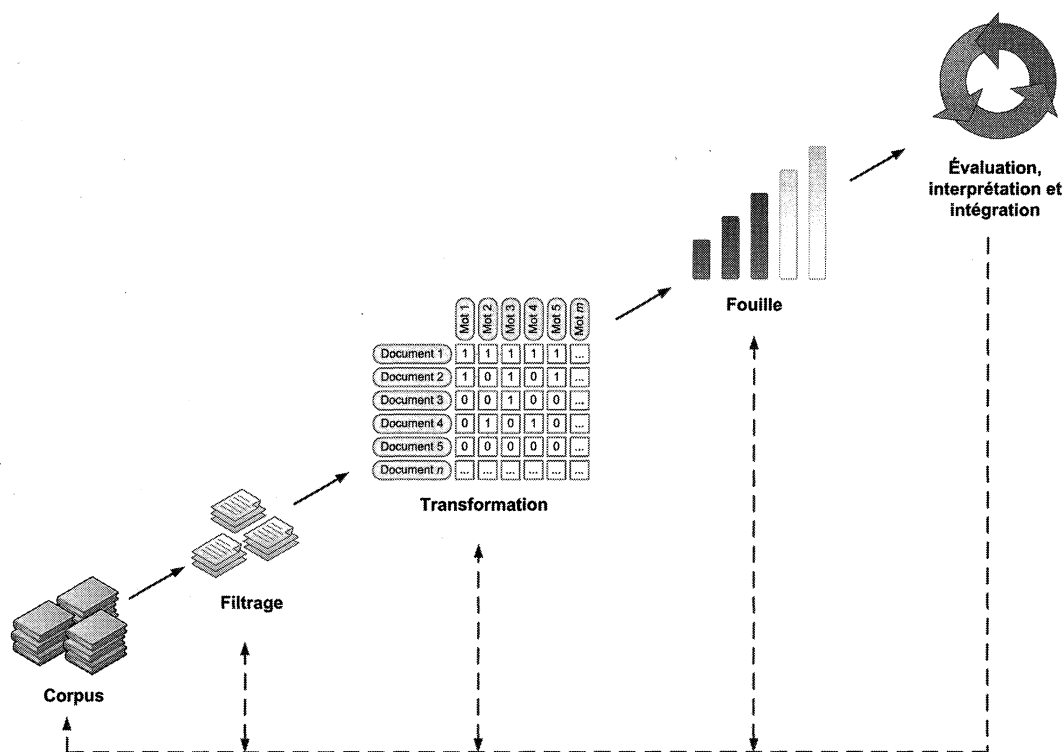
La première étape de tout processus de fouille de données réside dans le développement ou la constitution d'un corpus de documents. Il est essentiel que le corpus de documents soit constitué en tenant compte des objectifs à atteindre par le processus de fouille. À l'étape de constitution du corpus, quatre grandes familles de caractéristiques doivent être évaluées et prises en considération. La constitution d'un corpus à des fins de fouille de textes implique certains choix en ce qui concerne les caractéristiques :

- générales (provenance, taille, date de création, etc.) ;
- technologiques (support, format, etc.) ;
- informationnelles (thématiques et sujets abordés) ;
- linguistiques des documents (langue, genre, registres, etc.).

Cette opération est fondamentale, car la qualité des résultats ultérieurs dépend directement de la justesse des choix effectués lors de cette première étape.

La seconde étape de la démarche générique consiste à filtrer et, le cas échéant, à normaliser le lexique (l'ensemble des mots) du corpus. L'opération de filtrage du lexique est composée traditionnellement de plusieurs

8. Cette section est inspirée de notre contribution à un ouvrage d'introduction aux sciences de l'information, sous la direction de Clément Arsenault et Jean-Michel Salaün (2009, à paraître).

Figure 1Méthodologie générique de la fouille de textes (d'après Fayyad *et al.* 1996)

ments traités, du contexte de réalisation des traitements, des objectifs poursuivis, etc.

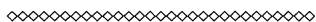
Les algorithmes supervisés peuvent être évalués selon les mesures classiques de rappel et de précision. En contrepartie, l'évaluation des résultats générés par les algorithmes non supervisés ne peut être accomplie en utilisant de telles mesures objectives. Les processus non supervisés permettent principalement d'identifier et de décrire certaines caractéristiques récurrentes observables statistiquement à l'intérieur du corpus de documents. Pour cette raison, il est fréquent de comparer les résultats obtenus par un algorithme non supervisé à ceux obtenus en utilisant plusieurs autres algorithmes comparables. Ceci permet de s'assurer de la stabilité des résultats. Idéalement, les mêmes patrons (*patterns*) devraient être observés, indépendamment de l'algorithme utilisé. Lorsque cela est possible, il peut aussi être utile de comparer les résultats à ceux obtenus par une analyse manuelle.

Par ailleurs, l'interprétation des résultats des algorithmes de fouille ne peut être dictée par aucun cadre théorique qui ferait abstraction du contexte dans lequel l'opération de fouille est réalisée. Finalement, les résultats doivent normalement faire l'objet d'un processus d'intégration à l'intérieur d'une application finale plus complexe dans laquelle le processus de fouille ne constitue qu'une étape bien précise. Les applications finales intégrant des processus de fouille de textes sont

de plus en plus nombreuses et variées. Parmi celles-ci, on trouve entre autres les applications de veille scientifique, de gestion électronique des documents et de recherche d'information. La figure 1, inspirée de Fayyad *et al.* (1996), présente les principales étapes de la méthodologie générique de fouille de textes.

Le développement d'outils d'analyse et de traitement de l'information numérique intégrant efficacement des fonctionnalités de fouille de textes est actuellement un important domaine de recherche dans les milieux académique et industriel. Certains domaines d'activités (dont, entre autres, ceux de la gestion des connaissances, de la bioinformatique et du génie biomédical) ont très rapidement perçu la pertinence d'intégrer des opérations de fouille de textes dans leurs pratiques. Depuis quelques années, plusieurs projets de recherche en sciences humaines et sociales explorent des modalités d'application de stratégies de fouille de textes en tenant compte des particularités d'analyse propres à chaque discipline (Forest, 2006). En sciences de l'information, les processus de fouille de textes sont de plus en plus souvent couplés aux outils de recherche d'information. Plus loin dans l'article, nous exposerons deux exemples d'outils de recherche d'information en ligne de nouvelle génération. Ces outils sont fondés sur l'application d'une opération de classification automatique des documents au sein du processus de recherche. Mais avant d'aborder plus précisément cette dimension applicative, il importe

La distinction entre les opérations de catégorisation et de classification se manifeste aussi dans le domaine de la fouille de textes.



d'expliciter les particularités de l'opération de classification automatique dans son acception dans le domaine de la fouille de textes. L'opération de classification automatique est au cœur des travaux de fouille de textes. Il s'agit d'une opération distincte de la catégorisation automatique, mais qui lui est complémentaire.

Catégorisation et classification automatiques

Les opérations de catégorisation et de classification automatiques font l'objet de travaux de recherche dans les secteurs de l'intelligence artificielle et de l'apprentissage machine. Elles s'inscrivent dans deux catégories de techniques informatiques. La catégorisation automatique figure dans la catégorie des méthodes dites supervisées ; les méthodes supervisées impliquent une intervention directe de l'utilisateur dans le processus réalisé. Cette intervention prend différentes formes : dans le cas de la catégorisation automatique, l'intervention consiste à guider le processus d'apprentissage en indiquant au système quelles catégories doivent être associées à certains documents. L'ensemble des documents catégorisés manuellement porte le nom d'ensemble d'apprentissage. Sur la base de cet ensemble, l'algorithme de catégorisation effectue un apprentissage (la nature de celui-ci est fonction de la nature de l'algorithme employé) et projette l'information apprise sur un ensemble de test, constitué de documents dont le système ne connaît pas *a priori* les catégories auxquelles ils appartiennent.

Quant à la classification automatique, elle figure traditionnellement dans la catégorie des méthodes dites non supervisées. Les différentes méthodes utilisées ici sont réalisées de manière autonome, sans aucune intervention de l'utilisateur. Ces méthodes non supervisées intègrent, tout comme les méthodes supervisées, un mécanisme d'apprentissage. Cependant, celui-ci opère uniquement sur la base de l'information soumise au système, sans aucun recours à des informations ou métadonnées spécifiées par l'utilisateur. L'algorithme peut nécessiter l'utilisation d'un ensemble d'apprentissage et d'un ensemble de test, mais ce n'est pas obligatoire. L'objectif de l'opération consiste à regrouper les données en ensembles homogènes en employant uniquement les traits caractéristiques ayant servi à décrire chaque élément.

La distinction entre les opérations de catégorisation et de classification se manifeste aussi dans le domaine de

la fouille de textes. Ces deux opérations correspondent à deux types d'analyse distincts. L'opération de catégorisation automatique, réalisée en employant certaines métadonnées préalablement connues, est très étroitement reliée à des processus d'analyse de nature prédictive ou normative. Par exemple, une telle opération semble des plus adaptées à l'identification du contenu thématique des documents. L'opération de catégorisation présuppose cependant que nous disposions *a priori* d'informations valides ou attestées pouvant être employées pour effectuer l'apprentissage. En ce sens, sur le plan théorique, l'opération de catégorisation automatique fait nécessairement appel à des informations externes aux données à traiter. En d'autres termes, la réalisation de cette opération implique l'application d'une structure d'informations sur les données textuelles. Dans une tâche de catégorisation thématique, cette structure externe prend la forme d'une taxinomie de catégories thématiques ou d'un plan de classification.

Par contre, l'opération de classification automatique est effectuée en ne considérant que les données provenant exclusivement des documents textuels soumis au système. L'opération de classification vise à regrouper des documents (ou des segments de documents) en ne considérant que leur contenu. Il s'agit d'une opération exploratoire et descriptive.

Dans son application au traitement automatique des documents, le processus de catégorisation implique plusieurs considérations. Premièrement, le processus de catégorisation présuppose l'élaboration *a priori* d'une taxinomie ou d'une hiérarchie de catégories thématiques adaptées au contenu et aux spécificités de la collection à traiter. Deuxièmement, ce processus implique non seulement l'identification, à partir de l'ensemble des documents, des caractéristiques (linguistiques et statistiques) qui serviront de base à la catégorisation, mais aussi le choix d'une méthode d'attribution des catégories aux documents.

Pour Charniak (1993), le défi de la classification automatique s'organise autour de trois axes. Le premier axe concerne la description des différents éléments à soumettre à l'opération de classification. En effet, il importe à cet égard de ne retenir que les traits caractéristiques discriminants, sur la base desquels les différents segments de documents seront comparés. Très souvent, cette description implique l'application d'une ou de plusieurs fonctions de filtrage et de nettoyage. Pour que le résultat soit pertinent, il importe que le processus de classification repose sur des descripteurs significatifs. Les décisions prises à cet égard doivent tenir compte aussi bien des objectifs de la classification que de la nature même des objets à classer. Le deuxième axe concerne la fonction discriminante, à la base de toute opération de classification. Cette fonction discriminante peut faire appel à plusieurs critères : l'identité, la similarité, l'homogénéité, l'équivalence, etc. Finalement, le dernier axe concerne le choix de l'algorithme employé pour réaliser

Figure 2
Représentation d'un regroupement plat

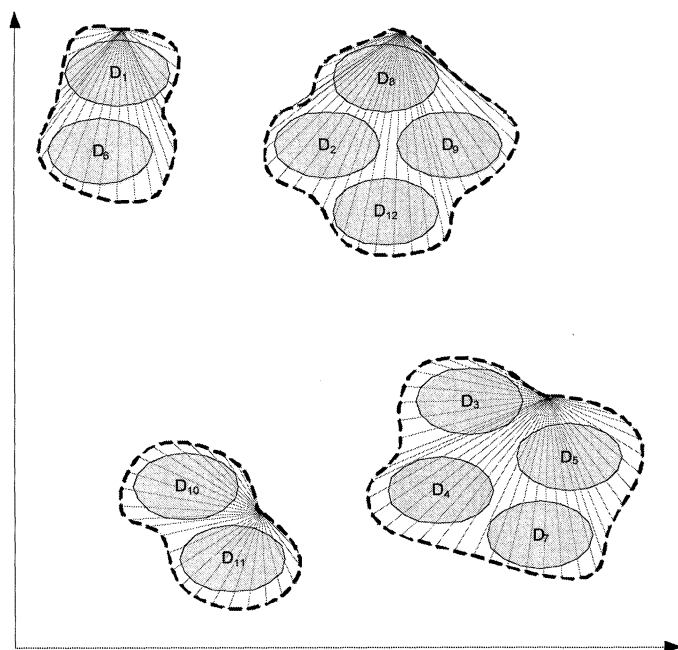
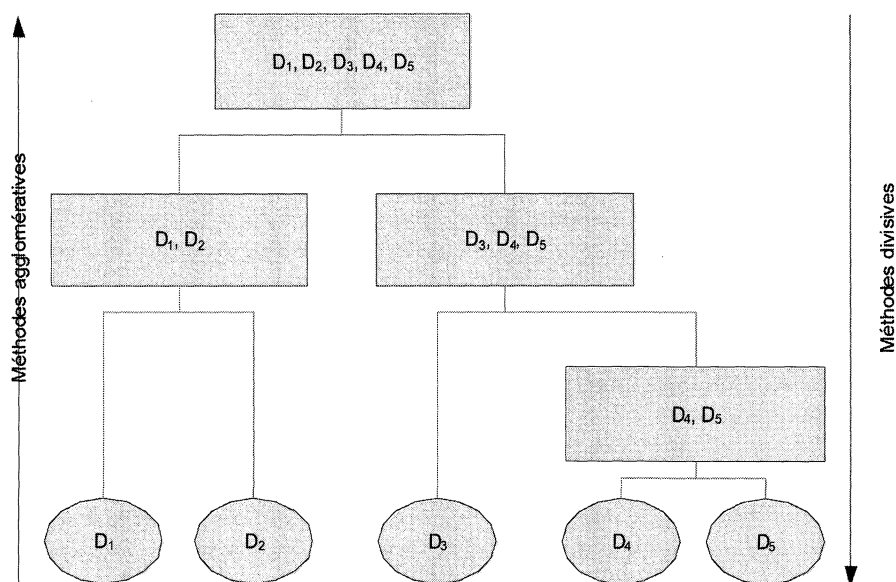
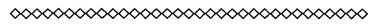


Figure 3
Représentation d'un regroupement hiérarchique



Actuellement, les processus de classification et de catégorisation automatiques ont été les plus exploités dans le contexte de la recherche d'information.



l'opération de classification. Ce choix implique des enjeux à la fois théoriques et pratiques qui relèvent tant de la nature des données à regrouper que des caractéristiques de la classification souhaitée.

Plusieurs techniques informatiques ont été explorées afin d'accomplir des tâches de classification des documents textuels (Jain, Murty and Flynn, 1999). On distingue les différentes méthodes de classification en fonction de la rigidité des regroupements effectués (*hard clustering* vs *soft clustering*) (Manning et Schütze, 1999). Les méthodes de *hard clustering* permettent de positionner un document dans un seul groupe. Par contre, certaines techniques de *soft clustering* permettent de situer le document dans plus d'un ensemble en précisant, pour le document, son degré d'appartenance à chacun de ceux-ci.

Par ailleurs, les techniques de *hard clustering* peuvent aussi faire l'objet de distinctions supplémentaires. Il est possible de distinguer, d'une part, les techniques effectuant des regroupements ou des partitionnements plats (*flat*) et, d'autre part, les techniques effectuant des regroupements hiérarchiques.

Les regroupements plats sont caractérisés par le fait qu'aucune relation précise n'est déterminée entre les différents regroupements générés (figure 2). Les méthodes de cette catégorie sont très souvent de nature itérative. En effet, elles procèdent d'abord en déterminant un nombre fini de regroupements, ensuite en raffinant chacun des regroupements effectués.

Quant aux techniques de classification permettant d'effectuer des regroupements hiérarchiques, elles se caractérisent par la production de nœuds qui représentent chacun une sous-classe d'un nœud de niveau supérieur (figure 3). Les méthodes les plus fréquemment employées pour la classification hiérarchique peuvent être regroupées en deux sous-catégories. La première englobe les méthodes dites « agglomératives » (*bottom-up*) ; celles-ci procèdent d'abord en identifiant tous les éléments à classer, ensuite en effectuant successivement plusieurs regroupements. La seconde sous-catégorie englobe les méthodes dites « divisives » (*top-down*) ; celles-ci procèdent d'abord en identifiant un seul regroupement, ensuite en divisant ou en fractionnant le groupe initial.

Application à la recherche d'information

Dans le domaine de la fouille de textes, les récents développements techniques (concernant, entre autres, le développement d'algorithmes d'extraction et d'organisation automatiques d'information) et méthodologiques (desquels relève principalement la modélisation des démarches à suivre afin d'appliquer efficacement les techniques) ont récemment été explorés au sein d'applications de recherche d'information en ligne. D'ailleurs, depuis quelques années, plusieurs prototypes d'applications industrielles d'analyse et de recherche d'information ont été proposés afin d'intégrer certaines techniques de fouille de textes davantage sensibles à la dimension sémantique des documents.

Actuellement, les processus de classification et de catégorisation automatiques ont été les plus exploités dans le contexte de la recherche d'information. Ces deux processus, bien qu'ils puissent être exploités afin d'assister différentes tâches liées à l'analyse de l'information numérique, sont principalement utilisés dans un contexte de recherche d'information afin de lever l'ambiguïté sémantique des termes de requête. Cette perspective d'application repose en partie sur l'hypothèse suivante : les mots avec lesquels les termes de la requête cooccurrent sont de bons indices pour identifier la signification précise des termes de la requête en question. Cette hypothèse est d'ailleurs à la base même de l'utilisation du modèle vectoriel pour la recherche d'information (Salton et McGill, 1983). Lors de l'utilisation de stratégies d'extraction et d'organisation automatiques d'information, l'objectif est d'exploiter au maximum cette idée en procédant, en plus, au regroupement automatique des documents. Comme nous l'avons souligné, le regroupement automatique des documents peut être accompli en utilisant soit un algorithme de classification, soit un algorithme de catégorisation. Dans les deux cas, au terme du processus, les documents sont regroupés sur la base de thématiques communes, permettant ainsi à l'utilisateur de porter son attention sur les documents inscrits dans des regroupements thématiques qui combleront plus adéquatement son besoin d'information.

Les principales applications de recherche d'information en ligne fondées sur l'utilisation de la classification ou de la catégorisation automatiques des documents sont *Grokker*¹⁰ et *Clusty*¹¹. Ces deux applications possèdent plusieurs caractéristiques communes, mais se distinguent l'une de l'autre à plusieurs égards.

La différence importante entre ces deux outils est qu'ils ne sont pas actuellement applicables sur des corpus de même type. *Clusty* interroge la partie visible du Web en entier, alors que *Grokker* n'offre actuellement que la

10. <http://www.grokker.com>

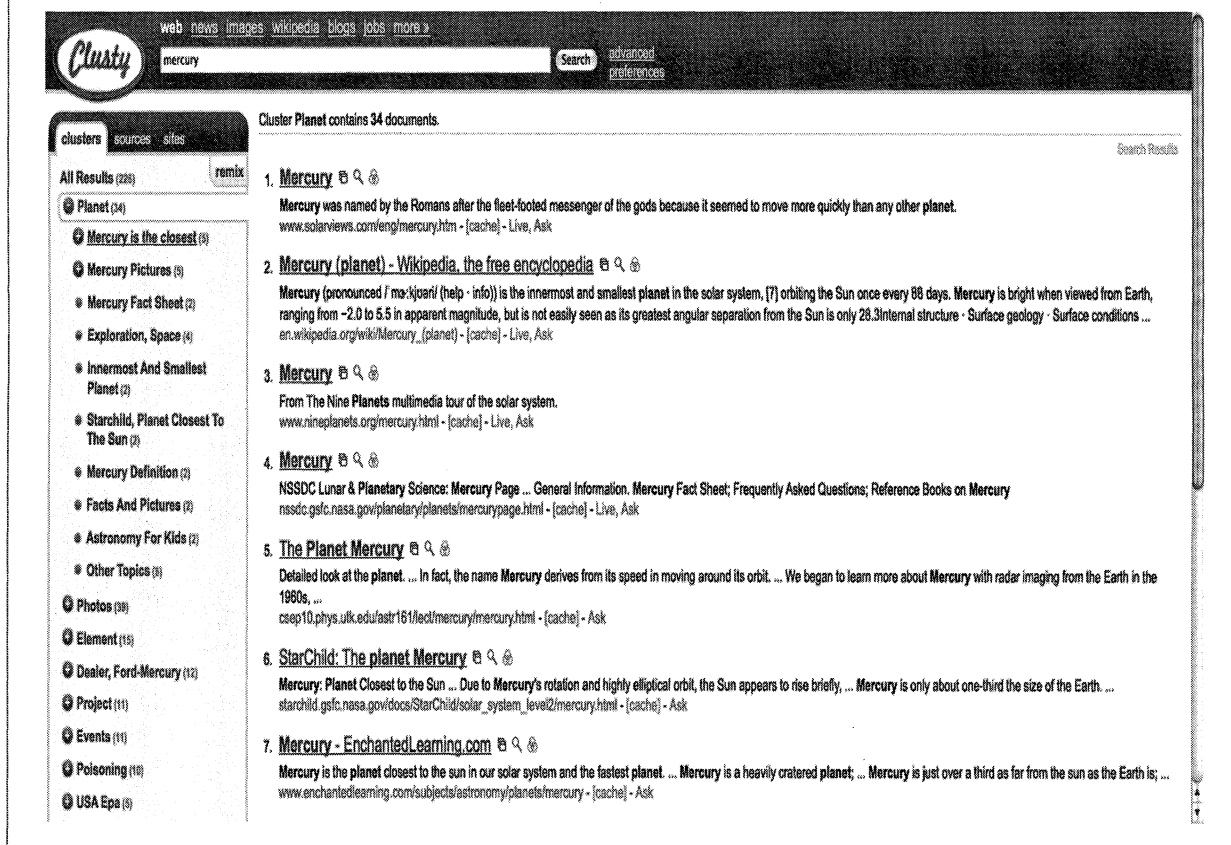
11. <http://www.clusty.com>

Les nouveaux outils d'analyse et de recherche d'information en ligne qui reposent sur des opérations de classification offrent un ensemble de fonctionnalités et de paramètres de base avec lesquels l'utilisateur peut interagir. Les figures 4 et 5 représentent les interfaces de recherche des deux moteurs mentionnés précédemment. Ces deux moteurs offrent des interfaces de recherche possédant de nombreuses caractéristiques communes. Ainsi, on constate que les deux moteurs fondés sur la classification automatique des documents numériques présentent, à l'instar des moteurs de recherche classiques, les résultats des requêtes sous forme de liste (à la droite de l'interface). Par défaut, les résultats sont présentés selon un tri de pertinence. Ce qui distingue les moteurs de recherche traditionnels des outils de recherche de

nouvelle génération repose en grande partie dans la liste des regroupements de documents présentée dans la partie gauche de l'interface. On y retrouve l'inventaire des regroupements hiérarchiques générés par l'algorithme de classification. Ainsi, tous les documents récupérés par le moteur de recherche sont d'abord regroupés selon des mots discriminants qu'ils partagent avant d'être affichés dans la composante de droite de l'interface. Comme nous l'avons expliqué précédemment, l'opération de classification automatique consiste à regrouper des documents partageant certains critères de similarité, lesquels permettent d'identifier des thématiques ou des univers lexicaux caractéristiques des documents. Ces informations extraites par l'algorithme de classification s'avèrent utiles pour assister l'utilisateur dans l'identification des documents lui permettant de mieux répondre à son besoin d'information. Cependant, *stricto sensu*, l'opération de classification n'implique pas une composante permettant de caractériser le contenu des classes ou des regroupements ainsi générés. C'est à cet égard qu'il est préférable de réaliser une opération de catégorisation lorsque les documents nous permettent de le faire, c'est-à-dire lorsque les documents de la base interrogée sont accompagnés de métadonnées permettant d'effectuer adéquatement l'opération de catégorisation. Une solution de rechange s'impose lorsqu'aucune métadonnée n'est disponible, ou lorsque les métadonnées sont disponibles en quantité insuffisante ou encore lorsque leur qualité ne peut être attestée – ce qui est précisément le cas des documents sur le Web, du moins dans son état actuel. Il est alors possible d'identifier la thématique du contenu des regroupements issus de l'étape de classification en procédant à l'extraction automatique des termes les plus discriminants de chaque regroupement. Dans nos travaux antérieurs (Forest, 2006), nous avons démontré la pertinence de cette démarche en utilisant la mesure de pondération TF-IDF¹² comme facteur de discrimination des termes de chaque regroupement. Il est cependant possible d'employer et de combiner plusieurs mesures statistiques afin d'extraire

DOCUMENTATION ET BIBLIOTHÈQUES | AVRIL • JUIN 2009 | 85

Figure 4
Interface de recherche du moteur *Clusty*



les principales catégories thématiques présentes dans les documents préalablement regroupés.

L'exemple suivant illustre bien la pertinence de recourir à des procédés de classification et de catégorisation automatiques des documents pour assister la recherche d'information. En utilisant le terme « *mercury* » comme requête, les deux outils de recherche offrent à l'utilisateur la possibilité de parcourir les documents récupérés, tout en lui indiquant qu'il s'agit d'un terme très polysémique. En effet, le terme anglais « *mercury* » peut être utilisé pour désigner tantôt un dieu romain, tantôt une planète, tantôt un élément chimique, ou encore une marque de voiture, et même d'autres types de moyen de transport, plusieurs lieux géographiques, un chanteur, etc¹³. Compte tenu de la polysémie du terme « *mercury* », il n'est donc pas surprenant qu'une majorité de documents récupérés par les moteurs de recherche traditionnels sur le Web présentent des informations non pertinentes pour les utilisateurs. Ce problème pourrait être corrigé, en partie, en spécifiant la requête. Mais cette solution, très imparfaite dans les faits, implique que l'utilisateur sache *a priori* quelle information il recherche et quels sont les meilleurs termes pour y accéder. Or, en pratique,

les utilisateurs parcourent (*browse*) le Web sans y chercher des documents précis. De plus, ils sont souvent peu compétents dans la formulation des requêtes. Pour ces raisons, la classification et la catégorisation automatiques des documents s'avèrent des plus fécondes pour assister la découverte d'information, car elles permettent de mettre en lumière certaines dimensions thématiques potentiellement inconnues de l'utilisateur, sans toutefois lui imposer de faire des choix inutiles.

Dans notre exemple, il n'est pas étonnant de constater le nombre élevé de regroupements proposés tant par *Clusty* que par *Grokker*. Ainsi, pour le moteur de recherche *Grokker*, on note les principaux regroupements suivants (la valeur entre parenthèses correspond au nombre de documents dans chaque regroupement) : *Planet Mercury* (39), *General* (28), *Mercury Information* (17), *Solar System* (15), *Project Mercury* (12), *Mercury in Fish* (9), *Mercury Element* (8), *Mercury Levels* (8), *Sun* (8), etc.

Voici les résultats du moteur de recherche *Clusty* (la valeur entre parenthèses correspond ici aussi au nombre de documents dans chaque regroupement) : *Planet* (34), *Photos* (39), *Element* (15), *Dealer Ford-Mercury* (12), *Project* (11), *Events* (11), *Poisoning* (10), etc.

Dans les deux applications, il importe de mentionner que les regroupements sont effectués par un algorithme

13. On retrouve dans la version anglaise de Wikipédia plus d'une trentaine d'entrées associées au terme « *mercury* ».

Figure 5
Interface de recherche du moteur *Grokker* (sans visualisation graphique)

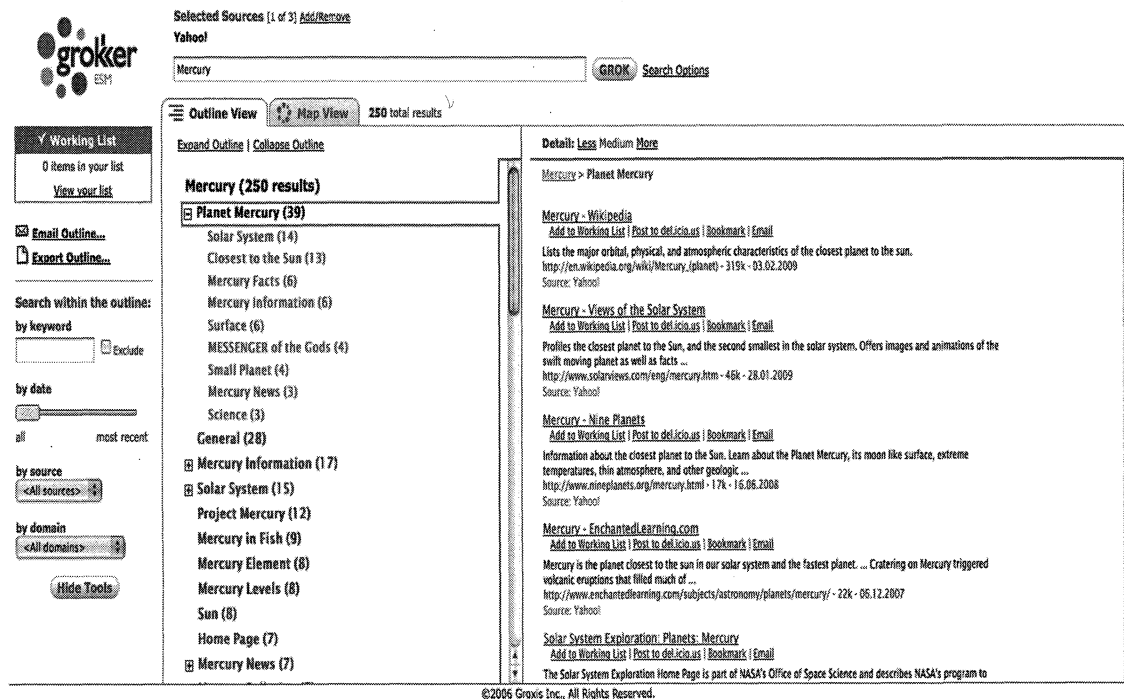
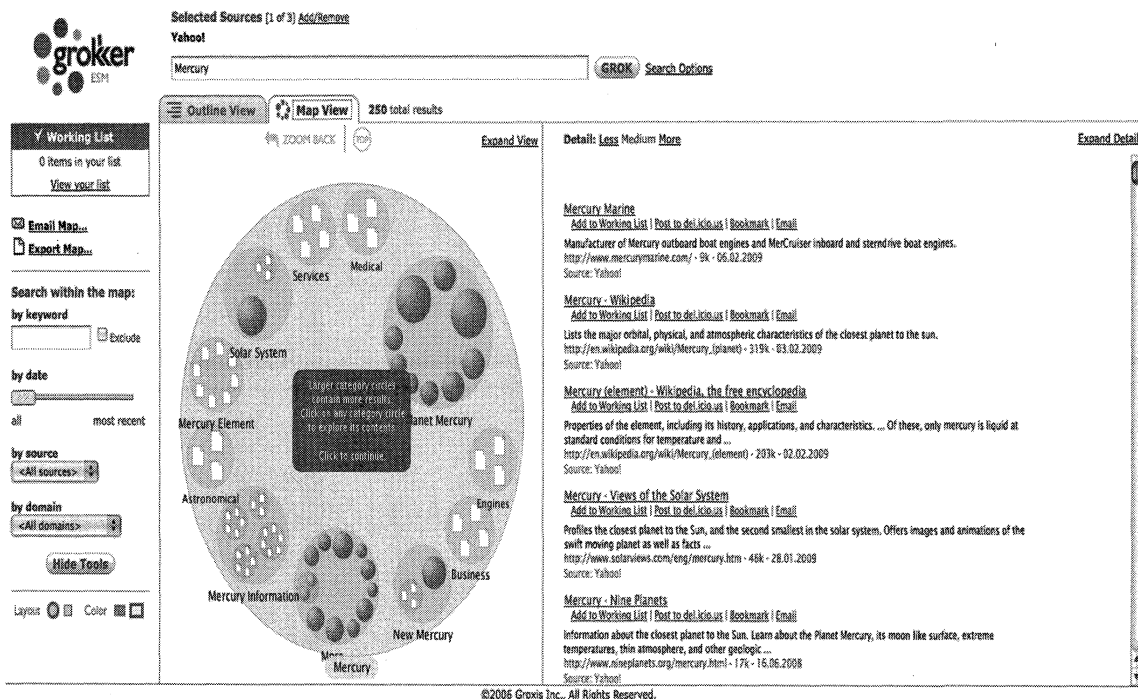


Figure 6
Interface de recherche du moteur *Grokker* (avec représentation graphique de la même structure présente dans la figure 5)



Les prochaines années seront caractérisées par une augmentation des contenus générés dynamiquement (pensons, à titre d'exemple, aux informations disponibles en continu dans les fils RSS). Ce phénomène sera aussi accompagné d'une augmentation des contenus générés grâce à une participation plus active des utilisateurs à ce qu'il convient désormais de nommer le Web 2.0. Il s'agit là de deux facteurs parmi d'autres qui nous portent à croire que la quantité d'information numérique ne fera qu'augmenter et qu'il est désormais essentiel d'explorer de nouvelles avenues pour en faciliter la recherche, l'analyse et la gestion. ☐

- DOCUMENTATION ET BIBLIOTHÈQUES | AVRIL • JUIN 2009 | 89