

網路自動分群搜尋引擎之 使用者評估研究

User Evaluation of Web Clustering Search Engines

卜小蝶

Hsiao-Tieh Pu

臺灣師範大學圖書資訊學研究所 教授

Professor

Graduate Institute of Library and Information Studies

National Taiwan Normal University

陳思穎

Sih-Ying Chen

臺灣師範大學圖書資訊學研究所 研究生

Graduate Student

Graduate Institute of Library and Information Studies

National Taiwan Normal University

【摘要 Abstract】

本文由使用者角度，分別採用實驗、觀察、問卷、訪談及檢索過程記錄分析等方法，就檢索效率、檢索效益及滿意度等範疇，評估比較分群與排序條列式搜尋引擎之檢索表現。整體而言，排序條列式搜尋引擎能在較快的時間內，找到品質相當且數量較多的相關網頁，使用者對檢索結果的滿意度也較高。而分群搜尋引擎則具有突顯重要概念、提供多元思考及降低資訊超載等特性。此外，使用者對於個別群集的接受度較整體群集架構來得高。分群搜尋引擎提供了嶄新的檢索經驗，但仍有不少問題需要進一步分析探討，本文最後則提供一些觀察與建議。

投稿日期：2007.10.01；接受日期：2007.11.21

email: 卜小蝶htpu@ntnu.edu.tw；陳思穎cindy-in@yahoo.com.tw

This study provides user evaluation of clustering and ranked list search engines by comparing their search efficiency, search effectiveness, and users' satisfaction level. The methods include experiment, observation, questionnaire, interview, and search log analysis. The results show that the ranked list search engine performs better in users' search speed, relevant pages retrieved, and the satisfaction level. On the other hand, the clustering search engine has the merits of demonstrating important concepts, providing clues for diverse thinking, and helping to reduce the information overload. Meanwhile, users' acceptance of individual clusters is higher than that of the cluster scheme. The paper also provides some discussions and suggestions on improving the clustering search engines.

[關鍵字 Keywords]

網路分群搜尋引擎；網路使用者研究；搜尋引擎評估

Web Clustering Search Engine; Web User Studies; Search Engine Evaluation

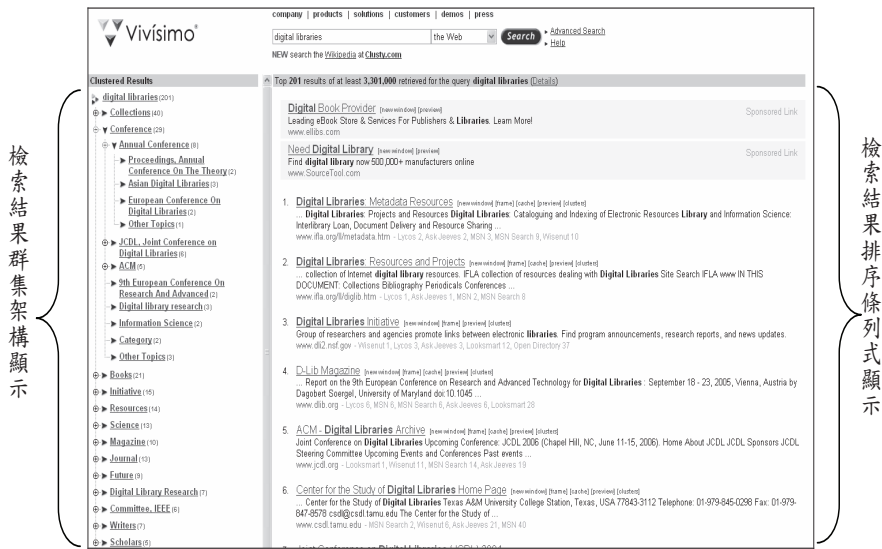
壹、前言

隨著網路的普及與資源的成長，使用搜尋引擎查詢資訊，已成為上網最重要的活動之一。但根據Delphi Group（2004）針對知識工作者的調查，59%使用者認為在網路搜尋環境中，取用資訊雖比以往更為快捷便利，但卻有68%的使用者認為查詢資訊仍相當困難，而62%的使用者對於檢索效率並不滿意。其主要原因在於，搜尋引擎多僅提供關鍵字檢索功能，面對龐雜的網路資源，使用者輸入任一關鍵字，常常就有動輒數十萬筆的查詢結果，無法快速有效過濾。同時，在利用關鍵字檢索時，使用者若不清楚應輸入哪些合適的檢索詞彙，檢索結果也不一定令人滿意。此外，在檢索過程中，有許多使用者並非一直在進行關鍵字檢索，而是花費更多時間在瀏覽檢索結果。更令人困擾的是同義詞（synonym）、同形異義（homonym）等語意混淆問題，例如查詢“chips”，會查到“chocolate chips（巧克力片）”，也會查到“micro chips（電腦微晶片）”等。由上述說明可知，僅依賴關鍵字檢索功能，是無法滿足使用者多元檢索需求，及解決資訊超載問題。

為了讓使用者能快速有效查詢或瀏覽資訊，將資訊做適當分類，是一可行的改善方向（Samler & Lewellen, 2004）。基本上，分類對檢索

的價值可由兩方面探討，一是分類提供了某種脈絡（context）資訊，可降低語意的模糊度（ambiguity），讓使用者避免查詢或瀏覽大量無關的資訊。舉例來說，若能將詞彙做適當分類，如前述chocolate chips屬食品類、micro chips屬電子類，如此一來，chips就不致於產生混淆。分類的另一種價值是觸類旁通，亦即使用者可藉由各種型的分類架構，刺激其聯想相關概念，並進一步發掘（discover）一些原本未預期到的資訊。例如由micro-chips所屬的電子類別，觀察到其中尚有processor、ASIC、memory chips等相關主題類別。簡言之，分類應有助於檢索效益的提升。

就目前的網路搜尋環境中，利用分類提昇檢索效益的作法主要有兩種，一是利用人工或半自動方式建立分類目錄，如Yahoo!的分類目錄；另一則是藉由大量資料關聯計算，將檢索結果進行分群，即自動分群（automatic clustering）技術，如Vivisimo的檢索結果分群應用。一般而言，人工作法多利用事先定義（pre-defined）的分類架構如分類表（classification scheme），將網頁進行分類；而自動化作法則依賴統計或演算法，自動產生群集（cluster）後，進一步形成一由下而上的類別架構（category scheme）。面對龐雜的網路資源，人工方式成本高、也緩不濟急，因此自動分群技術就變得相當重要。以目前最著名的網路自動分群搜尋引擎Vivisimo為例，當輸入關鍵字“digital libraries”，其檢索結果如圖一所示，右方為一般排序條列式（ranked list）的呈現方式，而左方則為一階層式的群集架構（cluster scheme）。此乃根據檢索結果進行相似性比對後，產生群集，並由各群集計算出適當的類別，再予以呈現。雖然這些主題類別（category）或群集類別（cluster）並無明確的邏輯結構（如Conference是資料類型，Digital Collections是意涵，Alexandria、Project是計畫名稱，三種類別並無概念上的邏輯關係），類別名稱也不盡然一致易懂，但這種呈現方式其實也有不少優點，例如可讓使用者快速瞭解所輸入關鍵字“digital libraries”的一些重要相關概念及關聯，同時也有一些預期之外的主題類別，提供聯想。換言之，有了類別架構，使用者較有機會瞭解、釐清檢索詞彙及檢索結果所呈現的概念，同時也能產生一些新的概念，進一步提升檢索效益。



圖一 網路自動分群搜尋引擎Vivisimo之檢索結果舉例

資料來源：http://vivisimo.com/

上述自動產生類別架構是近年興起的研究主題（Krishnapuram & Kummamuru, 2003），預期也將是改善網路搜尋中資訊超載問題的重要利器。然而，對使用者而言，究竟這些自動產生的類別架構及群集所呈現的意義為何？是否能如預期有助於檢索效益的提升？其與一般排序條列式、不分群的呈現方式比較，又有何優劣異同？都是值得探討的問題。因此本文希望能由使用者角度，評估自動分群搜尋引擎對檢索效益的幫助與限制。評估範疇主要包括檢索效率（efficiency）、效益（effectiveness）及滿意度（satisfaction）等面向，分別採用實驗、觀察、問卷、訪談、檢索過程記錄分析等方法，分析探討使用者對網路自動分群搜尋引擎的實際使用情形及看法，以供後續相關系統發展之參考。

貳、相關研究

一、分類架構之意涵

分類架構的概念相當複雜多元，相關概念包括由知識觀點探討知識本體（ontology）與分類學（taxonomy），也包括由方法或技術觀點，探

討分類（classification）與分群（clustering）的異同。分類架構在本文是指根據某些物件屬性所訂定之分類系統，如圖書資訊學所探討的分類表（classification scheme），網站資訊架構設計中有關各元素組織安排的分類架構（website taxonomy），或企業導入知識管理時所建構的知識分類表（knowledge taxonomy）等。而本文所分析的分類架構則是以網路搜尋引擎中所使用的群集架構（cluster scheme）為主。雖然上述架構均以組織資訊為目的，偶爾也會交互使用，但其概念不盡相同，進一步說明如下。

分類學（taxonomy）的概念源起於生命科學，可追溯至18世紀的瑞典學者Carl Linnaeus，其建置了組織與描述動植物關係的階層架構（Hackos, 2005）。Taxonomy通常應用於特定機構或主題領域，不同於圖書館的分類表，是以整理一般性資源為目的。分類學早期被認為是劃分生物種類的理論與實務（Mayr, 1982），近年來則廣為應用於資訊檢索系統，主要是藉由結合分類表與索引典技術（Gilchrist, 2003），以提供主題階層清單。其目的在提升資訊檢索效率，因其有助使用者依資訊資源在分類架構中的「情境」（context）位置，選擇相關主題，進一步縮小查詢範圍（Cisco & Jackson, 2005）。進一步分析，一個有用的taxonomy其基本功能有二，一是讓使用者透過瀏覽類目的階層關係，建立資訊的關聯、甚至形成新的看法（Delphi Group, 2002），即扮演「外部認知鷹架」（Clark, 1997; Jacob, 2001）角色；二是幫助使用者發掘查詢時未聯想到的相關概念（Samler & Lewellen, 2004），即作為「偶遇指引」（serendipitous guidance）（Bruno & Richmond, 2003）的角色。

而分類（classification）則是圖書資訊學的核心研究主題之一。Kwasnik（1999）認為「分類是表達與瞭解知識的一種程序（process）與方法（approach）」，我們無時無地不在從事分類，吸收知識、表達知識、建立理論都需要分類。而一個好的分類架構（classification scheme）即是能以有用的結構（structure），將概念作有意義的呈現與串連。Kwasnik並提出四種分類結構：階層（hierarchies）、樹狀（tree）、典範（paradigm）及分面（facet）。其中依知識的階層關係呈現是最常見的結構，例如圖書分類表；樹狀結構其實與階層相似，所不同在於後者的類別並無嚴謹的繼承關係；而典範又稱矩陣維度（matrix），是將兩種以上屬性並列呈現，資料庫的表格（table）或詮釋資料（metadata）即是類似結構；而分面其實與矩陣結構類似，只是分面的屬性較具通用性，且不侷限於二維結構。

二、分類架構於網路資源組織與檢索之應用

將分類架構應用於網路資源組織整理已行之有年（Vizine-Goetz, 1999）。其分類架構的訂定主要有採取標準的分類表（如Dewey Decimal Classification, DDC），與自訂的分類表（如Yahoo!）。二者結構多為階層式或樹狀式。Chan（2001）曾指出，雖然目前已有許多網路資源的分類架構，但多數仍缺乏如DDC與LCC（Library of Congress Classification, LCC）般嚴謹的階層與概念架構。Vizine-Goetz（1999）也曾評估以DDC與LCC整理網路資源的可行性，他認為DDC與LCC的階層架構，在廣度與深度上足以支援瀏覽，而其類號也能協助資訊檢索。Vizine-Goetz（2002）更進一步比較DDC、Yahoo!與LookSmart在分類架構的最上層標目、資源分佈情況、瀏覽以及主題樹是否可應用於多語環境等，他發現DDC的類目結構與網路資源主題目錄，皆可提供階層與字順瀏覽方式，同樣也都可以應用於多語環境，而且有近七成的資源均可歸入第五層以上的類目，因此其認為以DDC作為瀏覽架構及整理大量資源是可行的作法。

上述研究多肯定圖書資訊分類架構的價值，但面對網路資源的龐雜與多元，仍存在不少問題。例如Assadi 和 Beauvisage（2002）就認為由於網路並非百科全書，亦非圖書館，網路涵蓋了不同主題與品質的服務與資源，且網路的使用情境、使用者興趣及需求皆具多樣性，因此並不適合套用某學科領域的分類架構來整理網路資源。而Schwartz（2001）認為使用者並不瞭解分類架構的設計與發展原則，因此隨著資源數量的增加，階層架構可能無法涵蓋所有資源，甚至讓使用者感到困惑。即使套用圖書分類架構來整理網路資源，困擾也不少。Schwartz提到，因圖書分類架構多以學科為基礎，未必能因應網路這樣跨學科或多學科的環境。此外，在非網路環境中，分類架構通常是與資源本身結合（如圖書的分類標籤與圖書是共同存在於特定實體），但在網路環境中，分類架構與資源則是分開儲存，其結構與設計似乎更需考量檢索上的效用，而不需侷限於學科知識體系或實體空間安排。此外，Schwartz也提出網路分類架構應具備以下特性：修訂容易、具有彈性、易於使用、具包容性及權威性等。

對於如何設計一適用於網路環境的分類架構，Harvey（1999）曾嘗試整理一些注意事項如下：

- (一)將分類架構的文獻保證原則（literary warrant）擴展至網路資源，必要時增加新類目。

- (二)修訂類別名稱，增強其表達性與通用性。
- (三)分解和標記類號的組成因素，以識別其所表達的特定主題。
- (四)持續增添新的術語作為索引詞彙。
- (五)擴展分類法與其他控制詞彙的關聯。
- (六)控制類別的深度，多數網路分類檢索系統只使用到前三層。

就上述討論，圖書資訊學的分類研究多集中於分類的知識內容及資訊組織作法的探討 (Mai, 2004; Bates, 2002; Ellis & Vasconcelos, 1999)。若進一步針對網路資源檢索角度，一般而言，查詢網路資源可區分為關鍵字檢索 (keyword search) 與主題瀏覽 (subject browsing) 兩種模式。其中主題瀏覽模式，除了利用主題詞表如標題表 (subject headings) 或索引典 (thesauri) 來標記資源內容主題，以提供檢索外，最主要作法仍是藉由分類架構對資源加以組織，提供使用者依不同類別及結構來瀏覽資源。分類架構不同於主題詞表，在於分類架構可以將相關資源整理於階層架構中，而主題詞表則通常只有單一概念排序，並未提供概念間的整體關聯架構。在網路搜尋環境中，分類架構除了被許多入口網站作為網路資源的主題瀏覽工具外，同時也被應用於跨語檢索的輔助工具。此外，對於同一資源若予以多重分類，也可提供使用者以不同角度瀏覽 (Cross, et al., 2000)。Koch 和 Day (1997) 提供了一些分類對於網路資源檢索的價值，說明如下：

- (一)便於瀏覽資訊：當使用者不清楚網頁所包含的內容時，分類架構可協助其快速瞭解。同時使用者也可藉由階層式架構逐一選取各項檢索結果。
- (二)擴大及縮小檢索範圍：藉由相關類別的擴展與限制，可有效提升求全率與求準率。此外，分類架構也有助於過濾大量的檢索結果。
- (三)降低詞彙語意問題：以類別查詢可避免因詞彙造成的同義詞、同形異義詞等問題。
- (四)提供脈絡資訊：檢索結果若具類別資訊，則將有助於使用者更清楚檢索結果的相關性。
- (五)跨語檢索的輔助：藉由類別檢索或瀏覽，使用者即使未具備語言背景，也能有機會進行檢索。
- (六)跨資料庫檢索的輔助：若網路資源檢索系統採用同一種分類架構，即使資料庫的類型、內容不同，使用者仍可藉由相關類別，進行跨資料庫瀏覽及檢索。

綜言之，圖書資訊學的分類研究已具有相當歷史，也累積不少成果。然而面對網路資源大量、異質、變動等特性，分類架構的設計與使用確實需要進一步調整與改善。上述研究提出一些編製上的指導原則，但對於使用者的需求與使用探討仍不多見，也是亟需努力的方向之一。

三、資訊架構中之分類架構設計

所謂的資訊架構（Information Architecture）是指對資訊系統的結構進行描述的一套方法（Barker, 2005），包括資訊如何被組織、資訊瀏覽的方式、詞彙或術語的選擇等。一個具有良好資訊架構的系統，使用者可輕易的使用系統、快速獲取所需資訊。資訊架構的應用相當廣泛，最常見於網站及企業內部知識管理系統的介面設計。其中，如何讓使用者能快速瀏覽內容，則有賴設計符合使用者需求與認知的分類架構（taxonomy design）。

根據Rosenfeld 和 Morville（2002）建議，分類架構可區分為精確性架構（exact organization scheme）與模糊性架構（ambiguous organization scheme），前者僅針對資訊定義明確的項目，分成集群，再依字母、地區、年代來排序，類似圖書資訊的編目作法；後者則是依資訊的主題內容，將相關資料彙集起來，類似圖書資訊的分類作法。模糊性架構的建構雖然成本較高，但對使用者的幫助較大，其建置方法主要有5種，分別為依主題（topic）、任務（task）、觀眾（audience）、隱喻（metaphor）及混用（hybrid）等。其中依主題來設計的架構最有用，但也最具挑戰性。Rosenfeld 和 Morville進一步對分類架構的結構設計提出一些建議：

- (一)階層式（由上往下）架構：對於成長快速的網站，宜採寬而淺的階層系統；類別之間必須彼此互斥（mutually exclusive）；單一的分類架構必須在排它性與包容性求取平衡；容許多重分類及具多階層（polyhierarchical）架構；注意寬度與深度的平衡；注意人類認知能力的限制。
 - (二)資料庫式（由下往上）架構：較適合大型分散式環境，如利用 Metadata將關聯式資料庫的結構與非結構化網頁做適當結合。
 - (三)超文字式架構：較有創意，彈性大，可彌補階層及資料庫式之不足，但需注意使用者認知負擔。
- 而在設計分類架構時，也有以下挑戰：
- (一)模糊性（ambiguity）：分類架構的基礎是語言，而語言具有模糊性，也因此造成分類架構不易保持穩定。

(二)異質性 (heterogeneity)：不同資訊系統或網站所包含的內容與期望達成的目的不同，因此不易形成一共通的分類架構。

(三)觀點差異 (differences in perspectives)：分類架構深受建構者觀點影響，同樣的主題內容，不同觀點也會產生不同的分類架構。

(四)內部政治因素 (internal politics)：分類架構多在某一機構組織中形成，為配合機構目標，分類架構常常會受機構組織的權力結構所影響。此外，分類架構也能反映顧客觀點，對整個機構組織影響深遠。

Vogel (2003a) 建議可依系統的設計目的來發展多個分類架構，而不應只發展一個單一多目的的分類架構。此外，他也提出一些編製良好分類架構的原則如下：

(一)深度 (depth)：以五層為宜，最多不超過七層。

(二)廣度 (width)：同義詞不應分割為層級，而應儘量合併。

(三)平衡度 (balance)：類別所含子類應儘量一致，避免有大小類差異。

(四)單一路徑 (single path progression)：避免尋找或瀏覽的路徑重複。

除此之外，Vogel (2003b) 也提到分類架構支援知識管理系統應注意的兩個層面，第一是管理層面，要關心分類架構的可維護性 (maintainability)、可靠性 (reliability)、通用性 (universality)、可理解性 (understandability) 等；另一則是以使用者為中心的設計原則，應著重機動性 (dynamic)、有力性 (powerful)、可擴充性 (scalability) 及彈性 (flexibility) 等。

Sravanapudi (2004) 認為所謂良好的分類架構可由兩個層面探討，即結構 (structure) 及完整性 (completeness)。前者是指分類架構應具邏輯性，以階層式結構呈現，讓企業組織能明瞭各類別的意義；而後者則指分類架構應廣納所有可以描述該領域的術語。Hackos (2005) 則提到符合使用者需求的分類架構應具備特色，包括：使用者認為分類架構相較全文檢索，是較有效率的檢索工具；最有效的分類架構應呈現「方形」結構，即每個層級的類別總數等於每個類別其子類別的總數；類別標籤若具鑑別度，使用者其實可以處理相當多的類目；分類架構在使用者瞭解該領域的術語時效果最好；每個類別都應包含完整的關鍵字；類別所使用的術語與命名方式應儘量一致等。

Chen 和 Dumais (2001) 是少數針對使用者介面進行評鑑的研究之一。其請使用者比較七種檢索結果的呈現介面，研究結果發現有分類的介面都比條列式介面的效果好。該研究也建議一些分類介面的改善建

議，例如分類介面最大的用處在幫助使用者對檢索結果進行相關判斷，不僅需要提供類別名稱，最好也能提供類名的說明；類目層級不要過多，以免使用者瀏覽時無所適從；同一類目中的網頁最好按照與該類目之間的相關程度排列等。

上述研究多以企業網站應用為範疇，並以分類架構的管理角度，提出一些分類架構編製時的注意事項。在這些研究中，可發現對使用者角度的重視，並強調分類架構的使用性（usability）或有用性（usefulness），這不僅是資訊架構相當重要的研究議題，也可做為網路資源檢索系統的分類架構評估參考。

四、自動分群技術及搜尋引擎應用

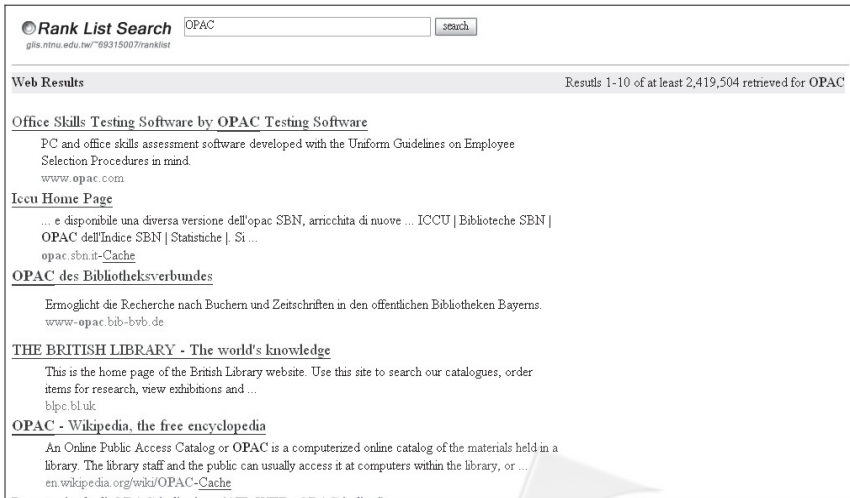
自動分群技術在資訊檢索（information retrieval）與文件探勘（document mining）領域應用相當普遍。基本上，分群是一種將資料不斷群聚，使得屬於同一群內的元素彼此之間的相似程度，高於屬於不同群元素之間相似程度的過程（Lewis & Croft, 1990）。根據Jain等人（1999）歸納，分群演算法主要區分為階層式（hierarchical）與分割式（partitional）兩大類。前者不僅將資料分群，同時也產生階層，這是本文欲探討的範圍。目前提供自動分群功能的網路搜尋引擎並不多見，主要原因可能在於其技術門檻相當高，也無明確的獲利模式所致。目前最主要系統包括Vivisimo（<http://www.vivisimo.com/>）、Grokker（<http://www.grokker.com/>）及Mooter（<http://www.mooter.com/>）。Vivisimo是由三位來自美國卡內基美隆大學的教授，於2000年所發起的研究計畫，並進而轉型為商業公司的搜尋引擎，其也是目前網路上最重要的自動分群搜尋引擎。該公司並於2004年另外成立了Clusty（<http://clusty.com>），專門提供檢索結果的自動分群功能，並以階層式類別架構呈現分群結果，且目前以文字式的階層結構呈現。而Grokker則將檢索結果的分群結果，同時提供綱要式（outline view）及圖形化的二維方式（map view）呈現。Mooter則與Grokker類似，皆相當重視資訊視覺化效果，只是二者選用的圖形化介面有所不同。Grokker採用圓圈，而Mooter採用放射線。除上述三種搜尋引擎外，其它還有KartOO、ujiko等系統，但多以資訊視覺化為應用重點，並不特別強調分群技術的差異。

有關自動分群系統評鑑已有不少研究，但多以求準率及求全率等傳統資訊檢索評估標準做為基礎（Zamir & Etzioni, 1998; Tonella, et al., 2003; Crabtree, et al., 2005），較少以使用者觀點來評估系統檢索效益。Zamir 和 Etzioni（1999）曾針對其所設計的Grouper自動分群搜尋引擎，

進行使用者評估。其主要目的在比較使用者使用前後版本（Grouper I及Grouper II）功能之差異。而Käki 和 Aula（2005）則針對其設計的Findex系統進行使用者評鑑，包括使用者檢索的速度、正確性、成功率及觀感等。研究結果發現，相較無分群功能的系統，有40%的使用者可以找得更快、更好；而使用者也傾向不要有太多的類別呈現；此外，習慣使用分群系統的使用者，有不小比例會利用群集所提供的資訊，做為其檢索語句的來源。

參、研究方法

本研究之評估對象Vivisimo為目前最具代表性的網路分群搜尋引擎。此外，為進一步瞭解分群與無分群對使用者檢索效益的影響，本研究將Vivisimo無分群的檢索結果，設計另一排序條列式介面（Ranked List Search，RLS）（圖二），進行比較。此外，由於Vivisimo是一以英文資源為主的搜尋引擎，因此參與評估者須兼具相當的英文能力及主題專業背景。本研究對象之選擇以具有英文文獻之閱讀需求及能力的碩博士生為主；同時，其多已進入論文完成階段，對所尋找主題皆具有一定的先備知識。研究採用的方法包括實驗、觀察、問卷調查、訪談及檢索歷程記錄分析等，簡述如下：



圖二 本研究實驗設計之排序條列式搜尋引擎介面RLS

一、實驗法

實驗分為指定及自定檢索任務，由受試者實際利用Vivisimo及RLS進行檢索。指定任務的受試者包括八位圖書資訊研究所學生，檢索任務依問題類型的難易度及開放程度設計四項任務，如表一所示，問題內容以資訊組織之相關主題為範疇。每項任務提供資訊需求及情境說明，並限定受試者使用本研究訂定之檢索詞彙，以便受試者能就一致的檢索結果進行評估，限時五分鐘以內完成。此外，為避免使用順序可能造成的影響，受試者分為A、B二組，前者先使用Vivisimo，再使用RLS；後者則先使用RLS、再使用Vivisimo。此外，受試者利用Vivisimo時，可自由運用其檢索結果的群集架構及條列式架構。

除指定任務外，本研究希望透過自定任務的自然情境，以便更完整觀察或避免遺漏使用者的重要行為特性。自定任務不區分問題類型及內容主題，也不設定檢索詞彙，完全由受試者依其實際需求及任務情境進行檢索，受試者包括人社、理工、商管、醫農等四種背景共八位研究生。此外，十六位受試者皆為網路重度使用者（上網經驗平均五年、每天上網平均三小時、自評檢索技巧皆為中等以上），且熟悉搜尋引擎的使用（最常使用依次為Google及Yahoo!），僅有三名使用過Vivisimo。

表一 本研究設計之問題類型及檢索任務

難易度 封閉性	簡	單	困	難
封閉	任務一、何謂書目記錄功能需求（Functional Requirements for Bibliographic Records，簡稱FRBR）？		任務三、何謂ontology？	
開放	任務二、數位圖書館（digital libraries）有哪些相關標準？		任務四、知識組織（knowledge organization）有哪些做法？	

二、觀察法

研究者於十六位受試者進行檢索時，從旁採不介入方式，觀察受試者在檢索過程中的情緒反應（如焦慮、困擾、滿意等表情）及特殊檢索行為（如停頓、反覆等動作），以便於訪談時，能進一步釐清受試者的實際想法。

三、問卷調查法

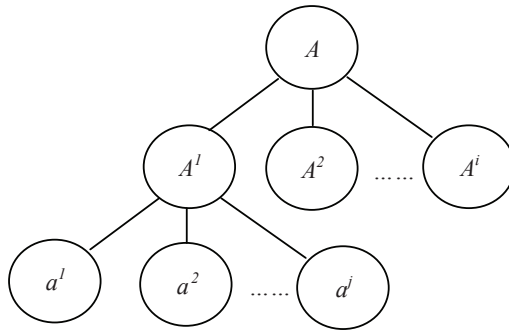
受試者在完成檢索任務實驗後，需填寫本研究依相關文獻、前導研究及專家審查等分析歸納結果，所設計出之網路分群搜尋引擎評估表（表二）。評估範疇包括檢索效率、檢索效益、及檢索滿意度（滿意度再區分為整體、結構、及內容三項子範疇）共二十五項評估問題，希望受試者依據個人使用觀感，以1-5分等級、區分為「非常不同意」至「非常同意」，分別評估Vivisimo及RLS的表現。其中檢索滿意度的結構及內容項目，為單獨針對Vivisimo的群集類別架構進行評估。

此外，本研究所指之群集類別架構是由檢索結果分群後，由群集所組成之階層式架構。如圖三所示，每一群集可視為一主題類別（ A^1 、 A^2 、 $\dots A^i$ ），相關群集再匯集成上層群集（如圖中之A），最底層之群集則包括實際的檢索結果網頁（ a^1 、 a^2 、 $\dots a^j$ ），目前Vivisimo所提供的群集層級約有1-5級不等。

表二 本研究設計之網路分群搜尋引擎評估表

評估範疇	評估項目	評估問題
效率 efficiency	加快瀏覽速度	我能迅速瀏覽檢索結果
	加快檢索速度	我能很快找到所需資訊
效益 effectiveness	瞭解檢索主題本身概念	我能更加瞭解檢索主題的意涵
	瞭解檢索結果整體概念	我能更加瞭解檢索結果的整體概念
	瞭解相關主題	我能更加瞭解與檢索主題相關的主題
	獲得先前不知的相關主題	我能找到先前不知道、但與檢索主題相關的主題
	獲得先前未預期的檢索結果	我能找到先前不知道、但與檢索主題相關的檢索結果
	獲得所有與檢索主題相關的資訊	我能找到更多與檢索主題相關的網頁
	過濾所有與檢索主題無關的資訊	我比較不會找到與檢索主題無關的網頁

滿意度 satisfaction	整體 integral	減輕認知負擔	我能輕鬆地找到所需資訊
		對整體檢索結果感到滿意	整體而言，我很滿意所找到的資訊
	架構 structure	群集架構的廣度適當	我認為每一層級所包含的群集數目合宜
		群集架構的深度適當	我認為層級的數目合宜
		群集架構的階層具邏輯性	我認為層級的階層排列具有邏輯性
	內容 content	群集具主題代表性	我認為每一群集皆能代表其包含檢索結果的主題概念
		群集主題意義明確	我認為每一群集的主題概念明確，不會與其他主題概念混淆
		群集主題容易理解	我認為每一群集的主題概念容易理解
		相關群集具相容性	我認為同一類別群集中的子群集皆相關
		不相關群集具互斥性	我認為不同類別群集中的子群集皆能有所區隔
		群集架構具完整性	我認為整體群集架構包含了所有與檢索主題相關的群集
		群集架構具相關性	我認為整體群集架構所包含的群集皆與檢索主題相關
		群集名稱代表其所包含的主題	我認為每一群集的名稱皆能代表其所包含檢索結果之主題概念
		群集名稱符合認知	我認為每一群集的名稱符合個人認知
		群集架構具有提示聯想作用	我認為群集的階層架構具有提示聯想作用
		群集名稱具有提示聯想作用	我認為群集的名稱具有提示聯想作用



圖三 分群類別架構示意圖

四、訪談法

在完成上述研究步驟後，研究者與受試者進行一對一訪談。除了釐清受試者檢索過程的一些疑點外，主要希望瞭解受試者使用自動分群搜尋引擎的實際感受及後續使用意願等。

五、檢索歷程記錄分析法

本研究利用Morae電腦螢幕錄製軟體，記錄受試者之檢索歷程，以便客觀分析受試者的檢索行為。分析項目以檢索詞彙（如主題內容、詞彙類型）及各種檢索動作（如點選連結、網頁、群集、群集層級的情形）的相關記錄為主。

本研究因尚屬探索階段，有諸多限制，簡述如下。首先，各搜尋引擎的自動分群技術差異很大，為降低系統規模及功能不足所可能造成的評估偏誤，本研究選擇目前最具代表性的分群搜尋引擎為研究對象，同時也僅採用其檢索結果，製作條列排序式的對照系統，進行比較。再者，由於網路資源的變動性及自動分群的動態特性，受試者在進行檢索時，每次所產生的檢索結果及群集可能略有差異，因此本研究已儘量安排受試者在一週內完成實驗，以減輕無法以固定分群結果進行實驗的影響。此外本研究之評估範疇以群集架構的結構與內容為主，針對不同分群呈現介面所可能造成的影響，不包括在本研究探討範圍。最後，本研究現階段包含十六名受試者，也建構一可行的評估架構與方法，希望未來能以此為基礎，進行較大規模的實驗，以提升研究結果之信效度。

肆、研究結果

一、網路分群搜尋引擎之檢索效率

根據指定任務之檢索歷程記錄分析結果，如表三所示，使用者利用條列排序式RLS時，相較於分群搜尋引擎Vivisimo，在檢索速度上有較佳的表現。RLS的任務平均完成時間略低於Vivisimo。若以使用者查獲每一筆相關網頁所需時間來看，利用RLS僅需58秒，但Vivisimo則需1分32秒。同時，RLS在1分7秒內就查獲第一筆相關網頁，Vivisimo則需1分39秒。此外，就每查獲一筆相關網頁平均所需的點擊次數及開啟網頁次數，RLS的次數都較Vivisimo來得低。簡言之，受試者在RLS查詢相關網頁的速度較Vivisimo來得快。

表三 指定任務之檢索效率分析結果

搜尋引擎	問題類型	平均每項任務完成時間	平均每人查獲相關網頁筆數	平均每筆相關網頁需花費時間	找到第一筆相關網頁時間	平均每筆相關網頁需點擊次數	平均每筆相關網頁需開啟網頁次數
RLS	簡單/封閉	4'17"	3.75	1'09"	59"	4.87	2.17
	簡單/開放	4'33"	6.00	46"	47"	3.65	1.50
	困難/封閉	3'41"	4.88	45"	57"	4.20	1.72
	困難/開放	4'30"	3.88	1'10"	1'44"	5.38	2.29
	平均值	4'15"	4.63	58"	1'07"	4.40	1.86
Vivisimo	簡單/封閉	4'23"	2.88	1'31"	1'37"	5.99	2.04
	簡單/開放	4'26"	3.00	1'29"	1'21"	6.50	1.63
	困難/封閉	4'37"	3.13	1'28"	1'32"	6.91	2.28
	困難/開放	4'47"	2.83	1'41"	2'05"	7.30	2.59
	平均值	4'33"	2.96	1'32"	1'39"	4.87	2.17

二、網路分群搜尋引擎之檢索效益

根據指定任務之檢索歷程記錄分析結果，如表四所示，RLS較Vivisimo找到較多的相關網頁，平均每人可查獲4.63筆，而Vivisimo則僅有2.96筆。同時，根據兩位專業人員逐一檢視其相關網頁，並依相關性及權威性判斷，二者所查獲的相關網頁品質相當。此外，就未完成任務

人數，RLS的完成率較高，特別在困難/封閉的問題類型上，RLS的完成人數明顯較Vivisimo來得高。簡言之，在固定實驗時間內，RLS在檢索效益上表現較佳。

但值得注意的是，由二者的檢索結果重複率來看（即同時出現在RLS及Vivisimo的查獲相關網頁之中），依問題類型分別為簡單/封閉（22.58%）、簡單/開放（15.79%）、困難/封閉（24.14%）、及困難/開放（21.88%），比例都不高，這也透露受試者的確由Vivisimo找到一些未在RLS找到的網頁。主要原因可能是Vivisimo讓原本在RLS排序在後的檢索結果藉由分群呈現方式，讓受試者有機會瀏覽得到。換言之，分群呈現介面多少影響到受試者的瀏覽行為。若就受試者在各別搜尋引擎查獲相關網頁的重複筆數來看，受試者在RLS的重複筆數相當高，有超過五成的檢索結果是一致的；而受試者在Vivisimo的重複筆數則僅有三成。顯示即使使用同一種分群介面，受試者查獲的相關網頁一致性也不高。究其原因，可能與檢索結果的多元性有關。由於分群提供更多、更豐富的資訊，受試者自然有較多的選擇，檢索結果也會趨於多元。

表四 指定任務之檢索效益分析結果

搜尋引擎	問題類型	查獲相關網頁總筆數	平均每人查獲相關網頁筆數	重複筆數 (%)	查獲相關網頁之平均品質	未完成任務人數
RLS	簡單/封閉	30	3.75	13 (43.33%)	3.44	4
	簡單/開放	48	6.00	29 (60.42%)	3.67	5
	困難/封閉	39	4.88	22 (56.41%)	3.34	3
	困難/開放	31	3.88	16 (51.61%)	3.46	5
	平均值	37	4.63	20 (54.05%)	3.48	4.25
Vivisimo	簡單/封閉	23	2.88	9 (39.13%)	3.76	5
	簡單/開放	25	3.00	6 (24%)	3.94	4
	困難/封閉	25	3.13	13 (52%)	2.90	6
	困難/開放	20	2.83	3 (15%)	3.27	6
	平均值	23.25	2.96	7.75 (33.33%)	3.48	5.25

三、網路分群搜尋引擎之檢索滿意度

就問卷主觀滿意度之分析結果（表五），RLS的整體滿意度略高於Vivisimo（3.29 vs. 3.10）。就降低認知負擔滿意度而言，受試者普遍

認為Vivisimo因同時提供排序條列檢索結果及分群相關資訊，需費較多心力瀏覽，自然對認知造成一些負擔；反之，RLS的使用則較為輕鬆。根據研究者觀察，當檢索需求明確時，受試者傾向認為使用RLS較為快速；但當檢索需求不明確或對檢索主題不熟悉時，Vivisimo或許能減少一點資訊超載，但卻也產生新的資訊焦慮，因為一些受試者可能在層層群集中迷失，反而無法聚焦思考，造成更多認知負擔。

就檢索結果滿意度而言，指定任務受試者認為二者差異不大，但自定任務受試者則認為RLS的檢索結果較佳。根據訪談結果，主要是有些受試者對於Vivisimo分群之檢索結果是否完整感到疑慮。換言之，因為各群集所包含的檢索結果數量不多，受試者有時擔心會遺漏重要資訊。同時，也有受試者表示平常使用RLS的檢索結果，在量及質皆已足夠，Vivisimo並未提供更多更好的檢索結果。此外，由於自定任務受試者是以Vivisimo與平常所使用的RLS做比較（如Google或Yahoo!），可能多少影響其主觀滿意度。例如自定任務中一名受試者因日常所使用的搜尋引擎，其檢索結果品質並不比Vivisimo好，因此不同於其它受試者觀感，反而較滿意於Vivisimo的檢索結果。最後，有關對檢索結果多重分類的看法，有些受試者認為多重分類可以避免遺漏資訊，但也有受試者認為反而容易造成混淆、且需費時過濾。

表五 受試者主觀滿意度分析結果

滿意度 任務類型	降低認知負擔滿意度		檢索結果滿意度	
	RLS	Vivisimo	RLS	Vivisimo
指定任務	3.38	3.25	3.38	3.38
自定任務	3.13	2.88	3.25	2.88
平均值	3.26	3.07	3.32	3.13

基本上，多數受試者對RLS的檢索速度及檢索結果的質量都感到滿意，同時也對其檢索結果的呈現方式感到符合直覺、易於使用。對於Vivisimo，十六名受試者普遍對其檢索方式感到新奇有趣，也對其所能產生的觸類旁通效果，感到相當滿意。換言之，不論對檢索主題熟悉與否，使用Vivisimo多少都能帶來一些刺激與聯想，讓受試者樂於更進一

步探索，並找到一些新的資訊。針對再使用意願，多數受試者表示，基於使用習慣，日後仍會以使用RLS搜尋引擎為主，但也會將Vivisimo視為延伸工具，樂於嘗試。同時受試者也一致表示有再次使用Vivisimo之意願，甚至有半數受試者在本次實驗後，將Vivisimo加入了個人書籤。

四、網路分群搜尋引擎之群集使用分析

針對群集及群集架構的使用情形，指定任務的受試者其使用群集的比例頗高（表六）。在檢索過程中，不論問題類型及受試者對問題的熟悉度，平均達到44.14%的群集點擊率。此外，受試者對問題愈不熟悉，點擊群集的比例也略高。就自定任務的受試者，其群集點擊率雖不若指定任務來得高，但也有21.51%的點擊率。整體而言，十六位受試者除了一位完全未使用群集外，皆有一定的使用率。顯見分群功能的重要性。

表六 指定任務之Vivisimo群集使用率

問題類型	問題熟悉度	總點擊數	群集點擊數	點擊率
簡單/封閉	3.13	17.25	7.38	42.78%
簡單/開放	2.25	19.5	10.13	51.95%
困難/封閉	2.88	21.63	8.63	39.90%
困難/開放	2.25	20.67	8.75	42.33%

註：問題熟悉度為受試者於實驗前，以1-5分自評對檢索問題的熟悉程度。
分數愈高表愈熟悉。

針對分群架構的層級使用情形，如表七所示，第一層的點擊率最高，到第三層則幾乎無人使用。此外，指定任務與自定任務受試者在層級點擊情形也略有差異。例如前者點擊第二層的比例較頂層高，而後者則是點擊頂層比例較第二層高。但整體而言，受試者的層級使用以1-2層為主；同時，根據訪談，受試者認為第一層的資訊最具參考價值，第二層以下因細節較多，有時反而會造成認知負擔。

表七 群集層級之使用分析結果

群級層級 任務類型	頂層	第一層	第二層	第三層 以上
指定任務	48 (14.12%)	193 (56.76%)	98 (28.82%)	1 (0.29%)
自定任務	30 (16.67%)	132 (73.33%)	18 (10.00%)	0 (0.00%)
總數 (%)	78 (12.38%)	325 (51.59%)	226 (35.87%)	1 (0.16%)

根據問卷調查結果，有關分群架構之層級結構滿意度，包括對廣度（3.25）、深度（3.50）、邏輯性（3.25）的平均滿意度為3.33分，顯示受試者對結構尚表滿意。而根據訪談，受試者進一步建議，廣度最好能以每層10個群集內的數目為宜，因較能輕鬆看出檢索結果的整體概念及架構。而在深度部分，受試者則認為其實2層就已足夠，這與前述受試者實際點選群集的分析結果一致。此外，就層級的邏輯性，受試者多數認為尚能符合認知；但也有受試者表示，當檢索問題是個人相當熟悉的主題，有時反而會對其邏輯性感到疑惑。

受試者對於分群架構的群集內容也感到滿意，如表八所示，平均達3.49分。就個別及整體群集的評估結果來看，受試者對個別群集的滿意度較高，其中對群集主題的理解性及名稱都感到相當滿意（4）。但對於整體群集的滿意度則相對較低，且除了提示性項目外，基本上，對各項目的滿意度都不太高，特別是對群集彼此間的相容及互斥性不表滿意；此外，對於整體群集的完整性與相關性滿意度也不高，呼應前述受試者對檢索結果是否能完整呈現於分群架構中的疑慮有關。簡言之，受試者對於個別群集相當滿意，但對整體群集或是群集關聯的呈現，顯然較有疑慮。



表八 分群架構之內容評估結果

評估範疇	評估項目	滿意度
個別群集	群集具主題代表性	3.35
	群集主題意義明確	3.75
	群集主題容易理解	4.00
	群集名稱代表其所包含的主題	3.75
	群集名稱符合認知	3.75
	群集名稱具有提示聯想作用	4.00
	平均值	3.77
整體群集	相關群集具相容性	2.75
	不相關群集具互斥性	2.75
	群集架構具完整性	3.13
	群集架構具相關性	3.25
	群集架構具有提示聯想作用	3.88
	平均值	3.15
總平均值		3.49

五、綜合討論

本文由使用者角度，就檢索效率（efficiency）、效益（effectiveness）及滿意度（satisfaction）等範疇，評估比較分群搜尋引擎與不分群的排序條列式搜尋引擎的檢索表現。整體而言，排序條列式搜尋引擎能在較快的時間內、找到較多的相關網頁，且品質與Vivisimo相當，同時主觀滿意度也較Vivisimo來得高。而根據訪談結果，受試者多表示基於使用習慣，未來仍會以排序條列式搜尋引擎為主，但對Vivisimo也頗具好感，再使用意願也很高。就研究者觀察，由於多數受試者皆為第一次使用Vivisimo，心態上抱著新奇有趣的成份頗高，未來是否持續使用，並對檢索實際產生效益，仍待觀察。

針對分群架構的看法，受試者的接受度頗高，也肯定分群架構的價值，但也有一些疑慮，整理簡述如下：

（一）突顯重要概念 vs. 遺漏重要結果

透過分群架構，使用者較有機會瞭解檢索結果所呈現的整體概念，以及其中所包含的重要主題。也同時讓一些原本排序在後的相關檢索結果，有機會在分群架構中呈現，讓使用者藉此避免遺漏重要資訊。但另

一方面，由於使用者的注意力有限，有時反而因此忽略了其它重要相關資訊。一得一失之間，不易評斷孰優孰劣，端視使用者檢索需求。例如若只希望獲得幾筆相關資料，二者其實都可以滿足；若希望多多益善，二者同時使用似乎對檢索結果的完整性較有幫助。

(二)提供多元思考 vs. 增加認知負擔

對使用者而言，分群架構的結構與內容提供了檢索結果以外的資訊。這些群集的名稱、結構都有助於釐清檢索主題概念、及發掘新資訊或新觀點。換言之，使用者藉由觸類旁通，有了更多元的思考方向。但另一方面，使用者相對也得付出較多的心力與時間來瀏覽。在使用不分群的排序條列搜尋引擎時，使用者多專注於與檢索問題相關的網頁判讀上；但在使用分群搜尋引擎時，常需在檢索結果與群集間跳躍，偶而也會被新資訊所吸引，而忽略原來的檢索問題。究竟分群所增加的資訊，是提供更多元的方向，還是造成更多的負擔，同樣也需視使用者檢索需求而定。例如使用者對檢索主題相當熟悉，分群或不分群可能差異不大；若對檢索主題較無法掌握，或許分群可以提供一些額外的訊息。

(三)降低資訊超載 vs. 增加資訊焦慮

分群的最重要目的之一是希望降低資訊超載，就研究結果觀察，分群的確有助於縮小瀏覽範圍。但也由於瀏覽範圍縮小，反而可能產生新的焦慮。例如由於每個群集所包含的檢索結果數量有限，一些使用者擔心會遺漏重要資訊，而不斷點選各個群集，以獲取較完整的檢索結果。但分群多半提供多重分類結構，因此有時又會重複看到相同的檢索結果，而需要過濾。也有時群集內的檢索結果太少（如無法與其它群集合併），與使用者的期望有很大落差。此外，當分群架構過於龐大，在不斷重複瀏覽點選中，使用者有時也會迷失在層層群集之中。上述問題都可能造成使用者不耐、茫然與挫折的檢索經驗。因此，分群架構是否適用於需求單純的一般網路使用者，值得進一步探究。

伍、結論

分群搜尋引擎的群集及群集架構提供豐富多元的資訊，也帶給使用者新的檢索經驗。分群本質上屬於一種資訊表徵的方法，如何與資訊檢索系統做良好的結合，還需要更多的研究與探討。本研究目前僅是初步，以下是一些未來研究建議，希望能有更多研究投入。首先，由研究結果發現，使用者對群集的認同度較高，對群集架構則較為保留；同

時，這些群集的主要價值多與檢索主題提示性有關，這與一般搜尋引擎所提供的相關詞提示（*relevant term suggestion*）功能，其實相當類似，分群與相關詞提示的比較需要進一步探討。再者，本研究目前以檢索結果為主要評估來源，但由檢索過程中觀察，使用者對分群的認知深受檢索介面影響，未來可進一步針對介面設計進行評估。最後，就一般認知，分群有助大量資訊過濾，相當適合網路資源檢索應用，但研究觀察，當使用者檢索需求相當明確，或僅需要少量資訊，簡單的排序條列式搜尋引擎其實已經足夠。加上使用者所需之分群功能多是具備語意層次的分類，就龐雜的網路資源環境，分群效益可能有限。使用者注意力是重要資源，若分群效益不高，似乎應該思考將空間釋放給其它功能。換言之，對於分群功能的合適應用範疇，尚需進一步分析探討。

誌謝

感謝國科會專題研究計畫之補助（計畫編號NSC95-2413-H-003-024）。

參考文獻

- Assadi, H. & Beauvisage, T. (2002). *A Comparative Study of Six French-language Web Directories*. Retrieved Sept. 28, 2007, from <http://thomas.beauvisage.free.fr/pubs/isko2002.pdf>
- Barker, I. (2005). *What is information architecture?* Retrieved Sept. 28, 2007, from http://www.steptwo.com.au/papers/kmc_whatisinfoarch/index.html
- Bates, M. (2002). After the dot-bomb: getting web information retrieval right this time. *First Monday*, 7(7). Retrieved Sept. 28, 2007, from http://www.firstmonday.org/issues/issue7_7/bates/
- Bruno, D. & Heather, R. (2003). The truth about taxonomies. *The Information Management Journal*, 37(2), 44-51.
- Chan, L. M. (2001). *Exploiting LCSH, LCC, and DDC to Retrieve Networked Resources: Issues and Challenges*. Retrieved Sept. 28, 2007, from http://www.loc.gov/catdir/bibcontrol/chan_paper.html
- Chen, H., & Dumais, S. (2001). Optimizing search by showing results in context. *SIGCHI'01* (pp. 277-284), Seattle, WA., USA.

- Cisco, S. L. & Jackson, W. K. (2005). Creating order out of chaos with taxonomies. *The Information Management Journal*, (May/June), 45-50.
- Clark, A. (1997). *Being there: Putting brain, body, and world together again*. Cambridge, MA: MIT Press.
- Crabtree, D., Gao, X., & Andreae, P. (2005). Standardized evaluation method for web clustering results. *Proceedings of the 2005 IEEE/WIC/ACM International Conference on Web Intelligence* (pp. 280-283).
- Cross, P., et al. (2000). Subject classification, browsing and searching. In M. Belcher, V. Knight, & E. Place (Eds.), *DESIRE Information Gateways Handbook*. Retrieved Sept. 28, 2007, from <http://www.carnet.hr/CUC/cuc2000/handbook/handbook.pdf>
- Delphi Group (2002). *Taxonomy and Content Classification*. Retrieved Sept. 28, 2007, from http://www.delphigroup.com/research/whitepaper_request_download.htm
- Delphi Group (2004). *Information Intelligence: Content Classification and the Enterprise Taxonomy Practice*. Retrieved Sept. 28, 2007, from <http://stratify.com/infocenter/download/DelphiResearchReport2004.pdf>
- Ellis, D. & Vasconcelos, A. (1999). Ranganathan and the net: using facet analysis to search and organize the World Wide Web. *Aslib Proceedings*, 51(1), 3-10.
- Gilchrist, A. (2003). Thesauri, taxonomies and ontologies – An etymological note. *Journal of Documentation*, 59(1), 7-18.
- Hackos, B. (2005). Taxonomies: lessons from users. *CIDM Information Management News*. Retrieved Sept. 28, 2007, from <http://www.infomanagementcenter.com/enewsletter/200510/fourth.html>
- Harvey, R. (1999). *Organising Knowledge in Australia*. New South Wales: Center for Information Studies.
- Jacob, E. K. (2001). The everyday world of work: two approaches to the investigation of classification in context. *Journal of Documentation*, 57(1), 76-99.
- Jain, K., Murty, M. N., & Flynn, P. J. (1999). Data clustering: a review. *ACM Computing Surveys*, 31(3), 264-323.
- Käki, M. & Aula, A. (2005). Findex: improving search result use through automatic filtering categories. *Interacting with Computers*, 17(2), 187-206.

- Koch, T., & Day, M. (1997). *The Role of Classification Schemes in Internet Resource Description and Discovery*. Retrieved Sept. 28, 2007, from <http://www.ukoln.ac.uk/metadata/desire/classification/classification.pdf>
- Krishnapuram, R., & Kummamuru, K. (2003). Automatic taxonomy generation: issues and possibilities. *Lecture Notes in Computer Science*, 2715, 52-63.
- Kwasnik, B. H. (1999). The role of classification in knowledge representation and discovery. *Library Trend*, 48(1), 22-47.
- Lewis, D. D., & Croft, W. B. (1990). Term clustering of syntactic phrases. *Proceedings of the 13th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 385-404.
- Mai, Jens-Erik (2004). Classification of the Web: challenges and inquiries. *Knowledge Organization*, 31(2), 92-97.
- Mayr, E. (1982). *The Growth of Biological Thought: Diversity, Evolution, and Inheritance*. Cambridge, MA: Harvard University Press.
- Rosenfeld, L., & Morville, P. (2002). *Information Architecture for the World Wide Web*. (2nd ed.). Sebastopol, CA.: O'Reilly.
- Samler, S. & Lewellen, K. (2004). Good taxonomy is key to successful searching. *EContent*, 27(7/8), S20.
- Schwartz, C. (2001). *Sorting out the Web: Approaches to Subject Access*. Stamford, Conn.: Ablex Pub.
- Sravanapudi, A. (2004). Categorization: it's all about context. *EContent*, 27(7/8), S23.
- Tonella, P., Ricca, F., Pianta, E., Girardi, C., Lucca, G. D., Fasolino, A. R., & Tramontana, P. (2003). Evaluation methods for web application clustering. *5th International Workshop on Web Site Evolution*, 33-40.
- Vizine-Goetz, D. (1999). *Using Library Classification Schemes for Internet Resources*. Retrieved Sept. 28, 2007, from <http://staff.oclc.org/~vizine/InterCat/vizine-goetz.htm>
- Vizine-Goetz, D. (2002). Classification schemes for Internet resources revisited. *Journal of Internet Cataloging*, 5(4), 5-18.
- Vogel, C. (2003a). A roadmap for proper taxonomy design. *Computer Technology Review*, 23(7), 42-44.
- Vogel, C. (2003b). Designing a knowledge discovery system. *Computer Technology Review*, 23(10), 42-43.

- Zamir, O. & Etzioni, O. (1998). Web document clustering:a feasibility demonstration. *Proceeding of the 21st Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, Melbourne, Australia, 46-54.
- Zamir O. & Etzioni, O. (1999). Grouper: a dynamic clustering interface to web search results. *Proceedings of the 8th International World Wide Web Conference (WWW8)*, Toronto, Canada. Retrieved Sept. 28, 2007, from <http://www8.org/w8-papers/3a-search-query/dynamic/dynamic.html>