

自動化多媒體電視節目檔案之資訊探勘

包曉天

國立交通大學
管理科學系所
副教授

httpao@mail.cc.nctu.edu.tw

徐永煜

國立交通大學
資訊工程系所
助理研究教授

yyxu@csie.nctu.edu.tw

傅心家

國立交通大學
資訊工程系所
教授

hcfu@cs.nctu.edu.tw

摘要

本文旨在描述多媒體電視節目資料檔案之資訊探勘。本資訊探勘綜合應用語音、影像以及視訊等領域的資訊擷取及分析技術於電視節目之關鍵字產生、演員偵測及場景分段，來建構精簡且有結構意義的多媒體摘要。透過在語音及音訊分析，可以將節目內容與廣告分離開來，應用人臉辨識與影像偵測技術，可以將節目內容中各個場景分段。再用光學文字辨識技術可自螢幕取得每段場景中的關鍵字組。這些關鍵文字還可以用來搜尋或連結來瀏覽更多全球資訊網上的相關資訊內容。為了驗證研發技術的實用性，曾自 2007 年 3 月到 2007 年 8 月錄製了 400 小時的電視新聞節目來實驗。其中關鍵字的擷取及辨識其正確率可達到 88.8%。新聞故事分段的正確率可以達到 89.4%。為了強化資料檔案系統與使用者間的互動，利用 Web 2.0 概念設計的網路檔案資料庫，可以即時記錄使用者上線瀏覽時的點閱活動及收集使用者上網發佈意見等資訊。這些資訊可幫助觀眾找到想看的新聞電視台及節目，讓新聞從業人員找到最常被訪談的人員，讓學者專家得以追蹤研究新聞故事的發展及探討相關後續效應等。也可以提供商業客戶購買新聞時段廣告的重要參考資訊。

關鍵字：多媒體，電視節目檔案，資訊探勘，電視新聞，資訊樹。

Abstract

This paper addresses an automated multimedia TV content archive construction and information mining techniques. The utilized techniques from the fields of acoustic, image, and video analysis, for text information on TV program story, actor tracking and scene identification. The goal is to construct a compact yet meaningful abstraction of broadcast TV program, allowing users to browse through large amounts

of data in a non-linear fashion with flexibility and efficiency. By using acoustic analysis, a TV program can be partitioned into content and commercial clips, with 88.8% accuracy on a data set of 400 hours TV program recorded off the air from March 2007 to August of 2007. By applying speaker identification and/or image detection techniques, each content scene can be segmented with an accuracy of 89.4%. On screen captions or subtitles are recognized by OCR techniques to produce the keywords associated with each TV program. The extracted key words can be used to link or to navigate more related contents on the WWW. By using the Web 2.0 based archive, the information browsing system can collect users' click activities and web post messages, which are useful information for TV business people and advertising commercial in addition to the Nelson Rating.

Keywords: Information Mining, Multimedia, TV Archives, TV News, Information tree.

1. 緒論

自從第一家商業電視台 (WNBT, 1941) 創立以來，電視節目很快的成為主流媒體。以新聞為例，在諸多新聞資訊的來源之中，電視新聞一直都是很多人的新聞資訊主要來源。但是要找歷史新聞資料雖然可以容易的取自公共圖書館的舊報紙資料，但要從圖書館找到電視新聞資料可以說是不可能的。而美國田納西州的 Paul C. Simpson 先生，一位保險業從業者出資於 1968 年成立范德堡電視新聞檔案 (Vanderbilt Television News Archive, <http://tvnews.vanderbilt.edu>)，如圖(1)所示。電視新聞資料庫在美國已經有四十年的歷史了，這也是全世界第一個電視節目的檔案資料庫。

如[9]及[1]所提及的，一個以網頁為平台的全自動電視新聞系統可達到下列目標：

- 學術及應用方面：有了節目檔案資料庫將大量提高電視新聞的品質。當研究學者與



圖(1)：范德堡電視新聞資料庫首頁。

公眾可在新聞檔案資料庫上做內容的分析時，新聞媒體工作者必會在新聞的報導上更小心更用心^{*1}。

- 保存及儲存方面：范德堡電視新聞檔案自1968年Betacam錄像帶的時代便開始了。保留在這些錄像帶的資料隨著時間變質將成為保存上的大問題。今日，我們可以將這些新聞資料以數位形態儲存在網際網路的WWW中。
- 時間及區域方面：網際網路上的使用者可以在任何時間、任何地點取的最豐富的資訊。

有關國際上電視節目檔案資料庫網頁的研究及應用已有很多，茲簡述一二如后。Infomedia [4] 是美國Carnegie Mellon大學所啟動的一個整合型計畫。使用先進的人工智慧技術來整理電視節目影片資料為其全面性的目標。VACE-II [8] 為Infomedia的子計畫。此計畫自新聞視訊的視覺內容中自動偵測、萃

取並編輯出高度令人感興趣的事物，以及故事發展。

本論文提出了一個自動化的電視節目檔案資料庫產生及資訊探勘。其所需的核技術是一個系統化的方法來自動產生電視節目的語意標籤(Semantic label)，並使用先進統計及人工智慧的方法來發掘其中隱藏的資訊。我們期望達成下述重要的目標：

- 雖然WWW網頁提供了非常廣泛的知識及娛樂內容，但工作之餘觀看電視新聞及各類節目已經成為許多人的生活習慣了。因此，我們將電視節目編輯成網頁內容，就可提供大眾一種隨時隨地都可享用的多媒體服務。
- 雖有越來越多的電視頻道在24小時播放節目。人們極需要結構化的資訊來幫助選擇符合需求的新聞及娛樂節目及頻道。
- 雖然每個電台都會宣告節目內容的準則規範及特質，例如新聞節目都強調中立性。但人工編輯作業的節目，要做到與預定規範毫無偏差也是不太容易的。利用統計及人工智慧的方法來評估節目的特性，可幫助電台主管瞭解節目

^{*1}美國CBS公司的前主播Dan Rather先生曾在接受德州時報記者訪問說：「他每天生活中有兩個重擔——(1)收視率與(2)Vanderbilt電視新聞資料庫」。

內容的個性 (characteristics), 也可提供觀眾選擇節目時所需的資訊。

本文後續的內容將包括：第二節是關於建構電視節目檔案資料庫的介紹。第三節將描述由錄製的新聞節目資料中產生所需之語意標籤的方法。第四節將介紹實做多媒體電視新聞檔案資料庫時，各類資訊擷取技術及成效評估。第五節則專注於描述自擷取之資訊中探勘資訊的方法。最後將摘要與結論詮釋在第六節。

2. 電視節目檔案資料庫之建構

一個架構在網際網路上全自動建構電視節目資料檔案系統包含了二個主要模組：(1) 電視節目錄製及內容擷取、分析及結構化檔案建置模組(TV Content Index Generator)；及(2) 提供節目內容搜尋與網路瀏覽的網頁伺服器(Web Server)，茲簡述如后。圖(2)描述此二個模組的整體架構及其與外界環境及使用者間之互動關係。(1) 電視節目錄製及內容擷取模組的主要工作是將電視節目先錄製成便於網路儲存及播放的視訊檔案。再將整段錄製的電視節目檔案切割為一則一則的故事場景單元，並自每個故事場景單元中擷取文字影像辨識成該單元之標籤及相關關鍵字組。再利用關鍵字組到網際網路搜尋來取得更多的相關網頁文字資訊來建構一個結構化節目資訊檔案。(2) 網頁伺服器提供一個友善的網路搜尋與網頁瀏覽介面 (Interface)，供使用者取得想要或有興趣的電視節目資料，甚至對節目發表個人意見，都是網頁伺服器最主要的工作。在 Web 2.0 時代，網頁也同時用來收錄使用者的個人資訊及上網行為，供電視及廣告業者在節目規劃及廣告行銷方面之重要參考。

以下簡述電視節目錄製及內容擷取與分析的原理及方法。首先，將電視新聞錄製成串流格式檔案，並依電視台及節目播出日期時間命名及儲存於資料庫中。隨後，利用場景偵測器來分割串流視訊檔案為一個個的場景，再利用人像辨識技術在場景中找到主角（演員）出現的畫面，每個場景中首次出現主角的畫面作為該場景之關鍵畫面。如場景中沒有主角畫面，則以首張有演員的畫面作為關鍵畫面。關鍵畫面中的文字影像可用光學文字辨識技術[6] 辨識出的文字作為每則場景單元的標題與關鍵字組的參考資訊。取出的關鍵字組接著可用來在網際網路上搜尋與電視節目相關的

資訊。這些資訊除了可以用來幫助建構下節要詳細介紹的電視節目資訊樹外，更可提供網頁瀏覽者豐富的相關資訊及有興趣的節目內容。

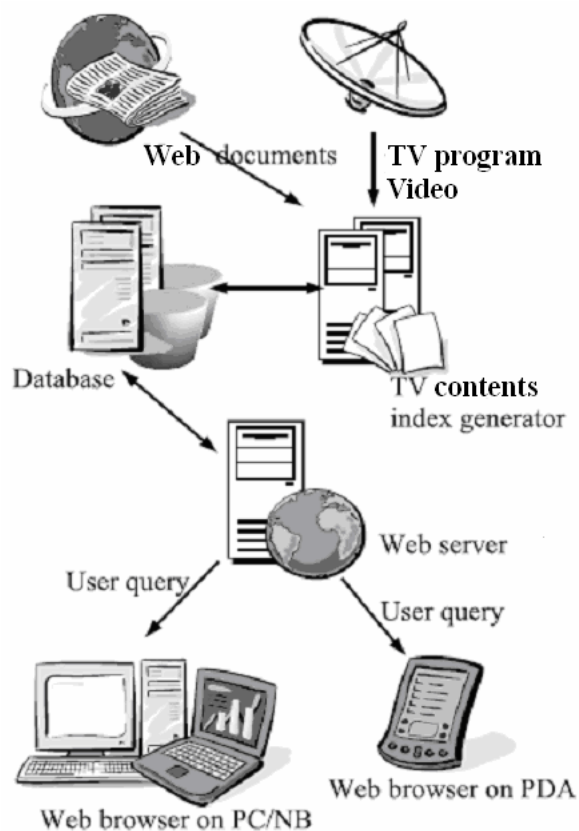


圖 (2)：電視新聞資料庫的整體架構。

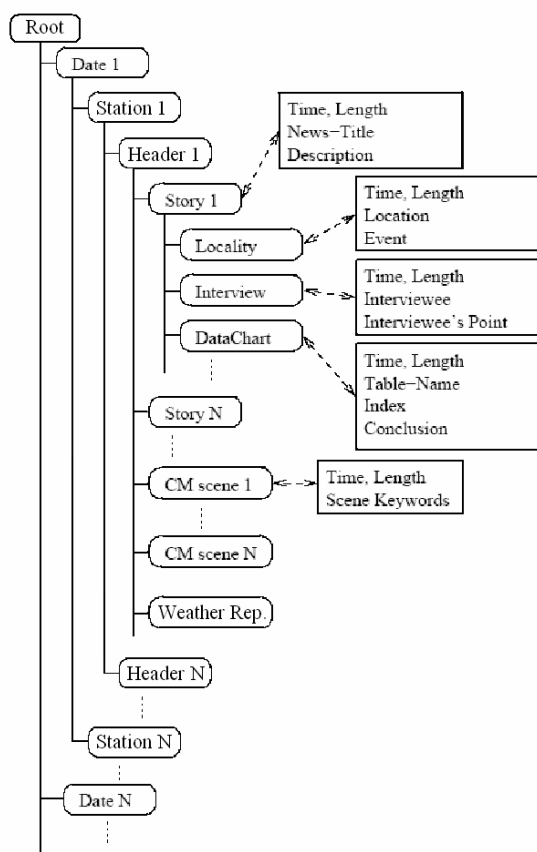
3. 資訊樹之建立

有關資訊樹的定義及學理介紹，可參考[2]所提供之嚴謹及完整的描述。由於電視節目種類繁多，如以通用化的概念形式來描述電視節目資訊樹，將與前述資訊樹的學理介紹相去不遠。因此本節將以新聞資訊樹為例的實務方式來介紹電視節目資訊樹之建構，既可避免內容過於抽象化。又因新聞資訊內容包羅萬象，比較容易推廣到各種類形電視節目資訊樹的建構。

如眾所周知，編寫電視新聞節目的結構化摘要時，宜把握五個 W 的原則：Who, What, When, Where, 以及 Why。這幾個原則在抓住讀者的注意力與簡介故事的要點時是非常有用的。因此在建構新聞資訊樹時，先從新聞事件中擷取五個 W 的資料，再用結構化的方式建構資訊樹。

3.1 新聞資訊樹

圖(3)展示了新聞資訊樹的階層式架構。這個階層式的樹狀結構有六種類型的資訊紀錄：(1)日期(date)，(2)電視台(Station)，(3)標頭(Header)，(4)新聞內容，(5)廣告，及(6)氣象報告。「標頭」紀錄包含了節目開始時間，長度及其所對應之視訊片段的簡潔說明。「新聞內容」紀錄可進一步的分成三類子紀錄型式：(a)現場，(b)訪談，以及(c)圖表。茲分別說明如后。

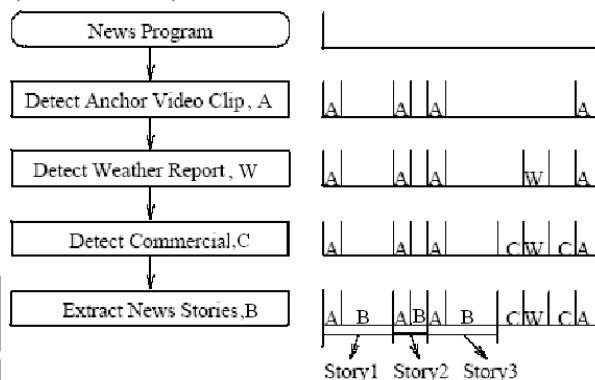


圖(3)：新聞資訊樹的階層式架構。

一般來說，一時段的電視新聞節目包含三種類型的片段：新聞故事，廣告以及氣象報導。這些片段的偵測與分割可用[5]所述之方法及技術來完成。圖(4)展示了用[5]所介紹的方法來切割電視新聞節目的流程圖，茲說明如下。在各式各樣的場景片段中，主播(Anchor, A)視訊片段是首先被偵測出來的。一般來說，主播片段是整個電視新聞時段中最常出現的片段，因此我們提出了使用[3]所建議的非監督式方法來將主播片段自其他片段中分割出來。在找到所有的主播片段後，再用一個以支援向量機器(Support Vector Machine, SVM)為核心技術的視訊分類器[7]

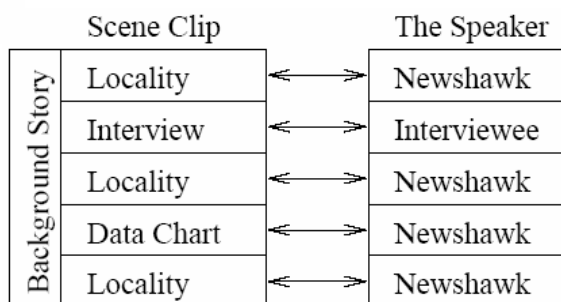
用來偵測氣象報告(W)的片段。最後，廣告內容(C)也可用[5]所述的方法與現場故事(B)分離出來。

針對每個現場故事(B)，如圖(5)所示，我們進一步的將場景切割成三類：現場(Locality)，訪談(Interview)，以及圖表(Data chart)。



圖(4)：電視新聞節目各項內容分割及萃取之流程圖。

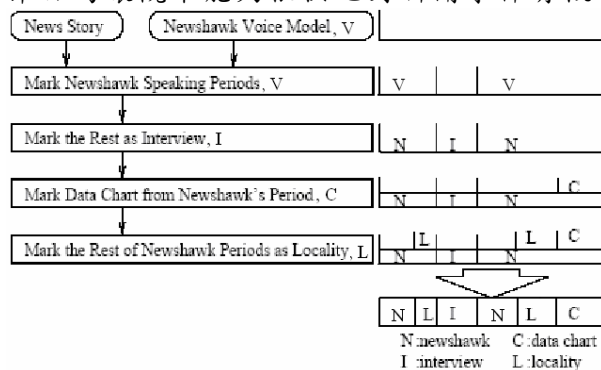
圖(6)展示了如何將現場故事切割成現場，訪談，以及圖表或引言。使用語音偵測技術可將現場記者引言及人員訪談場景偵測(Voice, V)測出來，再用人像辨識將現場記者引言場景(Narrative, N)與人員訪談場景(Interview, I)分離。偵測畫面上文字出現的區域能夠用來找到圖表的場景(Data Chart, C)。最後，如還有沒有標誌場景便應是地區現場(Locality, L)場景。



圖(5)：新聞現場故事內容包含三種資訊：事件現場(Locality)，訪談(Interview)，以及圖表(Data chart)及對應相關人員：現場記者(Newshawk)與現場訪談人員(Interviewee)。

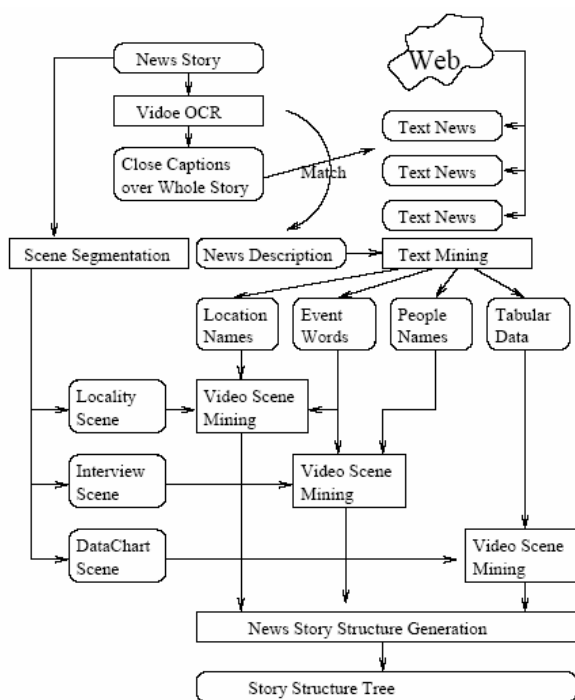
3.2 各單元之語意標籤

本節將描述如何在每個從新聞節目中切割出來的單元上指定語意標籤 (Semantic label)。基本上，標籤上的文字是由對應的關鍵影像中辨識擷取出的。一般來說，電視新聞節目為讓觀眾能夠很快地對新聞事件有概



圖(6)：新聞事件現場場景切割之流程圖。

念，會在螢幕上顯示簡介新聞故事的字幕，例如地名，人名，以及事件的關鍵字組等等。通常，這些文字給各分割單元之標籤提供了相當足夠的資訊。圖(7)展示了新聞資訊樹所需的文字資訊的產生及語意標籤的建立流程。新聞資訊樹的建構流程包括：(1) 新聞故事的文字資訊的產生與 (2) 場景之語意標籤的建立等兩階段。

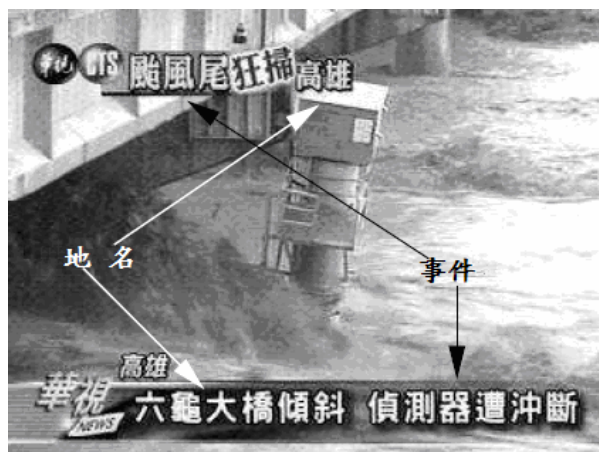


圖(7)：新聞資訊樹所需 (1) 新聞故事的文字資訊的產生與 (2) 場景之語意標籤的建立等兩階段的建立流程。

在新聞故事的文字資訊的產生階段，先從前節所得的關鍵畫面上用文字辨識技術擷取文字影像並將文字辨識出來。辨識出來的文字將作為搜尋關鍵字組，在網際網路上檢索相關文字新聞稿。在找到的文稿中，最相符的文字新聞稿之標題與內容，可用來當作資訊樹中該新聞事件的重要文字內容參考資料。另外，還可在其中取得語意標籤所需的候選字詞，包括 地點，人名，事件名稱，引用文句，以及圖表等資料。

在場景之語意標籤的建立階段，將自前述候選字詞中選取合適的語意標籤給新聞資訊樹中每個場景。如圖(8)所示，新聞事件現場場景的字幕中通常會含有事件現場的地名與事件的簡單描述。因此，在現場場景中，我們可由螢幕字幕所產生的候選字中尋找語意標籤所需與地名或事件名稱相關的文字詞。

圖(9)展示訪談場景的一個範例。在訪談場景中，受訪者的姓名與他們的觀點常常會在螢幕字幕上出現。因此，在訪談場景中，我們可以搜尋候選詞中的人名與事件名稱作為場景的語意標籤。



圖(8)：新聞事件現場場景的畫面範例。現場場景常被用來展示「何事」於「何地」發生。所以，常常能夠自現場場景的字幕中取得「地名」與「事件」描述的資訊。

資料圖表場景的範例如圖(10)所示。一般來說，螢幕字幕常會充滿在整個圖表畫面上。這些字幕表示的可能是新聞故事中所引用的句子或表格資料。因此，自圖表場景中尋找引用的句子或圖表資料可為這類場景的語意標籤。



圖(9)：一個訪談場景的範例。訪談場景常被用來展示現場人們對於新聞事件的觀點。所以，自現場場景的螢幕文字中能夠取得受訪者的名字與他們的觀點。

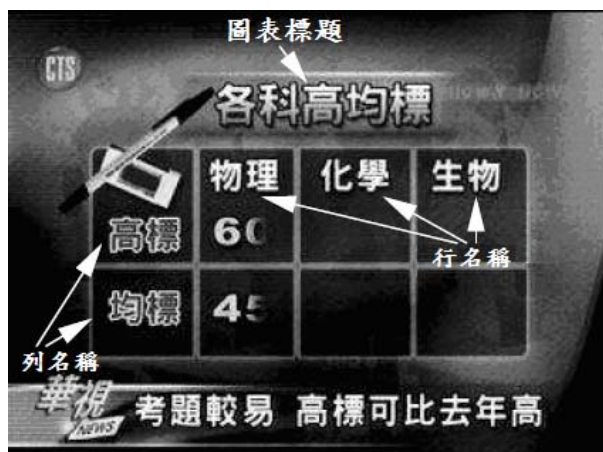


圖 (10)：資料圖表場景的範例。資料圖表常被用來有組織的顯示資訊。詳盡的資訊可以自其螢幕字幕中取得。

4. 電視新聞資訊檔案實驗平台及資訊探勘結果及討論

為了充分了解前述的電視節目分析方法及技術的可行及實用性，我們建構了一個包括 6 台個人電腦 (PC) 的實驗平台。實驗平台中，每台 PC 均以區域網路相連，並與網際網路相通以便各地的使用者上線來利用資料檔案庫中的資訊。各台 PC 的功能分別為：(1) 電視節目串流錄影及場景分割 PC，(2) 螢幕文字擷取辨識、分析 PC，(3) 影像分析及辨識 PC，(4) 內容整合 (建構電視新聞資訊樹) PC，(5) 資料庫 PC，及 (6) Web 伺服器 PC。

我們將此平台應用於全自動電視新聞節目多媒體資料檔案之建構及資訊探勘上。自 2003 年 7 月開始，每天中午及晚間各錄一整段新聞。這個電腦化的系統除了停電外，每週七天，每年 365 天不曾間斷的建構了完整的電視新聞多媒體檔案。為了便於瀏覽，也建置了一個網頁 (<http://nn.csie.nctu.edu.tw/3-1.htm>) 供大眾上線試用、測試及建議指教。這個電視新聞服務網頁可同時展示每天新聞事件的標題、關鍵影像及視訊片段。使用者亦可用「關鍵字」來查特定或相關新聞。圖(11)展示了 2007 年 11 月 30 日晚間新聞節目的部分標題及關鍵畫面。

在這個網頁上，使用者可以指定日期來查閱當天的全部多媒體新聞摘要，包括：(1) 新聞事件的文字標題及 (2) 關鍵畫面，更可以在標題上點閱該則新聞的關鍵畫面及視訊片段如圖 12。

圖(13)展示以「陳總統」為關鍵字查詢相關新聞。查詢結果顯示實驗平台自 2003 年 7 月到 2007 年 11 月間，收錄了 6900 餘則「陳總統」的相關新聞。使用者尚可在每列最尾端的視訊按鈕上點選所想看的視訊片段。

為確實了解我們使用的各種分析、技術及工具的效果，自 2006 年 7 月~2007 年 8 月，將錄製的 400 小時新聞節目內容做了一次資訊探勘效率評估，其結果列述於表 1。

表 1: 電視新聞節目檔案資料庫資訊探勘實驗結果

量測項目 測量值	場景 切割	事件 分割	廣告 辨識	氣象 辨識	標題 辨識	文數字 辨識
正確率	89%	87%	85%	87%	88%	86%
召回率	93%	92%	88%	91%	90%	92%
平均效率	90.9%	89.4%	86.5%	88.9%	88.9%	88.8%

我們以正確率 (precision) 與召回率 (recall) 做為評量辨識偵測效果的標準。其中正確率與召回率的定義如式 (4.1)、(4.2) 所示，並且以式 (4.3) 中所定義的平均效率 (Performance) 來評量平均的辨識偵測結果。



圖(11)：2007 年 11 月 30 日晚間新聞節目的部分標題及關鍵畫面。

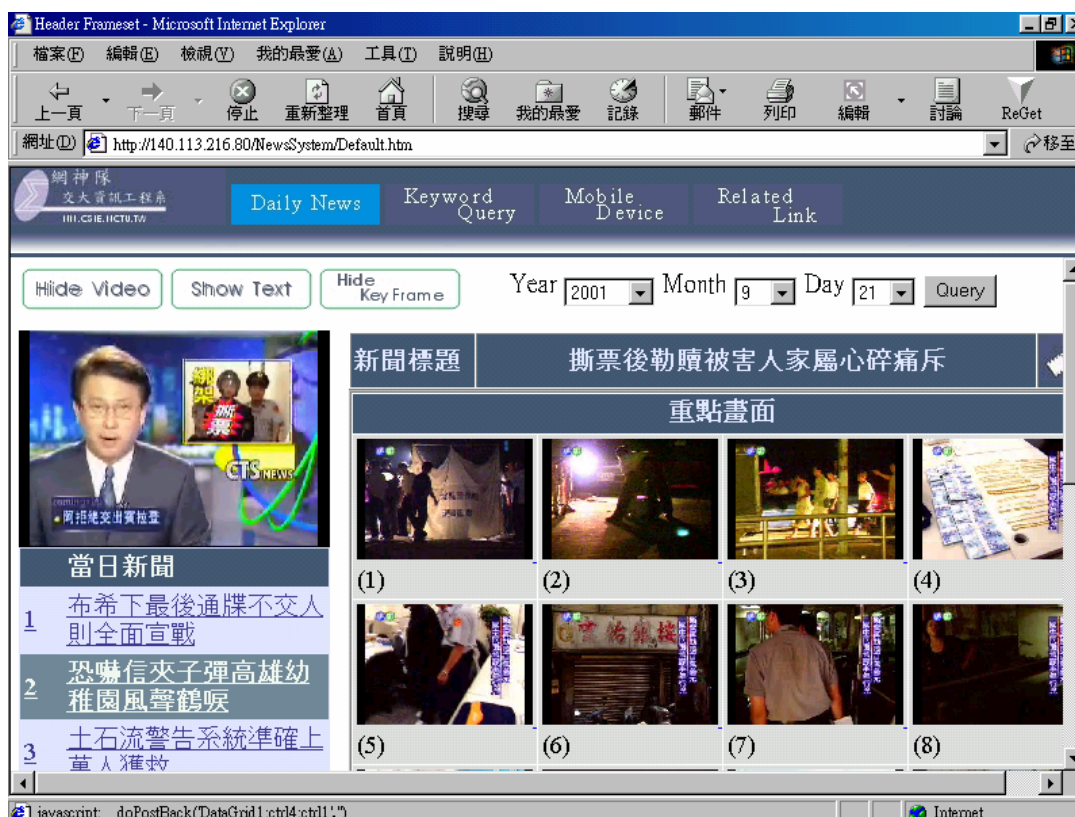


圖 (12)：使用者點選一則新聞後，視窗展示該新聞的關鍵畫面及全部視訊片段。



圖 (13): 以「陳總統」為關鍵字查詢所得之相關新聞。

$$Precision = \frac{Correctly\ detected\ items}{All\ detected\ items} \quad (4.1)$$

$$Recall = \frac{Correctly\ detected\ items}{All\ items} \quad (4.2)$$

$$Performance = \frac{2 * Precision * Recall}{Precision + Recall} \quad (4.3)$$

當正確率越高時，代表辨識偵測結果中錯誤辨識偵測 (false detection) 越少。而召回率越高時，則代表辨識偵測結果中辨識偵測的遺漏 (miss detection) 越少。而平均效率的計算為正確率與召回率的調和平均數 (harmonic mean)，此為一般在資訊檢索 (information retrieval) 中常見的計算方式，稱為 F-measure。

實驗平台自 2003 年 7 月啟用以來，歷經多次改進，系統穩定性已經相當高，目前一年平均當機 2~3 次，執行速度大約是每錄 1 小時節目需 3 小時來作內容分析及資料檔案的建構。

5. 新聞資料樹之資料探勘

本節介紹自新聞資料樹探勘資訊的方法與所能取得的資料。在第(一)小節，我們將展示探勘電視台於不同題材之喜好程度的方法。新聞資料樹還能夠用來追蹤一系列新聞的發展 (看第(二)小節)。除此之外，自新聞資料樹探勘出來的結果以及即時的評比能夠合併來提供有用的導引給購買電視新聞廣告的客戶。

5.1 探勘電視台對新聞的愛好

一般來說，電視台依據新聞的重要性與其對觀眾的吸引力來安排每個新聞故事在節目中播放的時間長短及順序。事實上，受喜好的新聞故事時常能分到較多的播放時間。透過分析新聞故事的時間長短及其播放的順序，我們便能大概估計或判斷電視台所偏好的新聞類別。自新聞資訊樹中抽取這樣得資訊將幫助觀眾尋找他們所想要的新聞電視台。茲簡介建議的新聞探勘方法如下：

給定 N 個對應到 N 個主題的關鍵字集合， $K_1, K_2, \dots, K_i, \dots, K_N$ 。使用下列的函式 $\delta(k, K_i)$ 來定義特定關鍵字 k 與關鍵字集合 K_i 間的關係：

$$\delta(k, k_i) = \begin{cases} 1 & \text{if keyword } k \in k_i \\ 0 & \text{otherwise.} \end{cases}$$

1. 自新聞節目的場景單元 s_j 中抽取關鍵字 $\{k_{s_j}^l; l=1, \dots, L_j\}$ 。
2. 對每個場景單元 s_j ，計算其對應於特定主題 K_i 的關聯頻率 $F(K_i|s_j)$ ，

$$F(K_i|s_j) = \frac{\sum_{l=1}^{L_j} \delta(k_{s_j}^l, k_i)}{\sum_{m=1}^N \sum_{l=1}^{L_j} \delta(k_{s_j}^l, k_m)}$$

3. 計算與日期 d 時間 t 之新聞的關聯頻率 $F_d(t|K_i)$:

$$F_d(t, k_i) = \begin{cases} F(K_i, s_1) & \text{for } s_1.start \leq t \leq s_1.end \\ F(K_i, s_2) & \text{for } s_2.start \leq t \leq s_2.end \\ \vdots & \vdots \\ F(K_i, s_M) & \text{for } s_M.start \leq t \leq s_M.end \end{cases}$$

在這裡 $s_j.start$ 和 $s_j.end$ 指的是場景單元 s_j 在 d 日 t 時之新聞節目中開始與結束的時間。

由於單一新聞節目中故事片段的關聯頻率分佈不足以代表全面性的趨勢，因此長時間統計是有必要的。透過一段長時間（例如說幾個月）對新聞主題之關聯頻率統計，新聞電視台的喜好就可以被發掘出來了。如圖 14 所示，透過選取與政治，社會，和體育相關的關鍵字集，並計算新聞主題出現的頻率。如圖 14 中的範例，該新聞電視台對政治新聞的愛好要高過社會與體育新聞。

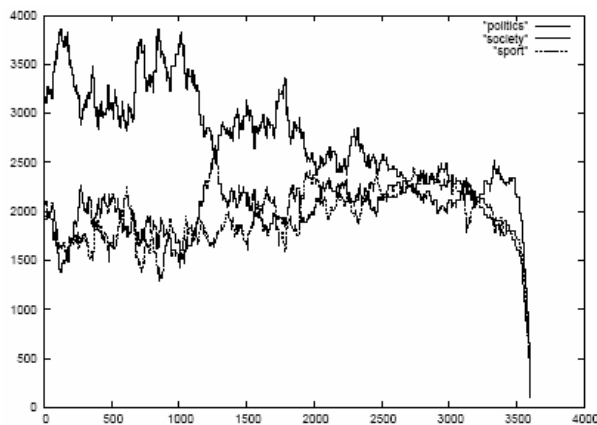
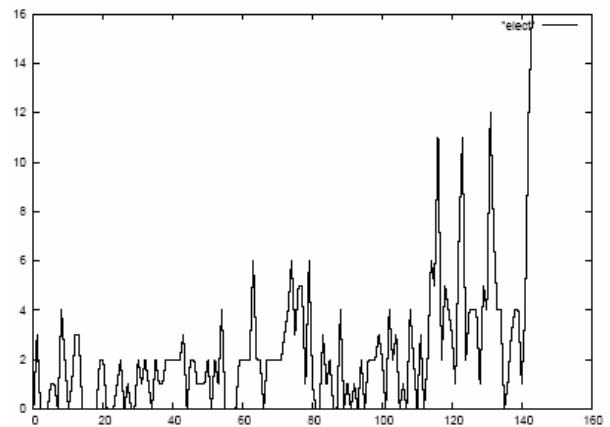


圖 (14)：政治，社會與體育三個代表性的關鍵詞集合被使用來偵測其在特定新聞電視台的出現頻率。

5.2 系列新聞事件的演化與追蹤

新聞故事的演化也可以自新聞資訊樹中探勘出來。透過與特定事件相關的關鍵字集來統計一段時間內相關的新聞場景，統計出來之符合場景單元之關聯頻率便表現了特定事件的發展過程。圖(15)展示了立委選舉前，選舉新聞出現的頻率變化。除此之外，特定事件之散佈地的情形，例如：城市，郡縣，國家等，也可以由比對到之場景單元相關連之地名得到。舉例來說，某人可以透過尋找使用特定的人名來自新聞資訊樹搜尋新聞相關的地名，來得到某人每日的行程。



圖(15)：立委選舉前三個月，選舉新聞出現的狀況。

5.3 電視廣告探勘

除了背景故事外，廣告紀錄也是很有價值的資訊。[5] 提出了一個電視視訊片段上廣告偵測與辨認的方法。當視訊畫面中的廣告的畫面包含了影像關鍵字，視訊光學文字辨識技術便能被使用來自相對應的視訊片段中抽取作為標籤用的關鍵字。在其他方面，以關鍵區塊為基礎的影像檢索的方法[10]可以很有用的來表現與辨認每個廣告片段。然而，關鍵區塊的標籤需要人工來標注。透過在新聞節目中收集這些標籤與關鍵字的統計資訊，電視廣告，即時評比以及新聞故事間交互的關係能夠被觀察與分析以完成一個有用的市場資料庫。接下來將討論「客戶模型」與「交叉銷售」兩個市場資料庫的範例。

- (1) **客戶模型** 顧客（即，電視廣告買家與觀眾）模型的基本構想是要透過鎖定特定廣告可能反應來預測最佳期望增進觀眾的反應率。這可以透過建立一個可能性預

測模型來將喜好同類型之新聞節目，在相同時間觀看新聞的觀眾聚集在一起。除此之外，透過更有效率的鎖定期望以及現有的廣告買家，電視台的營運者能夠增進並強化顧客關係。

- (2) **交叉銷售** 交叉銷售的基本概念是透過槓桿作用於現存有顧客來銷售額外的產品（即，廣告時段）或提供更多的新聞服務。透過分析一起購買之產品或服務的群組，並使用歷史資料預測每個顧客對不同產品的共同點，電視台可以最大化他的銷售潛力至現存顧客。交叉銷售是近年來**資料庫市場**中一個重要的領域，基本上可利用預測資料探勘技術來達到這個目的。使用顧客資料庫中不同產品的歷史購買資料（如觀眾喜好的**新聞型態**，**觀看時間與電台頻道**），廣告買家能夠辨識出他們之產品最能夠吸引那類新聞的觀眾。同樣地，對每種產品型態（即，廣告或廣告群），一個排序過的不同「新聞型態」或「觀眾群組」的列表，可顯示出最可能被該產品吸引的人群。接著，安排伴隨著符合該類型新聞的廣告來達到一個高可能性的觀眾回應率。

5.4 Web 2.0 時代的資料檔案

隨著 2001 年秋天互聯網公司（dot-com）泡沫的破滅，標誌著互聯網的一個轉捩點——基於 Web 2.0 概念的互聯網公司漸漸的冒出台面，蒸蒸日上的開始佔領互聯網舞台的中央。例如 Google 公司以搜尋來服務互聯網上的客戶，互聯網上的資料愈多，Google 的服務便顯得越好——資料愈多 Google 越可以在其中找到使用者想要的資料。同樣的 YouTube 網站公司提出視訊分享的概念，要大家上傳視訊短片來供大家瀏覽。在很多人瀏覽的情況下，熱門短片便很容易的冒出台面，許多人也樂於將自己個人的或喜歡的短片上傳到 YouTube。為了便於搜索 YouTube 提供用戶標誌短片的機制，因此到 YouTube 便可以很容易的搜尋到想找的短片。YouTube 無意間，將多年來許多專家學者想解決的視訊短片搜尋的難題，利用眾人合作分享的機制輕易的達到了。

本系統也將利用 Web 2.0 的概念，將已初步整理標誌過的電視節目的資料檔案，讓上線的大眾來分享的同時，也請眾人幫助做一些目前難以自動化的資訊處理或分析工作。同時

也可從眾多上線的使用者，分析出許多寶貴的資訊，如觀眾的喜好、年齡層、性別、教育程度、甚至收入水平。從而取得許多在 Nelsan「收視率」中從沒能提供的資訊，讓電視台節目製作人、業務經理人、廣告購買者、代理商獲得更多的節目與廣告間互動及關連資訊，產生更多的商機。

6. 結論

本論文提出於多媒體電視節目資料庫之資料探勘與可能的應用。本文以電視新聞為例，來解說可有的資料探勘包括：（1）將錄製的電視節目錄影訊分割成一則則的場景片段，（2）使用視訊文字辨識技術來自視訊取出字幕與影像文字以供每個場景產生關鍵字之用，（3）使用關鍵字來幫每個場景產生語意標籤，以及（4）自節目片段中分離出廣告。這些標籤以及場景資訊（例如：場景的開始與結束時間）都儲存在提出的節目資訊樹中。利用資訊探勘技術將節目資訊樹中許多的隱藏資訊，例如 電視頻道喜好的節目類型以及節目故事的演變等資訊發掘出來。這些資訊可幫助觀眾找到想看的電視台及節目，讓從業人員瞭解最常被觀眾注意的新聞人物及節目演員，讓學者專家得以追蹤研究新聞事件或節目故事的發展，來探討相關後續效應等。也可以提供商業客戶在購買新聞時段或相關節目廣告時的重要參考資訊。

參考文獻

- [1] 賴柏伸、曾政龍、陳岳宏、包曉天、傅心家。電視新聞多媒體資料庫之自動化資訊探勘。**資訊傳播與圖書館學**，11，1-10，（2005）。
- [2] Feinstein, C.D., & Morris, P.A. (1988, May/June). *Information tree: a model of information flow in complex organizations*. *Systems, Man and Cybernetics, IEEE Transactions*, 18(3):390-401.
- [3] Fraley, C., & Raftery, A. E. (1998). *How many clusters? which clustering method? answers via model-based cluster analysis*. In *Computer Journal*. vol. 41, pp. 578-588.
- [4] *Informedia*. [Online]. Available: <http://www.informedia.cs.cmu.edu/>
- [5] Lai, P.-S., Huang, Tzu-Yang, & Fu, H.-C.

- A shot-based video clip search method.* in Proceedings of CVGIP2004, Hualien, ROC. August, 2004.
- [6] Lai, P.S., Lai, L.Y., Tseng, T.C., Chen, Y.H., & Fu, H.-C. *A fully automated web-based tv-news system.* in Proceedings of PCM2004, Tokyo, Japan, Dec 2004.
- [7] Sun, S.-Y., Tseng, C.L., Tseng, Y.H., Chuang, S.C., & Fu, H.C. (2004). *Cluster-based support vector machine in text-independent speaker identification.* in Proceedings of International Joint Conference on Neural Networks, Budapest, Hungary.
- [8] *Vace-II.* [Online]. Available: <http://www.informedia.cs.cmu.edu/arda/vaceII.html>
- [9] Xu, Y., Chen, Y., Tseng, C., Lai, P., Hsieh, R., Lu, Y., Shen, Y., & Fu, H.-C. *Multimedia TV news browsing system.* in Proceedings IEEE International Conference on Multimedia and Expo, Taipei, Taiwan, ROC, June 2004.
- [10] Zhu, L., Rao, A., & Zhang, A. *Theory of keyblock-based image retrieval.* ACM Trans. on Information Systems, vol. 20, no. 2, pp. 224–257, April 2002.