

一個以教材結構為基礎之自動化試題分類測驗系統

馮琪惠[†] 蔡崇煒[‡] 楊竹星* 江明朝[§]

^{†,§} 國立中山大學資訊工程學系

*國立成功大學電腦通訊研究所

[†]m943040059@student.nsysu.edu.tw' [‡]cwtsai87@gmail.com

*csyang@cse.nsysu.edu.tw' [§]mcchiang@cse.nsysu.edu.tw

摘要

在本篇論文中，我們將提出一個自動化試題分類測驗系統，稱之為自動試題分類系統 (Automatic Item Classification System; AICS)。此系統依照教學者提供的教材結構或特定的教材結構，並將其建置成為教材樹，將大量試題與教材內容做相關性連結。系統的主要工作在為試題在教材結構的分類，並將相關分類的教材內容計算其相關性與記錄連結。當系統找出試題與教材內容的相關性之後，試題難易度計算、試卷整體困難度等，系統可分析試題並自動選題，且產生的試題內容。此研究的主要工作主要可分為：1.根據連結教材的結果，系統可呈現試題關連的教材，方便教師快速瞭解此自動出題系統所產生的試卷之範圍及難易程度，以輔助教師在建立試卷時，減少其編輯試卷所花費的時間。2.在另一方面，並且在學生測驗結束後，系統可立即利用試題與教材的關連性，將學生錯誤的試題，回饋相關教材以供學習，以提升學生的學習成效。

關鍵詞：測驗系統、文件分類、分類系統

Abstract

In this paper, we present an automatic item classification system, called AICS. This system uses the content structures provided by the teacher to create a content tree. The content tree correlates items with contents. The main work of AICS is to classify the items and find the most similar contents. After the system computes the relationship between the items and contents, AICS can automatically compute the difficulty of items and examinations. The main tasks of this research can be divided into two categories: (1) The system can show the contents that are related to the items and help the teacher understand the difficulty of the examination quickly. (2) After the examination, the system can provide the contents for student to understand the irrelevant items.

Keywords: examination system, document classification and classification system

1. 前言

近年來，由於網際網路的技術及電腦科技的進步，學習的行為，正逐漸由傳統紙本及教室的學

習，逐漸轉變成為由數位內容所構成的線上學習方式。對於學生的學習進度及狀態的掌握，也因由數位的學習方式，變得更佳容易。在數位學習系統之中，『評量測驗』在各種類型的教學活動中，扮演著相當重要的角色。經由測驗的進行，學習者可以瞭解自己在學習過程中，是否已經瞭解及掌握教學內容的要點。教師也可藉此，深入了解學習者的學習狀況，並評估學生者是否達到標準，作為補救教學與加強學習的依據。在評量測驗的這項學習活動中，構成一份測驗試卷的基本單位即為試題，有良好的試題才會有公平的測驗，透過個別試題的檢視，包含質的分析與量的分析，可以得知試題的優劣。質的分析乃針對試題的內容與形式，從取材的適切性與編制試題的技術方面加以評鑑，量的分析則是根據試題實際施測後，收集受試者的答題結果，分析其試題優劣的指標，其中包含難易度、鑑別度、猜測度、受試者答題的選項分析等，以作為修改或者選擇試題的依據。在這項研究中，我們將在中山大學資工系網際網路實驗室先前開發的GLMS數位學習管理系統之上，設計及實做一個測驗系統，目的在解決以下問題，以及其他延伸方向：1. 試題分類，2. 試題與相關性教材，3. 題庫修正回饋 及 4. 教學與學習回饋。

在此系統下，命題的題庫由學校教師提供，將所累積的題目匯入系統中，成為系統的基礎題庫。其後，康軒及翰林出版社相繼提供支援，擴大題庫現有數量，到目前為止，系統共有三萬多條題目，涵蓋是非、選擇及簡答題，教師選題時的多樣性適足。選題系統是完成考卷的輔助，系統提供困難度、重要度及引用率的統計，便於教師控制選題的難度與重要度。此評量系統提供線上測驗與立即評分，並以長條圖、派圖及四分位圖，表示評量後的量化數據結果。教師可得結果回饋，做為選題修正與補救教學依據。學生可得回饋，以修正讀書方式，補正不足。在教材部份，本系統在教材結構取捨上多了許多選擇，採取的方法是以教育部公布的課程大綱為主體架構，接著依各版本提供的教材結構與內容為系統教材分類，以統一過多的分類雜項。期望提供一個穩定、安全且適用的學習平台，以提升學習者的使用率。

2. 相關研究

在本章節中，我們將藉由介紹數位學習、適性測驗、試題分類及文件相似度計算的研究，說明本研究所建立的系統之相關性研究。

2.1 數位學習與適性測驗

郭生玉[1]指出網路測驗能提供多項優點:1.提升施測與批閱給分效率, 2.便於教師蒐集受試反映資訊; 提供一致性的測驗準備與試題製作, 3.有利於適性測驗的發展, 4.易於反覆練習, 迅速提供受測結果與正確性的回饋導引資訊, 5.易於建立競爭式環境, 提供學習動機與學習成效。網路測驗系統是可提供開放性、自主性、自發性的評量工具, 並提供學習建議[2]給教師與學習者。除了網路學習平台本身應具備的功能外, 陳昭秀[3]整理了Apple(1987), Shneiderman(1992)、Hiltz&Turoff(1993)、Browne(1994)和Leavens(1994)等學者的研究成果, 顯示系統應對於使用者的每個動作, 都該提供明確且有意義的訊息回饋。林姿妙[4]在探討國小教師對於兒童學習網站品質評鑑, 提到優良的兒童學習網站最重要的前三項評鑑準則, 依據為: 網站的設計針對各個學習者學習的不足, 能提供符合個別能力程度的加強或補救練習、網站設計具有線上求助的功能、網頁內容能提供線上測驗區或問題導向的學習活動。最多數國小老師強調網站要能給予學生學習的回饋, 即利用網站教學或幫助學生自我學習。

與傳統的學習環境相較, 數位學習雖然減少面對面的互動機會, 但卻提供了更多其他方式的回饋功能, 增加了更多的互動機會與回饋。例如測驗評量可利用回饋及時得知測驗結果, 並且延伸提供適當的補救學習策略給予學生或者教師。而以電腦化適性測驗 (Computerized Adaptive test; CAT)為目的的選題策略, 有許多研究參照試題反應理論 (Item Response Theorem; IRT)[5], IRT將試題視為測量受試能力的基本單位, 並以試題參數來描述試題的特質, 例如在三參數模式中即使用鑑別度、難度、和猜測度等三種參數來表示。

2.2 試題分類

試題以其文字內容當主軸做為分類的依據, 因此我們可採用文件分類方法當作試題分類的概念。文件分類通常是將一群文件進行分類的工作, 然而分類的方法會隨著分類的目的而有所不同。隨著文件的增加, 必須要以有效率的方法來對文件自動分類。分類後找出文件中使用者真正所需要的資訊, 稱為資訊檢索 (information retrieval; IR)。在此先利用文件分類方式來分類試題, 接著以資訊檢索的方法來尋找試題相對應的教材文件。最早的文件分類由Maron於1961年提出[6], 分類方法是由文件中萃取關鍵字當分類的依據, 假設電腦可以由文件中自動萃取出關鍵字, 便可以自動文件分類。先取文章的摘要作描述樣本, 再使用其他的文件作訓練資料, 接著以測試資料測試, 得到一群詞彙, 去掉出現頻率最高者, 以及出現過一兩次的詞彙。剩下的詞藉由Entropy公式, 計算這些詞彙分佈的情形, 將分佈平均者去掉, 保留分佈不平均者, 才當作分類依據的關鍵字。在文件分類中詞彙權重(term weight)十分重要, 目的是為了由權重來取得當作分

類依據的特徵詞彙 (feature word)。TF(Term Frequency)&IDF(Inverse Document Frequency)是Rocchio在1971年所提出計算詞彙權重的演算法。目前在文件分類與資訊擷取是很重要的理論[7], 其中TF (Term Frequency)計算詞彙出現頻率, 指的是某詞彙在某個文件中出現的次數。IDF (Inverse Document Frequency): 反文件頻率, 出現某一詞彙的文件數量稱為反文件頻率。如果使用IDF來表達詞彙的特徵, 則可以提高回歸率[8]。在目前已有多學者利用 TF&IDF來計算如何取得適當的權重來讓詞彙提高他的特徵性[8][9]。

2.3 文件相似度計算

Maron提出想要對文件作分類, 可由文件中某些重要詞彙作為關鍵字(keywords), 以當作分類的依據。當兩者文件的關鍵字相同性極高的時候, 可稱為這兩文件相似度高。下列列舉兩類研究中最常使用的相似度的計算方式:

1. 機率模式: 機率模式由1976年開始發展, 機率法為探討文件與蒐集的文件是否有相關, 若文件剛好在所蒐集的文件中出現, 表示相關反之沒出現代表著不相關, 接著計算相對應的類別文件與他的相關性。機率模式是求文件中關鍵字屬於某類別的機率, 根據貝式定理(Bayesian theorem), 可算出某文件屬於各類別的機率[10], 機率最大的類別即為文件的分類所在。

2. 向量模式: Vector Space Model (VSM) 1975由Salton 等人提出[11], 文件的內容被化成多維度空間中的點, 並以向量來表示一個文件(圖1), 透過計算和比較向量間的距離, 以前述介紹之多種文件分類方法, 許多人將其應用在試題的分類上, 呂盈輝[12]考慮字詞的鑑別度外, 還加入了字詞的屬性、特徵的位置、特殊符號的個數及語句長度等特徵, 這些試題特徵最後再由類神經網路學習並過濾, 形成試題的分類。張翔宇[13]利用物件導向的模式來表現出本體論的架構, 並依據此模式的建構程序來發展自動建立本體論技術, 從中文操作型題庫中計算字詞權重函數選出重要的特徵詞。再藉由類神經網路、資料採擷等相關技術來達成自動化建構本體論架構。

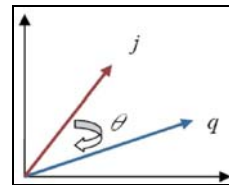


圖 1 VSM 文件相似性計算

而在中文的文件分析中, 中文句子是由一連串的詞所組成, 若要對中文文件做分類或者相似度比較, 勢必要對中文文件做前置處理, 以取得中文關鍵字。斷詞的目的是將中文句子拆成數個有意義的詞, 以方便後續語意分析處理。語意分析部份的目的是經過斷詞後, 藉由自訂且建立的資料庫剖析句子, 以達成自行需要的語意分析目標。中文斷詞的方法有很多種, 基本為建構一個大型且分類好的語

料庫，以用來比對與詞類分析。斷詞的演算法有長詞優先法、經驗法則、統計式模型[14]。教師以概念知識為基礎做教學，概念通常隱藏於文章、教科書或課程指引之中，教師要將它們找出來，並判斷什麼樣的專有名詞是理解重要觀念所必須存在的。將教科書的文章中文斷詞並抽取關鍵詞，去除無用的冗詞，篩選出代表概念的關鍵詞。便能利用此關鍵詞做為試題分類的基礎。

3. 系統設計

在本章節中，我們將介紹本系統的系統架構、系統流程及各項模組的設計。

3.1 系統架構

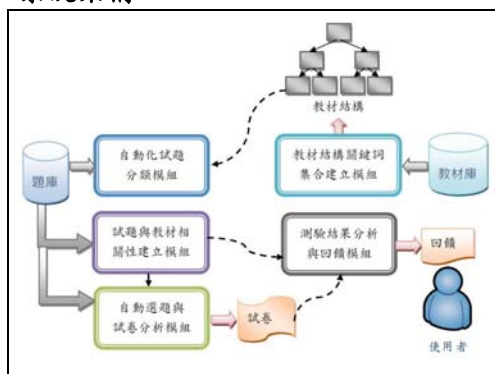


圖 2 系統架構圖

圖 2 為本研究所建置的自動化試題分類測驗系統，由五個模組形成，第一個模組『教材結構關鍵詞集合建立模組』主要在建構一個教材結構，並且利用教材內容替每個節點建立代表其屬性的關鍵詞集合，以此結構為基礎進行其他模組的功能。第二個模組『自動化試題分類模組』，利用前一個模組所建立的教材結構，類似於搜尋引擎的分類目錄，在此將試題歸類到某葉節點(單元)底下，接著第三個『試題與教材相關性建立模組』來比較試題與該單元的多段教材相關性，類似於將查詢的字串與某目錄底下的文件比對，找出相關性最高者與其他次相關性者。第四個『自動選題與試卷分析模組』提供一個自動選題機制來形成測驗用試卷，接著利用前面幾個模組建立的資訊來分析試卷，給予施測者和受測者試題與教材之間的交互訊息。最後第五個『測驗結果分析與回饋模組』將前述四個模組的結果交叉應用，在試題庫部分給予試題難易度修正回饋，在使用者分為教師與學生兩種不同回饋模式。

圖 3 說明本系統的執行流程，其中先建立教材庫，經由中文斷詞及內容分析後建立教材的 Domain Tree。當題庫產生試卷後，試卷的內容以試題為單位，將試題與教材 Domain Tree 進行比對，找出其試題與教材內容的相關性。在學生受測後，回饋相對應於試題的教材內容給予學生，以提升學習效率。

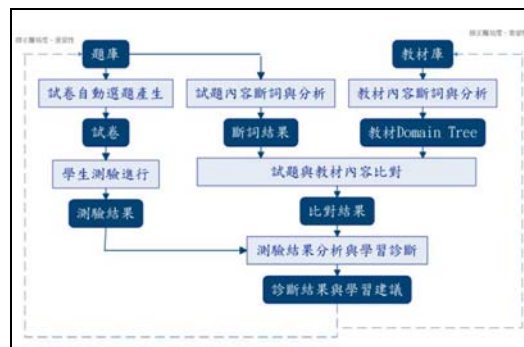


圖 3 系統流程圖

3.2 各項模組的設計

在此我們將詳細說明此系統的五項模組的設計方式。

1. 教材結構關鍵詞集合建立模組

為了在使系統可以整合不同出版社的課程教材，本系統以樹狀結構最能完整的包含不同學制的教材結構，根底下依照學制建立第一層節點，往下一層一層依不同學制的分科目或者科系建立子節點，直到建立課程章節的最小單位，並採用 M-tree 結構(圖 4)來建立課程的教材樹。

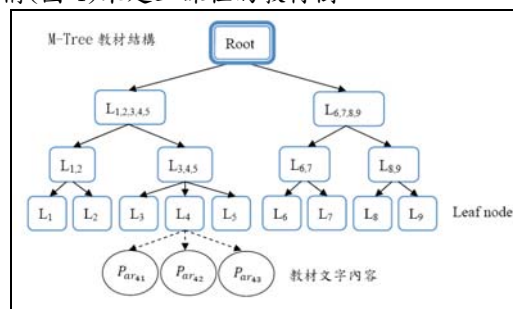


圖 4 M-tree 課程教材樹

為使得每一個節點都有其代表性的屬性，採用分類法則的關鍵詞彙當作屬性，關鍵詞彙的擷取法如下：將教材結構樹葉節點所記錄的課文內容做中文斷詞分析，採用中研院研發之中文斷詞系統[15]，將課文內容中文自動分詞，分詞後去除冗詞，找出符合課程能力指標與代表課文內容的概念關鍵詞，將這些關鍵詞集合起來，建立 domain-based concept 教材(圖 5)，由這些概念關鍵詞組成一個概念關鍵詞庫，此詞庫可以當作知識庫來管理，讓教材課文內容對應並記錄到代表該課文內容的關鍵詞。依據此概念關鍵詞庫，可讓關聯到此詞庫的任何資料表內容互相連結，進一步分析他們的相關性與影響性。

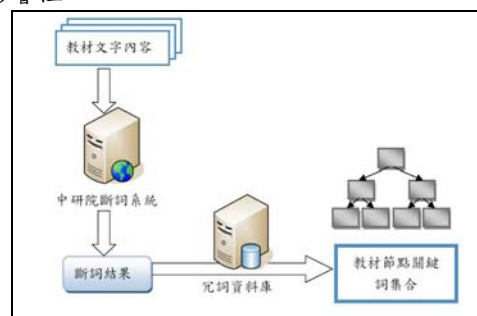


圖 5 關鍵詞集合建立流程

而經由去除冗詞及去除重複並計算頻率(出現頻率及去除重複)兩個程序後,系統可計算出這個教材樹中的各個教材單元的所屬關鍵字出現頻率。圖6說明這顆教材樹中其教材單元中的關鍵字集合。

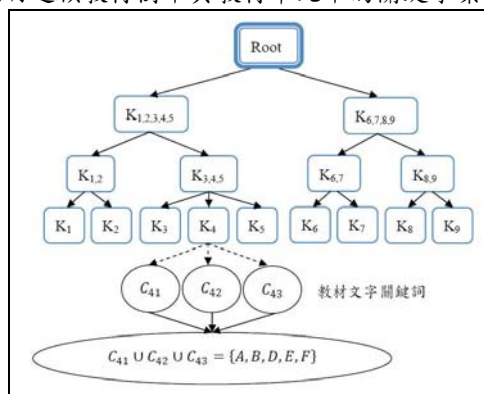


圖6 教材結構與關鍵字集合

2. 自動化試題分類模組

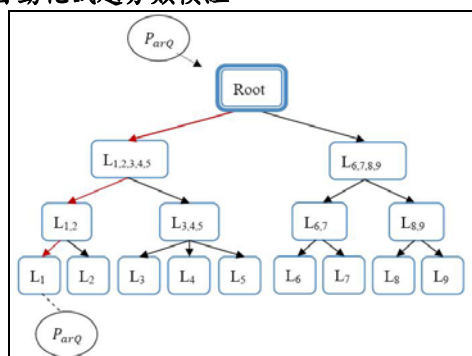


圖7 試題 P_{arQ} 與教材結構的分類結果

圖7表示教材樹狀結構的各節點儲存能代表該節點之關鍵字，將試題由第一層開始，往下一層一層依試題的關鍵字與節點關鍵字之比較，直到最底層。以國小課程為例，第二層為年級的上下學期，第三層可以為科目，以此類推，在最底層為課，課包含著多段不同教材文字內容。 P_{arQ} 代表一道試題進入課程教材樹，其中 $L_{1,2,3,4,5}$ 代表root以下的第一個父親節點，與 $L_{6,7,8,9}$ 所代表的課程為不同單元，利用這樣的教材樹與試題的分類方式，可以把試題快速的對應到相關的教材內容。

3. 試題與教材相關性建立模組

試題的編制必須依據雙項細目表，包括教學目標與教材內容兩個向度。教學目標(以橫軸表示)以Bloom所提的認知領域六個教學目標為依據：知識、理解、應用、分析、綜合、評鑑。教材內容(以縱軸表示)以選題的範圍，表示出包含幾個不同的單元。單一試題可能包含了一個或者以上不同的單元，試題的內容敘述通常會和教材內容相似或者相同，也必定包含了代表該單元的部份知識概念關鍵詞，要將此單一試題歸納在哪一個教材單元，可以視為一段文字要在多篇文件中歸類，視為中文文件分類的議題。中文文件分類方式有幾種，類神經網路模式，向量統計模式，基因演算法模式，資料探

勘模式。將試題歸類在某一單元的教材後，更細分的而單一試題與教材內容的比對，較必須以Segmented Words Comparing來處理。以期達到比對片段文字內容敘述。能將相對應的教材內容顯示出來。本研究針對中文文件分類方式，提出一個較合適的相似度計算方式，稱之為，Intersection Term Frequency (ITF)。此方式將試題與幾段教材內容文字比較，相同的詞彙越多，代表兩者的相關性越高。圖8說明ITF與課程教材的分類過程，ITF為類似向量統計方式來計算試題與教材的相關性，採用試題與教材相同的詞彙(term)加上該詞彙在教材出現的頻率(frequency)，當作相關性的權重計算。

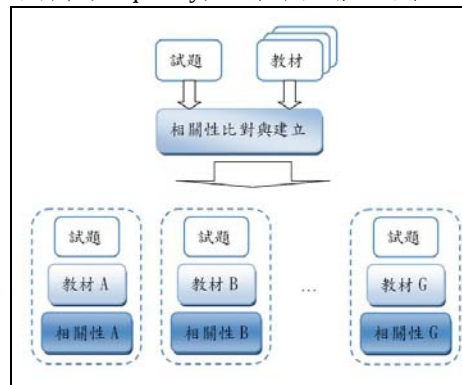


圖8 試題在教材結構分類流程

其方式假設有一個input C_i 和試題詞彙形成的集合 Q ，output為 $R_i = C_i \cap Q$ ， $W_{R_i} = f(R_i)$ ， $P_{R_i} = W_{R_i} / f(C_i)$ 。假設在同一單元有兩段課文的文字內容為 P_{ar1} 及 P_{ar2} ，經由中文斷詞及去除冗詞後，形成兩個集合，分別為 C_1 及 C_2 ，其中 $C_i = \{A, B, D, E, F\}$ for P_{ar1} ，各詞彙出現的頻率對應為 $TF_1 = \{1, 1, 2, 1, 1\}$ ， $C_2 = \{B, C, G, H, K, L\}$ ， $TF_2 = \{1, 1, 1, 2, 1, 2\}$ for P_{ar2} ，要歸類的試題內容為 P_{arQ} ，所形成的集合為 $Q = \{A, B, D\}$ 。首先取 Q 與 C_1 的交集 $R_1 = \{A, B, D\}$ ， $W_{R_1} = 1+1+2=4$ ， $P_{R_1} = 4/6$ ，而 Q 與 C_2 的交集 $R_2 = \{B\}$ ， $W_{R_2} = 1$ ， $P_{R_2} = 1/8$ 。因此 $P_{R_1} > P_{R_2}$ ，故 C_1 與 Q 的相關性大於 C_2 ， P_{ar1} 的相關性大於 P_{ar2} 。

4. 自動選題與試卷分析模組

依據測驗的目的與受試者，依據不同的條件需建立不同程度的測驗，如何從題庫中選取適當的試題，來組成份適合的測驗，自動選題系統必須要有適當的選題策略。依據測驗理論，試題包含三種參數：試題的鑑別度、難易度及猜測度。假設現有一題庫具有專業事先命題者給予適當的鑑別度與難易度，來設計選題策略。而教師最常使用的常模參照測驗，是將選取出來的試題在試卷必須符合常態分佈。而本系統依據這樣的方式從資料庫中選取適當的試題來建立試卷。

步驟1. 選出範圍內所有試題，依照難易度分群。

步驟2. 教師選擇試卷的難易度(圖9)及難易度分配分散或集中。

步驟3. 依照教師選擇試卷的難易度及難易度分配，求出常態分佈之p. d. f。

步驟4. 依照此p. d. f計算出各難易度題目所佔的比例。

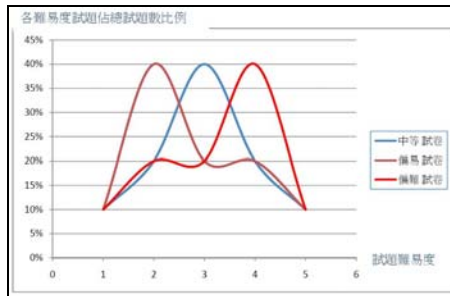


圖9 試卷難易度分佈曲線

自動選題系統所篩選出來的試題，無論是依照預設的條件，或者教師設定的特定條件，最後試題在試卷中的配置方式或者試卷的屬性(包括教材涵蓋度與難易度等)，若能讓教師快速的預覽，便能節省教師出考卷的時間。在教師設定的試卷教材範圍內，依照教材知識分析結果，可得知該設定的教材範圍內，包含哪些該範圍學生該熟知的概念關鍵詞，一份試卷包含了多筆試題，由試題知識分析的結果延伸，對照教材的知識概念，可分析出此份試卷包含了哪些教材的概念關鍵詞，並且可統計該概念關鍵詞出先的頻率，教師進而預覽是否合乎所期待的試卷。每單一試題能對應到教材的某個部份內容，整份試卷所集合的試題對應的教材內容，可以表示該份試卷所代表的測驗內容。在自動選題後若能及時提供這項資訊，可讓選題教師在選題結束時，預覽整份試卷所包含的教材內容。並且在學生測驗完時，各題都有對照的教材內容可讓學生複習與加強。

5. 測驗結果分析與回饋模組

試題的分析，可由兩個方面來看，就試題的內容與形式，主觀的編輯試題技術是否良好來評鑑試卷內容的專業性。另一方面，由測驗結果來分析各試題的難易度、鑑別度與受試者的反應，可以提供命題者瞭解試題的品質。試題與教材章節建立了完整的關連性，可以提供教師作為實施補救教學的依據，並提醒學生(受試者)應加強學習的單元。試題與教材的相關性決定後，在測驗使用的試題即可應用相關性與對應教材的結果。在學生介面(圖10)，我們可以依據測驗的結果分析，答題的錯誤部分試題依其相關的教材，顯示給學生閱讀，是加強學習的回饋，若在教材有所謂的延伸學習教材或補充教材，亦可對應到答題正確的試題部分，是正向回饋。在教師部分，出完試卷後，可依據試卷內試題對應的教材瞭解到考題的教材涵蓋範圍，提供教師出題資訊，並且觀看學生錯誤比例較高的試題所對應的教材範圍，也瞭解到學生在哪部分學習情況不好，提供教師加強教學建議的範圍。

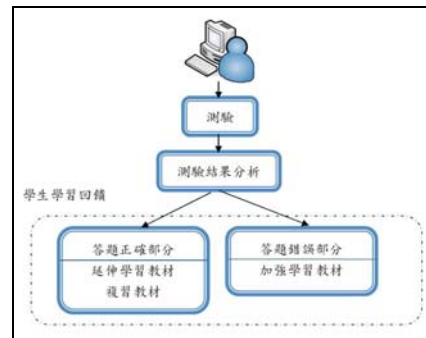


圖10 試題在教材結構分類流程

4. 系統成果及實驗分析

在本章節中，我們將介紹系統的設計及實驗分析結果。

4.1 系統設計

此系統係由先前中山資工所，IT實驗室所開發的GLMS[16]數位學習系統基礎，再針對中小學校園環境的需求，匯整學習教材資源，建立題庫與測驗評量系統，提供教材內容與授課內容影音檔的分享，經由格網的技術將資源整合應用，在IPv6環境中提供一個穩定、安全且適用的教學平台，以提升學習者的使用率。本計畫邀集中山大學、正修科大等幾所大專校院與台南市教育網路中心和國民小學共同參與。利用TANet網路環境與數位學習系統(Learning Management System)為基礎，整合並建立跨校數位學習系統環境，完成數位學習格網(e-Learning Grid)之規劃與建置。

4.2 實驗分析

在實驗分析中，我們利用國小的教材當作測試的資料，其中共有950題試題，共有13課的課程內容。實驗分為兩個部分，實驗室分析ITF方法的在下列課程中，將試題與課程內容的關連正確度。其中表1說明實驗中，所有課程內容中題目的數量。

表1 實驗的課程內容中所含的題目數量

單元名稱	總題數
#1 家鄉人口的分佈	88
#2 家鄉人口的組成	58
#3 家鄉人口的變化	81
#4 行業與生活	75
#5 行行出狀元	82
#6 生活大不同	86
#7 生活中的外來文化	67
#8 為民服務的機構	72
#9 家鄉機構的利用	76
#10 家鄉新建設	54
#11 家鄉建設與問題	65
#12 鄉民的覺醒	79
#13 家鄉的永續發展	67
總計	950

實驗1(表2)說明ITF在不同的教材上的正確率

平均為74%，而錯誤判斷的情況絕大部分在於題目的內容。其中正確率(A)表示在進行分析時，僅利用課程及教材的文字內容。正確率(B)則是加入圖片文字敘述內容進行分析。部分的題目由於包含單元的內容，因此在分類判斷時，會造成系統的誤判。因此，本系統除了找出與試題最為相近的課程內容，並將次相近的課程內容呈現在系統的介面上，以避免ITF判斷錯誤時，提供錯誤的資訊給教師或學員。

表2 實驗1-ITF在不同的教材上的正確率

單元名稱	正確率 (A)	正確率 (B)
#1	82%	72%
#2	72%	72%
#3	65%	73%
#4	72%	73%
#5	65%	70%
#6	69%	79%
#7	60%	70%
#8	57%	68%
#9	76%	95%
#10	65%	74%
#11	55%	63%
#12	49%	58%
#13	46%	94%
總計	65%	74%

實驗2(表3)則是ITF與VSM這個方法的分析及比較。在這個實驗中，我們利用230題試題進行實驗，其中ITF可獲得非常高的正確率，並優於傳統的分類方式VSM。在計算時間的分析上，VSM總共需要0.39ms，而ITF總共需要0.20ms。

表3 實驗2-ITF與VSM的比較

科目	總題數	VSM 正確率	ITF 正確率
家鄉新建設	52	44%	98%
家鄉建設與問題	58	70%	94%
社會鄉民的覺醒	64	70%	95%
社會家鄉的永續發展	56	68%	96%
總計	230	63%	96%

5. 結論

此研究方法可讓大量的試題，依據不同的教材結構來分配試題適當的位置，對應到適當的教材內容，並且提供幫助學習的回饋。在教材結構每個節點，建立了代表該節點屬性的關鍵詞集合，這些關鍵詞除了拿來與試題作相似性比較，亦可當作各不同教材內容關連的建立之用。經由實驗的分析，可以明顯的發現ITF可以提供較高的正確率進行課程與教材的關連。此外，ITF可以用1/2 VSM的計算時間來獲得更好的效果。而本系統在課程教材內容的研究上，可以提供不同出版社的課程內容合併。除了上述的功能，這個系統提供教師更方便的出題環境。在學生的考試方面，可以快速提供學生錯誤題目的教材，以提升學生的學習效率。在另一方面，教師也可以透過系統的統計介面，瞭解學生在

課程的學習中，較不易瞭解的課程內容。

參考文獻

- [1] 郭生玉 (民 78)。“心理與教育測驗”，精華書局，台北。
- [2] 吳錫修(87) 電腦輔助測驗之動態命題與試卷驗證模式研究。第七屆國際電腦輔助教學研討會，國立高雄師範大學。
- [3] 陳昭秀，網路化電腦輔助合作學習之使用者介面設計研究，國立交通大學傳播研究所碩士論文，1995。
- [4] 林姿妙，兒童學習網站品質評鑑準則之發展研究，台南師範學院國民教育研究所碩士論文，2001。
- [5] Hambleton, R. K., and Swaminathan, H., "Item response theory: Principles and applications." Boston: Kluwer-Nijhoff., 1983.
- [6] M.E. Maron, "Automatic Indexing: An Experimental Inquiry," Journal of the ACM, vol. 8, 1961, pp. 404-417.
- [7] Joachims, T.(1996), "A Probabilistic Analysis of the Rocchio Algorithm with TFIDF for Text Categorization", Proceedings of the Fourteenth International Conference on Machine Learning, pp.143-151.
- [8] Salton, G. and Buckley, C., "Term weighting approaches in automatic text retrieval", Information Processing & Management, 24(5), pp. 513-523.
- [9] 林頌華，新聞標題自動分類，清華大學資訊工程學系碩士論文，1999。
- [10] Robertson, S. E. and Sparck Jones, K.(1976), "Relevance weighting of search terms" Journal of the American Society for Information Sciences, Vol 27, No.3, pp.129-146.
- [11] G. Salton, A. Wong, and C.S. Yang, "A Vector Space Model for Automatic Indexing", Communications of the ACM, Nov. 1975, Vol. 18, No. 11, pp. 613-620.
- [12] 呂盈輝，應用於數位學習之中文試題分類技術，中正大學資訊工程學系碩士論文，2002。
- [13] 張翔宇，應用本體論建構自動化試題分類機制之研究，國立臺灣師範大學工業教育研究所碩士論文，2005。
- [14] 杜海倫，「以標題進行新聞自動分類」，清華大學資訊工程學系碩士論文，1999。
- [15] <http://ckipsvr.iis.sinica.edu.tw/> 中研院資訊科學所詞庫
- [16] GLMS, <http://140.117.170.15/glms/11>.