

建構財務危機預測之最佳資料探勘模式

Development of Optimal Financial Distress Prediction Model based on Data Mining Techniques

蔡志豐 (Chih-Fong Tsai)

國立中正大學

會計與資訊科技學系專任助理教授

e-mail : actcft@ccu.edu.tw

邱偉明 (Wei-Ming Chiu)

國立中正大學

會計與資訊科技學系碩士生

e-mail : chiu0324@hotmail.com

摘要

近年來，企業經營環境隨著資訊全球化時代的來臨而有了重大的轉變，在整體經濟情勢的愈加困難下，財務危機發生的可能性已隨之逐年的增加。對企業的投資者而言，企業能否繼續經營將是他們是否願意將資金投入資本市場的主因；企業財務狀況同時也是金融機構決定是否將貸款授予客戶所必須考慮的關鍵因素之一，對金融機構來說，對財務危機做出準確的預測，將使他們更能提高資金分配的效率、減少呆帳的可能。

本研究的主要目的為建構一套能做為財務危機預測的混合模式，並以單一分類器、多重分類器來比較最後的預測效果。經由實證結果發現，本研究所建構的混合模式中，以自我組織映射所結合的方式效果最佳，其平均正確率都比單一分類器、多重分類器高，且測試模糊樣本與全部樣本後，其分類效果也都是最好、最穩定的。

關鍵詞：財務危機、資料探勘、混合架構、多重分類器、

1. 前言

財務危機預測是近年來會計與財務相關領域研究的主要重點之一，發展財務危機預測模式更是被學術界與商業界視為一個重要而且被廣泛研究的議題，因為財務預測對授信決策、金融機構的獲利、投資大眾均有顯著的影響。因此，準確的預測企業的財務危機已經是一個重要而且關鍵的議題[39]。

以往針對財務危機預測的研究，主要是以傳統的統計方法來進行，例如：羅吉斯迴歸、迴歸分析等。然而，截至目前為止，機器學習和人工智慧技術亦被廣泛的應用在發展財務危機預測的模型，例如：類神經網路、決策樹等，這些分類工具已被發展出來。而相關的研

究也均顯示，機器學習技術的效果與預測準確率都比傳統的統計方法還要佳[24,25,27]。

早期機器學習技術在財務危機預測的應用，主要著重於：在不同的技術中找出一個準確率最高的模型（Single Best）來做為預測，但是僅依賴單一模型，還是有可能會有所誤差，導致最近出現一個新的趨勢，就是整合多重分類器（multiple classifiers or ensemble classifiers）來改善資料分類結果，多重分類器為改善傳統單一方式的一種分類器建構模式方法，可提供更穩定且具效能的分類器，所以在財務危機預測領域，啟發了使用多重分類器來求得更佳效能的研究，而相關研究也均顯示，多重分類器的效果確實比單一分類器來的佳[9,43]。

隨著多重分類器的盛行，在財務危機預測領域已有學者已經開始使用較為進階的混合式架構，雖然，混合模式的概念，率先在[26]的研究中早已被提出，但應用在財務危機預測卻是近年來的事，且在相關文獻中，學者所建構的混合模式，僅是著重在分群的結果[13,14]，並無法用來做為準確的預測、也沒有一個標準的建構方法來建制混合式財務危機預測模型。

有鑑於對財務危機做出準確預測的重要性與各預測模式（單一、多重、混合）的開發使用，本研究主要將針對債務信用不良此種資料集進行資料挖掘，主要目的如下：

- 一、建構一套能做為財務危機預測的混合架構並提供相關實驗流程配置。
- 二、比較單一分類器、多重分類器、與本研究所建立的混合模式以找出能提供最準確的財務危機預測模式。

以下，本研究將進行文獻探討，第三部份主要為研究之實驗設計與架構，第四部份為實驗結果與分析，第五部份將針對研究發現進行結論匯總，同時提供後續研究的相關建議。

2. 文獻探討

2.1 財務危機

企業管理的主要理念是希望能讓企業永續經營，而一個企業的生存與否，常受到許多內在與外在因素的影響，諸如：錯誤的投資決策、不良的投資環境、偏低的現金流量等都有可能導致企業的破產。在銀行授信上也一樣，評估授信的人員希望藉由先前收集的資訊，來降低逾期放款率的發生；近年來信用核准的高成長和大量放款業務的管理，無不希望能謹慎評估企業經營品質、顧客信用狀況，減少呆帳之可能[6]。造成財務危機發生的原因很多，各學者專家的認定標準也亦不同，本研究茲將有關公司財務危機定義彙整如表 1 所示：

表 1 財務危機之定義

年份	學者	定義
1992	Bahnson and Bartley	財務危機為債信不足或債權到期仍無法償還本息。
1986	Pastena and Ruland	財務危機為企業資產淨值為負、且沒有能力償還債務的企業。
1972	Deakin	將財務危機公司定義為經歷倒閉、無償債能力及清算的廠商。
1968	Beaver	將企業發生財務危機定義為企業宣告破產、公司債違約、銀行帳目透支、無法支付優先股股息等四類。

2.2 資料探勘

資料探勘是近年來隨著人工智慧和資料庫技術發展，而崛起的一種應用技術，它是利用各種資料分析工具，能夠自動的在資料庫中找尋並分析資料，最後將獲得的結果呈現出來，由於資料探勘能有效地幫助企業處理龐大的資料，並從中萃取出關鍵的資訊以作為管理者在其決策過程中的重要參考，所以，資料探勘的技術在近年來已被大量的使用在不同的領域中[24]。

由於目標的不同，所使用的資料探勘方式也不相同。資料探勘主要劃分兩類：監督式學習與非監督式學習，根據[4]將資料探勘的方式分為六種，分別是分類、估計、預測、分群、關聯、視覺化。

2.2 資料探勘與相關技術

2.2.1 分群分析

（一）自我組織映射網路-SOM

自我組織映射網路（Self-Organizing Map, SOM）為[22]於 1980 年所提出的非監督式學習網路模式，迄今仍是非監督式學習網路模式的典範，在財務危機預測的領域，也被廣泛應用[13,14,21]，下圖 1 為一自我組織映射網路的結構：

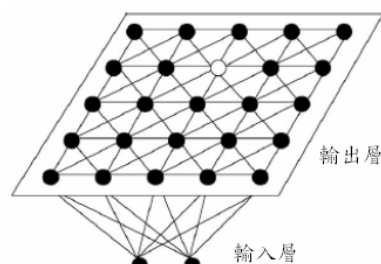


圖 1 自我組織映射網路架構

SOM 的架構包括輸入層、輸出層、網路連結，圖 1 中的輸入層是輸入用來訓練的資料，而輸出層是拓撲座標，映射圖上的類神經元是以矩陣的方式排列，並且根據目前的輸入向量，神經元間彼此相互競爭以爭取勝出。輸出層的神經元會根據輸入向量的「特徵」以有意義的「拓撲結構」將具有相似的特徵，並依網路設定值進行聚類展現在輸出空間中。

（二）K-means

K-means，為 MacQueen 於 1967 年提出的演算法，在許多分群分析中，K-means 演算法為最早的分群計算方式，同時也是最典型、最常被使用的方法之一，K-means 能夠這樣廣泛的被運用最主要的原因就是它的演算法架構簡單、可以處理大量的數值型資料、運算容易且快速的優點能有效的應用在不同領域[17,37,32,41]，所以相較於其他分群方式，K-means 是一種簡單而且有效的統計分群技巧，不過 K-means 也有其缺點：集群的數目（即 K 值）必須由使用者輸入，因為初始中心點的

選擇會影響集群的分割，進而造成所得的分割不能保證為全域最佳，而影響分群效能，可能將因此降低分群可靠度[32]。

2.2.2 分類分析

(一) 類神經網路-Neural Network

類神經網路，即是一種類似、模仿生物神經網路的資訊處理系統，輸入的資料經由學習與記憶的過程，發展出對目標值近似（Approximation）能力和聯想能力，其由許多個別的神經元（Neuron）構成一個網路，使用大量簡單的相連神經元來做運算，並透過這個網路來處理並解決問題。

類神經網路的架構大至可分成兩大類，一為遞迴式網路（Recurrent Network）另一類為前饋式網路（Feed-forward Network）。在類神經網路中最常使用的類型便是多層感知器（Multilayer perceptron, MLP），它主要是用倒傳遞（Back Propagation）的方式來進行演算[7]。MLP 的架構如圖 2 所示，就是將多個感知器連接在一起，形成一串函數連結網路，有輸入層、隱藏層與輸出層，輸入訊號以前送（feed forward）的方式由輸入層一直遞傳計算，經過隱藏層後到達輸出層。

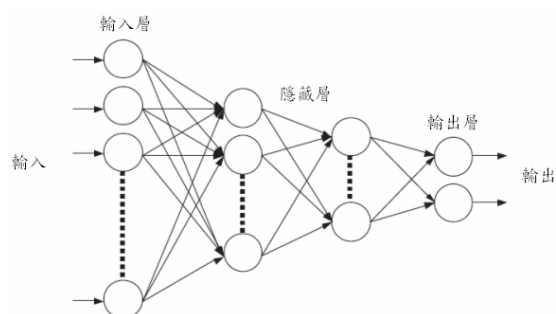


圖 2 MLP 架構

(二) 分類迴歸樹-CART

分類迴歸樹（Classification and Regression Tree, CART）為決策樹的一種，是[5]在 1984 年所發展出來的一種新型演算法，它是以遞迴二元分類的觀念將資料分類，依照預測變數與其相對的各項指標將既有的訓練樣本劃分成數個已知的類別，並將其劃分程序彙整成一連串的規則（Rule），在尋找規則過程中，預測新樣本的類別或目標值的迴歸分析預測，圖 3 所表示的即為一個 CART 二元分類樹之模型[8]。

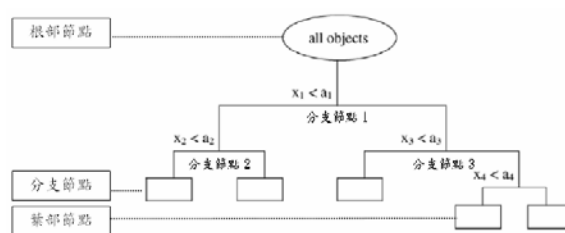


圖 3 CART 二元分類樹之模型

(三) 羅吉斯迴歸-Logistic Regression

羅吉斯迴歸（LR）模型是一個應用非常廣泛的統計技術，係用來預測一個二分類或次序變項的值[12,34]。建立羅吉斯迴歸模型，可以用來描述解釋係數與反應係數之間的關係，進而預測反應係數。所以，當資料是二元反應係數時(Binary Response Variable)，在信用評等、財務危機、破產預測等二分類法的領域經常被應用[16,23,44,45]。

2.2.3 多重分類器

多重分類器就是整合多個分類器的個別決策，一般整合分類器的方式可分為以下幾種[11,19]：

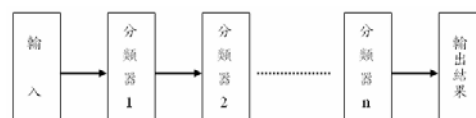


圖 4 串接式整合

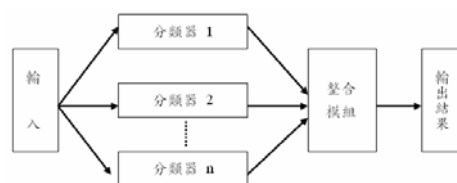


圖 5 並接式整合

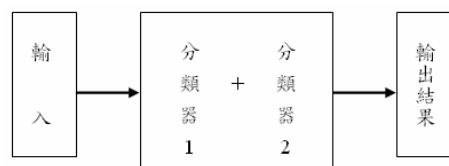


圖 6 結合式整合

(1) 串接式整合(Serial combination)：

數個分類器以序列的方式安排，如圖 4 所示，第一個分類器的結果會成為第二個分類器

的輸入，所以分類器的先後順序將會決定此系統的效能。

(2)並接式整合(Parallel combination)：

將分類平行式地接在一起，當輸入一個值時，多層分類器同步地處理這些輸入值，這些分類器產生的結果再交由一個整合模組（combining module）來整合，並做出一致的決定，一般較為常見的整合方式為多數表決。

(3)結合式整合(Mixed combination)：

結合式即為整合二種不同分類器（合二為一）的混合方法，通常都是進行演算法的修改，其主要目標是希望藉此來提升模型的學習效率。

通常整合分類器的方式會影響多重分類器的效能[40]，在過去有許多研究都指出，使用單一分類技術存在許多限制，一般來說，數個分類器的整合比單獨使用一種分類器的情況下會有更好的分類正確率，這是因為不同的分類器之間可能提供了互補的資訊，因此整合數個分類器對於做出最後的決策將會有所助益[10,38]。

2.2.4 混合模式

混合模式的概念，率先在[26]的研究中被提出，而混合模式的整合方式，是利用兩種不同技術來建構，在財務危機預測中，已經被提出的模式有下列二種。如下圖所示：

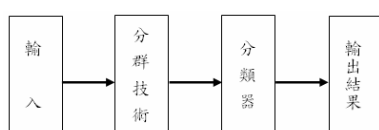


圖 7 混合架構(I)



圖 8 混合架構(II)

圖 7，混合模式（I）：主要是先將資料進行分群，資料經過分群後的結果，做為下一階段分類器的輸入[13]。圖 8，混合模式（II）：由於分群技術屬於非監督式學習，並沒有像監

督式學習演算法如此的準確區分資料，所以事先訓練好分類器，將分類器的輸出做為分群技術的輸入[14]。

2.3 相關研究與討論

財務危機預測是會計與財務領域一直持續不斷的重要研究議題同時也越來越受到重視，而從 1990 年代開始，機器學習技術也漸漸地被廣泛的應用來發展財務危機預測的模型。本研究整理了從 2000—2006 年間、24 篇期刊、32 個實驗，有關資料探勘技術於財務危機預測應用之相關文獻，包括了：採用預測模式的類型（單一、多重、混合式）、資料資訊（資料集來源、交叉驗證），以下表 2 就各文獻依年份整理並敘述討論如下：

表 2 2000-2006 年
資料探勘技術於財務危機預測之相關應用

年份	學者	預測模式類型			資料資訊	
		單一 (single)	多重 (ensemble)	混合 (hybrid)	資料來源	交叉驗證
2006	Gestel et al.		✓		Benelux	✓
	Huysmans et al.			✓	Benel (I)	
	Lensberg et al.	✓			Benel (II)	
	Lee et al.	✓			Norwegian	
	Min et al.		✓		Taiwan	
	Tsakonas et al.		✓		Korea	
	Wu et al.		✓		Greek	
2005	Hsieh			✓	Taiwan (I)	
	Lee et al.	✓			Taiwan (II)	
	Min and Lee	✓			Australian	✓
	Ong et al.	✓			German	✓
	Shin et al.	✓			Korea	✓
	West et al.		✓		Australian	
					German	
2004	Huang et al.	✓			Bankruptcy	
	Li	✓			Taiwan	✓
2003	Kim and Han	✓			US	
	Lee et al.		✓		China	
2002	Malhotra and Malhotra		✓		Korea	
	McKee and Lensberg		✓		Taiwan	
	Pang et al.	✓			US	
	Shin and Lee	✓			US	
	Atiya	✓			China	
2001	Kao et al.		✓		Korea	
	West		✓		US	
2000					Taiwan	
					Australian	✓

（一）預測模式類型

(1)單一分類器

從上述文獻來看，一般而言，財務危機預測：諸如破產、信用等問題（包括其他樣本分類問題）可直接地藉由設計一個簡單的單一分類器來達成，而機器學習與人工智慧技術（例如：類神經、決策樹、基因演算法等）也被廣泛的使用做為分類器。在相關的研究結果中也顯示，機器學習技術所能達到的效能都比傳統的統計方法（例如：邏輯斯迴歸、區別分析等）還要好。

往年在單一分類器應用上以 MLP 為多，但近年來有漸漸減少的趨勢，這可能是因為 MLP 從 1990 年代，早已被廣泛的使用在財務危機預測，導致 MLP 可能傾向於被做為建構其他分類器時比較預測準確度的基準(Baseline Classifier)，而 CART 在近年來也被應用在財務危機預測的領域。

(2)多重分類器

[20]提出多個分類器的整合會比使用單一分類器所得到的效能更佳，他們認為不同的分類器之間可以提供互補的資訊，以利產出更準確的結果。過去，基於以設計多重分類器為目的的文獻裡，他們各別設計使用的多重分類器有三種，分別是串接、並接、結合式。

結合式的整合方式在分類器演算法的選擇上都以 SVM 和 GA 為主，此種結合式的整合方式即是以修改演算法為主。而並接式的設計方式以最常被使用的 MLP 設計一個同質性的多重分類器，但在此並無相關文獻指出，採用異質性的多重分類器效果和同質性之比較。在近二年內，串接式的整合方式沒再度被使用，可能是其準確率無法有效提升。

(3)混合模式

從表顯示，2005 年才開始有學者使用較為進階的混合式架構。[13]先將資料進行分群，由於分群會自動產生多個群集，接著藉由 SOM 產生視覺化的圖表將最不明顯的群集分割出來，再透過分類器進行分類，其目的主要是想了解，因為什麼樣的顧客因素導致在信用狀態良好與不良下難以清楚區分，藉此探討原因，但其研究並不著重在分類的準確與否。

在[14]的研究裡，主要目的是在顧客信用評等資料中，藉由 SOM 來清楚區分信用良好與信用不好，利用 SOM 進行分群時，並無法清楚的區分這兩群，因此藉由提出一個能提升 SOM 在分群時能更準確的方法，由於 SOM 屬於非監督式學習，並沒有像分類器這種監督式學習演算法如此的準確區分，所以藉由事先訓練好分類器，將分類的輸出做為 SOM 的輸入，以提高 SOM 在分群時能更清楚，再藉由 SOM 產生的視覺圖表來進行資料的解讀，在此著重的也是分群的績效。

因此，在這二篇的文獻裡，其目的並不是用來做準確的財務危機預測模式建立，他們的研究皆著重在分群效果，而不是分類，故目前

在財務危機預測領域中，並沒有文獻提出建立一套能做為預測使用的混合架構及實驗的流程與配置。

(二) 資料資訊

(1)資料來源

各個研究所採用的資料集來源皆不同，然而從表中各文獻的資料集來源分析，32 個實驗中，僅有 8 個是屬於公開使用的資料 (Australian、German)。意味著，有 3/4 的實驗所使用的都是自行蒐集而來的資料，這樣的實驗結果將會欠缺公平性與可驗證性，各不同實驗所採用的不同技術或模式類型很難去比較何者的準確率高，因此較受爭議。

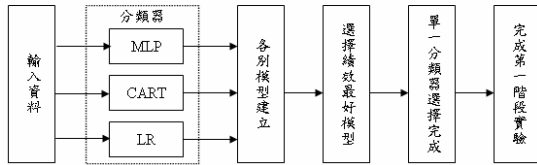
(2)交叉驗證

為了訓練與測試所設計的分類器，所選擇的資料集必須區分成訓練樣本與測試樣本，且為了提供更可靠的評估與訓練模型時能減少資料的相依性，又可避免訓練樣本與測試樣本量的不一致，最好的方式即是進行交叉驗證，但從文獻可以發現，有進行交叉驗證的實驗僅有 6 篇 (佔全部的 1/4)。

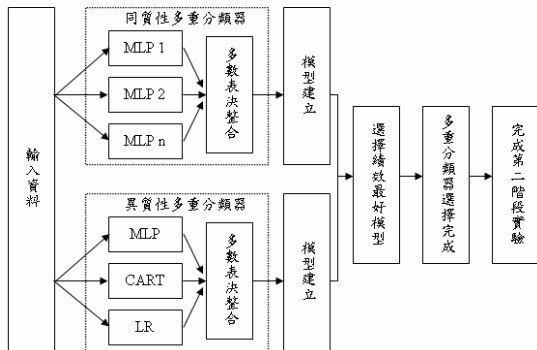
3.研究方法

本研究目的係以資料探勘技術建構一混合架構，並與單一分類器、多重分類器二者做比較，建構出財務危機預測之最佳模式。本研究以類神經網路 (MLP)、決策樹 (CART)、羅吉斯迴歸 (LR) 建構三種不同的單一分類器。在多重分類器部分，藉由上述三種不同單一分類器參數的設定，分別建立其同質性與異質性之多重分類器，再藉由多數表決投票 (Majority vote) 的方式加以結合。在混合模式的建立主要以自組織映射網路 (SOM) 與 K-means 兩種分群技術做為資料的前處理，再結合單一分類器，最後比較單一分類器、多重分類器、混合模式之預測績效。實驗主要分為四個階段，如圖 9 所示：

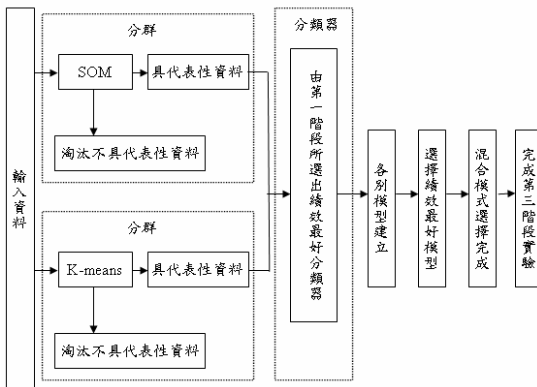
第一階段：單一分類器之建立



第二階段：多重分類器之建立



第三階段：混合模式之建立



第四階段：實驗結果比較

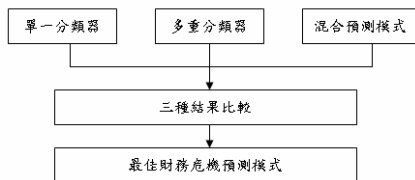


圖 9 實驗流程配置

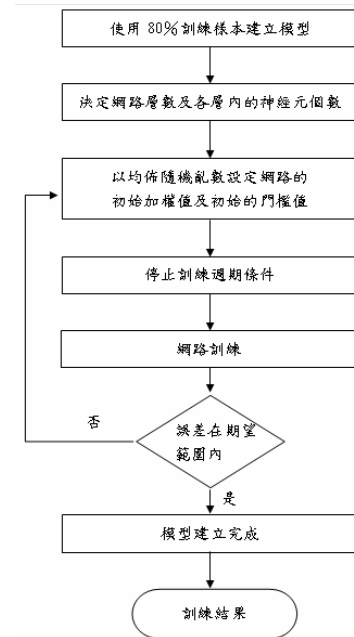


圖 10 類神經網路訓練流程

在建立類神經網路訓練樣本階段，為了提供更可靠的評估與訓練模型時能減少資料的相依性，又可避免訓練樣本與測試樣本量的一致，本研究將採用交叉驗證的方式。交叉驗證：將全部資料樣本劃分為 5 個子集，拿出其中的一個子集當測試樣本，剩下的 4 個子集當做訓練，找出模型中的參數，然後用測試樣本子集來測試由 4 個子集訓練出來的模型，得出準確率，重複 5 次的交叉驗證，故本研究利用 80% 資料做為訓練樣本；20% 資料做為測試樣本，反覆進行交叉驗證。

在決定網路層數及各隱藏層內的神經元個數、停止訓練週期條件設定此兩個階段，由於僅有相關文獻指出，最適的類神經網路層數為一層外，其他設定沒有理論根據指出，何種的設定能產生最好的預測績效（亦即 problem dependent），而相關文獻也只是選定某一種設定來建構他們的分類器，為了避免只由單一的設定來推論最後結果，本研究以 Try-and-Error 的方式，設定不同的參數，再輔以上述交叉驗證的方式找出最好設定。本研究將採用以下 20 種不同設定來進行訓練：

表 3 類神經網路訓練參數設定值

學習週期	50					100					200					300				
隱藏層神經元	8	12	16	24	32	8	12	16	24	32	8	12	16	24	32	8	12	16	24	32

3.1.1 單一分類器之建立

(一) 類神經網路 (MLP)

(二) 分類迴歸樹 (CART)

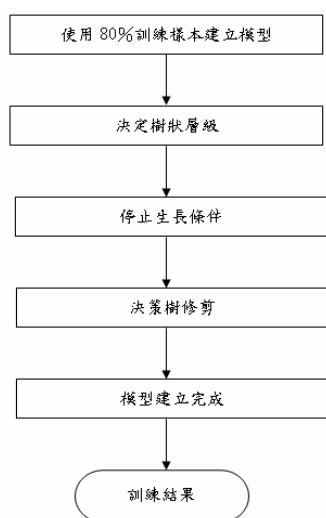


圖 11 分類迴歸樹訓練流程

由於分類迴歸樹的功能強大，使用簡單，僅需要設定預測輸出欄位即可。基本上，樹狀層級不需要設定，因為在整個訓練過程中，會最佳自動化判斷停止生長條件，並進行決策樹的修剪，接著完成訓練，形成最佳樹木層級，產出預測結果。

(三) 羅吉斯迴歸 (LR)

羅吉斯迴歸的建構流程與上述分類迴歸樹類似，僅需要設定輸入的變數欄位即可，在訓練過程中，會最佳自動化形成迴歸式，產出預測結果。

3.1.2 多重分類器之建立

(一) 同質性多重分類器

同質性多重分類器的主要概念是：在相同的演算法下，因為設定參數的不同，形成不同的結果，因為每一種結果的不同，造成預測能力的不一，所以藉由整合多個單一的結果形成最後分類器的結果。本研究以採用 20 種不同的類神經網路參數設定，找出最好的結果。因此，在同質性多重分類器的建立，將從 20 種設定中選出 3 個績效最好的模型來建構同質性分類器。

(二) 異質性多重分類器

有別於同質性多重分類器，異質性多重分類器便是整合不同演算法的單一分類器結果，形成最後分類器的結果。本研究將從實驗

第一階段中所產出的各別單一分類器 (MLP、CART 和 LR) 選出預測績效最好的模型來建構異質性分類器。

通常整合分類器結果的方式則是利用簡單的多數表決法，將奇數個分類器整合，將它們各自的結果集合投票，產出最後分類器的結果。多數決投票，由[15]提出，多數決投票方法所產生的決定乃根據是否全場一致或是至少超過一半分類器同意而定，所得結果以多數決投票進行，此法可提高不穩定演算法的分類正確率。

3.1.3 混合模式之建立

(一) SOM 分群法+分類器的建立

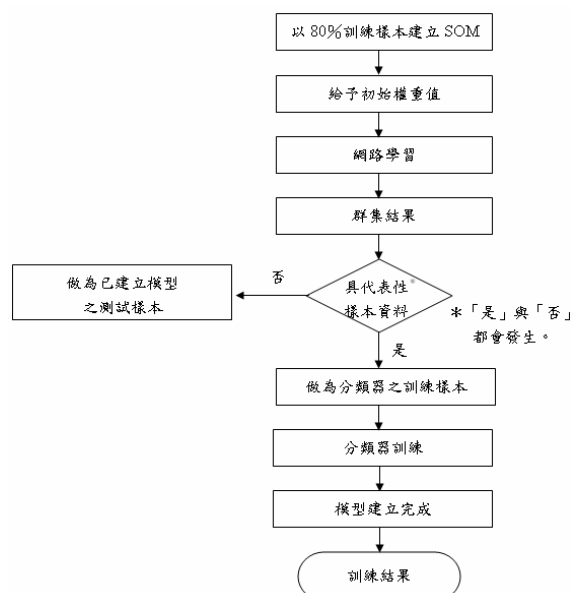


圖 12 混合模式訓練流程

在 SOM 初始權重值設定上，由於 SOM 屬非監督式學習，資料依照本身自我特徵進行集群，在設定拓撲座標時，並無一定的設定值，而是得根據不同設定所產生的不同集群效果來判定，故，本研究將採用 4 種座標的設定來試驗，測試何者集群效果最為明顯，分別為：2x2、3x3、4x4、5x5。

接著，我們利用分群產生的結果，將資料區分為具代表性樣本與不具代表性樣本，由於不具代表性樣本將會影響分類器在訓練模型時的干擾與確準度，故將之淘汰，利用較具代表性的樣本進行訓練，最終產出一分類模型。而不具代表性樣本資料，本研究將它區分出來做為測試模型時的樣本，在此本研究將用來測

試單一分類器、多重分類器、混合模式，因為訓練出來的模型如果能提高預測這些不具代表性或模糊的樣本，那麼才能證明此模型是非常具有可靠的預測能力。

當模型訓練完成後，接著進行模型的測試，在此的測試樣本，將區分成三種不同的樣本資料，除了在分群階段所淘汰的模糊樣本外，還包括：資料集 100% 的全部資料、20% 測試樣本，其目的都是用來比較並檢試混合模式之預測可靠度。

(二) K-means 分群法+分類器的建立

K-means 的建立流程，大至上於 SOM 方法類似，而差別在於，K-means 需設定集群的個數值 (K)，由於為了比較 SOM 與 K-means 分別的結果中尋求最好的模式，以利兩個能有比較的公平基準，[13]在其研究中，便以 SOM 集群的結果來決定 K-means 的個數，故，本研究也將由上述 4 個 SOM 的最佳集群結果來設定 K-means 的 K 值。

3.2 資料來源

本研究採用現今在財務危機預測真實世界情況的資料集來進行實驗，分別是：Australian¹、German²。本研究採用此二個公開的資料集來進行實驗，一來可以免除只用單一資料集的結果來推論最後結論，除了能使實驗結果更為可靠外，還兼具可驗證性。

註：

¹<http://www.niaad.liacc.up.pt/old/statlog/datasets/australian/australian.doc.html>

²<http://www.niaad.liacc.up.pt/old/statlog/datasets/german/german.doc.html>

4.實證結果

4.1 類神經網路參數設定

在進行 MLP 分類器設計時，為了找出最佳參數設定、並且減少資料相依性的問題，本研究採用五個資料子集的交叉驗證方式，進行 20 種參數設定訓練，其結果如下表：

表 4 Australian 參數設定訓練結果

Australian										
週期	50					100				
神經元	8	12	16	24	32	8	12	16	24	32
子集 1	0.840	0.855	0.833	0.826	0.870	0.855	0.833	0.819	0.833	0.855
子集 2	0.826	0.804	0.833	0.797	0.812	0.841	0.819	0.841	0.804	0.812
子集 3	0.877	0.884	0.891	0.891	0.891	0.906	0.906	0.891	0.899	0.891
子集 4	0.833	0.841	0.833	0.841	0.833	0.826	0.833	0.826	0.833	0.812
子集 5	0.832	0.839	0.825	0.839	0.861	0.847	0.839	0.832	0.839	0.832
平均	0.842	0.845	0.843	0.839	0.853	0.855	0.846	0.842	0.842	0.840
週期	200					300				
神經元	8	12	16	24	32	8	12	16	24	32
子集 1	0.841	0.804	0.804	0.804	0.797	0.833	0.746	0.797	0.812	0.804
子集 2	0.819	0.812	0.833	0.804	0.775	0.812	0.812	0.804	0.812	0.790
子集 3	0.920	0.920	0.913	0.913	0.906	0.891	0.906	0.906	0.906	0.877
子集 4	0.812	0.833	0.848	0.833	0.804	0.826	0.826	0.833	0.812	0.790
子集 5	0.876	0.847	0.854	0.854	0.861	0.869	0.839	0.839	0.839	0.854
平均	0.853	0.843	0.850	0.842	0.829	0.846	0.826	0.836	0.836	0.823

表 5 German 參數設定訓練結果

German										
週期	50					100				
神經元	8	12	16	24	32	8	12	16	24	32
子集 1	0.716	0.716	0.741	0.726	0.726	0.731	0.731	0.736	0.741	0.741
子集 2	0.780	0.775	0.790	0.800	0.795	0.770	0.745	0.785	0.770	0.785
子集 3	0.695	0.685	0.690	0.685	0.685	0.700	0.675	0.695	0.700	0.690
子集 4	0.705	0.695	0.725	0.710	0.715	0.725	0.705	0.710	0.705	0.705
子集 5	0.729	0.724	0.714	0.734	0.719	0.734	0.729	0.744	0.739	0.739
平均	0.725	0.719	0.732	0.731	0.728	0.732	0.717	0.734	0.731	0.732
週期	200					300				
神經元	8	12	16	24	32	8	12	16	24	32
子集 1	0.731	0.736	0.746	0.741	0.741	0.731	0.736	0.736	0.731	0.746
子集 2	0.745	0.770	0.745	0.765	0.780	0.755	0.720	0.745	0.690	0.735
子集 3	0.680	0.685	0.695	0.675	0.695	0.660	0.680	0.700	0.670	0.700
子集 4	0.695	0.710	0.660	0.700	0.725	0.685	0.705	0.675	0.710	0.700
子集 5	0.744	0.729	0.724	0.724	0.724	0.734	0.724	0.724	0.724	0.724
平均	0.719	0.726	0.714	0.721	0.733	0.711	0.713	0.716	0.705	0.721

經由五次的交叉驗證後，我們將五個子集的訓練結果平均，以最高平均準確率之參數為最佳設定值，除做為單一分類器設定的參數外，每個資料集各選出前三個最高的設定，以做為同質性多重分類器的設定值。最後選定之參數設定值分別如下表 6 所示：

表 6 參數值選定結果

Datasets	名次	週期	神經元個數
Australian	1	100	8
	2	50	32
	3	200	8
German	1	100	16
	2	300	32
	3	100	8

4.2 各預測模式之結果與比較

經由前述之各預測模式之建立與實驗後，本研究將各資料集實驗結果與比較彙總如下表 7、表 8 所示：

表 7 Australian 實驗結果與比較

	訓練	測試	測試	測試
	80%樣本	20%樣本	100%樣本	2%模糊樣本
單一	NN	0.91957 (3)	0.91098 (3)	0.90364 (3)
	CART	0.90792 (6)	0.90890 (4)	0.89754 (5)
	REG	0.87634 (7)	0.87452 (7)	0.87460 (7)
多重	同質性	0.92113 (4)	0.90678 (5)	0.90566 (2)
	異質性	0.91077 (5)	0.91270 (2)	0.90160 (4)
混合	SOM	0.92567 (1)	0.92002 (1)	0.90628 (1)
	Kmeans	0.92382 (2)	0.89468 (6)	0.88622 (6)

表 8 German 實驗結果與比較

	訓練	測試	測試	測試
	80%樣本	20%樣本	100%樣本	2%模糊樣本
單一	NN	0.89549 (3)	0.85512 (3)	0.86048 (3)
	CART	0.78904 (6)	0.76632 (6)	0.77220 (6)
	REG	0.77682 (7)	0.75594 (7)	0.76980 (7)
多重	同質性	0.91177 (1)	0.87111 (2)	0.87588 (2)
	異質性	0.83386 (5)	0.80935 (5)	0.81461 (5)
混合	SOM	0.90340 (2)	0.88958 (1)	0.87650 (1)
	Kmeans	0.89278 (4)	0.83346 (4)	0.84542 (4)

註：各數值後面的括號為總體排名的名次

首先，在第一階段實驗：單一分類器的建構上，類神經網路的預測結果都比分類迴歸樹、羅吉斯迴歸還要佳，這也證實了過去研究結果所發現，亦表示在單一分類器的建構上以 MLP 最為準確、也較為穩定，其預測平均正確率分別為：Australian 91.098%、German 85.512%。

在第二階段實驗：多重分類器的建構上，我們利用先前由交叉驗證所選出的前三名參數來做為同質性多重分類器的設定值，接著利用類神經網路、分類迴歸樹、羅吉斯迴歸三者來建構異質性多重分類器，經由實驗的結果顯示，多重分類器的預測結果比單一分類器來得好，其預測平均正確率分別為：Australian 異質性多重分類器 91.27%、German 同質性多重分類器 87.111%。雖然在不同資料集所呈現出來的多重分類器結合的效果均不同，但本研究藉由客觀的測試三種不同的樣本（測試樣本、全部樣本、模糊樣本），我們可以發現，不管在任一資料集的多重分類器結合方式所預測的結果都比較穩定，而且效果也都較單一分類器來得佳，也再次證實了藉由多個分類器的整合會比使用單一分類器所得到的效能更佳，因為不同的分類器之間可以相互提供互補的資訊。

在第三階段實驗：混合模式的建構，藉由先前在第一階段實驗中所選出最佳的單一分類器一類神經網路，來做為混合模式的建立。首先 SOM+MLP 的建構上，在分群階段，我們經由嘗試 4 種：2x2、3x3、4x4、5x5 不同的座標設定結果，最後找出以 5x5 為最佳的集群效果，接著將模糊樣本區隔開來，將具代表性的樣本投入分類器做訓練。在 K-means+MLP 的

建構上，我們利用 SOM 的集群個數值來做為 K-means 的 K 值設定，藉由上述 SOM 的集群結果，我們以 10 做為 K 值輸入，亦將分群後之模糊樣本區隔，接著做分類器的訓練。經由實驗的結果顯示，SOM+MLP 的效果比 K-means+MLP 的效果佳，平均正確率分別為：Australian 92.002%、German 88.958%。

在混合模式 K-means+MLP 的效果上，經由實驗結果我們發現其整體預測效果並不佳，這可能是因為 K-means 集群的數目（即 K 值）通常必須由使用者輸入，因為初始中心點的選擇會影響集群的分割，進而造成所得的分割不能保證為全域最佳，而影響分群效能，因此降低分群可靠度[32]。而本研究為了能在 SOM+MLP、K-means+MLP 的建構上能有一致的比較基準，故以[13]在其研究中提及：以 SOM 集群的結果來決定 K-means 的個數，雖然這樣的方式能比較有公平的比較基礎，但也因此可能造成了 K-means 無法有最佳的分群效果，導致最後預測結果的不佳。

在實驗的第四階段，本研究將各階段所產出的結果彙總於表做比較。由表 7、表 8 中，我們可以明顯的發現，不管是在 Australian 或 German 的資料集中，本研究所建構的混合模式（SOM+MLP）平均預測正確率在這七種不同的預測模式中都是最高的。除了測試樣本外，本研究將各預測模式進行另外二種不同樣本測試，經由實驗結果顯示，本研究所建構的混合模式不管測試何種樣本所產出的預測效果也都是最高的，在測試全部樣本上分別是：Australian 90.628%、German 87.650%；測試模糊樣本分別是：Australian 74.651%、German 89.333%。特別的是，經由測試淘汰的模糊樣本後，本研究所建構的混合模式仍然具有較高的預測結果，這也意味著，本研究所建構的模型是非常具有可靠的預測能力、而且穩定度也較高。

4.3 穩定度分析

由於過去相關文獻僅提供其實驗所建構的分類器經由測試「測試樣本」後的平均正確率來做為模型的評估報導，本研究為了加強模型的可靠度與穩定度，另外採用了「全部樣本」、「模糊樣本」來做為測試，藉由不同客觀的角度來驗證模型的穩定度，以下圖 13、14 為本研究將實驗所設計的七種模式中，各取出在單一分類器、多重分類器、混合模式三種模式中，總體預測正確率最高來比較模型的穩定

度：

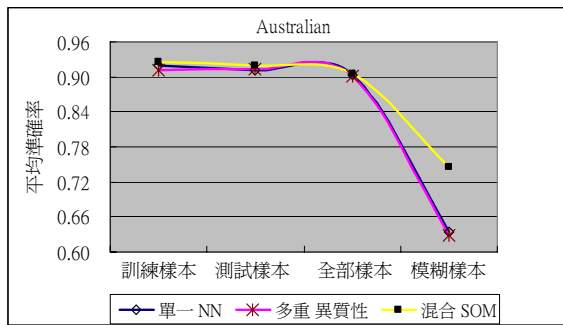


圖 13 Australian 穩定度分析

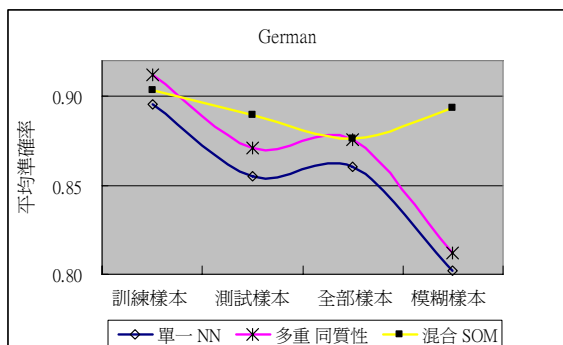


圖 14 German 穩定度分析

藉由圖表中的結果顯示，本研究所建構的混合模式，不管在 Australian、German 資料集的測試結果下，從訓練樣本到測試模糊樣本的平均準確率，其線條呈現較為平緩穩定與一致性，而其他二者分類器所呈現的結果較為波動不穩，且當進行模糊樣本測試後，單一分類器、多重分類器的平均準確率大幅度的往下降，因此，從穩定度的角度來看，本研究所建構的混合模式也是較具有高度的預測能力。

5. 結論

企業財務危機預測在這近三十年來一直是會計與財務領域的一個重要研究議題，企業的破產倒閉會影響每個國家的經濟和政策制訂者為了促進經濟成長的意圖，同時也造成了社會界、工商業界巨額的損失，在資本市場也會對於每個不同個體公司的繼續經營情況做出不同的反應，所以企業應該要不斷的對經營體質的健康度做出危機預測，因此，財務危機預測對企業而言，是非常重要的，而過去機器學習技術應用於財務危機預測的相關研究亦紛紛的建構出不同的預測模式，希望能提供實務界一個降低損失成本的機會。

本研究基於對財務危機做出準確預測的重要性與各預測模式的開發使用，藉由嘗試多種資料探勘技術的建構與實驗、並相互比較最後預測結果，經由研究結果指出，本研究所建構的混合模式 (SOM+MLP) 在七種不同模式裡，確實能提供較高的預測準確率，平均正確率分別為：Australian 92.002%、German 88.958%，而且面臨不同樣本的測試後，依然能有較高的穩定度。

雖然本研究所建構的預測模式能有較高與較穩定的預測結果，但未來研究還是有很多地方可繼續延伸與探討：

1. 本研究所建構的混合模式主要係以分群的概念來做為「資料樣本」的前處理，將模糊樣本汰淘，以利分類器訓練樣本時能更為精確，因此，建議未來研究可加入特徵擷取 (Feature selection) 將「資料變數」進行前置處理。
2. 多重分類器的整合方法眾多，而本研究僅採用較為簡單的多數表決法來做整合，因此，未來相關研究可以嘗試不同的整合方法來建構多重分類器。
3. 在模型評估上，本研究除了首先採用不同的樣本進行準確率測試外，尚可再檢試模型的型 I、型 II 錯誤率、t-test 等。
4. 混合模式 K-means+MLP 的建構中，本研究為了能有比較的基礎，以 SOM 集群的結果來決定 K-means 的個數，因此導致預測效果並不是非常良好，所以對於 K 值的設定還是可以再深入研究。

參考文獻

- [1] Atiya, A.F., "Bankruptcy prediction for credit risk using neural networks: a survey and new results," *IEEE Transactions on Neural Networks*, Vol. 12, No. 4, 929-935, 2001.
- [2] Bahnson, P. and Bartley, J., "The Sensitivity of Failure Prediction Models to Alternative Definitions of Failure," *Advances in Accounting*, 55-64, 1992.
- [3] Beaver, W.H., "Market prices, financial ratios and the prediction of failure," *Journal of Accounting Research*, 179-192, 1968.
- [4] Berry, M.J. and Linoff, G., "Data Mining Techniques: For Marketing Sale and Customer Support," *New York: John Wiley & Sons, Inc.*, 1997.
- [5] Breiman, L., Friedman, J. H., Olshen, R. A.

- and Stone, C. J., "Classification and Regression Trees", *The Wadsworth Statistics/Probability Series*, Belmont, CA, USA, 1984.
- [6] Chen, M.C. and Huang, S.H., "Credit scoring and rejected instances reassigning through evolutionary computation techniques," *Expert Systems with Applications*, Vol. 24, 433-441, 2003.
- [7] Chiu, C.-C., Tien, C.-C. and Chou, Y.-C., "Construction of Clustering and Classification Models by Integrating Fuzzy Art, CART and Neural Network Approaches," *Journal of the Chinese Institute of Industrial Engineers*, Vol. 22, No. 2, 171-188, 2005.
- [8] Deconinck, E., Hancock, T., Commans, D. Massart, D.L. and Heyden, Y.V., "Classification of drugs in absorption classes using the classification and regression trees (CART) methodology," *Journal of Pharmaceutical and Biomedical Analysis*, Vol 39, 91-103, 2005.
- [9] Gestel, T.V., Baesens, B., Suykens, J.A.K., Van den Poel, D., Baestaens, D.-E. and Willekens, M., "Bayesian kernel based classification for financial distress detectoin," *European Journal of Operational Research*, Vol. 172, 979-1003, 2006.
- [10] Ghosh, J., "Multiclassifier systems: Back to the future," in *Proceeding 3rd International Workshop on Multiple Classifier Systems*, 1-15, 2002.
- [11] Giacinto, G., Roli, F. and Fumera, G., "Design of Effective Multiple Classifier Systems by Clustering of Classifiers", *Proceedings of ICPR 2000 15th International Conference on Pattern Recognition*, Barcelona, Spain, 160-163, 2000.
- [12] Hosmer, D.W. and Lemeshow, S., "Applied logistic regression," *John Wiley & Sons, New York*, 2000.
- [13] Hsieh, N.-C., "Hybrid mining approach in the design of credit scoring models," *Expert Systems with Applications*, Vol. 28, 655-665, 2005.
- [14] Huysmans, J., Baesens, B., Vanthienen, J. and van Gestel, T., "Failure prediction with self organizing maps," *Expert Systems with Applications*, Vol. 30, 479-487, 2006.
- [15] James, G., "Majority Vote Classifiers: Theory and Applications," PhD Thesis, Department of Statistics, University of Stanford, 1998.
- [16] Joanes, D.N., "Rejecting inference applied to logistic regression for credit scoring," *IMA Journal of Mathematics Applied in Business and Industry*. Vol 5, 35-43, 1993.
- [17] Kanungo, T., Mount, D. M., Netanyahu, N. S., Piatko, C. D., Silverman, R. and Wu, A. Y., "An efficient k-means clustering algorithm: analysis and implementation," *IEEE Transactions on Pattern and Machine Intelligence*, Vol 24, No. 7, 881-892, 2002.
- [18] Kao, L.J. and Chiu, C.C., "Mining the Customer Credit by Using the Neural Network Model with Classification and Regression Tree Approach," *IEEE International Conference on Systems, Man and Cybernetics*, 923-928, 2001.
- [19] Kim, Eunju., Kim, Wooju. and Lee, Yillbyung., "Combination of multiple classifiers for the customer's purchase behavior prediction," *Decision Support Systems*, Vol 34, No. 2, 167-175, 2003.
- [20] Kittler, J., Hatef, M., Duin, R.P.W. and Matas, J., "On combining classifiers," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 20, No. 3, 226-239, 1998.
- [21] Kiviluoto, K., "Predicting bankruptcies with the self-organizing map," *Neurocomputing*, Vol 21, No.1-3, 191-201, 1998.
- [22] Kohonen, T., "Self Organized formation of topologically correct feature maps," *Biological Cybernetics*, Vol 43, 59-69, 1982.
- [23] Laitinen, E.K., "Predicting a corporate credit analyst's risk estimate by logistic and linear models," *International Review of Financial Analysis*, Vol 8, Issue 2, 97-121, 1999.
- [24] Lee, K., Booth, D. and Alam, P., "A comparison of supervised and unsupervised neural networks in predicting bankruptcy of Korean firms," *Expert Systems with Applications*, Vol. 29, 1-16, 2005.
- [25] Lee, T.-S., Chiu, C.-C., Chou, Y.-C. and Lu, C.-J., "Mining the customer credit using classification and regression tree and multivariate adaptive regression splines," *Computational Statistics and Data Analysis*, Vol. 50, 1113-1130, 2006.
- [26] Lenard, M.J., Madey, G.R. and Alam, P., "The design and validation of a hybrid information system for the auditor's going concern decision," *Journal of Management Information Systems*, Vol. 14, No. 4, 219-237, 1998.
- [27] Lensberg, T., Eilifsen, A. and McKee, T.E.,

- "Bankruptcy theory development and classification via genetic programming," *European Journal of Operational Research*, Vol. 169, 677-697, 2006.
- [28] Malhotra, R. and Malhotra, D.K., "Differentiating between good credits and bad credits using neuro-fuzzy systems," *European Journal of Operational Research*, Vol. 136, 190-211, 2002.
- [29] McKee, T.E. and Lensberg, T., "Genetic programming and rough sets: a hybrid approach to bankruptcy classification," *European Journal of Operational Research*, Vol. 138, 436-451, 2002.
- [30] Min, J.H. and Lee, Y.-C., "Bankruptcy prediction using support vector machine with optimal choice of kernel function parameters," *Expert Systems with Applications*, Vol. 28, 603-614, 2005.
- [31] Min, S.-H., Lee, J. and Han, I., "Hybrid genetic algorithms and support vector machines for bankruptcy prediction," *Expert Systems with Applications* (in press), 2006.
- [32] Mingoti, S.A. and Lima, J.O., "Comparing SOM neural network with Fuzzy c-means, K-means and traditional hierarchical clustering algorithms," *European Journal of Operational Research*, Vol. 174, No. 3, 1742-1759, 2006.
- [33] Ong, C.-S., Huang, J.-J. and Tzeng, G.-H., "Building credit scoring models using genetic programming," *Expert Systems with Applications*, Vol. 29, 41-47, 2005.
- [34] Ozdamar, K., "Paket programlarla istatistiksel veri analizi—1", *Kaan Kitabevi*, Eskisehir, 2004.
- [35] Pang, S.-L., Wang, Y.-M. and Bai, Y.-H., "Credit scoring model based on neural network," *IEEE International Conference on Machine Learning and Cybernetics*, Beijing, Vol. 4/5, 1742-1745, 2002.
- [36] Pastena, V. and Ruland, W., "The merger/bankruptcy alternative," *The Accounting Review*, Vol. 61, No. 2, 288-302, 1986.
- [37] San, O.M., Huynh, V. and Nakamori, Y., "An alternative extension of the K-means algorithm for clustering categorical data," *International Journal of Applied Mathematics and Computer Science*, Vol. 14, No. 2, 241-247, 2004.
- [38] Schiele, B., "How many classifiers do I need?" *International Conference on Pattern Recognition ICPR 2002*, Quebec, Canada, 2002.
- [39] Shin, K.-S., Lee, T.S. and Kim, H.-J., "An application of support vector machines in bankruptcy prediction model," *Expert Systems with Applications*, Vol. 28, 127-135, 2005.
- [40] Skurichina, M. and Duin, R. P. W., "Bagging, Boosting and the Random Subspace Method for Linear Classifiers," *Pattern Analysis and Applications*, Vol. 5, No. 8, 121-135, 2002.
- [41] Tian, J., Zhu, L., Zhang, S. and Liu, L., "Improvement and parallelism of K-means clustering algorithm," *Tsinghua Science and Technology*, Vol. 10, No. 3, 277-281, 2005.
- [42] Tsakonas, A., Dounias, G., Doumpos, M. and Zopounidis, C., "Bankruptcy prediction with neural logic networks by means of grammar-guided genetic programming," *Expert Systems with Applications*, Vol. 30, 449-461, 2006.
- [43] West, D., Dellana, S. and Qian, J., "Neural network ensemble strategies for financial decision applications," *Computers and Operations Research*, Vol. 32, 2543-2559, 2005.
- [44] Westgaard, S. and van derWijst, N., "Default probabilities in a corporate bank portfolio: a logistic model approach," *European Journal of Operational Research*, Vol. 135, No. 2, 338-349, 2001.
- [45] Wiginton, J.C., "A note on the comparison of logit and discriminant models of consumer credit behavior," *Journal of Financial Quantitative Analysis*, Vol. 15, 757-770, 1980.
- [46] Wu, C.-H., Tzeng, G.-H., Goo, Y.-J. and Fang, W.-C., "A real-valued genetic algorithm to optimize the parameters of support vector machine for predicting bankruptcy," *Expert Systems with Applications* (in press), 2006.