

探討支持向量機器在發行人信用評等分類模式 之應用

施人英

長庚大學工商管理學系

陳文華

國立臺灣大學商學研究所

吳壽山

長庚大學管理學院

摘要

信用評等制度在金融市場已行之有年，其在企業籌資、投資人資訊取得、銀行授信參考，以及規範一般機構投資標的上，均扮演著相當重要的角色。信用評等的主要目的乃在評量債券、票券發行機構或存款機構信用品質的良窳，以利投資人做出合理的決策。過去信用評等的研究大多針對一般產業的公司債建立分類模式，較少針對發行機構本身的信用評等進行研究。早期的研究方法大多採用統計方法，近期乃有以人工智慧為基礎之各種演算法，例如類神經網路及 case-based reasoning 等。本研究嘗試應用一項新近發展且已獲得相當高分類正確率的人工智慧方法 support vector machines (SVM) 來建構發行人信用評等分類模式。為驗證 SVM 的可用性，我們將以標準普爾公司(Standard and Poor's，以下簡稱 S&P)所發佈的一般產業發行人信用評等資料樣本為例，選擇 S&P 評等時所考慮之相關重要財務變數及國家風險因素作為模式的輸入變數，並以類神經網路方法為基準與 SVM 方法進行比較，實證結果顯示 SVM 模式優於類神經網路模式。

關鍵詞：信用評等、支持向量機器、倒傳遞類神經網路



A Study of SVM Classification Models in Issuers' Credit Ratings

Jen-Ying Shih

Department of Business Administration, Chang Gung University

Wun-Hwa Chen

Graduate Institute of Business Administration, National Taiwan University

Soushan Wu

College of Management, Chang Gung University

Abstract

Credit rating systems have existed for a long time in most financial markets and played a major role in corporate capital raising, providing investment information for both individual investors and institutional investors, and credit granting in banks. The purpose of credit ratings is to measure the credit worthiness of credit securities' issuers so as to provide investors valuable information in making financial decisions. Due to the fact that the subordination of bonds has a great impact on the bond's rating (hence render the rating problem much easier to solve), most of the early researches have focused on industrial bond ratings rather than issuers' credit rating. In terms of classification approaches, early researches relied on conventional statistic methods, while recent studies tended to apply artificial intelligence based techniques, such as artificial neural networks and case-based reasoning. The main objective of this research is to propose a classification model for the issuers' credit ratings based on support vector machines, a novel classification algorithm famous for dealing with high dimension classification.

To verify the capability of the proposed model, a set of Standard and Poor's issuers' credit rating data was used as the test bed. To construct our classification models, the ten key financial variables used by Standard and Poor's (S&P), and country risk were chosen as the input variables. An artificial neural network based classification model was selected as the benchmark. Our empirical results showed the superiority of the support vector machine model over the neural artificial network model.

Key words: credit ratings, support vector machines, backpropagation neural networks

壹、導論

發行人信用評等(issuers' credit rating)乃是信用評等機構(rating agency)所授與的交易對手信用評等(counterparty credit rating)，此評等係評等機構對受評發行人償還債務整體能力之意見，主要係著重於評估發行人是否有準時履行其財務承諾之能力及意願，但並非反映任何發行人對各債務的優先償還順序或偏好¹。發行人的信用風險愈高，所獲得的信用評等等級愈低，相對地，其所發行的公司債等籌資工具所能獲得的債信評等也將受發行人信用評等的影響。因此，在金融市場中，信用評等扮演傳遞信用風險資訊的功能，資金需求者可以相對應的信用等級在資本市場籌資，有效降低籌資成本與資金缺乏的風險；資金供給者在資訊透明的投資環境中，可依本身對風險偏好的程度選擇適當的投資債權工具，此資訊尤其對受到高度管制投資標的的行業而言，更具指標性的意義，例如有些法人投資機構或基金被限定僅能投資於 BBB（或 Baa）級以上的投資工具²，使得信用評等資訊的提供更具重要性。綜上所述，信用評等的價值在於能透過簡易的符號提供投資人企業信用風險的資訊，以利其依據本身的風險偏好做出適合的投資決策（Molinero et. al. 1996）。

基本上，評等機構藉由蒐集與訪談發行人相關的各種資訊，包括發行人的國家風險、企業營運風險和財務風險資訊，經由分析師團隊的專業評估，給予發行人整體的信用風險等級。雖然授予信用評等為一系統化的分類決策過程，但是分析師團隊的專業評估技術迄今則未被系統化的說明，意即無法量化各評等要素與評等等級間的關係；除此之外，許多企業尚未被評等機構評等，但投資人仍對這些企業的信用風險資訊存在很大的需求，故許多學者嘗試運用各種統計方法(Belkai 1980; Ederington 1985; Pinches & Mingo 1975)與人工智慧的方法 (Dutta & Shekhar 1988; Surkan & Singleton 1993; Shin & Han 1999)來解決此分類問題，並獲得一些有用的決策資訊與分類效果。但前人的研究多著重於債信評等的分類，往往可運用一些特定的發行條件（例如債權的從屬性）發展分類準確率較高的模式，但幾乎鮮少對更基本的議題——發行人信用評等去發展分類模式，在實務上，市場人士往往對發行人信用評等的資訊需求更為殷切，潛在的債券發行人可預先評估本身的信用等級，以便考慮以何種方式籌資；投資人則可事先對尚未獲得信用評等的發行人風險有初步的瞭解。因此，本研究將探究發行人信用評等分類模式。此外，過去在債信評等研究議題上，已有學者提出人工智慧方法的優異性(Dutta & Shekhar 1988; Maher & Sen 1997)，並應用類神經網路(artificial neural networks, 簡稱 ANN)、基因演算和 case-based reasoning（簡稱 CBR）等方法解決債信評等的問題，但較少探究近年來新發展的分類器 support vector machines(SVM)的適用性。

¹ 以上發行人信用評等的說明係根據 S&P 公司網站資料 <http://www2.standardandpoors.com>。

² S&P 對長期信用投資工具的評等，共分為九個等級，其中屬於投資等級的債券為 AAA、AA、A 及 BBB；投機等級的債券則包括 BB、B、CCC、CC、C 與 D。Moody's 對長期信用投資工具的評等，亦分為九個等級，其中屬於投資等級的債券為 Aaa、Aa、A、Baa；投機等級的債券則包括 Ba、B、Caa、Ca、C 與 D。

SVM 多運用於型態分類問題，並獲得高度的分類正確率，應用領域包括製藥問題(Burbidge et. al. 2001)、蛋白質結構分類(Cai & Lin 2002)及生物體分類(Morris & Autret 2001)等。在財務問題上則多為運用 support vector regression (SVR)解決預測各種金融商品價格及報酬率(Tay & Cao 2001)的問題，較少分類問題(support vector classification, 簡稱 SVC)的研究，故本文將探討 SVM 在財務分類問題——信用評等的適用性，並與許多學者常用的類神經網路(ANN)比較。

以下文章將探討與整理過去信用評等研究的概況，其次，將簡介 SVM 與 ANN 方法，由於 SVM 為新方法，故我們將針對 SVM 做一詳細的介紹，第三部份將說明研究設計，第四部分將說明並討論發行人信用評等分類之研究結果。最後，總結本研究的結論，並提出未來可進一步研究的方向。

貳、信用評等相關研究

一、探究問題之類型

早期信用評等的研究多專注於長期公司債信評等分類問題，近代則有一些短期票債評等的研究，因此，在輸入變數的選取上，通常將該公司債的債權順位(subordination)、發行規模或償債期限納入(Horrigan 1966; West 1970; Pinches & Mingo 1973, 1975)，故在分類的正確率上，往往比不納入考量來得高，依據 Pinches & Mingo (1975)，Moody's A 級以上的公司債幾乎為非從屬的公司債(nonsubordinated bonds)，Ba 級以下的公司債幾乎為從屬的公司債，故此因素是債信評等的重要預測變數。但由於實務上，評等機構往往先產生發行人評等後，才針對其發行的信用投資工具進行評等，故我們希望能發展更基礎及適用範圍更廣泛的發行人信用評等分類模式。

二、評等方法

在研究方法上，傳統統計多變量分析是最常被應用的方法，包括多元線性區別分析(Belkai 1980; Ederington 1985; Pinches & Mingo 1975)、多元非線性區別分析(Pinches & Mingo 1977)、直線迴歸法(Horrigan 1966; West 1970; Ederington, 1985)、Probit regression (Ederington 1985)、Logit regression (Ederington 1985)及 multidimensional scaling (Molinero et. al. 1996)。但這些統計方法往往需滿足特定的統計假設(例如資料符合常態分配)才能適用，因此，不需針對資料組合有任何統計假設的人工智慧方法因運而起，再加上資訊科技的進步可滿足大量運算的需求，因此，學術上及實務上愈來愈多的專家學者嘗試採用人工智慧方法解決問題，目前已有學者採用倒傳遞類神經網路(Dutta & Shekhar 1988; Surkan & Singleton 1990)、genetic algorithm 與 CBR 整合方法(Shin & Han 1999)、類神經網路與 CBR 整合方法(Kim & Han 2001)及 CBR (Shin & Han 2001)解決債信評等問題，並獲得相當不錯的分類正確

率。近年來，許多研究提出 SVM 方法為一個良好的分類方法，可跳脫 ANN 可能遭遇的局部最小化(local minimum)問題，以求得全域最佳解，故被用以解決許多分類問題，因此，本研究亦嘗試運用 SVM 方法解決發行人信用評等分類問題。

三、輸入變數選取

在輸入變數的選取上，有些研究採用多變量統計方法的主成份分析、因素分析及逐步選取的方式選取重要的輸入變數(Pinches & Mingo 1975, 1977; Shin & Han 1999; Kim & Han 2001; Shin & Han 2001)，有些學者以經濟合理性主觀選取適當的輸入變數(Horrigan 1966; Belkai 1980)，這些變數往往因應資料的可取得性而多著重於企業財務資訊（尤其是財務比率資訊）和信用投資工具的發行條件（公司債的從屬性），鮮少使用營運面或市場面的資訊，故 Pinches & Mingo(1975)認為這是正確率無法大幅提昇的原因，也是信評機構主張分析師主觀判斷的價值所在。

本研究認為既然要分類標準普爾公司(Standard and Poor's，以下簡稱 S&P)所給定的評等，因應資料的可取得性，可嘗試僅利用 S&P 評等技術所著重的財務變數來作為輸入變數，一般而言，S&P 宣稱其主要考量之財務面影響因素包括槓桿(leverage)、保障(coverage)、獲利能力(profitability)及現金流量(cash flow)四個層面（S&P, 1996），以探究模式分類的效果。

四、資料期間

由於評等機構在進行評等時，並非單憑一個年度的資料來決定，而前人的研究除了少數的獲利成長性因素涉及多個年度的資料(West 1970; Pinches & Mingo 1973, 1975)，例如過去九年來的盈餘變異數，多半僅蒐集過去一年的資料去分類債信評等，抑或以過去五個年度的平均數作為輸入變數(Pinches & Mingo 1973, 1975)，這與實務分析上有很大的差別，故本研究將運用三個完整年度的歷史資料去分類信用評等。

五、分類輸出的設計

過去，在運用類神經網路方法建構公司債信用評等模式的研究中(Dutta & Shekhar 1988; Surkan & Singleton 1990)，並非像運用傳統統計方法的學者一般作多等級的分類（意即分類等級大於二級），而只是二分法，意即區分某公司債是否為某一信用評等等級，在此二分法的情況下，類神經網路的分類效果普遍皆很好，正確率最高達 88% 左右；但遺憾的是並沒有作多等級的分類。在這不同的比較基礎下，僅能聲稱類神經網路方法在區分是否為某一信用評等等級債券的議題上，其正確率高於傳統統計方法；至於在多等級的分類情形中，則有待進一步的探究。

以下我們就依此五個構面，將過去的研究彙總如表 1 所示。



表 1：信用評等研究比較表

研究 構面	問題類型	研究方法	輸入變數	資料 期間	輸出變數	最高正確率
Horrigan 1966	Moody's 及 S&P 所發 佈的債信 評等	多元迴歸模 型	債權順位、總資 產、營運資金佔銷 貨比率、淨資產佔 總負債比率、銷貨 額佔淨資產比率、 淨利佔銷貨額比率	一年	Moody's 及 S&P 的六個信 評等級	58% (Moody's) 52%(S&P)
West 1970	Moody's 所 發佈的債 信評等	多元迴歸模 型	根據 Fisher (1959) 的研究，輸入變數 為九年盈餘的變異 數、償債期間、負 債佔權益比率及流 通在外的負債	一年	六個信評 等級(Aaa, Aa, A, Baa, Ba, B)	62%
Pinches & Mingo 1973	Moody's 所發佈的 債信評等	線性多元區 別分析(以下 簡稱 MDS)	運用因素分析 (factor analysis) 從 35 個變數選擇 6 個 輸入變數，包括連 續不斷發股利的年 數、發行金額、淨 利加計利息費用/利 息費用、長期負債 佔總資產的比率、 淨利佔總資產的比 率及債權順位	一年， 再加上 過去五 年各輸 入變數 的平均 數	五個信評 等級(Aa, A, Baa, Ba, B)	71.5 %
Pinches & Mingo 1975	Moody's 所發佈的 債信評等	依據債券順 位分別利用 quadratic MDS 建立 兩個模式	運用因素分析選擇 輸入變數，包括連 續不斷發股利的年 數、發行金額、淨 利加計利息費用/利 息費用、長期負債 佔總資產的比率、 淨利佔總資產的比 率及債權順位	一年， 再加上 過去五 年各輸 入變數 的平均 數	五個信評 等級(Aa, A, Baa, Ba, B)	75.4%
Belkaoui 1980	S&P 發佈 的債信評 等	逐步多元區 別分析 (Stepwise MDS)	以經濟的合理性選 擇輸入變數，包括 總資產、總負債、 長期負債佔總投資 資本的比率(總投資 資本包括總負債、 特別股和普通股權	一年	六個信用 等級 (AAA, AA, A, BBB, BB, B)	62.8%

			益)、短期負債佔總投資資本的比率、流動比率、利息與特別股股息保障倍數、(淨利+稅後利息費用)/(稅後利息費用+特別股股息)、每股股價/普通股每股權益，以及是否為第一順位的公司債(0-1)。			
Ederington 1985	Moody's 所發佈的債信評等	直線迴歸模型(LR); ordered probit (OP); unordered logit(UL); 線性 MDA(LM); quadratic MDA(QM)	財務保障的預測值、獲利能力的時間序列預測值、獲利能力預測值的估計標準誤(estimated standard error)、(現金流量/長期債務)比率的時間序列預測值、前項比率的標準誤。	一年	六個分類等級(Aaa, Aa, A, Baa, Ba, B)	LR=65%, OP=78%, UL=73%, LM=69%, QM=72%
Dutta & Shekhar 1988	S&P 所發佈的債信評等	倒傳遞類神經網路	依據對債信評等的影響力及資料的可取得性，設計 10 個輸入變數，包括負債/(現金+固定資產)、債務比例、銷貨金額/淨資產、營業利益/銷貨金額、財務強度、銷貨淨利/固定資產、過去五年來的營收成長率、預測未來五年的營收成長率、營運資金/銷貨金額、對企業的主觀展望(subjective prospect of company)	除了採用五年的收入成長率外，其餘皆採用一年的資料	兩個等級，包括是否為某一信用評等等級(例如 AA)	83.3%
Surkan & Singleton 1990	Moody's 對從 AT&T 分割出去的貝爾電話營運公司之債信評等	倒傳遞類神經網路	依據 Peavy and Scott 的研究選擇 7 個變數，包括債務/總資本、稅前利息費用/淨利、股東權益報酬率(ROE)、過	除了 ROE 採用過去五年的變異係數外，	兩個等級，包括 Aaa and (A1,A2,A3)兩級	88%

			去五年來 ROE 的變異係數、log (總資產)、固定設備建構成本/總現金流入、長途電話收入比例	其餘則採用一年期的資料		
Molinero et. al. 1996	S&P 對西班牙銀行的短期債信評等	多元尺度分析 (MS)	選取衡量獲利能力、資本結構、財務成本和風險結構的 24 個相關財務比率。	一年	產生一張群集地圖以瞭解各銀行的分佈	未明確說明
Maher & Sen 1997	Moody's 債信評等	倒傳遞類神經網路及 Logistics 兩種方法進行比較	7 個變數，包括總資產、長期債務/總資產、淨利/總資產、發行債務的償還順位、企業普通股 β 值、退休金負債淨額、來自於停業部門及非常項目的淨利	除了總資產、長期債務 / 總資產、淨利 / 總資產以五年平均數計算外，其餘皆為一年期資料	六個分類等級(Aaa, Aa, A, Baa, Ba, B)	70%
Shin & Han 1999	韓國的商業本票評等	用基因演算法 (GA) 找 weight vector 以作為 CBR 方法的屬性。	從 168 個財務比率中進行因素分析和 ANOVA 檢定，篩選出 27 個財務比率，再以逐步(stepwise)選取方法篩選出 12 個財務比率。	一年	5 個分類等級	各等級的加權平均正確率為 75.5%
Kim & Han 2001	韓國短期債券評等	整合運用類神經網路之 SOM 及 LVQ 的 CBR 方法	針對 129 個變數(包括 4 個分類變數及 125 個財務比率)進行因素分析篩選 26 個財務比率，再以逐步(stepwise)選取方法篩選出 13 個財務比率。	一年	5 個分類等級	各等級的加權平均正確率為 69.1%
Shin & Han 2001	韓國的商業本票評等	使用 inductive indexing 之 CBR	運用因素分析、ANOVA 檢定及 Kruskal-Wallis 檢定篩選出 27 個變數(包	一年	5 個分類等級	各等級的加權平均正確率為 70.0%

			括 23 個量化變數及 4 個質化變數，再以逐步 (stepwise) 選取方法篩選出 12 個財務變數。			
--	--	--	---	--	--	--

參、分類模式

本節將介紹本研究所採用的兩種方法--SVM 方法及 ANN 方法。由於 SVM 方法在管理領域上的應用情形尚屬新概念，故以下我們將從 SVM 發展的理論基礎詳細介紹，若讀者已熟悉該方法，則可跳過本節，將不會對閱讀整篇文章造成影響。

一、Support vector machines

(一) SVM 觀念

SVM 為一奠基於統計學習理論的學習機器(learning machine)，其基本的運作概念為將輸入向量以線性或非線性的核心函式(kernel function)映射到一個高維度的特徵空間(feature space)，在特徵空間中找到最適的超平面(hyperplane)以區別各個分類。如此一來，原本在低維度中不能用線性求解的問題，就能在高維度中進行分類；這個由高維度所組成的特徵空間，甚至可以是無限維度的，因為在計算上，權重(weights)這個參數是不需被計算出來的。藉由選擇適當的核心函數，非線性的映射可以使決策函數在這個新的特徵空間將問題求解。此性質讓 Vapnik 可以運用最小化結構性風險 (structural risk minimization, 簡稱 SRM) 在非線性問題上，並且還是用一樣的最適化技巧。而 SVM 所決定的決策函數是由一群特殊的向量所組成的，而這群向量是由訓練的資料中挑選出來的，稱為支持向量(support vectors)，也因此整個演算法稱為 support vector machines (Vapnik 1995)。

以圖 1 為例，資料主要可分為兩群“o”和“+”。從幾何上來解釋，SVM 就是要找出一個最適切割超平面（或一條決策函數）將這兩群資料分開，如圖中的實線。而這條決策函數有許多很好的統計性質。經由 SVM 計算後，圖中圈起來的點就是支持向量，有四個在虛線上的點主要是用來決定決策函數；另一個不在虛線上的點是違反分割限制的點，因為無法被分類，所以也納入支持向量中，此為考慮成本函數的 soft margin 分類器的代表例子。



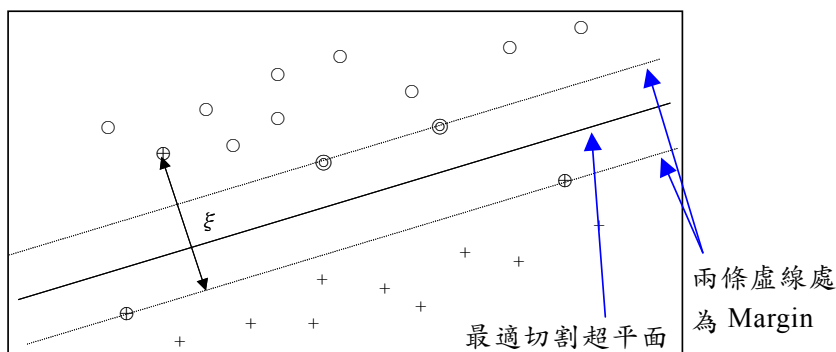


圖 1：運用線性函數發展 SVM 分類器的範例

SVM 由於在最適化時不會遭遇局部最佳化(local optimum)的問題，因此，有不少的實證研究中指出，SVM 的表現比 ANN 還好(Morris & Autret 2001; Tay & Cao 2001; Cai & Lin 2002)。

(二) Structural Risk Minimization

ANN 存在難以一般化的問題，在訓練學習時，往往產生過度配適(over fitting)的問題，主要原因在於 ANN 係採用 empirical risk minimization (ERM)原則產生最佳的分類模式，而 SVM 則是運用 structural risk minimization (SRM)原則發展最佳模式，並已被 Gunn(1998)證實優於 ANN 所採用之 ERM 原則。簡而言之，SRM 原則是最小化期望風險的上界(upper bound)，而非 ERM 所採用最小化訓練組資料樣本的誤差。此差異使得 SVM 擁有較佳的一般化能力，而此能力正是統計學習理論的目標 (Vapnik 1995)。

(三) Support Vector Classification

以下將描述 SVC 如何藉由現有範例所產生的超平面以作為最適分割超平面(optimal separating hyperplane)分類器，其基本精神為找出一組最大化 margin(即最大化超平面與每個分類的最近資料點之間的距離)。以下分三種型態的例子說明。

1、線性可分割範例

假設訓練範例資料為 $(x_1, y_1), \dots, (x_l, y_l), x \in R^n, y \in \{+1, -1\}$ ，其中 x 為輸入向量，該組資料可被一個超平面分割為兩類，一類為 $+1$ ，另一類為 -1 ，若該組資料可被正確無誤地分割，而且每個分類的最近向量(nearest vector)與該超平面間的距離最大化，則我們可說該資料可被一組超平面最適分割。我們可以下列形式表示該超平面：

$$y_i[\langle w, x \rangle + b] \geq 1, \quad i = 1, \dots, l \quad (1)$$

所找到兩個最近該超平面的向量與該平面間的距離 margin 如下：

$$\begin{aligned}
\rho(w, b) &= \min_{x_i, y_i = -1} d(w, b; x_i) + \min_{x_i, y_i = 1} d(w, b; x_i) \\
&= \min_{x_i, y_i = -1} \frac{|\langle w, x_i \rangle + b|}{\|w\|} + \min_{x_i, y_i = 1} \frac{|\langle w, x_i \rangle + b|}{\|w\|} \\
&= \frac{2}{\|w\|}
\end{aligned} \tag{2}$$

最大化 (2) 可表示為達成最小化 $\Phi(w) = \frac{1}{2}\|w\|^2$ ，採用拉式鬆弛法 (Lagrange relaxation) 將該問題表達為：

$$\min_{w, b} \Phi(w, b, \alpha) = \frac{1}{2}\|w\|^2 - \sum_{i=1}^l \alpha_i (y_i [\langle w, x_i \rangle + b] - 1) \tag{3}$$

為利求解，可將該原生問題 (primal problem) 轉化為對偶問題 (dual problem)：

$$\max_{\alpha} W(\alpha) = \max_{\alpha} \left(\min_{w, b} \Phi(w, b, \alpha) \right) \tag{4}$$

根據 (3) 對 b 及 w 偏微分可將 (4) 表達為：

$$\max_{\alpha} W(\alpha) = \max_{\alpha} -\frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle + \sum_{k=1}^l \alpha_k \tag{5}$$

可求解下式 (6)：

$$\begin{aligned}
\alpha^* &= \arg \min_{\alpha} \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle - \sum_{k=1}^l \alpha_k \\
s.t. \quad &\alpha_i \geq 0 \quad i = 1, \dots, l \quad \sum_{j=1}^l \alpha_j y_j = 0
\end{aligned} \tag{6}$$

求得最適分割超平面為：

$$w^* = \sum_{i=1}^l \alpha_i y_i x_i \quad b^* = -\frac{1}{2} \langle w^*, x_r + x_s \rangle \tag{7}$$

其中， x_r, x_s 是每個分類可滿足 $\alpha_r, \alpha_s > 0$ ， $y_r = -1, y_s = 1$ 的任何支持向量。我們可得到 hard classifier 如下：

$$f(x) = \text{sgn}(\langle w^*, x \rangle + b) \tag{8}$$

若考量無法完全被分類的情況時，其 soft classifier 為下式：

$$f(x) = h(\langle w^*, x \rangle + b) \quad \text{where } h(z) = \begin{cases} -1 & : z < -1 \\ z & : -1 \leq z \leq 1 \\ +1 & : z > 1 \end{cases} \tag{9}$$

根據 Kuhn-Tucker condition：

$$\alpha_i (y_i [\langle w, x_i \rangle + b] - 1) = 0, \quad i = 1, \dots, l \tag{10}$$

因此只有可滿足 $y_i [\langle w, x_i \rangle + b] = 1$ 的 x_i 點有非零的拉式乘數 (Lagrange multipliers)，我們稱這些點為支持向量 (簡稱 SV)，若資料可被線性分割，則所有的 SV 點將落在

margin 上，也因此可求得很少量的 SV。結果，超平面將由一組少量的訓練資料點決定，其他點可從訓練組中移除，即使重新計算超平面，依然可得到同樣的答案。故 SVM 可用以摘要隱含在 SV 所產生的資料組合中的資訊。

2、線性不可分割範例—Soft margin technique

除了運用尋找上界的方法以找出可正確分割資料群的超平面外，有時真實世界因為無法完全正確將資料分割，故 Vapnik (1995) 又另外引進與錯誤分類有關的額外成本函數(cost function)的觀念，以求得此分界，進而達到一般化的效果。前述的最適分割超平面可表達為：

$$\min \Phi(w, \xi) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^l \xi_i \quad (11)$$

s.t.

$$y_i [\langle w, x_i \rangle + b] \geq 1 - \xi_i, \quad i = 1, \dots, l.$$

$$\text{where } \xi_i \geq 0$$

其中 ξ_i 是錯誤分類的誤差衡量， C 是一個給定的參數值，根據 Minoux (1986) 可得到拉式鬆弛法如下式：

$$\Phi(w, \xi) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^l \xi_i - \sum_{i=1}^l \alpha_i (y_i [w^T x_i + b] - 1 + \xi_i) - \sum_{i=1}^l \beta_i \xi_i \quad (12)$$

其中 α, β 是拉式乘數。如前所述，可找到其對偶問題如下：

$$\max_{\alpha} W(\alpha) = \max_{\alpha} -\frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle + \sum_{k=1}^l \alpha_k \quad (13)$$

求解式(14)：

$$\alpha^* = \arg \min_{\alpha} \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle - \sum_{k=1}^l \alpha_k \quad (14)$$

s.t.

$$0 \leq \alpha_i \leq C \quad i = 1, \dots, l$$

$$\sum_{i=1}^l \alpha_i y_i = 0.$$

其中，係數 C 必須被決定，這個參數為該分類器的另外一個 capacity control。

3、高維度特徵空間的一般化—替代核心函式

當線性的分界(boundary)不適當時，SVM 可以將輸入向量(x)映射到高維度的特徵空間(z)，藉由選擇一個非線性的映射，SVM 可以在特徵空間中建構一個最適的分割超平面。此映射包括常用的 polynomial $K(x, x') = (\langle x, x' \rangle + 1)^d$ 和 Gaussian radial basis functions (RBF) $K(x, x') = \exp(-\gamma \|x - x'\|^2)$ 。因此，(14)式可表達為：

$$\alpha^* = \arg \min_{\alpha} \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j K(x_i, x_j) - \sum_{k=1}^l \alpha_k \quad (15)$$

其中 $K(x_i, x_j)$ 為核心函式，用以執行非線性映射到特徵空間的任務。

二、類神經網路

類神經網路模擬生物神經網路執行分散式運算功能，是一個由許多簡單但以高度複雜的方式互連的處理單元(processing unit)所構成的網路，介於處理單元間的訊號傳遞路徑稱為連結，並且不同的拓樸(topology)與演算法(algorithm)可以組成各種網路模式(Lippmann 1987)，例如多層式倒傳遞網路(backpropagation)、Hopfield Networks、Self-Organizing Maps (簡稱 SOM) Networks 等。類神經網路已經成功地運用在許多領域，例如解決行銷、零售、銀行、財務、保險及電信等問題(Smith & Gupta 2000)。

倒傳遞演算法是類神經網路演算法中適合用來解決預測與分類問題的演算法。其主要網路架構可分為輸入層、隱藏層及輸出層，其中隱藏層可分為零層或若干層，每一層均由數個處理單元排列組成，每一層的輸入資料為前一層的輸出資料，各層間的連結具有加權值(weight)，藉由加權值的強弱控制前一層輸入資料的影響程度。它的基本原理是利用「坡降法 (gradient descent)」的觀念，將表達網路實際輸出與目標輸出之差異的誤差函數最小化（即 ERM），並透過連結加權值的不斷調整，來達成網路的訓練。當輸入每一筆學習範例資料時，在輸出層會得到預測的輸出值，此時藉由比較目標輸出值和預測值之間的差異，可得一誤差函數。接著將誤差函數予以微分求其最小化，網路則利用微分產生的結果調整層與層之間的權值，不斷地改變而達到學習的效果。這個由輸出的誤差結果向後傳遞至隱藏層和輸入層以調整權值的過程，就是它之所以被稱為倒傳遞的原因。訓練多層網路架構的能力是建立一個智慧型應用程式的重要步驟，而選擇適當的類神經網路參數則是獲得較佳預測能力的重要因素。詳細演算法可參考(Lippmann 1987)。

肆、研究設計

一、研究樣本

本研究係根據 S&P 公司出版的刊物—*Global Sector Review* 和 CARD 光碟片上所記載的信用評等和財務變數資料，選出在 1996 年接受 S&P 評等的美國和紐澳地區的企業為研究樣本。這些樣本所處產業包括消費性產品業、高科技產業、零售業、化工業、建築材料業、汽車業、機器設備業和鋼鐵業，樣本總數共計 429 個。在訓練組及測試組的樣本分佈上（參見表 3），係依據所有資料在各等級的分佈情形分配，本研究共蒐集 429 個受評企業樣本，大約依照 75% 及 25% 的比例分層（信用評等等級）抽取

為 325 個訓練組與 104 個測試組樣本。為研究模式的信度與效度考量，本研究另以不重複抽樣及 0.75 Bootstrapping 兩種方式(Witten & Frank 2000)各抽出 100 組訓練樣本與測試樣本，以進行 SVM 與 ANN 兩種方法的比較。

表 3：樣本分佈

信用評等等級	樣本分佈	百分比	訓練組	測試組
AA 級以上	41	9.6%	32	9
A	100	23.3%	75	25
BBB	93	21.7%	71	22
BB	111	25.9%	84	27
B 級以下	84	19.6%	63	21
	429	100.0%	325	104

本研究在自變數的選取上，採用 S&P's 在評估財務風險時考量的主要因素(表 4)，包括獲利能力、利息保障、資本結構、現金週轉能力(現金流量面的短期償債能力)四個層面的十個輸入變數，以及一個考慮國家風險因素(以 0 和 1 區分 AAA 級和 AA 級國家風險)的輸入變數。至於營運風險，由於資料搜集上的困難度較高且主觀成份過重，故暫不予考慮，但這有可能會對模型的效果產生不利的影響。在輸入變數資料蒐集期間方面，本研究取評等年度往前推三個年度的財務資料(1993、1994 及 1995 年)，故共計有 31 個自變數(10 個財務變數*3 年+1 個國家風險變數)，但由於這些變數的值域並非分佈一致，若未經資料轉化或正規化處理，將導致值域較小者的變數之重要性無法顯現，而由值域較大者控制整個網路的學習過程，進而左右了學習成果。故有必要針對輸入變數值作轉化或正規化的處理，本研究將以最大最小對映法將資料值域轉換為介於-1 與 1 之間的值。

在輸出變數的設計上，由於 AAA 級樣本數較少，故與 AA 級合併視為同一組，CCC 級以下的樣本數也較少，故與 B 級合併為同一分類。所以模式輸出變數共分為五個等級，即從 AA 級以上、A、BBB、BB 到 B 級以下。

二、SVM

本研究運用以 C++ 所開發的 LIBSVM³軟體程式，核心函式為 RBF，核心函式的 γ 及成本函式的 C 為待決定的參數，我們依據 Hsu et. al. (2003)所提出的"simple grid search"方法，並採用 ten-fold cross validation 去尋找適合的 γ 及 C 參數值，請參考圖 2，找到的兩個參數值分別為 $\gamma=0.5$ 及 $C=8$ ，以得到 SVM 分類模式。

³ <http://www.csie.ntu.edu.tw/~cjlin/libsvm/>



表 4：輸入變數定義

變數	意義
X_1	國家風險變數：由於美國、澳洲的國家風險等級平均為 AAA，紐西蘭等國家的評等是 AA，故設定一虛擬變數以為區別，AAA 級國家變數值為 1，AA 級國家變數值為 0。
$X_{2i}, i=1,2,3$ 年	持續營業稅前息前之利息保障 = $\frac{\text{持續營業稅前純益} + \text{利息費用}}{\text{利息費用} + \text{資本化利息}}$
$X_{3i}, i=1,2,3$ 年	持續營業稅前息前折舊與折耗前之利息保障 = $\frac{\text{持續營業稅前息前折舊與折耗前純益}}{\text{利息費用與當期資本化利息}}$
$X_{4i}, i=1,2,3$ 年	永久性資本稅前報酬率 = $\frac{\text{持續營業部門稅前純益} + \text{利息費用}}{\text{平均一年內到期之長期負債} + \text{長期債務} + \text{非流動遞延所得稅} + \text{股東權益} + \text{平均短期借貸}} * 100\%$
$X_{5i}, i=1,2,3$ 年	營業利益佔銷貨百分比 = $\frac{\text{營運利益}}{\text{銷貨}} * 100\%$ 營業利益 = 銷貨淨額 - 銷貨成本 (在提列折舊與折耗前) - 銷管費用 - 研發成本
$X_{6i}, i=1,2,3$ 年	營業產生的現金流量對總債務的比例 = $\frac{\text{營運產生的現金流量}}{\text{總債務}} * 100\%$ 營業產生的現金流量 = 繼續營業部門稅後淨利 + 折舊與折耗 + 遞延所得稅 + 其他非現金的項目
$X_{7i}, i=1,2,3$ 年	$\text{自由現金流量對總債務的比例} = \frac{\text{營運產生的現金流量} - \text{資本支出} - (+)\text{營運資金的增額(減少)}}{\text{總債務}}$ 註：營運資金的增額(減少)並不包括現金和短期投資的變動
$X_{8i}, i=1,2,3$ 年	總債務佔資本的比例 = $\frac{\text{總債務}}{\text{總債務} + \text{股東權益}}$
$X_{9i}, i=1,2,3$ 年	銷貨淨額
$X_{10i}, i=1,2,3$ 年	股東權益
$X_{11i}, I=1,2,3$ 年	總資產

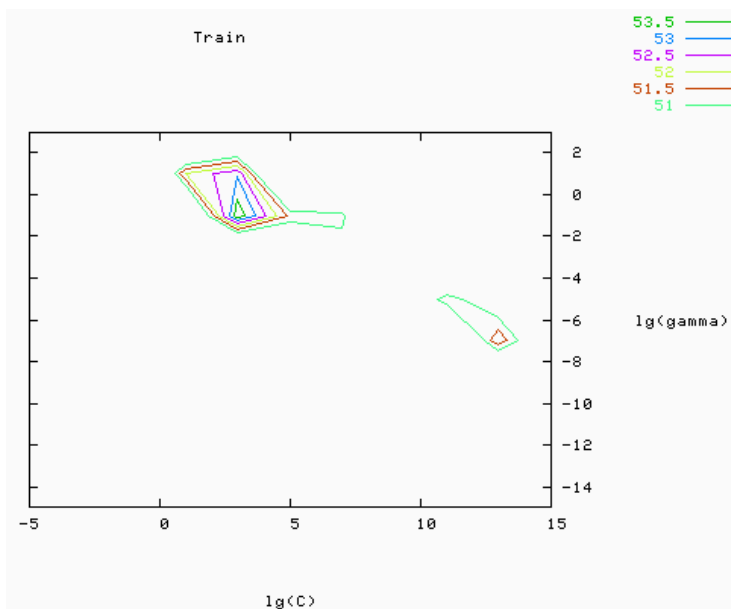


圖 2： γ 及 C 參數值等高線圖

(註：lg 為以 2 為底的對數，右上角的等高線數值為 validation set 的正確率)

三、ANN

本研究以 Matlab 6.1 軟體去設計 ANN 模式，採用倒傳遞類神經網路模式，運用 Levenberg-Marquardt algorithm，輸入層由 31 個輸入單元所組成；在隱藏層的設計上，處理分類問題時，隱藏層的處理單元數宜少於輸入層的輸入單元個數(Diamantaras & Kung 2000)，然最佳的隱藏層單元數仍未有一般性的定論，故本研究測試了 1~31 個隱藏層處理單元數，採用 hyperbolic tangent sigmoid (即 tansig function in Matlab) 轉換函數，輸出層包含 5 個處理單元，分別對應五類信用評等。在學習率與訓練次數方面，本研究測試了學習率 $\in \{0.001, 0.005, 0.01, 0.05, 0.1\}$ ，以及訓練次數 $\in \{10, 100, 1000, 1500, 2000\}$ ，有關各個組合的訓練組正確率數據，請參考表 5。為了避免模式過度配適的問題，我們從訓練組中分出一組驗證組(validation)，並採用 early stopping 的策略去克服該問題。

經過 $5 \times 5 \times 31$ 個不同的 ANN 模式設計後，能夠獲得最高的訓練組正確率的學習率為 0.005，訓練次數為 10 次，隱藏層單元數目為 22 個單元數，故選擇以 31-22-5 的 ANN 架構模式。

表 5：各種 ANN 模式的訓練組正確率

隱藏層 單元數	epoch=10 lr=0.001 lr=0.005 lr=0.01 lr=0.05 lr=0.1	epoch=100 lr=0.001 lr=0.005 lr=0.01 lr=0.05 lr=0.1	epoch=1000 lr=0.001 lr=0.005 lr=0.01 lr=0.05 lr=0.1	epoch=1500 lr=0.001 lr=0.005 lr=0.01 lr=0.05 lr=0.1	epoch=2000 lr=0.001 lr=0.005 lr=0.01 lr=0.05 lr=0.1
1	31.90% 38.96% 25.77% 42.02% 43.87%	27.30% 43.87% 48.16% 36.50% 19.94%	10.12% 48.77% 40.48% 43.25% 23.01%	29.14% 19.63% 44.17% 46.93% 19.63%	25.46% 28.53% 22.09% 25.77% 45.09%
2	53.68% 60.12% 53.07% 60.43% 40.18%	52.45% 40.18% 25.77% 51.84% 52.15%	23.31% 25.77% 63.80% 63.19%	21.47% 27.30% 40.49% 32.52% 38.65%	25.77% 43.56% 55.83% 36.50% 49.69%
3	29.45% 61.35% 61.66% 20.25% 43.56%	60.12% 45.71% 62.27% 43.25% 23.31%	24.23% 24.85% 63.19% 25.77% 35.28%	23.01% 22.39% 34.05% 21.78% 36.50%	47.85% 61.66% 62.58% 9.82% 26.07%
4	47.55% 60.43% 38.28% 57.98% 68.10%	44.79% 38.96% 43.25% 30.61% 68.40%	61.66% 62.58% 50.92% 25.77% 64.72%	31.29% 39.57% 34.05% 37.42% 68.71%	61.35% 58.28% 62.27% 46.32% 50.31%
5	61.66% 41.41% 9.82% 65.95% 62.58%	64.11% 30.37% 62.27% 46.01% 59.82%	16.26% 45.71% 63.80% 65.95% 28.22%	37.73% 62.58% 58.28% 19.94% 52.15%	60.12% 64.72% 29.75% 68.10% 51.23%
6	64.72% 25.77% 57.98% 58.90% 36.81%	58.28% 44.17% 67.79% 69.02% 53.99%	7.36% 64.72% 18.71% 15.03% 40.18%	65.64% 68.71% 38.04% 49.69% 39.57%	61.66% 51.23% 51.53% 66.87% 57.98%
7	24.85% 66.26% 18.10% 55.83% 28.53%	53.68% 20.25% 58.90% 51.53% 56.44%	41.41% 55.52% 67.48% 55.21% 59.51%	57.98% 66.87% 65.03% 49.39% 66.26%	55.83% 52.45% 25.15% 27.61% 65.95%
8	59.82% 63.50% 67.79% 57.06% 60.12%	58.28% 18.71% 27.91% 59.51% 57.06%	61.35% 39.88% 69.33% 64.72% 22.70%	65.95% 58.28% 30.98% 56.44% 67.48%	57.06% 63.50% 63.64% 32.52% 65.95%
9	39.57% 49.08% 67.79% 69.63% 60.74%	27.91% 58.90% 42.02% 51.23% 69.02%	59.82% 65.95% 30.08% 60.12% 63.19%	40.18% 28.53% 30.98% 68.71% 71.47%	54.91% 65.95% 69.33% 73.31% 50.00%
10	33.44% 44.17% 69.02% 71.47% 52.76%	24.54% 74.23% 49.39% 68.71% 48.77%	60.43% 28.83% 65.03% 27.30% 23.31%	29.14% 52.76% 49.08% 44.79% 74.54%	19.94% 66.56% 20.55% 62.58% 53.68%
11	61.35% 70.25% 46.01% 69.02% 71.47%	64.72% 63.50% 63.80% 71.78% 56.75%	52.76% 61.66% 64.42% 52.15% 79.45%	34.66% 26.38% 50.61% 50.31% 64.72%	57.36% 26.07% 67.79% 73.01% 22.39%
12	52.76% 67.48% 66.87% 59.51% 53.37%	58.28% 69.33% 70.25% 62.88% 23.31%	53.07% 67.18% 62.58% 67.18% 72.39%	57.36% 74.54% 23.62% 68.71% 31.29%	21.78% 68.40% 56.13% 58.28% 71.78%
13	65.03% 35.58% 50.31% 67.18% 9.51%	45.09% 61.96% 76.07% 19.02% 55.52%	64.72% 73.93% 55.83% 46.93% 50.00%	46.63% 64.11% 71.78% 36.50% 68.40%	62.58% 73.62% 17.48% 73.93% 62.27%
14	56.75% 42.33% 52.45% 70.25% 68.40%	49.69% 52.15% 57.06% 38.96% 17.48%	45.09% 57.98% 65.64% 67.18% 45.71%	41.41% 65.95% 72.70% 68.71% 23.31%	18.10% 59.51% 39.26% 59.20% 52.76%
15	27.61% 65.03% 66.26% 20.25% 69.33%	24.23% 75.46% 68.71% 20.55% 61.96%	23.31% 68.71% 64.11% 73.01% 69.94%	28.53% 74.85% 71.78% 50.00% 61.96%	37.73% 69.94% 30.06% 59.20% 38.96%
16	61.35% 56.13% 32.52% 58.90% 53.99%	66.56% 61.96% 22.70% 44.48% 67.79%	48.47% 32.21% 66.87% 74.54% 64.72%	53.68% 66.26% 73.62% 26.69% 74.85%	24.54% 72.09% 25.46% 72.09% 69.63%
17	51.84% 69.94% 73.01% 65.64% 70.55%	61.04% 19.94% 61.96% 20.25% 58.90%	25.46% 69.02% 66.56% 25.77% 23.93%	55.52% 68.10% 59.20% 63.80% 66.26%	48.77% 71.47% 68.10% 73.93% 23.01%
18	65.95% 70.55% 39.57% 64.11% 45.40%	25.46% 66.87% 60.43% 27.30% 26.07%	10.43% 51.84% 62.27% 69.94% 66.26%	21.78% 40.80% 62.58% 24.85% 69.02%	51.53% 64.72% 21.47% 44.17% 56.75%
19	46.32% 63.19% 68.71% 60.74% 69.94%	67.18% 66.56% 69.63% 45.71% 71.17%	45.09% 25.77% 55.52% 56.44% 69.33%	55.83% 34.05% 67.79% 73.01% 64.72%	60.12% 72.70% 60.12% 69.02% 71.47%
20	57.98% 65.34% 72.39% 74.54% 65.64%	46.32% 76.99% 58.28% 69.33% 55.83%	46.63% 48.16% 40.18% 69.02% 74.23%	35.58% 51.84% 44.79% 52.45% 65.95%	55.21% 58.59% 55.83% 60.43% 60.74%
21	57.98% 20.25% 64.42% 34.97% 74.85%	56.44% 52.15% 64.42% 20.55% 40.18%	62.58% 74.85% 71.78% 69.33% 57.36%	21.47% 75.15% 76.69% 24.23% 71.17%	60.74% 57.67% 51.53% 58.59% 73.31%
22	51.84% 80.06% 61.35% 31.90% 71.17%	58.90% 76.38% 69.63% 41.41% 64.11%	30.37% 69.02% 70.25% 73.62% 49.69%	63.19% 25.77% 26.69% 19.02% 26.99%	52.76% 73.01% 71.17% 67.48% 47.55%
23	56.44% 62.88% 59.82% 40.49% 56.44%	51.84% 18.10% 60.43% 73.93% 66.26%	57.98% 61.96% 69.02% 71.47% 66.26%	53.68% 46.63% 63.19% 67.18% 58.59%	58.90% 36.50% 54.29% 70.25% 33.44%
24	49.39% 45.40% 43.87% 73.31% 44.79%	17.48% 28.22% 67.18% 75.77% 71.78%	25.77% 64.11% 70.55% 77.30% 73.31%	32.82% 45.09% 66.87% 71.78% 10.43%	64.11% 56.13% 74.23% 62.58% 25.77%
25	63.50% 26.07% 60.43% 68.10% 68.10%	46.01% 66.87% 73.01% 69.63% 69.63%	65.34% 65.64% 37.73% 49.08% 45.40%	64.72% 30.37% 15.34% 79.75% 35.89%	52.76% 69.33% 30.06% 72.09% 59.82%
26	52.45% 35.58% 67.18% 37.73% 25.77%	53.99% 70.55% 46.01% 52.45% 54.91%	53.99% 25.77% 71.47% 56.13%	21.78% 57.67% 70.55% 58.59% 67.79%	37.12% 69.63% 46.01% 63.80% 68.10%

表 5：各種 ANN 模式的訓練組正確率

隱藏層 單元數	epoch=10			epoch=100			epoch=1000			epoch=1500			epoch=2000		
	lr=0.001	lr=0.005	lr=0.01	lr=0.05	lr=0.1	lr=0.001	lr=0.005	lr=0.01	lr=0.05	lr=0.1	lr=0.001	lr=0.005	lr=0.01	lr=0.05	lr=0.1
27	48.77%	28.22%	68.10%	23.31%	62.27%	57.06%	76.38%	53.99%	58.28%	73.93%	48.47%	71.78%	45.09%	62.27%	66.26%
28	68.10%	59.82%	76.69%	70.86%	66.26%	27.61%	66.26%	62.58%	70.55%	65.64%	55.52%	39.88%	71.47%	73.01%	59.51%
29	51.23%	19.94%	72.70%	61.66%	70.55%	23.01%	70.55%	47.85%	48.16%	31.29%	57.06%	46.93%	71.78%	71.47%	66.87%
30	53.99%	66.26%	60.74%	46.93%	68.40%	23.31%	59.82%	63.80%	75.46%	70.86%	56.44%	22.39%	76.38%	73.93%	72.09%
31	56.13%	22.09%	70.86%	67.18%	69.94%	49.39%	66.87%	32.82%	25.77%	69.33%	54.29%	19.63%	69.63%	63.80%	68.10%

註：epoch 為訓練次數，lr 為學習率



伍、結果與討論

一、SVM 模式分類結果

SVM 模式訓練組與測試組的分類結果混亂矩陣(confusion matrix)如表 6 及表 7 所示，其整體分類正確率分別為 70.77%及 60.58%。雖然正確率與 100%仍有一段差距，但若以隨機分類五個信用評等等級所計算的正確率來比較，仍是高於 20%。在協助分類的目的上，仍具有分類模式的資訊價值。

以個別等級的分類正確率觀察結果（訓練組請參看表 6，測試組請參看表 7），在訓練組中，A 級的正確率最高(80%)，其次為 BB 級(75%)，第三為 B 級以下(68.25%)，第四為 BBB 級(63.38%)，最低為 AA 級以上(59.38%)；而在測試組中，則是 BBB 級的正確率最高(72.73%)，其次才是 A 級(68%)，第三為 BB 級(62.96%)，第四為 B 級以下(57.14%)，最低為 AA 級以上(11.11%)。資料結果顯示以下現象：

1. A 級信用評等雖然在訓練組中分類正確率最高，在測試組中卻非最高，但仍是正確率次高的信用評等等級。根據測試組的混亂矩陣，發現分類錯誤的原因在於分類模式把錯誤分類的 8 個測試樣本低估一級（誤判為 BBB 級）。儘管如此，該等級相對上仍是本研究中 SVM 模式分類正確較高的信用評等等級。
2. 測試組中分類正確率最高的 BBB 級，在訓練組中分類正確率並不高，推其原因，可能在於訓練組中已學習該等級的一般特性，故即使在測試組中，分類正確率不降反升。同時，其錯誤分類的情形亦多發生於低估一級的情況（誤判為 BB 級）。
3. AA 級以上的分類正確率在訓練組及測試組中皆是最底的，而且其錯誤分類的樣本多發生於低估一級的情況（意即 A 級）。BB 級和 B 級以下，則發生的錯誤分類多為鄰近等級（意即低估或高估一個等級）。兩個極端等級「AA 級以上」及「B 級以下」的分類正確率在測試組中是較低的，推測其原因可能為這兩個極端等級的樣本數較少，分類模式尚未能充分學習這兩個等級的特性。

表 6：SVM 分類結果混亂矩陣——訓練樣本

Predicted Target	AA 級 以上	A	BBB	BB	B 級以下	正確率	低估一級 誤差正確率	上下一級 正確率
AA 級以上	19	11	2	0	0	59.38%	93.75%	93.75%
A	0	60	13	2	0	80.00%	97.33%	97.33%
BBB	1	11	45	13	1	63.38%	81.69%	97.18%
BB	1	5	5	63	10	75.00%	86.90%	92.86%
B 級以下	0	2	1	17	43	68.25%	68.25%	95.24%
整體模式						70.77%	85.23%	95.38%

表 7：SVM 分類結果混亂矩陣—測試樣本

Predicted Target	AA 級 以上	A	BBB	BB	B 級 以下	正確率	低估一級誤 差正確率	上下一級 正確率
AA 級以上	1	8	0	0	0	11.11%	100.00%	100.00%
A	0	17	8	0	0	68.00%	100.00%	100.00%
BBB	0	1	16	5	0	72.73%	95.45%	100.00%
BB	0	0	5	17	5	62.96%	81.48%	100.00%
B 級以下	0	0	1	8	12	57.14%	57.14%	95.24%
整體模式						60.58%	85.58%	99.04%

根據模式結果可得到以下對發展信用評等分類模式的決策意涵：

在本研究 SVM 分類模式中，BBB 級以上的分類等級誤判情況大部分發生於低估一級，推究其原因，我們知道根據許多國家的資本市場管制措施規定，某些法人投資機構僅能投資於投資等級以上（意即 BBB 級以上）的投資標的，基本上，這些發行 BBB 級以上投資工具的發行人之信用風險相對較低，信評機構在評等對這些企業時，除了財務面的因素外，營運管理面的因素可能是決定其等級的關鍵考量，而本研究無法蒐集營運管理面的資料可能是造成低估一級的原因。

雖然，本研究的分類正確率僅有 60.58%，若以輔助自動化信用評等決策的觀點來看，仍具有決策資訊價值，與隨機決策的基準(20%)比較，仍高出 40.58%。以錯誤分類所可能造成的損失成本來分析，SVM 分類模式的誤判多發生於鄰近等級的情況，尤其是低估一級信用評等等級的情形最普遍，故若我們以容忍低估一級的正確率來觀察，則模式的正確率為 85.58%，若以容忍上下一級誤判來看，則模式的正確率為 99.04%。實務上，上下一級的誤判對企業發行債券雖有影響，但並非嚴重的錯誤，因此，本研究仍可對資本市場提供初步的信用評等資訊，那些尚未接受信用評等的企業可先行瞭解本身的信用風險等級落點，投資人可進一步瞭解缺乏信用評等的投資工具所對應的信用風險。

二、與 ANN 比較討論

為瞭解 SVM 的分類效果，本研究另以 ANN 模式作為基準比較，由於 ANN 模式有可能無法求得全域最佳解，遭遇局部最佳化的情形，以及考慮抽樣樣本的差異性，故我們各以不重複抽樣與 0.75 Bootstrapping 重複抽樣兩種方式，各形成 100 組樣本，用以比較 SVM 與 ANN 模式的效果，每次的正確率如表 8 所示，用以比較 SVM 模式的正確率是否顯著高於 ANN 模式，並採用 paired t-test 檢定(Witten & Frank, 2000) 以瞭解本研究中兩個模式正確率的優劣。

H0: SVM 測試組分類正確率-ANN 測試組分類正確率 ≤ 0

H1: SVM 測試組分類正確率-ANN 測試組分類正確率 > 0

100 次不重複抽樣的實驗得到檢定值 $t = 10.863 > t(0.01; 99) = 2.3646$ ，100 次 0.75 Bootstrapping 的實驗得到檢定值 $t = 13.5318 > t(0.01; 99)$ ，皆為拒絕 H_0 ，檢定結果為 SVM 測試組的正確率顯著高於 ANN 測試組的正確率。

表 8：兩種方法測試組正確率實驗 100 次表

	不重複抽樣		0.75 Bootstrapping			不重複抽樣		0.75 Bootstrapping			不重複抽樣		0.75 Bootstrapping	
	SVM	ANN	SVM	ANN		SVM	ANN	SVM	ANN		SVM	ANN	SVM	ANN
1	60.58%	55.34%	52.88%	33.33%	34	48.08%	41.35%	50%	44.34%	67	56.73%	6.73%	44.23%	41.90%
2	52.88%	44.23%	55.77%	37.38%	35	50%	23.08%	56.73%	49.57%	68	44.23%	44.23%	45.19%	48.08%
3	55.77%	50.00%	58.65%	29.73%	36	56.73%	45.19%	60.58%	45.19%	69	45.19%	39.42%	49.04%	44.66%
4	58.65%	23.08%	46.15%	32.63%	37	60.58%	38.46%	49.04%	50.88%	70	49.04%	25.00%	55.77%	46.59%
5	46.15%	29.81%	50.96%	30.10%	38	49.04%	46.15%	63.46%	48.54%	71	55.77%	30.77%	48.08%	44.64%
6	50.96%	24.04%	54.81%	39.60%	39	63.46%	47.12%	56.73%	46.08%	72	48.08%	48.08%	57.69%	33.33%
7	54.81%	46.15%	56.73%	13.68%	40	56.73%	41.35%	51.92%	43.12%	73	57.69%	53.85%	52.88%	45.76%
8	56.73%	48.08%	51.92%	47.00%	41	51.92%	38.46%	53.85%	46.73%	74	52.88%	39.42%	54.81%	55.56%
9	51.92%	46.15%	55.77%	36.61%	42	53.85%	35.58%	55.77%	40.00%	75	54.81%	35.58%	58.65%	47.01%
10	55.77%	46.15%	44.23%	40.74%	43	55.77%	44.23%	50.96%	52.48%	76	58.65%	39.42%	57.69%	37.37%
11	44.23%	54.81%	52.88%	36.52%	44	50.96%	46.15%	52.88%	47.12%	77	57.69%	49.04%	58.65%	52.78%
12	52.88%	30.77%	57.69%	43.75%	45	52.88%	36.54%	53.85%	43.93%	78	58.65%	16.35%	58.65%	47.96%
13	57.69%	40.38%	49.04%	19.61%	46	53.85%	36.54%	42.31%	36.36%	79	58.65%	15.38%	46.15%	46.43%
14	49.04%	25.00%	49.04%	8.25%	47	42.31%	46.15%	58.65%	49.55%	80	46.15%	25.96%	45.19%	45.87%
15	49.04%	46.15%	56.73%	53.54%	48	58.65%	21.15%	52.88%	39.42%	81	45.19%	23.08%	53.85%	37.86%
16	56.73%	60.58%	52.88%	41.35%	49	52.88%	48.08%	59.62%	33.64%	82	53.85%	43.27%	50%	40.21%
17	52.88%	37.50%	56.73%	45.71%	50	59.62%	55.77%	60.58%	42.86%	83	50%	43.27%	51.92%	40.78%
18	56.73%	50.96%	51.92%	32.35%	51	60.58%	38.46%	54.81%	31.13%	84	51.92%	47.12%	58.65%	26.36%
19	51.92%	45.19%	50.96%	49.57%	52	54.81%	51.92%	51.92%	42.42%	85	58.65%	48.08%	52.88%	31.13%
20	50.96%	44.23%	55.77%	46.61%	53	51.92%	39.42%	58.65%	51.69%	86	52.88%	35.58%	59.62%	53.77%
21	55.77%	42.31%	59.62%	21.93%	54	58.65%	49.04%	52.88%	41.90%	87	59.62%	5.77%	55.77%	50.00%
22	59.62%	37.50%	50%	38.89%	55	52.88%	36.54%	61.54%	35.14%	88	55.77%	36.54%	58.65%	48.15%
23	50%	50.00%	47.12%	43.36%	56	61.54%	15.38%	45.19%	39.36%	89	58.65%	45.19%	50%	42.59%
24	47.12%	45.19%	51.92%	21.43%	57	45.19%	46.15%	49.04%	44.23%	90	50%	48.08%	49.04%	43.48%
25	51.92%	48.08%	47.12%	24.56%	58	49.04%	35.58%	55.77%	40.00%	91	49.04%	49.04%	52.88%	40.20%
26	47.12%	33.65%	51.92%	29.46%	59	55.77%	47.12%	57.69%	42.06%	92	52.88%	26.92%	53.85%	31.96%
27	51.92%	42.31%	52.88%	36.04%	60	57.69%	46.15%	50.96%	47.92%	93	53.85%	50.96%	60.58%	47.06%
28	52.88%	48.08%	44.23%	24.30%	61	50.96%	20.19%	54.81%	32.11%	94	60.58%	44.23%	47.12%	51.35%
29	44.23%	22.12%	53.85%	48.39%	62	54.81%	44.23%	54.81%	32.76%	95	47.12%	38.46%	50.96%	18.18%
30	53.85%	42.31%	50.96%	40.18%	63	54.81%	53.85%	52.88%	28.18%	96	50.96%	45.19%	51.92%	17.86%
31	50.96%	49.04%	49.04%	53.21%	64	52.88%	42.31%	59.62%	44.04%	97	51.92%	49.04%	49.04%	35.04%
32	49.04%	45.19%	56.73%	41.35%	65	59.62%	51.92%	51.92%	29.70%	98	49.04%	41.35%	52.88%	23.53%
33	56.73%	45.19%	48.08%	43.12%	66	51.92%	16.35%	56.73%	41.12%	99	52.88%	50.00%	47.12%	28.33%
										100	47.12%	50.96%	51.92%	44.17%

為便於與 SVM 模式進行分類比較，基於使用同樣一組抽樣樣本的情況下，我們另行分析 31-22-5 的 ANN 架構模式，訓練組請參看表 9，測試組請參看表 10，與 SVM 進行比較，在測試組部分，可得到以下結論：

1. 以個別信評等級的正確率觀察，SVM 模式的正確率在 A、BBB 及 BB 級三個類別高於 ANN 模式，在 AA 級以上及 B 級以下則低於 ANN 模式，顯示 SVM 模式在極端等級的分類上是較差的。
2. 以整體正確率觀察，SVM 模式高於 ANN 模式。
3. 以容忍上下一級誤差的正確率衡量觀察，SVM 模式幾乎可達 100% 的整體模式正確率及個別信評等級的正確率（除了 B 級以下為 95.24%）。

表 9：ANN 分類結果混亂矩陣——訓練樣本

Predicted Target	AA 級 以上	A	BBB	BB	B 級 以下	正確率	低估一級誤 差正確率	上下一級 正確率
AA 級以上	21	9	0	2	0	65.63%	93.75%	93.75%
A	7	52	12	3	2	68.42%	84.21%	93.42%
BBB	1	7	57	4	1	81.43%	87.14%	97.14%
BB	1	3	1	72	7	85.71%	94.05%	95.24%
B 級以下	0	1	1	3	59	92.19%	92.19%	96.88%
整體模式						80.06%	89.88%	86.05%

表 10：ANN 分類結果混亂矩陣——測試樣本

Predicted Target	AA 級 以上	A	BBB	BB	B 級 以下	正確率	低估一級 誤差正確率	上下一級 正確率
AA 級以上	4	3	2	0	0	44.44%	77.78%	77.78%
A	0	13	9	2	0	54.17%	91.67%	91.67%
BBB	1	2	14	4	2	60.87%	78.26%	86.96%
BB	0	1	8	12	6	44.44%	66.67%	96.30%
B 級以下	0	0	1	5	14	70.00%	70.00%	95.00%
整體模式						55.34%	76.70%	86.05%

陸、結論

本研究的主要貢獻在於探討過去研究較少提及的發行人信用評等議題，以及應用新近的人工智慧分類工具 SVM 方法來建構發行人信用評等分類模式。由於發行人信用評等不像發行評等可因考量債權順位等發行條件而較易於分類信評等級，故為一更複雜的分類決策問題。發行人信用評等為各種發行工具信用評等的基礎，而且目前台灣的信用評等資訊亦多屬於發行人的信用評等，本研究以 SVM 方法發展信用評等決策模式可為資本市場提供初步的發行人信用評等資訊。

SVM 方法為近幾年來在處理分類問題上獲得良好結果的分析決策工具，並已運用於醫學及工程等問題解決，在管理科學方面，亦同樣須面臨許多分類的問題，本研究以發行人信用評等為例去探討其在管理領域的適用性，並以另一人工智慧方法 ANN 模式為基準與 SVM 模式進行比較，發現 SVM 方法在本研究中的分類正確率高於 ANN 方法。

本研究的 SVM 模式測試組分類正確率僅達 60.58%，且多肇因於低估一個信用評等等級，建議未來的研究方向可朝向蒐集營運管理面的資料，讓模式本身可學習較完整的信用評等屬性，以消除低估一個信用評等等級的錯誤，進一步提升正確率。SVM 方法在發行人信用評等分類問題上，與本研究所設計的比較基準 ANN 方法比較後，初步得到較佳的分類正確率，建議未來可運用該方法解決其他的管理問題。

誌謝

作者感謝兩位匿名審查委員之寶貴意見。本研究係由國科會研究計畫所支持，計畫編號 NSC 95-2416 H-182-008，僅此誌謝。

參考文獻

1. Belkai, A. "Industrial Bond Ratings: A New Look," *Financial Management* (Autumn) 1980, pp: 44-51
2. Burbidge, R., Trotter, M., Buxton B. and Holden, S. "Drug Design by Machine Learning: Support Vector Machines for Pharmaceutical Data Analysis," *Computers and Chemistry* (26) 2001, pp: 5-14
3. Cai, Y.-D. and Lin, X.-J. "Prediction of Protein Structural Classes by Support Vector Machines," *Computers and Chemistry* (26) 2002, pp: 293-296
4. Diamantaras, K.I. and Kung, S.Y. *Principal Component Neural Networks: Theory and Applications*, John Wiley, New York, 1996
5. Dutta, S. and Shekhar, S. "Bond Rating: A Non-Conservative Application of Neural Networks," *Proceedings of the IEEE International Conference on Neural Networks* (II) 1988, pp: 443-450
6. Ederington, L.H., "Classification Models and Bond Ratings," *The Financial Review* (20:4) 1985, pp: 237-262
7. Fisher, L. "Determinants of Risk Premiums on Corporate Bonds," *Journal of Political Economy* (June) 1959, pp: 217-237
8. Gunn, S.R. "Support Vector Machines for Classification and Regression," unpublished manuscript, Faculty of Engineering and Applied Science Department of Electronics and Computer Science, University of Southampton, 1998, pp: 1-54
9. Horrigan, J.O. "The Determination of Long Term Credit Standing with Financial Ratios," *Journal of Accounting Research* (Supplement) 1966, pp: 44-62
10. Hsu, C.-W., Chang, C.-C. and Lin, C.-J. "A Practical Guide to Support Vector Classification," Department of Computer Science and Information Engineering, National Taiwan University, 2003
11. Kim, K.-S. and Han, I. "The Cluster-indexing Method for Case-based Reasoning Using Self-organizing Maps and Learning Vector Quantization for Bond Rating Cases," *Expert Systems with Applications* (21) 2001, pp: 147-156
12. Lippmann, R.P. "An Introduction to Computing with Neural Nets," *IEEE ASSP Magazine* (April) 1987, pp: 36-54

13. Maher, J.J. and Sen, T.K. "Predicting Bond Ratings Using Neural Networks: A Comparison with Logistic Regression," *Intelligent systems in accounting, finance and management* (6) 1997, pp: 59-72
14. Minoux, M. *Mathematical Programming: Theory and Algorithms*, John Wiley and Sons, 1986
15. Molinero, C.M., Gomez, C.A. and Cinca, C.S. "A Multivariate Study of Spanish Bond Ratings," *Omega* (24:4) 1996, pp: 451-462
16. Morris, C.W. and Autret, A. "Support Vector Machines for Identifying Organisms - A Comparison with Strongly Partitioned Radial Basis Function Networks," *Ecological modeling* (146) 2001, pp: 57-67
17. Pinches, E. and Mingo, K.A. "A Multivariate Analysis of Industrial Bond Ratings," *Journal of Finance* (March) 1973, pp: 1-18
18. Pinches, E. and Mingo, K.A. "The Role of Subordination and Industrial Bond Ratings," *Journal of Finance* (March) 1975, pp: 201-206
19. Shin, K.-S. and Han, I. "Case-based Reasoning Supported by Genetic Algorithms for Corporate Bond Rating," *Expert Systems with Application* (16) 1999, pp: 85-95
20. Shin, K.-S. and Han, I. "A Case-based Approach Using Inductive Indexing for Corporate Bond Rating," *Decision Support Systems* (32) 2001, pp: 41-52
21. Smith, K.A. and Gupta, J.N.D. "Neural Networks in Business: Techniques and Applications for the Operations Research," *Computers and Operations Research* (27) 2000, pp: 1023-1044
22. Standard & Poor's Corporation *Standard & Poor's Corporate Ratings Criteria*, McGraw Hill Book Company, 1996
23. Surkan, A.J. and Singleton, J.C. "Neural Networks for Bond Rating Improved by Multiple Hidden Layers," *Proceedings of the IEEE International Conference on Neural Networks* (2) 1990, pp: 163-168
24. Tay, F.E.H. and Cao, L. "Application of Support Vector Machines in Financial Time Series Forecasting," *Omega* (29) 2001, pp: 309-317
25. Vapnik, V. N. *The Nature of Statistical Learning Theory*, New York, Springer-Verlag, 1995
26. West, R.R. "An Alternative Approach to Predicting Corporate Bond Ratings," *Journal of Accounting Research* (Spring) 1970, pp: 118-127
27. Witten, I.H., Frank E. *Data Mining: Practical Machine Learning Tools and Techniques with Java Implementations*, San Francisco, Morgan Kaufmann Publishers, 2000

