

動態廣義式變精準灰粗集合預測模型結合類神經模糊系統 -香港之中國概念股實證分析

簡春娟

嶺東科技大學 財務金融系 講師

張洋秝

嶺東科技大學 財務金融系 學生

Abstract

This research proposes a trend filters the system, that combines rough set theory, modifiable neural fuzzy system and grey theory, which construct multi-level broad sense type become perfect rough set to build, abbreviated as GVP-Rough Sets. This model will have timeliness, seasonality or regular materials that will utilize dust to predict that the theory presents the concept with dynamic trend first, and then via assembling and screening the system thickly, explore a group of materials of group which provide trend value, utilize horses to link the theory now and apply to study in dynamic investment combination of the stock market finally. This text is in charge of the angle from the wealth and collect general H of share in Hong Kong, raise the financial rate index of every listed company of the type stock red, and do the condition, decision attribute and important discrimination of degree in accordance with master human relations Buffett of China's theory, and then use the dust to predict that carry on trend prediction to every wealth newspaper attribute, then utilize K-means to make mathematical calculations, hive off to every attribute , then in the result of hiving off of every listed company, use the thick set to the classification of its hiding , uncertain , insufficient materials , excavate ability to filter hiving off, run in order to screen finely, the company with sound physique, and offer the weight of investment, in order to make up the best investment combination.

Our research performance comes from simulating the investment operation model. Objects of study are Honk Kong enterprises in mainland subsidiaries and Hang Seng Index. Data collected from September to December in 2007. The simulation results show that Return On Investment (ROI) of six selected companies is larger than that of market index return. And comparing with return of investing weight allocation earnings and market index return, the supreme result is 32.55% from Markowitz's method, the second is 28.53% from grey theory, then 28.53% from equally weight, and the minimum is 15.96% from Hang Seng Index in Hong Kong. Through our investment combination filter model, the investment achievements is superior to market index return obviously.

Keyword : Grey Theory 、 Rough Sets Theory 、 K-means Clustering 、 Efficient Frontier 、 Portfolio Theory.

壹、前言

二十一世紀的經濟正邁向自由、國際與多元化的社會。而金融環境市場更是面臨全球化、證券化、信用增加與財務工程等趨勢的到來。因此，現今投資管道富於彈性多元，相形於投資理財規劃的觀念也日與俱增。由於香港歷經至今，對政權的平穩過渡、經濟產業結構改善、證券市場的發展以及加快大陸改革開放都起了積極的作用。H、紅籌公司於國際間之快速發展，不僅提高大陸和香港的經濟合作、優勢互補，以及香港的經濟繁榮、社會穩定，更改善了香港股市的產業結構，對香港保持國際金融中心地位有了正面的影響。同時，經過長期的發展，香港已形成了健全的市場機制，為高效率地配置和利用外部環境經濟資源提供了必要的條件，已成為中國走向世界的開端，也是國際資本進入中國大陸市場的橋樑，相對 H、紅籌類股勢必成為各國外資、散戶鎖定投資的對象。

隨著金融環境的蓬勃發展，金融工具的多元化，可藉由分析企業過去財務報表歷史資料，作為未來較精確的評估預測，相對於投資策略上，如何妥善資金管理亦甚是重要，有學者指出，資產配置可解釋絕大部分資產組合報酬的差異；主要訴求於長期的投資，其中以投資報酬率與資產組合的整體風險為最主要的考量因素。投資組合為一種規範性（normative）的理論，主要目的在達到資產及負債管理上特定的目標，及探討投資人該如何制定決策，才能於固定風險下，使報酬率達到最大；或報酬率固定情況下，使風險降到最低的投資組合。基於上述觀點在決定投資比重的方法中，馬克維茲(Markowitz)所提出的均數變異數最適化（Mean-Variance Optimization）的方法，乃成為投資組合理論於資產負債管理上一個重要的工具。

然而，在投資預測研究領域上，人們絞盡腦汁、抽絲剝繭的想從大量的歷史資訊中看出端倪，試圖窺探投資的規律性，用以預測未來。

資料探勘的各種分類預測方法應用，通常是限定某精確度上以設立門檻值。在此，本研究不求預測完全精確，望求能得知未來正確的趨勢。本文所欲給予的是一個灰色系統的趨勢動態概念，結合新興起的知識挖掘工具：「粗集合理論」，提出趨勢過濾投資組合模型。本研究目的有七：

- 一、本研究欲給予的是一個趨勢動態過濾系統，先將具有時間、季節或規律性的資料轉變為趨勢動向的數據，並結合粗集合理論，建構成趨勢粗集合（Trend Rough Sets）；於財報資料上，運用灰預測理論對其未來值做有效趨勢預測，使原屬落後的財報資料轉置成領先指標，以篩選出具未來趨勢動態預測性的集合。
- 二、粗集合理論中，屬性的選擇，採用台灣經濟新報資料庫(TEJ)，擷取 2005 年 12 月～2007 年 06 月 H、紅籌類股上市公司的財務比率數據，經粗集合理論層層過濾後可降低避免整體模型誤選地雷股，而篩選出具有永續經營條件，為股東創造高額報酬的優質企業。
- 三、使用 K-means 分群理論作為數值轉換的工具，將灰預測後的各項財報資料進

行數值不同等級區隔。將繁雜的數據給於適當歸類與等級區隔後，使用粗集合理論篩選出模糊性與難以辨識性的集合歸屬。以 K-means 分群理論克服財務報表數據中產業、規模不同不宜比較之限制。

四、以馬克維茲 (Markowitz)(1952) 的平均數—變異數前緣 (Mean-variance Portfolio Model) 風險與報酬組合的選擇作為資金權重配置。

五、比較均權資金配置、灰關聯係數排序的資金配置與馬克維茲風險與報酬組合的選擇，作為探討投資資金規模大小、投資標的多寡與投資標的之權重配置三種不同資金配置下投資組合報酬率之差異程度。

六、擷取過去香港中概股之 H、紅籌類股股市歷史資料，實際驗證本研究的投資組合篩選模型，測試模型其真正的效能與穩定度。

七、比較香港股票市場在利多利空行情下投資本模型與投資大盤績效。

由於股票市場至今仍存在不合法交易的發生，如內線交易、公司大股東炒作股票、媒體不實報導、做假帳…等，故難免影響研究結果之推論。另外由於資料庫建構所需資金龐大，本研究僅能以現有的台灣經濟新報資料庫，資料較為齊全的資料庫，作為研究樣本資料來源。

貳、文獻回顧

灰預測理論是根據所收集到的數據，建立適合於事件所需之預測模型，以達到預測的目的。Kung-Hsiung Chang & Chin-Shun Wu (1998) 研究探討農曆新年效應，亦即於股票市場中藉由在新年前、後分別買進及賣出股票來獲得異常報酬。他們並建立一灰色時間序列模型來預測新年效應。研究發現：他們認為農曆新年效應應存在，且應用灰預測所得之績效優於移動平均法及最小平方法。

施宜協 (2003) 此研究以在台灣證券交易所上市之電子產業為研究對象，研究期間為 1999 年 1 月至 2003 年 12 月，共 60 個月。利用不同的人工智慧模型對八組電子產業（其中每組產業挑選二家公司），進行股價比之預測，建構市場中立型避險基金，並評估不同模型之績效。其研究之實證結果在平均誤差方面，灰預測擁有最小之預測誤差、灰色馬可夫次之，而適應性模糊推論系統最差。

K-Means 分群 (McQueen, 1967; Dash et al., 2001) 是由 MacQueen 於 1967 年所提出來的分群演算法，也是最早提出來的分群演算法。由於 K-means 演算的邏輯簡單易懂、可以接受的時間複雜度的特性，文獻上已出現廣泛的應用：

Tou and Gonzalez (1974) 裡提到 K-Means 分群法，是一種被廣泛應用的分群法，相較於其他分群方式，其運算比較簡單，在分群速度上也較為快速，但分群的結果容易受到初始群心點所影響，其分群方式是設立初始群心點，然後利用距離的概念，將相近的資料點歸類到同一群，然後再重新計算群心點，再重覆以上的步驟直到群心點收斂為止。

吳柏林 (1996) 應用 K-means method 決定群落中心，利用此群落中心來辨別非線性時間數列的結構；並應用模糊理論，找出此時間數列的模糊轉折點，決定

模糊轉折點的歸屬準則與轉折區間認定，並將此方法應用在台灣匯率的資料分析與預測。其結果發現，以模糊分類法配適模式來預測所得到的值是最接近真實值的，且優於傳統的 ARIMA 模式及 ICSS 演算法。

類神經模糊系統的應用結合模糊與類神經的優點；模糊邏輯在解決案例可以很容易進行驗證和最佳化，類神經網路可以利用數據資料庫進行訓練學習。

Kimoto and Asakawa(1990)其研究內容為：以模組化的類神經網路，找出未來一個月的最佳買賣點。輸入變數為基本面的經濟指標為主，訓練測試資料有 1987 年 1 月至 1989 年 9 月的日經股價指數，其系統模擬時採用 moving simulation 的方式。其模組化指以多個神經網路分別學習單一變數，在經處理整合。其目的在使類神經網路產生可被解釋的輸出，了解各輸入變數相互關係。

Jang(1991)其研究內容為：以兩個模組(dual-module)的類神經網路，預測短期股價趨勢，並以台灣股票市場為標的，發展一智慧型股票投資管理系統，給予投資人相關建議。兩模組分別為選取特性回顧與趨勢預測，再針對長短兩種天期的技術指標做學習，其研究結果與多重線性迴歸分析比較，績效都優於迴歸的結果。

Bouqata 等人(2000)利用 Quick-pop 類神經網路、Adaptive-Network-Based Fuzzy Inference System(ANFIS)、模糊回歸與決策樹(ADRI)、傳統時間序列模式四種方法，應用於預測之上，並利用 Root Mean Absolute Squared Estimate Error(RMASE)比較四種方法之優劣。而 ANFIS 則是最佳的方法。

粗集合理論模型(Rough Sets Theory Model)應用層面廣泛，涵蓋醫學工程、製成管理、財務工程等，而目前主要大量應用於企業破產預警、資料庫行銷與金融投資預測三大領域。主要的原因：粗集合理論特徵主要是可以藉由歷史資料庫，挖掘資訊中隱藏的模型藉以預測未來。而這三領域的歷史資料庫皆由多種屬性資訊表建立，且擁有相似的預測特徵。

在金融投資預測方面，目前有兩個研究主題，其一，是在各種投資市場中依據交易行為建立交易系統；學者們根據各市場交易系統記錄的細部資訊，通常是以粗集合理論基礎或結合類神經網路，簡化系統資訊、窺探市場波動規則，進行短期或長期投資應用。在諸多研究顯示：以粗集合理論建構的模型和傳統的統計預測模型相較下，使用粗集合模型的投資績效遠勝於傳統。其二，是以投資組合偏好為另一項主要應用，學者們利用粗集合模型尋找投資組合偏好屬性，觀測偏好屬性變化進行投資決策，提高投資效度。

粗集合理論的應用通常是結合其他理論進行應用。以粗集合理論為基礎所衍生的混和模型，經多方理論的結合，對於粗集合的屬性、不可辨識與結果檢測進行探討與改良，使粗集合理論更具深度與廣度。(如表 2.1)

Rough sets models	Business failure prediction	Database marketing	Financial investment
RSES			Bazan et al. (1994) Baltzersen (1996)
LEERS		Poel (1998)	
DataLogic	Szladow and Mills (1993)	Mills (1993) Mrozek and Skabek (1998)	Ziarko et al. (1993) Golan (1995) Golan and Edwards (1993) Ruggiero (1994a,b,c) Skalko (1996) Lin and Tremba (2000)
TRANCE		Eiben et al. (1998) Kowalczyk and Slisser (1997) Kowalczyk and Piasta (1998) Kowalczyk (1998a) Poel (1998) Poel and Piasta (1998)	
ProbRough			
Dominance relation	Greco et al. (1998)		
RoughDas and ProFit	Slowinski and Zopounidis (1994, 1995) Dimitras et al. (1999) Slowinski et al. (1997) Slowinski et al. (1999) Ahn et al. (2000) Hashemi et al. (1998)		Susmaga et al. (1997)
Hybrid model			

表 2.1 粗集合理相關文 F. E.H. Tay& L. Shen 整理

李慧慈（2003）利用粗集合論預測網路銀行使用意願，研究發現：在網路銀行使用意願方面，粗集合模式突破統計模型對資料的限制，且獲得異於迴歸模型的預測結果，顯示粗集合分析的確可挖掘隱匿於資料背後的重要訊息。

許展碩（2003）提出一般化的基因演算法，並整合由粗集合所產生的知識規則。研究發現：只要能將一個問題分解成多個功能需求，就可以用一般化的基因演算法加以解決，而且整合由粗集合理論所發掘出來的知識，確實能提高此基因演算法的效能以及加速收斂，特別是用在簡化母體或限制交配這二方面。

劉淑賢（2003）首先，透過以價值流呈現目前的製造過程，及採用粗集合理論來找出被視為重點的流程類型，以精簡控制所最需要的部分。而後再以一般化的 label-correcting (GLC) 方法來決定在精簡製造中所需的流程範圍。研究發現：這方法論適用於有混合型態的反覆製造環境，如零工型或是流程型生產，並達到全面減少在製品存量的目標，這亦是精簡生產的目標之一；這個方法相當新穎因不同型態的問題能以一個演算法來解決，歸納出的規則能有效地找出重點流程類型，並能以系統思考的觀點來減少浪費。

參、研究方法

一、灰預測理論

在灰色系統理論中，灰預測是根據所收集到的各數據，利用生成建模方法，建立適合於事件所需之預測模型，以達到預測的目的。本研究採用灰預測中的 GM(1,1) 模型來做預測。

(一)AGO 方式數據處理

假設： $X^{(0)}$ 為一 n 期非負時間序列，其原始數列為

$$X^{(0)} = (X^{(0)}(1), X^{(0)}(2), \dots, X^{(0)}(n)) \quad (3.1.1)$$

累加生成（AGO）數列為

$$X^{(1)} = (X^{(1)}(1), X^{(1)}(2), \dots, X^{(1)}(n)) \quad (3.1.2)$$

其中

$$x^{(1)}(k) = \sum_{m=1}^k x^{(0)}(m) \quad (3.1.3)$$

，則

$$\begin{cases} x^{(0)}(k) = x^{(1)}(k) - x^{(1)}(k-1), & \text{for } k \geq 2 \\ x^{(0)}(1) = x^{(1)}(1), & \text{for } k=1 \end{cases} \quad (3.1.4)$$

$\underline{x}^{(0)}$ 與 $\underline{x}^{(1)}$ 的幾何位置關係圖：

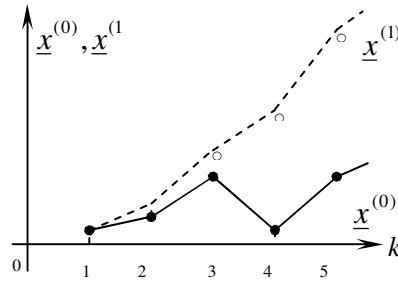


圖 3.1 幾何位置關係圖

(二)GM(1,1)模型

具有一階一個變量的灰模型稱為 GM(1,1)模型，a 稱為發展係數，b 稱為灰作用量。

$$X^{(0)}(k) + aZ^{(1)}(k) = b, k = 1, 2, \dots, n \quad (3.1.5)$$

$$Z^{(1)}(k) = 0.5[X^{(1)}(k) + X^{(1)}(k-1)], k = 2, 3, \dots, n \quad (3.1.6)$$

$$x^{(1)}(k) = \left[\sum_{m=1}^k x^{(0)}(m) \right], k = 1, 2, \dots, n \quad (3.1.7)$$

(三)灰色模式建構

GM(1,1)建模首先要計算 GM(1,1)模型中之參數 a、b，對參數 a、b 之估計一般使用最小平方法(Least square method)。

令 GM(1,1) 模型

$$X^{(0)}(k) + aZ^{(1)}(k) = b \quad (3.1.8)$$

，滿足序列 $X^{(1)}$ 與 $X^{(0)}$

$$X^{(0)} = (X^{(0)}(1), X^{(0)}(2), \dots, X^{(0)}(n)), k = 1, 2, \dots, n \quad (3.1.9)$$

$$X^{(1)} = (X^{(1)}(1), X^{(1)}(2), \dots, X^{(1)}(n)), k = 1, 2, \dots, n \quad (3.1.10)$$

$$x^{(1)}(k) = \sum_{m=1}^k x^{(0)}(m) \quad (3.1.11)$$

則有

$$X^{(0)}(j) + aZ^{(1)}(j) = b, \quad j = 2, 3, \dots, n \quad (3.1.12)$$

，即

$$\begin{bmatrix} x^{(0)}(2) \\ x^{(0)}(3) \\ \dots \\ x^{(0)}(n) \end{bmatrix} = \begin{bmatrix} -z^{(1)}(2) & 1 \\ -z^{(1)}(3) & 1 \\ \dots & \dots \\ -z^{(1)}(n) & 1 \end{bmatrix} \cdot \begin{bmatrix} a \\ b \end{bmatrix} \quad (3.1.13)$$

其中

$$Z^{(1)}(k) = 0.5[X^{(1)}(k) + X^{(1)}(k-1)], k = 2, 3, \dots, n \quad (3.1.14)$$

按最小平方方式，得 GM(1,1) 參數 a、b 的矩陣算式為

$$\hat{\theta} = \begin{bmatrix} a \\ b \end{bmatrix} = (B^T B)^{-1} B^T Y_N \quad (3.1.15)$$

$$B = \begin{bmatrix} -z^{(1)}(2) & 1 \\ -z^{(1)}(3) & 1 \\ \dots & \dots \\ -z^{(1)}(n) & 1 \end{bmatrix}, Y_N = \begin{bmatrix} x^{(0)}(2) \\ x^{(0)}(3) \\ \dots \\ x^{(0)}(n) \end{bmatrix} \quad (3.1.16)$$

(四) 灰預測

$$\text{將 } Z^{(1)}(k) = 0.5[X^{(1)}(k) + X^{(1)}(k-1)] \quad (3.1.17)$$

代入 GM(1,1) 模型

$$X^{(0)}(k) + aZ^{(1)}(k) = b \quad (3.1.18)$$

則可得

$$X^{(0)}(k) + 0.5a[X^{(1)}(k) + X^{(1)}(k-1)] = b \quad (3.1.19)$$

，其中

$$X^{(1)}(k) = X^{(1)}(k-1) + X^{(0)}(k) \quad (3.1.20)$$

整理可得

$$x^{(0)}(k) = \frac{b - a x^{(1)}(k-1)}{1 + 0.5a} \quad (3.1.21)$$

故

$$\hat{x}^{(0)}(n+1) = \frac{b - a x^{(1)}(n)}{1 + 0.5a} \quad (3.1.22)$$

其中，“ $\wedge$ ” 表某特徵的預測值。

(五) GM(1,1) 模型之發展係數的可容區

令 a 為 GM(1,1) 模型之發展係數，S0 為 a 的可容區，其中 S0=(-2,+2)。若

$a \in (-2, 0)$ ，則必須滿足

$$\frac{x^{(0)}(k)}{\sum_{m=2}^{k-1} x^{(0)}(k)} > \frac{|a|}{1 + 0.5a} \quad (3.1.23)$$

之條件。

二、K-means 分群法則

資料挖掘是一種可以從資料庫中分析原始資料，將隱含、先前並不知道的資訊從資料庫中萃取出來的技術。而空間資料挖掘是資料挖掘中的一個分支，它主要用在處理空間資料。在空間資料挖掘中，資料分群（data clustering）乃是從原始資料中發掘出我們有興趣的資料，是一種非常有用的技術。

將所有樣本點分割為 K 個原始群集，此 K 個群集重心稱為種子點（SEED POINT）。重複針對每一樣本，計算比較該樣本在該群中之距離平方和以及失去該樣本後該群之離差平方和。若失去該樣本後可使離差平方和降低，便將該樣本改為指定至其他群集中，重新上述之計算，目的是要使各群中之離差平方和最小化。重複動作直到其收斂，各樣本不需重新指派到其他群集中，便可完成分群，見圖 3.2、圖 3.3 與圖 3.4 的說明。而目前決定初始群數的方法，主要的有下列四點[鍾文杰,2001]：

（一）Anderberg 隨機方法

這是最常用的方法，目的就是將資料透過隨機方式，分割成 K 個群組。

（二）Forgy 方法

由 Forgy 於 1965 年所提出的，參考 Anderberg 的方法，先從資料庫中隨機選取 K 個種子，並藉著最近的種子分配到所代表的聚集。

（三）Macqueen 方法

由 Macqueen 於 1967 年所提出，先從資料庫中隨機選取 K 個種子，根據其順序指派剩餘的種子到最接近群心的聚集。經過每一個指派，所有群心重新計算完成。

（四）Kaufman 方法

由 Kaufman 和 Rousseeuw 於 1990 年所提出，最初群數的決定來自於連續代表種子的選取，直到 K 種群數都被找到。

一般來說，我們可以將資料分群的演算法分成四大類：分割式、階層式、密度基礎和格子基礎的分群方法。 K -means 演算法是屬於分割式分群方法，主要是以重心點或是中心點為基礎的方式，將資料群體進行分群，因為各群體的代表點不一定是群聚中的一點，所以可以在多數的情況下找到最佳群聚，它是一個非常普遍且廣泛被使用的資料分群方法。

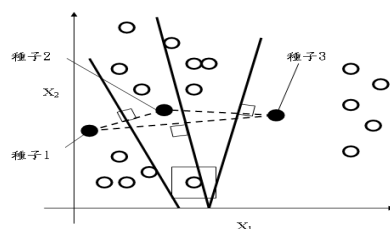


圖 3.2 初始種子決定了初始的群集方界。

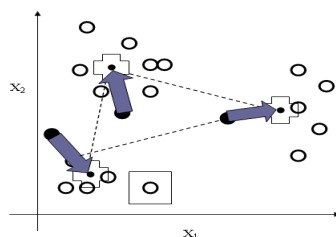


圖 3.3 計算新群集的質心。

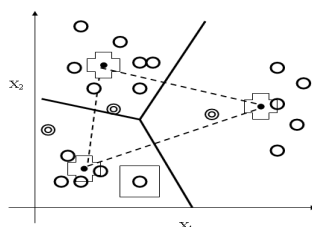


圖 3.4 每一次的重複，所有群集分配都須重新計算一次。

三、可調整式類神經模糊推論系統(ANFIS)

(一)ANFIS 簡介

可調整式類神經網路的結構與學習規則在以前的文獻當中已被完整的描述，功能上來說，除了一些片段的可辨性必要條件外，可調整式類神經網路的節點函數幾乎沒有什麼限制。而在結構上而言，類神經結構唯一的限制，是它必須為前饋控制的類型，如果我們不想使用更複雜的非同步操作模型的話。因為這些極微的限制，可調整式類神經網路可以被直接的應用在建模、決策、信號處理與控制等，多種的應用過程中。

在本文當中，我們將詳細的介紹可調整式類神經網路，它在功能上等同於模糊推論系統；『ANFIS』代表是“可調式以類神經網路為基礎的模糊推論系統”，或是“可調式類神經網路模糊推論系統”。本文中，將敘述如何去分解參數集，來幫助建立 ANFIS 結構中“Sugeno”與“Tsukamoto”模糊模型的混合學習規則；也將驗證在確定的主要限制下，徑向基類神經網路(the Radial Basis Function Network；RBFN)在功能上等價於 ANFIS 中的 Sugeno 模糊理論。ANFIS 中混合學習規則的廣泛可應用性，可透過以下四個模擬的例子來獲得驗證：

1. 建立一個二唯的 sinc 函數。
2. 建立一個三個輸入的非線性函數，這是其他模糊模型應用時的一個基本模型。
3. 解釋如何去辨識線上控制系統中非線性組成。
4. 預測 Mackey-Glass 的混沌時間序列。

驗證結果顯示 ANFIS 比常見的統計模型可以更廣泛的被應用，許多學者也相繼的提出了相似的類神經網路結構。

(二)ANFIS 結構說明

假設在模糊推論系統下，考慮有兩個輸入 x 與 y 且有一個輸出 z ，對一階的 Sugeno 模糊模型來說，在一般的規則集合中，有下列兩個 if-then 的模糊規則：

1. 若 x 為 A_1 且 y 為 B_1 ，

$$\text{則 } f_1 = p_1x + q_1y + r_1 \quad (3.3.1)$$

2.若 x 為 A_2 且 y 為 B_2 ，

$$\text{則 } f_2 = p_2x + q_2y + r_2 \quad (3.3.2)$$

在圖 3.5(a)中，以圖解方式來建構這個 Sugeno 模型，等價的 ANFIS 結構則表示在圖 3.5(b)中，在每一個階層的節點中有相似的函數，描述如下：

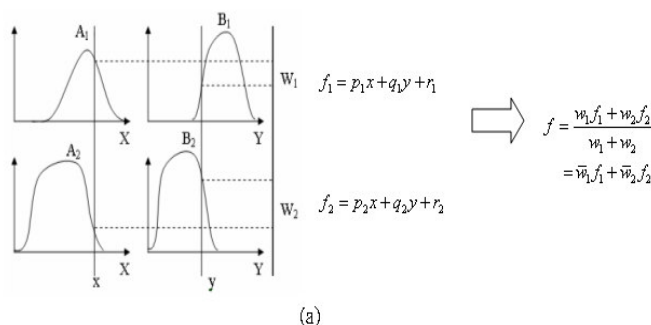


圖 3.5 (a)兩個規則兩個輸入一階 Sugeno 模型

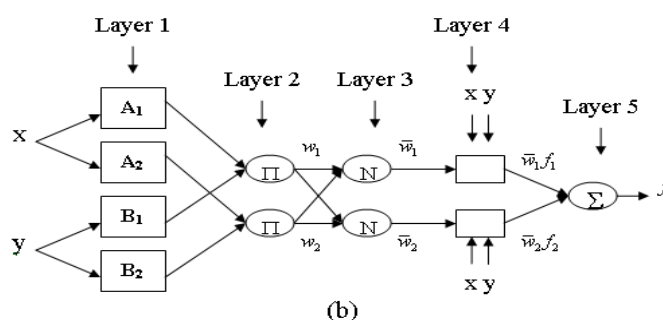


圖 3.5 (b)等價的 ANFIS 結構

(1)在此階層的每個節點 i 為一個含有節點函數的可調整式節點

$$O_{l,i} = \mu_{A_i}(x), \text{ for } i=1,2, \text{ or } \quad (3.3.5)$$

$$O_{l,i} = \mu_{B_{i-2}}(y), \text{ for } i=3,4 \quad (3.3.5)$$

其中， x 或 y 為節點 i 的兩個輸入變數，且 A_i 或 B_{i-2} 為連結此節點的語言符號

(例如”大”或”小”)，換句話說， $O_{l,i}$ 為模糊集合 $A(=A_1, A_2, B_1 \text{ or } B_2)$ 的歸屬層級，它具體指定在給定的輸入 x 或 y 且滿足語言符號 A 的階層，在此 A 的歸屬函數可以被定義入下，例如一般的鐘型函數為：

$$\mu_A(x) = \frac{1}{1 + \left| \frac{x - c_i}{a_i} \right|^{2b}}, \quad (3.3.6)$$

其中， $\{a_i, b_i, c_i\}$ 為參數集合，當這些參數值改變時，則鐘型函數會跟著相對應改變，因此，對於模糊集合 A 來說，這些參數代表歸屬函數的變化型式，在此階層的參數代表『前置參數』。

(2) 在此階層的每個節點為固定的節點符號 Π ，它的輸出為所有輸入信號所產生的結果：

$$O_{2,i} = w_i = \mu_{A_i}(x) \mu_{B_i}(y), \quad i=1,2. \quad (3.3.7)$$

每個節點的輸出代表一個規則的權重，一般而言，在此階層中任何其他 T -類型執行模糊 AND 的操作時，可以被用來當成一個節點函數。

(3) 在此階層的每個節點為固定的節點符號 N ，在第 i 個節點，計算第 i 個節點規則的權重相對於全部規則權重總合之比率：

$$O_{3,i} = \bar{w}_i = \frac{w_i}{w_1 + w_2}, \quad i=1,2. \quad (3.3.8)$$

為了方便說明，在此階層的輸出我們稱為『規則權重的正規化』。

(4) 在此階層的每個節點 i 為含有節點函數的可調整式節點：

$$O_{4,i} = \bar{w}_i f_i = \bar{w}_i (p_i x + q_i y + r_i), \quad (3.3.9)$$

其中， $\bar{w}_i$ 為階層 3 所計算出來的正規化規則權重，且 $\{p_i, q_i, r_i\}$ 為這個節點的參數集合，在這個階層的參數表示為『結果參數』。

(5) 在此階層的信號節點為固定的節點符號 Σ ，這裡計算所有輸入信號的總合作為全部的輸出：

$$\text{overall output} = O_{5,1} = \sum_i \bar{w}_i f_i = \frac{\sum_i w_i f_i}{\sum_i w_i} \quad (3.3.10)$$

因此，我們所建構的可調式類神經網路，在功能上同等於 Sugeno 模糊模型，而且這個可調式類神經網路的結構並不是唯一的；我們可以結合階層(3)與階層(4)來獲得只有四個階層的可調式類神經網路，圖 3.6 為此類型的 ANFIS 結構，在比較極端的例子中，我們甚至可以將全部的類神經網路，縮減為含有相同參數集的單一可調整式節點；很明顯地，節點的配置與每個階層的執行為有意義且有標準化的功能性。

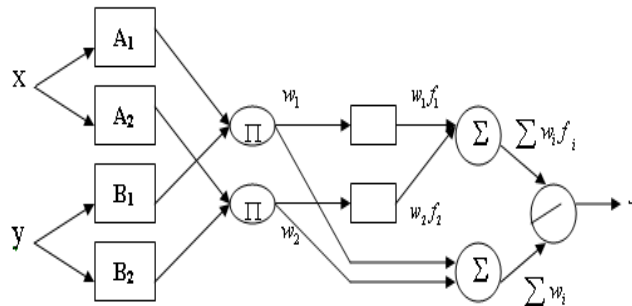


圖 3.6 Sugeno 模糊模型的 ANFIS 結構，在此權重的標準化在最後一個階層執行。

從 Sugeno ANFIS 到 Tsukamoto ANFIS 的延伸是很簡單的，如圖 3.7 所示，其中每個規則的輸出 ($f_i, i=1,2$) 是藉由歸屬函數與規則權重共同的來歸納出來，對於由 MAX-MIN 所構成的 Mamdani 模糊推論系統來說，若離散的近似值被用來代替在距心的積分，則相同的 ANFIS 可以被建構。然而，結果顯示 Mamdani ANFIS 是比 Sugeno ANFIS 或 Tsukamoto ANFIS 更加複雜的，但是，以 Max-Min 組成的 Mamdani ANFIS 之結構與計算所增加的額複雜性，不需包含更佳的學習能力或近似力。若我們採納以輸出-加總組成與距心解模糊化的 Mamdani 模糊模型，相同的 ANFIS 可以更容易的被構成。

透過這個章節的說明，因為 Sugeno 模糊模型的透明度與效力，我們將集中在一階 Sugeno 模糊模型的 ANFIS 結構。

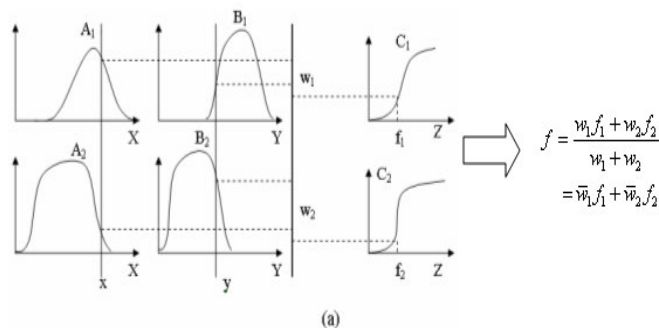


圖 3.7 (a)兩輸入與兩規則的 Tsukamoto 模糊模型

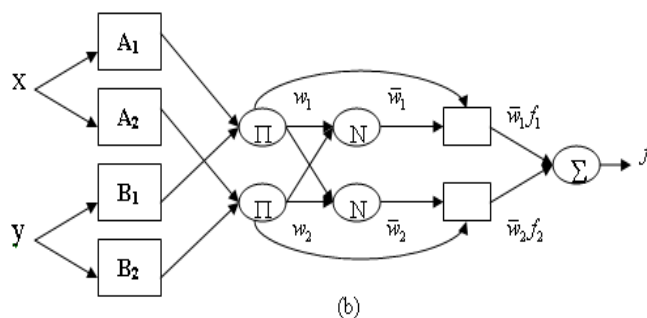


圖 3.7 (b)等價的 ANFIS 結構

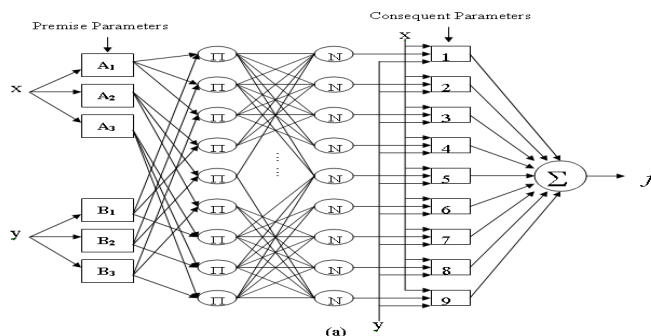


圖 3.8 (a)有九個規則的兩輸入一階 Sugeno 模糊模型之 ANFIS 結構

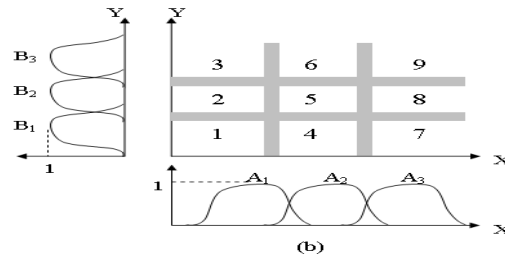


圖 3.8 (b)被化分為九個模糊區域的輸入空間

圖 3.8(a)為含有九個規則的兩輸入一階 Sugeno 模糊模型之 ANFIS 結構，其中，每個輸入假設有三個聯合的歸屬函數(MFs)，圖 3.8(b)說明如何將二維的輸入空間分割為九個重疊的模糊區域，每個區域是藉由模糊的 if-then 規則來控制；換句話說，一個規則的假設就是一個模糊區域，而接下輸出的部分則指定在此區域內。接下來我們將說明，如何應用混合學習規則系統來確定 ANFIS 的參數。

(三)混合學習規則系統(Hybrid Learning Algorithms)

從圖 3.5(b)的 ANFIS 結構中，我們可以觀察到當前置參數的數值固定時，全部的輸出可以被表示為結果參數的線性組合，這象徵在圖 3.5(b)的輸出 f 可以被

$$f = \frac{w_1}{w_1 + w_2} f_1 + \frac{w_2}{w_1 + w_2} f_2$$

改寫為：

$$\begin{aligned} &= \bar{w}_1(p_1x + q_1y + r_1) + \bar{w}_2(p_2x + q_2y + r_2) \\ &= (\bar{w}_1p_1 + \bar{w}_2p_2)x + (\bar{w}_1q_1 + \bar{w}_2q_2)y + (\bar{w}_1r_1 + \bar{w}_2r_2) \end{aligned} \quad (3.3.11)$$

其中，結果參數 $p_1, q_1, r_1, p_2, q_2, r_2$ 為線性組合，由這些觀察值我們可以得到：

S = 全部參數的集合，

S_1 = 前置(非線性)參數的集合，

S_2 = 結果(線性)參數的集合。

更明確的來說，混合學習規則系統的順向路徑，節點輸出一步步向前推論直到第四階層與結果參數都是藉由最小平方方法來計算；在倒傳遞路徑中，誤差訊號向後傳導與前置參數，皆是藉由梯度下降學習法來更新，表 3.1 說明了每種傳遞的詳細內容。

表 3.1 ANFIS 的混合學習系統之兩種傳遞方式

	順向路徑	倒傳遞路徑
前置參數	固定	梯度下降學習法
結果參數	最小平方估計法	固定
傳遞訊號	節點輸出	誤差訊號

結果參數是在前置參數固定的條件下所推論出來，因此，ANFIS 的混合式學習法則，若將其與多層類神經網路廣泛使用的標準倒傳遞方法作比較，可得知前者的收斂速度遠較後者快。因此，我們必須找尋分解參數集的可能性，對 Tsukamoto ANFIS 來說，若每個規則在結構部份的歸屬函數，可以用含有兩個結

果參數的分段線性近似來代替的話，分解參數集是可被完成的。

四、粗集合理論

(一)粗集合理論的基礎假設

基礎假設：論域（ U ）中每一個元素（ X ）和某些資訊是有關聯的。舉例而言，假設我們想要探討財務資料庫中某些財務比率與公司價值方面的關係，則財務資料庫中的財務比率資料，我們稱之為論域（ U ），其相關資訊是由各家公司的財務行為與價值特徵組合而成；論域（ U ）中各家公司即是我們所稱的元素（ X ）。

(二)粗集合之元素間不可辨識關係理論

通常在相同可使用的資訊下，各元素（ X ）的特徵存在許多相似性，而在資訊、工具不足下讓人無法辨識元素（ X ）間的差異性。然而，在粗集合的理論基礎下對於各元素的不可辨識關係（Indiscernibility），給予一個合理的集群關係，利用決策表給予一個適當的分類。藉此理論可發現論域（ U ）中元素集合的屬性性質與屬性間依賴程度並刪減多餘的屬性。

(三)資訊系統

資訊系統 $S = (U, R, V_q, f_q)$ 定義如下：

U ：定義為論域集合；

R ：代表整體屬性的集合；

$R = C \cup D$ ， C 定義為集合中條件屬性， D 定義為集合中的決策屬性；

$q \in R$ ， V_q 稱為 q 的定義域；

f_q 為資訊函數； $U \rightarrow V_q$

元素（ X ）：可被解釋為個案、狀況、製程、型樣、觀察資料…等；

屬性（ $C \& D$ ）：可被解釋為特色、變數與特徵條件。一件特別的資訊系統個案分別稱為決策表或屬性價值表；在決策表中列與行分別代表元素與屬性。

(四)上近似、下近似集合、邊界集合

現實生活中的資料通常存有某些不精確、含糊的部分，使得各元素（ X ）間相互矛盾而產生不可辨識（ $IND(B)$ ）。不可辨識是指在條件屬性集合下兩個或兩個以上屬於不同決策種類的元素難以辨識其關係或差異稱之；這裡的決策表被稱為不可辨識決策表。然而，粗集合理論中的近似集合，其主要的功能即是在處理元素間的不可辨識性。

$$S = (U, R, V_q, f_q) \quad , \quad X \subseteq U \quad (3.4.1)$$

而 R^+ 與 R_- 分別定義為 X 的上近似與下近似：

$$R^+(X) = \bigcup \{Y \in U / IND(B) : Y \cap X \neq \emptyset\}, \quad (3.4.2)$$

$$R_-(X) = \bigcup \{Y \in U / IND(B) : Y \subseteq X\}, \quad (3.4.3)$$

這裡的 $U / IND(B)$ 表示 R 的等價類； $IND(B)$ 稱為 R 的不可辨識關係，其定義如下：

$$IND(B) = \{(x, y) \in U^2 \\ : \text{for every } a \in R, a(x) = a(y)\} \quad (3.4.4)$$

上、下近似集合分別表示：使用屬性集合 R 時， R^+X 代表在 Y 決策下“完全”被分類為等類的 X 元素集合， R^-X 代表在所有 Y 決策下“有可能”被分類為等類的 X 元素集合。

$$Bnr(X) = R^+(X) - R^-(X) \quad (3.4.1)$$

稱為 X 的邊界集合。

(五)上、下近似集合的準確度

準確度主要是在解釋以 R 等價類分配的近似集合，集合中 X 元素是否被正確的分類。如果 $\alpha_B(R)$ 接近 1 代表分類與決策表是相當一致的，屬性集合 R 運用得當。

$$\alpha_B(R) = \frac{\sum \text{card}(R^-(X_i))}{\sum \text{card}(R^+(X_i))} \quad (3.4.1)$$

五、灰關聯分析(Grey Relational Analysis)

灰關聯分析，即對灰色系統因素之間的發展動態進行定量比較分析，它是一種根據因素與因素之間發展趨勢的相似或相異程度（關聯度）來衡量因素間關聯程度的方法，把系統有關因素間的各種關係，一一呈現出來，當做系統決策、預測控制提供有用信息和比較可靠的依據，這種分析模式可將灰色系統內各因素間灰關係清晰化。灰關聯分析程序可以歸納如下幾個步驟：

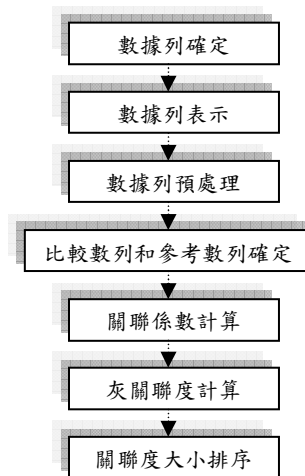


圖 3.9 灰關聯分析程序

(一)數據列確定

對抽象系統進行灰關聯分析時，須先確定能表現系統特徵之數據列，數據列的取得方法有兩種：1.直接法、2.間接法。

(二)數據列表示

數據列：1.時間序列 2.指標序列 3.空間序列。

(三)數據列預處理

進行灰關聯分析時常需對數據列進行處理，常用之處理方式包括：

- 1.時間數列：初值化、最小值化、最大值化、平均值化、區間值化。
- 2.非時間序列：指標區間值化、規一化。

(四)比較數列和參考數列確定

進行灰關聯分析時必須先確定參考數列，再比較其他數列與參考數列之接近程度。如此，才能對比較數列進行比較，進而作出判斷。

(五)灰關聯係數(Grey Relational Coefficient)計算

以灰關聯係數計算，得到的是各比較數列與參考數列在各點之灰關聯係數值，而鄧聚龍所定義之灰關聯係數為：

$$\gamma(x_i(k), x_j(k)) = \frac{\Delta_{\min.} + \zeta \Delta_{\max.}}{\Delta_{oi}(k) + \zeta \Delta_{\max.}}. \quad (3.5.1)$$

其中， $i=1,2,\dots,m; k=1,2,\dots,n \quad j \in I$ ， x_0 為參考序列， x_i 為一特定之比較序列。

$$\Delta_{oi} = \|x_0(k) - x_i(k)\| \quad (3.5.2)$$

為 x_0 和 x_i 之間第 k 個差的絕對值，

$$\Delta_{\min} = \min_{j \in I} \min_k \|x_0(k) - x_j(k)\|, \quad (3.5.3)$$

$$\Delta_{\max} = \max_{j \in I} \max_k \|x_0(k) - x_j(k)\|, \quad (3.5.4)$$

ζ 為辨識係數且 $\zeta \in [0,1]$ 。

(六)灰關聯度 (Grey Relational Grade) 計算

灰關聯係數計算，得到之資訊過於分散。因此，有必要將關聯係數集中表現在一個數值上，即灰關聯度。當求得灰關聯係數後，傳統方式（鄧聚龍）是取灰關聯係數的平均值為灰關聯度。

$$\gamma(x_i, x_j) = \frac{1}{n} \sum_{k=1}^n \gamma(x_i(k), x_j(k)). \quad (3.5.5)$$

然而在實際的系統上，各個因子對系統的重要程度並不見得完全相同，因此我們正視各個因子的權重不相等的實際情形，延伸上式中的關聯度的定義為

$$\gamma(x_i, x_j) = \sum_{k=1}^n \beta_k * \gamma(x_i(k), x_j(k)). \quad (3.5.6)$$

其中 β_k 表示因子 k 的常態化權重，由使用者決定，必須滿足 $\sum_{k=1}^n \beta_k = 1$ 。

(七)灰關聯序(Grey Relational Ordinal)

對參考數列 x_0 與比較數列 $x_i (i=1,2,\dots,m)$ 其關聯度分別為 $\gamma_i (i=1,2,\dots,m)$ ，按大小進行排序，即得灰關聯序(Grey Relational Ordinal)。

若 $\gamma(x_0, x_i) \geq \gamma(x_0, x_j)$ ，則稱 x_i 對 x_0 的關聯度大於 x_j 對 x_0 的關聯度，且表示為

$$x_i \succ x_j \circ$$

六、馬克維茲(Markowitz)

投資組合理論起始於諾貝爾經濟學獎得主馬克維茲Markowitz(1952)的平均數一變異數前緣(Mean-variance Portfolio Model)，(吳啓銘，民89)這個前緣代表在一個已知的投資組合預期報酬之下，可以達到的最低變異數之投資組合。所有落在最低變異數前緣線上的投資組合，可供作最佳的風險與報酬組合的選擇。整體最低變異數的投資組合的連線，稱之為效率前緣 (efficient frontier) (陳姿媚，2000)。其理論主要乃基於下列五項假設發展而成：

- (一)投資者希望財富愈多愈好，且其效用為財富的函數。但其財富的邊際效用遞減，一般通常假定投資期間為一期，因為多期的分析情況，在分析上相當麻煩。此外，由於財富與投資報酬率的關係相當密切，因此我們可以說投資者的效用為報酬率的函數。同理，其對報酬率的邊際效用遞減。
- (二)投資者能事先知道投資報酬率的機率分配為常態分配。
- (三)投資風險以報酬率之變異數或標準差來表示。
- (四)投資者希望其效用之期望值為最大，而此一效用之期望值是期望報酬率與風險之函數。因此影響投資決策的主要變數為期望報酬率與風險兩項。
- (五)遵守主宰原則的指導：即在同一風險水準下投資者希望報酬率愈高愈好，而在同一投資報酬率水準下，投資者希望風險愈小愈好。

上述五項為馬克維茲(Markowitz)投資組合理論的主要假設(王淑芬，民86；呂安悌，民89，pp6-7)。利用馬克維茲(Markowitz)之二次規劃法求出效率投資組合中表權數之 W_i 值，從而決定出效率前緣。基本的關係模式其使用到的變數參數及數學模式如下：

Max

$$L = (1 - \lambda)E(R_p) - \lambda\sigma_p^2 ; 0 \leq \lambda \leq 1 \quad (3.6.2)$$

S.T.

$$Var(R_p) = \sum_{i=1}^n \sum_{j=1}^n W_i W_j \sigma_{ij} \quad (3.6.2)$$

$$E(R_p) = \sum_{i=1}^n W_i E(R_i) \quad (3.6.3)$$

$$\sum_{i=1}^n W_i = 1 \quad (3.6.4)$$

n 種證券之投資組合，若由 $i=1,2,\dots,n$ 種證券組成，一個投資組合 R_i 表示第 i 個證券的報酬率、 W_i 表示第 i 個證券的持有百分比，則投資組合的報酬率為 $R_p = \sum_{i=1}^n W_i R_i$ 。當 λ 愈大，表示 $E(R_p)$ 之重要性減少，而 σ_p^2 重要性增加，也就

是投資者重視風險，故目標函數主要係用以求取 σ_p^2 極小值；若當 λ 愈小，表示

$E(R_p)$ 之重要性增加，也就是投資者重視報酬，而 σ_p^2 重要性減少，故目標函數主要係用以求取 $E(R_p)$ 之極大值。因此，只要不斷變動 λ 值，即可求得效率前緣上各項投資組合的投資比例。

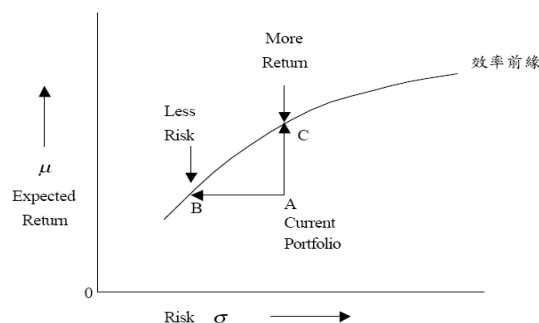


圖3.10 效率前緣

從上圖可以發現，位於效率前緣上的組合如B，與A有相同的預期報酬率，但是風險卻比A低；以C而言，預期報酬率比A高，但是風險卻與A相同。

肆、動態灰粗集合預測模型之建構

一、模型建構流程圖

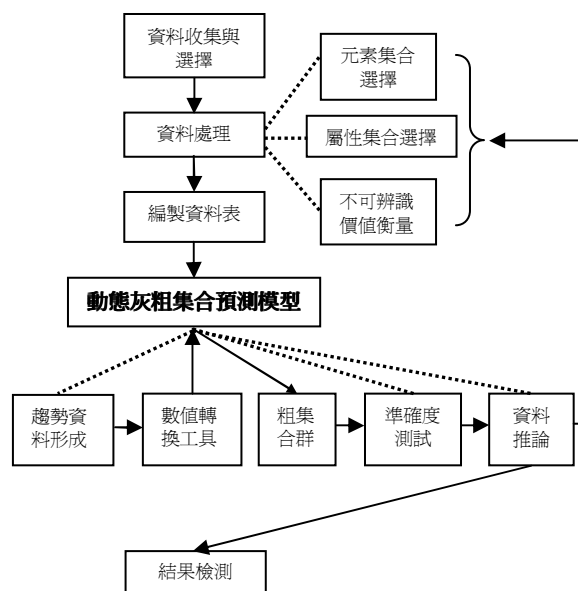


圖 4.1 動態灰粗集合預測模型建構流程圖

二、模型建構流程說明

- (一)決定探討的主題，收集相關資訊並選擇適當的資料庫。
- (二)確定主題的範圍與相關屬性，初步評估預期的結果與假設實用價值。
- (三)將主題與相關屬性編製成決策資料表。

- (四)將編製的決策資料表輸入動態灰粗集合模型。
- (五)使用 **K-Means** 分群工具將動態灰粗集合模型輸出的動態趨勢資料進行數值轉換，將轉換後資料表進行粗集合篩選。
- (六)利用重要性(**Significance**)刪減條件屬性。
- (七)計算相對分類誤差。
- (八)合併決策屬性。
- (九)篩選 **DGVPRS-Model** 之近似集。
- (十)以現期營收成長率與上期比較-增減(%)及去年同時期比較-增減(%)皆需大於 0 作為進一步篩選。
- (十一)以該股前四期(兩年)股價報酬率做馬克維茲(1952)的平均數一變異數前緣。

三、動態灰粗集合預測模型程式之撰寫

在動態灰粗預測模型建構後，依序依據 **K-means** 分群規則與粗集合原始理論，使用工程計算套裝軟體 **MATLAB** 中建立計算模型的程式，建構一套以基本面為基礎之投資組合初步篩選模型。先對各項灰預測後財報資料分群，再用粗集合理論篩選出模糊性與難以辨識性的集合歸屬。並以粗集合所篩選之下近似集合作為投資組合初步篩選標的，下近似集合中之子集合完全符合決策屬性之嚴格條件，由下近似集合所篩選之投資標的因而具有穩定度與安全性。

伍、實證分析

動態灰粗集合預測模型實證部分，是以香港 **H**、紅籌類股上市公司進行模型實證應用(如圖 5.1)。本研究以現有的台灣經濟新報資料庫(**TEJ**)，作為研究樣本資料來源。

實證的主要目的：測試動態灰粗集合預測模型之實用價值，是以動態灰粗集合中的下近似集合所篩選出之投資標的進行投資，以投資組合績效評估值作為模型實用價值衡量。

一、實證流程

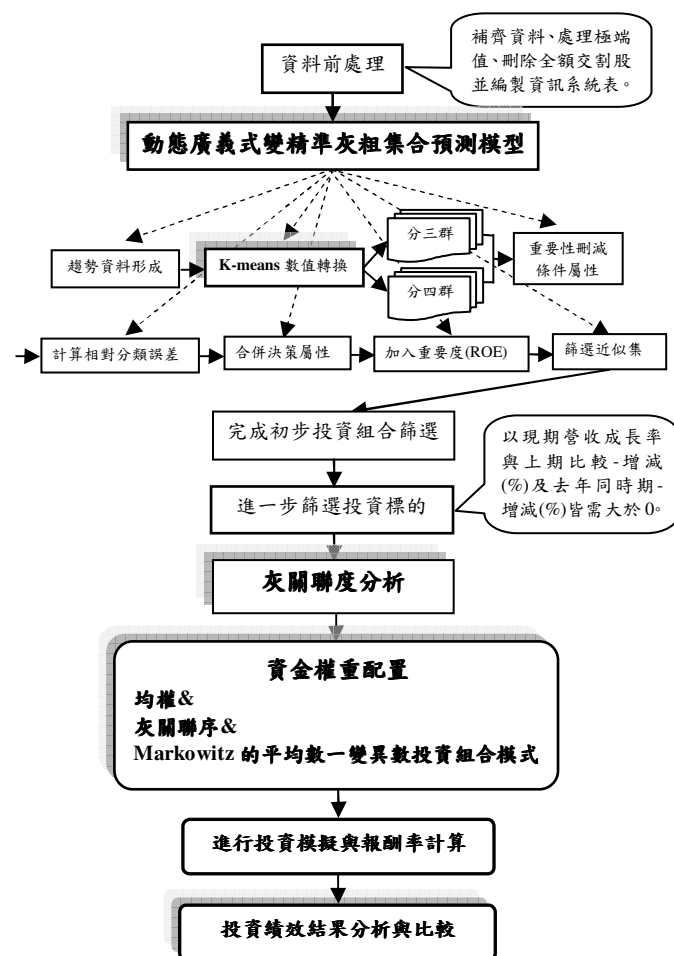


圖 5.1 動態灰粗集預測模型建構流程圖

二、實證步驟說明

(一)將2005~2007年之公司財務比率指標半年報資料（17項條件屬性、4項決策屬性、1項重要性屬性）轉入TG-Rough Sets Model；針對2005/12、2006/06、2006/12、2007/12共4期半年報績效指標進行趨勢動態預測。本模型灰預測是以過往4期半年報資料實際值對未來第5期半年報進行預測(如表5.1)。

表 5.1 趨勢資料形成之時間點配置圖

原始資料	期	預測之時間點
2005/07~2005/12	1	}
2006/01~2006/06	2	
2006/07~2006/12	3	
2007/01~2007/06	4	
2007/07~2007/12	5	預測出 2007/12 半年報數據

(二)將未來第 5 期半年報趨勢資料中的 17 項條件屬性特徵值利用 K-means 分群工具進行數值轉換(如圖 5.2)。

- (三)知識(屬性)約簡是粗集合理論的核心內容之一，本研究以決策屬性依賴條件屬性的程度，來作為刪減條件屬性的依據，也就是先計算出每個條件與決策屬性間的重要性，之後再將重要性較低的屬性給予刪除。研究以重要性的門檻為 0.2，經過重要度刪減條件屬性從 17 項降為 12 項(如圖 5.3)。
- (四)傳統粗集合理論在篩選近似集時，常因存在干擾(Noise)而使具有潛力的股票，被錯誤歸類到邊界集合中，為了解決此問題本研究在模型中引入了相對分類誤差的概念(Relative classification error)，且可分為正的相對分類誤差與負的相對分類誤差。
- (五)使用類神經模糊理論中的可調式類神經模糊推論系統(ANFIS)，來作為合併決策屬性的工具，試圖將不確定資訊系統(UIS)中的多個決策屬性，合併為一個最重要的決策屬性，主要目的是想要解決傳統粗集合理論中，決策屬性需設立門檻值的缺點(如圖 5.4)。
- (六)利用本研究建構的模型，篩選出正的下近似集合(POSp)、負的下近似集合(NEGn)、正的上近似集合(UPPp)、負的上近似集合(UPPn) 與邊界集合，而本研究以正的下近似集合(POSp)為投資標的(如圖 5.5)。
- (七)進行趨勢粗集合的準確度與資料推論，分別 對每一季預測值 K-means 分三群 (K=3)、四群 (K=4) 資料所產生之上近似集合、下 近似集合與邊界集合分別進行資料之推論。
- (八)將趨勢粗集合中的下近似集合，以現期營收和上期比較-增減(%)及去年同期比較-增減(%)需同時大於 0 進行進一步篩選。
- (九)以該股前四期(兩年)股價報酬率作馬可維茲效率前緣風險與報酬組合的選擇，並與灰關聯系數排序的資金權重配置以及均權的資金配置，作為探討投資資金規模大小、投資標的多寡與投資標的之權重配置三種不同資金配置下投資組合報酬率之差異程度。
- (十)決定買賣時點，並進行實際模擬投資。



圖 5.2 K-means 分群工具數值轉換

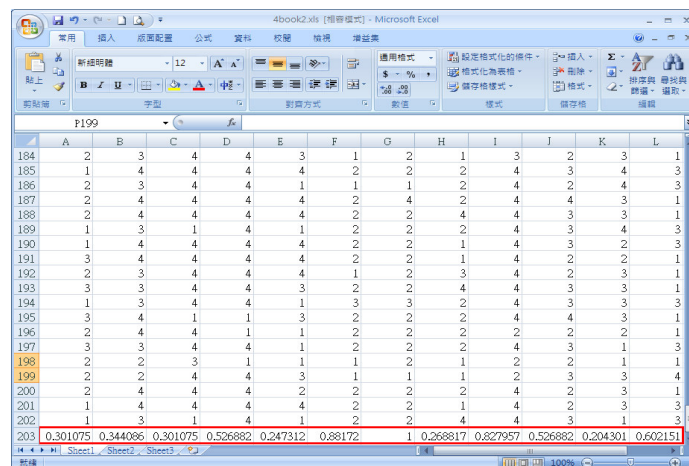


圖 5.3 各條件屬性之重要性

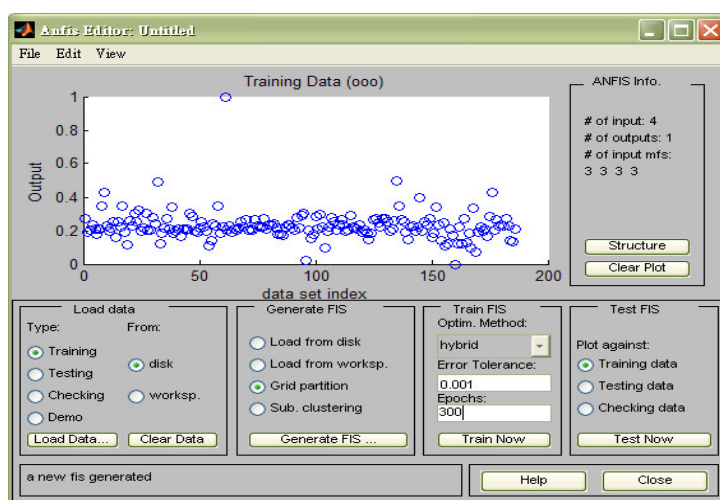


圖 5.4 可調式類神經模糊推論系統



圖 5.5 廣義式變精準 Rough Set 理論運算結果

三、實證結果

(一)資料推論

實證的結果：特徵值（R）在 K-Means 分爲四群時呈現較良好的粗集合準確度(如表 5.2)，所以在資料分群上是採取 K=4 之分群結果進行粗集合篩選。

表 5.2 灰粗集合模型之準確度

2007 年 12 月	K-Means 分群	
	K=3	K=4
準確度	0.28261	1
模糊度	0.71739	0

(二)模擬投資組合與大盤之投資績效比較

實證結果發現，於實際投資模擬期間，本研究動態灰粗集預測模型篩選出六家公司，其投資報酬率分別爲滬杭甬高速20.23%、華潤電力22.5%、中國海洋石油39.35%、中國移動30.09%、中國神華37.87%、中海油田服務18.67%皆高於大盤，且在投資權重配置所得的報酬與大盤相較下，馬可維茲最高32.55%(如表 5.5)，灰關聯序次之28.53%(如表5.4)，其次爲均權28.53%(如表5.3)，最低爲香港恒生指數15.96%。

表5.3 均權權重配置之投資組合績效比較

公 司 (代 號)	實際模擬投資時間：2007 年 9 月～12 月			
	權重	報酬率	加權報酬率	總報酬率
滬杭甬高速(0576)	16.67%	20.23%	3.37%	28.12%
華潤電力(0836)	16.67%	22.50%	3.75%	
中國海洋石油(0883)	16.67%	39.35%	6.56%	
中國移動(0941)	16.67%	30.09%	5.02%	
中國神華(1088)	16.67%	37.87%	6.31%	
中海油田服務(2883)	16.67%	18.67%	3.11%	
恒 生 指 數 (H S I)	100%	15.96%	15.96%	15.96%

表5.4 灰關聯序權重配置之投資組合績效比較

公 司 (代 號)	實際模擬投資時間：2007 年 9 月～12 月			
	權重	報酬率	加權報酬率	總報酬率
滬杭甬高速(0576)	9.52%	20.23%	1.93%	28.53%
華潤電力(0836)	19.05%	22.50%	4.29%	
中國海洋石油(0883)	14.29%	39.35%	5.62%	
中國移動(0941)	4.76%	30.09%	1.43%	
中國神華(1088)	28.57%	37.87%	10.82%	
中海油田服務(2883)	23.81%	18.67%	4.44%	
恒 生 指 數 (H S I)	100%	15.96%	15.96%	15.96%

表5.5 馬可維茲權重配置之投資組合績效比較

公 司 (代 號)	實際模擬投資時間：2007 年 9 月～12 月			
	權重	報酬率	加權報酬率	總報酬率
滬杭甬高速(0576)	0.00%	20.23%	0.00%	32.55%
華潤電力(0836)	5.35%	22.50%	1.20%	
中國海洋石油(0883)	10.58%	39.35%	4.16%	
中國移動(0941)	59.73%	30.09%	17.97%	
中國神華(1088)	24.34%	37.87%	9.22%	
中海油田服務(2883)	0.00%	18.67%	0.00%	15.96%
恒生指數(HSI)	100%	15.96%	15.96%	

綜合上述之差異比較，動態灰粗集合預測模型與平均數一變異數前緣的資金配置方式，擁有較佳投資組合績效以作為香港股市最後實證結果(如圖5.6)。

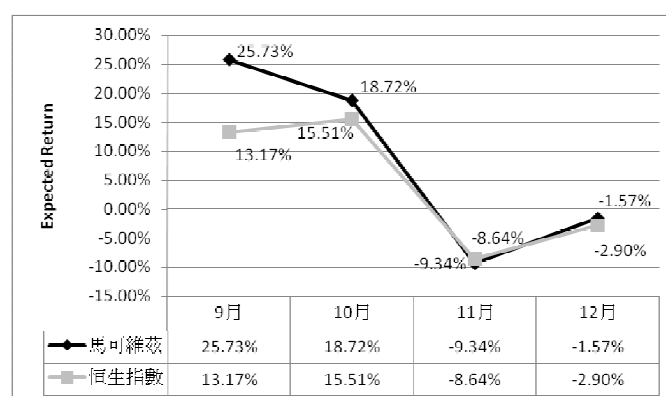


圖5.6 馬可維茲之投資績效比較

陸、結論

運用廣義式變精準動態灰粗集合預測模型可快速預測出趨勢性的三大集合與變化，投資者可適時更新投資標的與避險提高投資決策品質與投資績效，並與平均數一變異數前緣的資金配置方式具有以下幾點特性：

- 一、動態灰粗集合預測模型同時兼具粗集合模型中屬性刪減、辨識與灰趨勢預測等功能。
- 二、本模型解決了傳統粗集合設立門檻值的問題，並經由重要度刪減不必要的雜訊，使模型更能精準篩選出優質公司。
- 三、在平均數一變異數前緣資金得知投資公司 家數越多越是增加規避風險性，且報酬率亦會跟著升高的可能。
- 四、於香港股市利多、利空之長期投資情勢下，本模型篩選之投資組合明顯優於大盤；經動態灰粗集合預測模型層層篩選下，結合平均數一變異數前緣的資金配置方式，擁有較佳的投資組合績效。
- 五、無須考量各子集間關聯性，適用於確定性或不確定性系統資料，兼具資料

探勘、資訊預警功能。

六、此模式可應用領域相當的廣泛，例如投資避險、企業動態安全存量模式建立、生態環境評估、特殊疾病爆發的預測與評估…等方面，資料分析迅速，模型建構簡單、易懂。

參考文獻

- 李慧慈 (2003) 「利用粗集合論預測網路銀行使用意願」，南台科技大學國際企業系研究所碩士論文。
- 吳漢雄、鄧聚龍、溫坤禮，1996，「灰色分析入門」，高立出版。
- 吳柏林 (1996)，非線性時間數列轉折區間之模糊分類研究，國科會研究報告。財務金融研究中心，「投資分析+Matlab 應用」，全華科技圖書。
- 許展碩 (2003) 「一個複合型的演化式計算：基因演算法結和約略集合理論」暨南國際大學資訊管理學系研究所碩士論文。
- 施宜協(2003)，「運用灰預測、灰色馬可夫與適應性模糊類神經推論系統於市場中立型避險基金建構之研究」，國立台灣科技大學資訊管理研究所碩士論文。
- 劉淑賢 (2003) 「一個新穎並結合約略集合理論的精簡控制」暨南國際大學資訊管理學系研究所碩士論文。
- 鄧聚龍，1987，「灰色系統基本方法」，華中理工大學出版社。
- 鄧聚龍、郭洪，1996，「灰色原理與應用」，全華科技圖書公司。
- Ahn, B.S., Cho, S., Kim, C., 2000. The integrated methodology of rough set theory and artificial neural network for business failure prediction. *Expert Systems with Applications* 18, 65–74.
- Bazan, J.G., Skowron, A., Synak, P., 1994. Market data analysis: A rough set approach. *ICSRsearch Reports* 6/94, Warsaw University of Technology.
- Brachman, R., Khabaza, T., Kloesgen, W., Piatetsky-Shapiro, G., Simoudis, E., 1996. Mining business databases. *Communications of ACM* 39 (11), 42–48.
- Baltzersen, J.K., 1996. An attempt to predict stock market data: a rough sets approach. Diploma Thesis, Knowledge Systems Group, Department of Computer Systems and Telematics, The Norwegian Institute of Technology, University of Trondheim.
- Bouqata, B.; Bensaid, A., Palliam, R., Gomez Skarmeta, A.F., “Time series prediction using crisp and fuzzy neural networks: a comparative study,” *The IEEE/IAFE/INFORMS 2000 Conference on Computational Intelligence for Financial Engineering*, pp.170 –173,2000.
- Dimitras, A.I., Slowinski, R., Susmaga, R., Zopounidis, C., 1999. Business failure prediction using rough sets. *European Journal of Operational Research* 114, 263–280.
- Eiben, A.E., Euverman, T.J., Kowalczyk, W., Slisser, F., 1998. Modelling customer

- retention with statistical techniques, rough data models and genetic programming. In: Skowron, A., Pal, S.K. (Eds.), *Rough-Fuzzy Hybridization: A New Trend in Decision-Making*. Springer, Berlin, pp. 330–348 (Chapter 15).
- Golan, R., Edwards, D., 1993. Temporal rules discovery using datalogic/R+ with stock market data. In: Ziarko, W. (Ed.), *Rough Sets, Fuzzy Sets and Knowledge Discovery, Proceedings of the International Workshop on Rough Sets and Knowledge Discovery (RSKD'93)*, Banff, Alberta, Canada, October 12–15. Springer, Berlin.
- Golan, R., 1995. Stock market analysis utilizing rough set theory. Ph.D. Thesis, Department of Computer Science, University of Regina, Canada.
- Greco, S., Matarazzo, B., Slowinski, R., 1998. A new rough set approach to evaluation of bankruptcy risk. In: Zopounidis, C. (Ed.), *Operational Tools in the Management of Financial Risks*. Kluwer Academic Publishers, Dordrecht, pp. 121–136.
- Hampton, J., 1998a. Rough set theory: The basics (part 2). *Journal of Computational Intelligence in Finance* 6 (1), 40–42.
- Jang, G. S. (1991). An intelligent trend prediction and reversal recognition system using dual-module neural networks. *The first International Conference on Artificial Intelligent Applications on Wall Street*, 26(24), 54–59.
- Kimoto, T., & Asakawa, K. (1990). Stock market prediction system with modular neural networks. *The Mit Press*, 25(23), 36–43.
- Kowalczyk, W., 1998a. Rough data modelling: A new technique for analyzing data. In: Polkowski, L., Skowron, A. (Eds.), *Rough Sets in Knowledge Discovery*, vol. 1. Physica-Verlag, Wurzburg, pp. 400–421 (Chapter 20).
- Kowalczyk, W., Slisser, F., 1997. Modelling customer retention with rough data models. In: Komorowski, J., Zytrowski, J. (Eds.), *Principles of Data Mining and Knowledge Discovery, Proceedings of the First European Symposium, PKDD'97*, Trondheim, Norway, June 24–27, pp. 4–13.
- Kowalczyk, W., Piasta, Z., 1998. Rough set-inspired approach to knowledge discovery in business databases. In: Wu, X., Kotagiri, R., Korb, K.B. (Eds.), *Research and Development in Knowledge Discovery and Data Mining: Proceedings of the Second Pacific-Asia Conference, PAKDD-98*, Melbourne, Australia, April 15–17, pp. 186–197.
- Lin, T.Y., Tremba, A.J., 2000. Attribute transformations on numeric databases and its applications to stock market and economic data. In: Terano, T., Liu, H., Chen, L.P. (Eds.), *Proceedings of the 4th Pacific-Asia Conference, PAKDD 2000*, Kyoto, Japan, April 18–20, 181–192.
- Mills, D., 1993. Finding the likely buyer using rough sets technology. *American Salesman* 38 (8), 3–5.

- Mrozek, A., Skabek, K., 1998. Rough sets in economic applications. In: Polkowski, L., Skowron, A. (Eds.), *Rough Sets in Knowledge Discovery*, vol. 2. Physica-Verlag, Wurzburg, pp. 238–271 (Chapter 13).
- Poel, D., Piasta, Z., 1998. Purchase prediction in database marketing with the probrough system. In: Polkowski, L., Skowron, A. (Eds.), *Rough Sets and Current Trends in Computing, Proceedings of the First International Conference, RSCTC'98*, Warsaw, Poland, June 22–26, pp. 593–600.
- Poel, D., 1998. Rough sets for database marketing. In: Polkowski, L., Skowron, A. (Eds.), *Rough Sets in Knowledge Discovery*, 2. Physica-Verlag, Wurzburg, pp. 324–335 (Chapter 17).
- Ruggiero, M., 1994a. Rules are made to be traded. *AI in Finance* (Fall), 35–40.
- Ruggiero, M., 1994b. How to build a system framework. *Futures* 23 (12), 50–56.
- Ruggiero, M., 1994c. Turning the key. *Futures* 23 (14), 38–40.
- Slowinski, R., Stefanowski, J., 1994. Rough classification with valued closeness relation. In: Diday, E. et al. (Eds.), *New Approaches in Classification and Data Analysis*. Springer, Berlin, pp. 482–488.
- Szladow, A., Mills, D., 1993. Tapping financial databases. *Business Credit* 95 (7), 8.
- Skalko, C., 1996. Rough sets help time the OEX. *Journal of Computational Intelligence in Finance* 4 (6), 20–27.
- Susmaga, R., Michalowski, W., Slowinski, R., 1997. Identifying regularities in stock portfolio tilting. Interim Report, IR-97-66, International Institute for Applied Systems Analysis.
- Slowinski, R., Zopounidis, C., Dimitras, A.I., 1997. Prediction of company acquisition in greece by means of the rough set approach. *European Journal of Operational Research* 100, 1–15.
- Slowinski, R., Zopounidis, C., Dimitras, A.I., Susmaga, R., 1999. Rough set predictor of business failure. In: Ribeiro, R.A., Zimmermann, H.J., Yager, R.R., Kacprzyk, J. (Eds.), *Soft Computing in Financial Engineering*. Physica-Verlag, Wurzburg, pp. 402–424.
- Tay, Francis E.H., Shen, Lixiang. Economic and .nancial prediction using rough sets model, *European Journal of Operational Research* 141 (2002) 641–659.
- Tou, J.T. and Gonzalez, R.C. (1974). *Pattern Recognition Principles*, Addison-Wesley, Reading, MA.
- Ziarko, W., Golan, R., Edwards, D., 1993. An application of datalogic/R knowledge discovery tool to identify strong predictive rules in stock market data. In: *Proceedings of AAAI Workshop on Knowledge Discovery in Databases*, Washington, DC, pp. 89–101.