

2010 TOPCO 崇越論文大賞

論文題目：

混合式及具權重變化之管制圖型樣辨識
-ICA與SVM之整合應用

報名編號： _____F0035_____

混合式及具權重變化之管制圖型樣辨識 -ICA與SVM之整合應用

Recognition of the Mixture Control Chart Patterns Using ICA/SVM Schemes

摘要

管制圖型樣辨識在現今的工業製程中已成為不可或缺的監控技術之一，其主要功能為補強監控管制圖所不易偵測之不尋常型樣，透過辨識不尋常型樣之種類，可探討出其製程失控之可歸屬原因，藉以改善製程減少不良品的產生。目前大多數的研究，大都假設製程上所監控的觀察值是管制圖型樣中的基本型樣，並且這些基本型樣的權重都是固定的；然而，一般實務製程所監控到的觀測值，卻可能是由兩個基本型樣混合組成，且這些基本型樣本身可能存在著權重變化。為了辨識這些混合型態及具權重變化的管制圖型樣，本研究整合獨立成份分析(independent component analysis, ICA)與支援向量機(support vector machine, SVM)，提出一個有效的管制圖型樣辨識架構。本研究所提之方法首先利用 ICA 將混合或具權重變化之管制圖型樣分解成兩個獨立成分(independent component, IC)，接著在獨立成份中找出代表基本型樣的 IC，最後將這個 IC 作為 SVM 之輸入變數，以建構型樣辨識模型。實驗結果顯示，本研究所提出的整合 ICA 與 SVM 模式，可以有效的辨識混合管制圖型樣及權重變化型樣。

關鍵詞：管制圖型樣、獨立成份分析、支援向量機、型樣辨識

混合式及具權重變化之管制圖型樣辨識 -ICA與SVM之整合應用

壹、緒論

如何在製程失控後，盡可能的越早察覺出問題所在，是製程監控中很重要的一環。在過去數十年中，有許多方法或工具成功的被用在製程監控上，其中，統計製程管制(statistical process control, SPC)為最常使用來偵測和改善製程品質的方法之一。管制圖在工業製程中可以幫助減少變異以及藉著改善產品品質來增加競爭力，是 SPC 中有效的工具之一。

當管制圖中的一個樣本點落在管制界線外，或者是管制圖中存在不尋常(abnormal)的趨勢型樣(pattern)時，即表示此製程有失控的可能。有效的辨識管制圖型樣(control chart patterns, CCPs)在 SPC 是很重要的一項議題，因為我們可以利用這些不尋常的管制圖型樣來尋找製程出錯的可歸屬原因。根據統計品質管制手冊(Western Electric, 1958)經常討論的型樣有 8 種(Wang and Kuo, 2007; Gauri and Chakraborty, 2006, 2009)，如圖 1 所示，包括(1)正常型樣(normal, NOR)，也就是正常無失控的圖形型樣；(2)變異數位移型樣(stratification, STA)；(3)系統性位移型樣(systematic, SYS)；(4)週期性位移型樣(cyclic, CYC)；(5)遞增型樣(increasing trend, UT)；(6)遞減型樣(decreasing trend, DT)；(7)向上位移型樣(upward shift, US)及(8)向下位移型樣(downward shift, DS)，若製程中存在上述(2)-(8)型樣，即表示此製程是處於失控狀態。

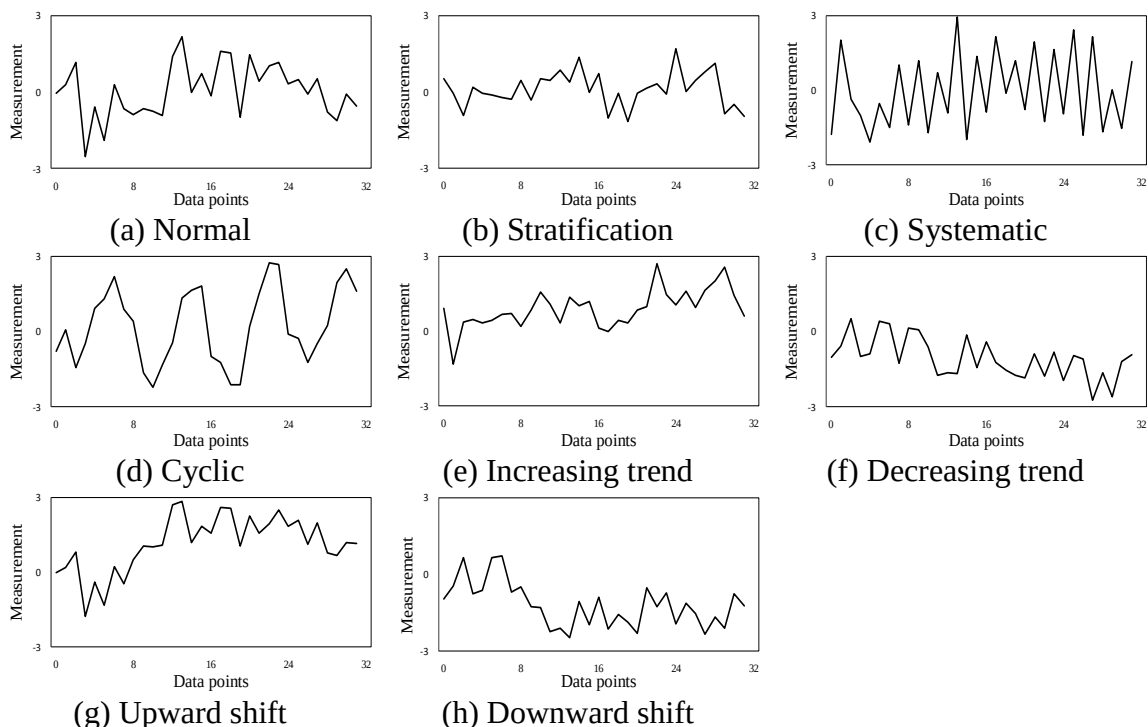


圖 1 八種常見的管制圖型樣

在辨識管制圖形樣式有相當多的研究文獻，部份研究(Guh, Zorriassatine, Tannock, and O'Brien, 1999; Perry, Sporre and Velasco, 2001; Guh and Shiue, 2005)應用類神經網路(artificial neural networks, ANNs)來辨識管制圖型樣，但類神經網路最大的缺點在於它必須具有大量的參數設定，較難有穩定的結果，而且模型可能會發生過度配適的情況；後續研究為了解決以上問題，提出結合兩種或兩種以上的人工智慧方法和預測模型，如：Wang and Kuo (2007)結合小波過濾(wavelet filter)和模糊分類(fuzzy clustering)用來辨識六種常見的管制圖型樣，相較於類神經網路的運算方法更有效率且較準確，但大多數的研究都只能分辨單一基本型態的不尋常管制圖型樣，原因在於這些研究均假設所偵察到的觀測值必須是常見 CCPs 中的基本型樣。然而，在實務上大多數所觀測到的管制圖可能不會只有簡單的單一基本型樣，而是會由兩個或兩個以上的基本型樣所結合而成的混合型樣，或者是單一基本型樣中具權重變化的權重變化型樣，如圖 2 所示，比較圖 1 和圖 2 的圖形樣式可知，在混合型樣和權重變化型樣中較不易辨識製程中存在哪些基本不尋常型樣。因此，如何有效地針對混合型樣及權重變化型樣進行辨識是個重要且具挑戰性的任務。

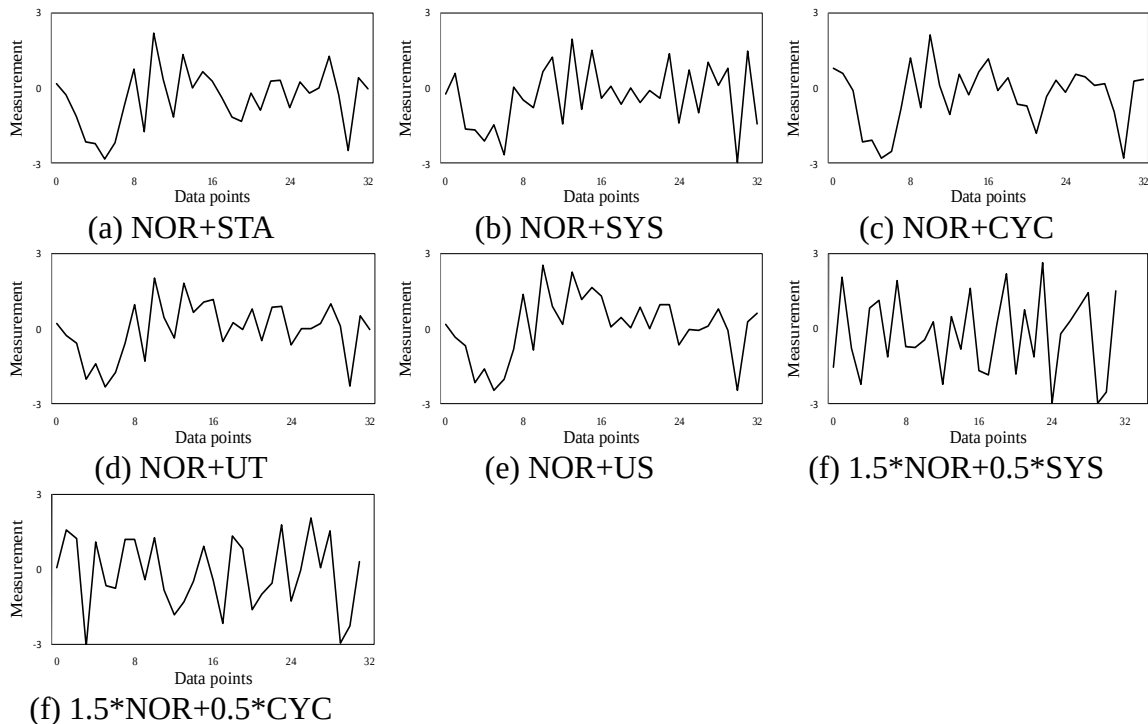


圖 2 混合型樣(a),(b),(c),(d),(e)；權重變化型樣(f),(g)

本研究提出一種結合兩種方法的模式，主要針對混合型樣和權重變化型樣，期望透過所提的方法能正確的辨識出所屬類型。此模式結合獨立成份分析(independent component analysis, ICA)和支援向量機(support vector machine, SVM)，本研究稱之為 ICA-SVM 模式。ICA 在未知來源訊號的特徵萃取上有強大的估計能力，它能在資料中從不同的觀察角度找出好的特徵值，這些特徵值經由

原始資料的線性轉換，用來使資料更加清楚的呈現出資料本身的特性。本研究選擇獨立成份分析做管制圖型樣辨識的前處理，將獨立成份分析應用在混合型樣和權重變化型樣的資料上，分離出近似統計獨立的獨立成份(independent component, IC)，透過這些估計的獨立成份找出混合型樣資料中的獨立來源，分解出隱藏在混合型樣中的基本管制圖型樣。

支援向量機是建立在統計學習理論(Statistical Learning Theory, SLT)的基礎上所發展出的一種嶄新的學習系統，本研究使用 SVM 的原因不僅在於支援向量機的運算速度較類神經網路、分類樹...等其他分類方法來的迅速之外，還可將問題轉化為一個二次型的最佳化問題，解得全域的最佳解，可以解決了在神經網路方法中無法避免的局部極值問題，並且具有多品項分類的能力，可在高維度的特徵空間中找出最佳分界線，建構多個超平面並將資料分成多類，剛好適合用於本研究中所提出的八種常見管制圖型樣。

ICA-SVM 模式目前尚未有其他研究嘗試使用在型樣辨識上，本研究主要利用獨立成份分析的前處理，估計出代表潛在基本型樣的獨立成份，以達到加快支援向量機的分類速度，比起直接將混合型樣和權重變化型樣放入支援向量機作分類的作法，更能提高其分類正確率，如此一來，可盡早發現失控製程的可歸屬原因，減少製程監控成本並提高生產品質。本篇論文架構組織依序是第 2 段：研究方法 ICA、SVM 的簡短回顧與研究流程；第 3 段：實驗結果的呈現；第 4 段：結論與建議。

貳、研究方法與設計

一、獨立成份分析

獨立成份分析是近來發展出的新興方法，用來尋找隨機變數中隱藏因子的統計方法(Lee, 2004; Hyvärinen and Oja, 2000)，目前被廣泛的運用在分離未知來源訊號和特徵萃取的議題上。一個基本的獨立成份分析模式可以表示為方程式(1)，如下所示：

$$\mathbf{X} = \mathbf{AS} \quad (1)$$

其中 $\mathbf{X}$ 為維度 $M \times N$ 的混合訊號矩陣， $\mathbf{A}$ 為維度 $M \times M$ 的未知混合矩陣(mixing matrix)， $\mathbf{S}$ 是維度 $M \times N$ 的未知來源訊號矩陣。獨立成份分析模式要解決的問題就是在只有觀察到的混合訊號矩陣 $\mathbf{X}$ 的情況下，估計出未知的混合矩陣 $\mathbf{A}$ 及未知來源訊號矩陣 $\mathbf{S}$ 。在假設個別的來源訊號 s_i (矩陣 $\mathbf{S}$ 之列向量)互為統計獨立的情況下，為了估計未知來源訊號矩陣 $\mathbf{S}$ ，獨立成份分析的方法是找到一個維度 $M \times M$ 的

解混合矩陣(demixing matrix) W 將所觀察到的混合訊號矩陣 X 進行轉換以產生維度為 $M \times N$ 的矩陣 Y ，亦即如方程式(2)所示(Hyvärinen and Oja, 2000; Hyvärinen et al., 2001)：

$$WX = Y \quad (2)$$

當解混合矩陣 W 為混合矩陣 A 的反矩陣時，即 $W = A^{-1}$ ，矩陣 Y 會等於未知來源訊號矩陣 S ，矩陣 Y 的列向量 y 稱為獨立成份，在互為統計獨立下，將可用來估計個別來源訊號 s 。

獨立成份分析方法在估計解混合矩陣 W 時，常利用各獨立成份間必須統計獨立的基本假設，將各獨立成份間的獨立性高低定義成一個目標函數，然後，藉由某一特定的最佳化技術或一非監督式的演算法，將解混合矩陣 W 估計出來(Hyvärinen and Oja, 2000; Hyvärinen et al., 2001)，而這個演算法的目標就是使得各獨立成份間的統計獨立性最大，利用此一性質，從混合訊號中分離出原始來源訊號，並藉由量化獨立性的衡量找出獨立成份。

獨立成份分析找尋獨立成分的過程可分為兩步驟，首先必須先選定衡量獨立性的準則作為目標函數，再利用最佳化的演算法求解目標函數，找出最佳解混合矩陣 W 。在衡量獨立性的準則方面，可利用獨立成份必須具有非高斯分配(non-gaussian distribution)的限制，定義獨立成份的非高斯特性(non-gaussianity)成目標函數來衡量各獨立成份間的獨立性。目前，已有許多成熟的技術被用來量測獨立成份分析中所謂獨立成份的非高斯性，例如峰態(kurtosis)、共同資訊法(mutual information)及負熵法(negentropy)等，而在這些方法之中，以負熵法最常被討論與使用(Hyvärinen et al., 2001)。

負熵法被用來測量隨機變數所能提供之資訊量與隨機變數的穩定性，表示隨機觀察值的資訊程度，當隨機變數的結構性越差越亂、越無法預測時，其熵值會越大，代表隨機變數越不穩定。

負熵 J 定義如下(Hyvärinen and Oja, 2000; Hyvärinen et al., 2001)：

$$J(y) = H(y_{\text{gauss}}) - H(y) \quad (3)$$

其中 y 為一個隨機向量，代表一個獨立成份， y_{gauss} 表示和 y 具有相同變異數的高斯變數向量，也就是具有相同之共變異數矩陣，因此負熵的值為非負值，即 $J(y) \geq 0$ ，只有當隨機向量 y 為高斯分配時，其值才會為零，所以目標函數可表示為： $\text{Max}_w J(y)$

但是負熵最大的問題就是計算太過困難。因此，Hyvärinen 提出下列的近似函數在其演算法中當目標函數：

$$J(y) \propto [E\{G(y)\} - E\{G(v)\}]^2 \quad (4)$$

其中 v 是平均值為零且變異數為 1 的高斯分佈之隨機變數。G 並沒有限制，只要是任何的二次方函數(non-quadratic function)都可以，因為 G 若為二次方函數， $J(y)$ 必為零，下列為本研究使用的函數 G：

$$G_1(y) = \frac{1}{a_1} \log \cosh(a_1 y)$$

$$G_2(y) = -\exp(-y^2/2)$$

$$G_3(y) = y^4 \quad (5)$$

在方程式(5)中， $1 \leq a_1 \leq 2$ ，通常將 a_1 設定為 1。

在獨立成份分析中有許多演算法可使目標函數最佳化，其中，FastICA 演算法是由 Hyvärinen 在 1999 年所提出且為目前最常用的演算法之一。因此，本研究使用 matlab 軟體中的 FastICA 套件，將 ICA 方法透過 FastICA 演算法作為估計資料中獨立成份的演算法，求得我們所需的獨立成份。詳細的 ICA 演算法可參照 (Hyvärinen and Oja, 2000; Hyvärinen et al., 2001)。

二、支援向量機

支援向量機是由 Vapnik 在 1995 年和 AT&T 實驗室團隊所提出的一個新方法，其主要的理論是來自統計學習理論的 VC 維理論和結構化風險最小誤差法 (structural risk minimization, SRM)，根據有限的樣本資訊在模型的複雜性(即對特定訓練樣本的學習精度, Accuracy)和學習能力(即無錯誤地識別任意樣本的能力)之間尋求最佳折衷點，以期獲得最好的推廣能力 (generalization ability)。支援向量機最初用來辨識管制圖型樣的概念可以簡述如下：第一，將輸入向量映射到一個特徵空間(可能具有高維度)，不論是線性或非線性，選擇一個相關性高的核函數，然後在特徵空間內找出一條最佳分界線，並建立一個超平面將資料分成兩類(也可分多類)。支援向量機主要是在訓練出一個全面的最佳解並且避免太多特徵值產生。

在典型的分類問題中，我們通常會定義以下基本的表示方法，首先令 $\{(x_i, y_i)\}_{i=1}^N$ ， $x_i \in R^d$ 、 $y_i \in \{-1, 1\}$ ， x_i 和 y_i 分別為訓練組的輸入向量和目標函數，

N 是觀測樣本， d 是維度。支援向量機的演算法主要是尋找一個超平面 $\mathbf{w}^T \cdot \mathbf{x}_i + b = 0$ ，其中 $\mathbf{w}$ 是超平面的權重向量， b 是偏差值，其分類器形式如方程式 (6) 所示 (Vapnik, 2000)：

$$y_i (\mathbf{w}^T \cdot \mathbf{x}_i + b) \geq 1 \quad \forall i \quad (6)$$

此超平面用最大邊際寬度 $2/\|\mathbf{w}\|$ ，來將資料分成兩類，並且在邊界的點皆稱為支援向量 (support vector)。為了尋找最適超平面，支援向量機的最佳化問題模式如下 (Vapnik, 2000)：

$$\begin{aligned} \text{Min } \Phi(\mathbf{x}) &= \frac{1}{2} \|\mathbf{w}\|^2 \\ \text{S.t. } y_i (\mathbf{w}^T \mathbf{x}_i + b) &\geq 1, \quad i = 1, 2, \dots, N \end{aligned} \quad (7)$$

要解出方程式 (7) 是不容易的，必須用 Lagrange 方法將其轉換為對偶問題來解較容易，轉換後則為最大化的模式，如下所示 (Vapnik, 2000)：

$$\begin{aligned} \text{Max } L_D &= \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i=1, j=1}^N \alpha_i \alpha_j y_i y_j \mathbf{x}_i^T \mathbf{x}_j \\ \text{S.t. } \sum_{j=1}^N \alpha_j y_j &= 0 \end{aligned} \quad (8)$$

為了精準的分出兩類，我們增加了一個寬鬆變數在 Lagrang 方程式中，使的方程式變成 $y_i (\mathbf{w}^T \mathbf{x}_i + b) \geq 1 - \xi_i$ ， $\xi_i \geq 0$ ，而寬鬆變數的目的在於增加有彈性的緩衝區界線。而方程式 (8) 的限制式中也加入了一映射函數 $K(*)$ ，稱之為核函數 (kernel function)，當對偶問題加上核函數後，方程式 (8) 會改變如下所示 (Vapnik, 2000)：

$$\begin{aligned} \text{Max } L_D &= \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i=1, j=1}^N \alpha_i \alpha_j y_i y_j K(\mathbf{x}_i, \mathbf{x}_j) \\ \text{S.t. } \sum_{j=1}^N \alpha_j y_j &= 0 \\ 0 \leq \alpha_i &\leq C, \quad i = 1, 2, \dots, N \end{aligned} \quad (9)$$

一般的核函數有線性 (linear)、多項 (polynomial)、輻射基底函數 (radial basis

function, RBF)和雙曲函數(sigmoid)，針對非線性問題所衍生的方法為後三種，但最被廣泛使用的是 RBF，且能解決大部分問題(Hsu, Lin and Lin, 2003; Cherkassky and Ma, 2004)，因此，本研究使用 RBF 作為支援向量機的核函數，其被定義為 $K(\mathbf{x}_i, \mathbf{x}_j) = \exp(-\gamma \|\mathbf{x}_i - \mathbf{x}_j\|^2)$, $\gamma \geq 0$ ，其中， γ 被標記為 RBF 的寬度，代表其設定參數，當以 RBF 作為 SVM 之核心函數時必須設定兩個參數值，除了 RBF 本身之參數 γ 外，另一個為懲罰係數 C ，其代表分類誤差的影響程度。本研究針對此兩個參數提出指數成長的序列搜尋方式作為尋找 SVM 的最佳參數方法(如：25,23,...,2-15)，藉此達到模型最佳分類效果(Hsu et al. 2003)。詳細的 SVM 演算法可參照(Hsu et al. 2003; Cherkassky and Ma, 2004)。

本研究使用臺灣大學林智仁(Chih-Jen Lin)博士等開發設計的 LIBSVM(a library for support vector machines)，其為一個通用的 SVM 套裝軟體，主要用來解決二元或多元分類的問題，因此本研究主要透過 LIBSVM 套件結合 Matlab 軟體，進行管制圖型樣分類辨識。

三、研究架構

本研究提出一個混合式型樣辨識架構，結合獨立成份分析和支援向量機，來對兩種實務製程上時常發生的混合型樣和具權重變化型樣資料進行辨識。研究流程的架構分為訓練階段和測試階段，分別重複進行十次實驗，取得平均正確率與標準差，觀察平均正確分類結果與變異大小，而訓練階段的流程如下圖 3 所示，其目的在於建構支援向量機模型來辨識管制圖型樣：

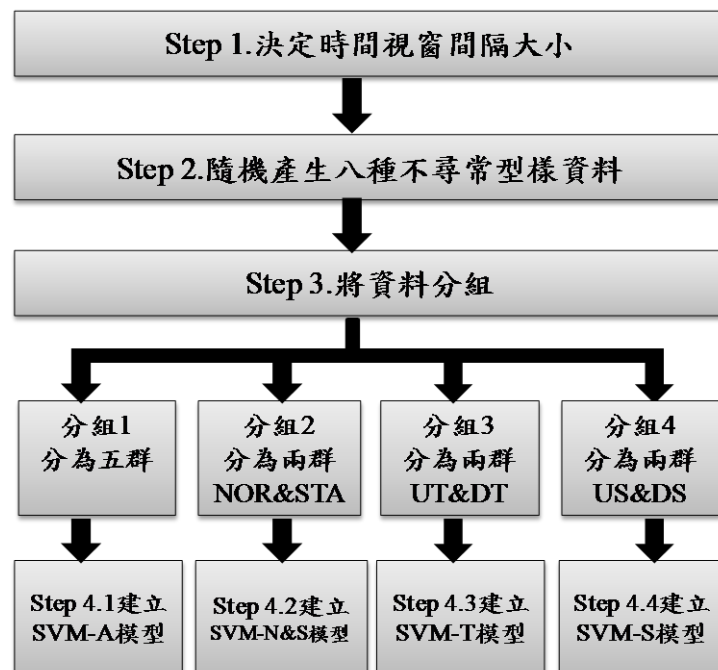


圖 3 訓練階段流程圖

Step 1.決定時間視窗間隔大小：

本研究利用時間移動窗口分析方法來模擬訓練和測試資料，根據文獻建議時間移動窗口以 32 個樣本點為一組觀測樣本，32 個樣本點能呈現所有不尋常型樣的週期，因此以 32 個樣本點為移動窗口的區間長度。(Yang and Yang, 2002, 2005; Assaleh and Al-assaf, 2005; Guh and Shiue, 2008; Gauri and Charkaborty, 2009; El-Midany et al., 2009)。

Step 2.隨機產生八種不尋常型樣資料：

透過表 1 的計算公式產生八種常見的不尋常型樣的訓練資料，分別為 NOR、STA、SYS、CYC、UT、DT、US 和 DS(Yang and Yang, 2002, 2005; Guh, 2005; Gauri and Charkaborty, 2006, 2009)，根據文獻建議參數變化在其限制範圍內透過均勻分配隨機產生(Gauri and Charkaborty, 2006, 2009)，各種型樣資料皆產生一組序列 2031 筆的製程觀測值，並利用向前遞增取樣以每 32 個樣本點為一筆模擬資料，每種型樣皆有 2000 筆，總共 16000 筆時間窗口為 32 的模擬資料。

表 1 八種型樣的計算公式

管制型樣	管制圖型樣方程式	管制圖型樣參數
NOR	$x_i = u + r_i \sigma$	平均數(u)=0 標準差(σ)=1
STA	$x_i = u + r_i \sigma'$	隨機雜訊(σ')=(0.2(σ) to 0.4(σ))
SYS	$x_i = u + r_i \sigma + d(-1)^i$	系統位移量(d)=(1(σ) to 3(σ))
Cyclic	$x_i = u + r_i \sigma + a \cdot \sin(2\pi i / t)$	振幅(a)=(1.5(σ) to 2.5(σ)) 週期(t)=(8 and 16)
Trend	$x_i = u + r_i \sigma \pm ig$	傾斜程度(g)=(0.05(σ) to 0.1(σ))
Shift	$x_i = u + r_i \sigma \pm ks$ $k=1$ if $i > P$, else $k=0$	位移量(s)=(1.5(σ) to 2.5(σ)) 位移起始點(P)=(7, 13, 19)

i 為型樣的觀測時間點($i=1 \dots 32$)， r_i 為標準常態分配在第 i 點的隨機觀測值， x_i 為型樣在第 i 點的觀測值。(本表參考文獻 Gauri & Charkaborty, 2009)

Step 3.將產生的八種型樣資料重新組合分群：

此步驟是因為測試階段時，會利用獨立成份作為支援向量機的輸入資料，但由於獨立成份分析在估計 IC 的過程中會使的各個輸入變數的平均數皆為 0，標準差皆為 1，如此導致 NOR 和 STA 透過 ICA 分解後的獨立成份會無法辨識，因為 NOR 和 STA 之間的差別就在於標準差的變化。此外，由於 ICA 所估計的 IC 有排序和正負號相反的問題產生，導致本研究估計後的兩個獨立成分有排序上的問題，例如：型樣 UT、DT 和 US、DS 無法辨識的狀況。為了解決上述的問題，因此在訓練階段時，先將隨機產生的八種常見的不尋常型樣資料分成下列四種類

型，並用以建構四個不同的 SVM 分類器，其分組方式如下：

分組 1：將八種型樣資料重新組合分成五群，分別為 NOR 和 STA 歸類為同一群，合稱 N&S 型樣；SYS 一群；CYC 一群；UT 和 DT 視為同一群，合稱 Trend 型樣；US 和 DS 視為同一群，合稱 Shift 型樣。分組的目的在於將相似的型樣合併，建構 SVM-A 的模型，先將差異較大的型樣分類出來，之後再分類其他型樣。

分組 2：將型樣資料中的 NOR 和 STA 獨立出來分成兩群，以建構用來辨識 NOR 和 STA 的 SVM-N&S 模型。

分組 3：將型樣資料中的 UT 和 DT 獨立出來分成兩群，以建構用來辨識 UT 和 DT 的 SVM-T 模型。

分組 4：將型樣資料中的 US 和 DS 獨立出來分成兩群，以建構用來辨識 US 和 DS 的 SVM-S 模型。

Step 4. 建立型樣辨識分類模型：

本研究利用支援向量機分類器來建立模型，核函數設定為非線性的 RBF，並討論在 RBF 模式下最具影響力的兩個參數 C 和 γ 。本研究參考文獻 Hsu et al. (2003) 所提出的參數設定規則，直接利用 C 和 γ 的指數遞增序列來定義參數值，例如： $C = 2^{-5}, 2^{-3}, 2^{-1}, \dots, 2^{15}$ ，最後依正確分類率來決定使用哪組最佳參數，然後再依最佳參數所建構的 SVM 模型丟入測試階段，用來辨識管制圖型樣。此階段主要建立上述四種 SVM 分類器，流程如下所述：

Step 4.1 建立 SVM-A 分類器：

將分組 1 的資料透過標準化處理，目的在於將資料整合、去除離群值的影響和凸顯型樣的型態走勢，另外，在測試階段時因為 ICA 會把測試資料的變異簡化為 1，因此為了訓練出可辨識 ICA 簡化後的資料，本研究將每個樣本點減去其平均數再除上標準差，其計算公式如下： $Z = \sum x_i - \mu / \sigma$ ，其中 x_i 代表第 i 筆觀測值， μ 為平均數， σ 為標準差。其餘 3 組資料則不需要標準化，因為其餘 3 組資料主要用來辨識相似的型樣，且其測試階段時的資料沒有透過 ICA 處理，因此，為了不失去太多重要訊息，不需要將其餘 3 組資料標準化。分組 1 的資料標準化後，將資料用來建立可同時分類全部型樣的模型，但建模時，為了不失去太多重要訊息，因此將標準化後的全數 32 個樣本點作為訓練模型的輸入變數，作為 SVM 的訓練樣本，並建立模型，命名為 SVM-A，適用於任何型樣的測試資料，因此最先用來辨識型樣資料，主要可將型樣辨識成五群，分別為 N&S、SYS、CYC、UT&DT 和 US&DS。

Step 4.2 建立 SVM-N&S 分類器：

分組 2 的資料為了分類 NOR 和 STA，文獻建議特徵值指標變數 ADIST(average distance between the consecutive points in terms of sd)，此指標可將 STA 的型樣資料凸顯出來，STA 的 ADIST 值會大於其他型樣，因此本研究將此指標作為辨識 NOR 和 STA 的輸入變數，可用來提高 NOR 和 STA 的正確分類率，並且有效的辨識兩種型樣間的差異，因此本研究利用 32 個樣本點算出指標變數 ADIST，其計算公式如下(Gauri and Charkaborty, 2006, 2009)：

$$ADIST = \left\{ \sum_{i=1}^N \left[(t_{i+1} - t_i)^2 + (x_{i+1} - x_i)^2 \right]^{1/2} / (N-1) \right\} / SD \quad (10)$$

其中 $t_i = ic$, ($i=1, 2, \dots, N$) 代表觀察值中第 i 個時間點與起始點的距離， c 為一常數，用來表示管制圖上採樣間隔的線性距離， x_i 為第 i 個時間點上所監察到的觀測值， N 代表的是觀察視窗的大小， SD 為標準差。最後，將產生的 NOR 和 STA 兩群資料以特徵值 ADIST 為輸入變數，作為 SVM 的訓練樣本，並建立模型，命名為 SVM-N&S，主要用來辨識 NOR 和 STA，解決 SVM-A 模型不能辨識 NOR 和 STA 的問題。

Step 4.3 建立 SVM-T 分類器：

因分組 3 的資料，其主要的差異僅有上跟下的正負號問題，因此，本研究直接以 32 個樣本點為輸入變數，作為 SVM 的訓練樣本，並建立模型，命名為 SVM-T，主要用來辨識 UT 和 DT，解決 SVM-A 模型不能辨識 UT 和 DT 的問題。

Step 4.4 建立 SVM-S 分類器：

因分組 4 的資料，其主要的差異僅有上跟下的正負號問題，因此，本研究直接以 32 個樣本點為輸入變數，作為 SVM 的訓練樣本，並建立模型，命名為 SVM-S，主要用來辨識 US 和 DS，解決 SVM-A 模型不能辨識 US 和 DS 的問題。

測試階段的研究流程如下圖 4 所示，主要用來測試混合型樣以及權重變化型樣資料用於本研究所提出的辨識模式之分類效果：

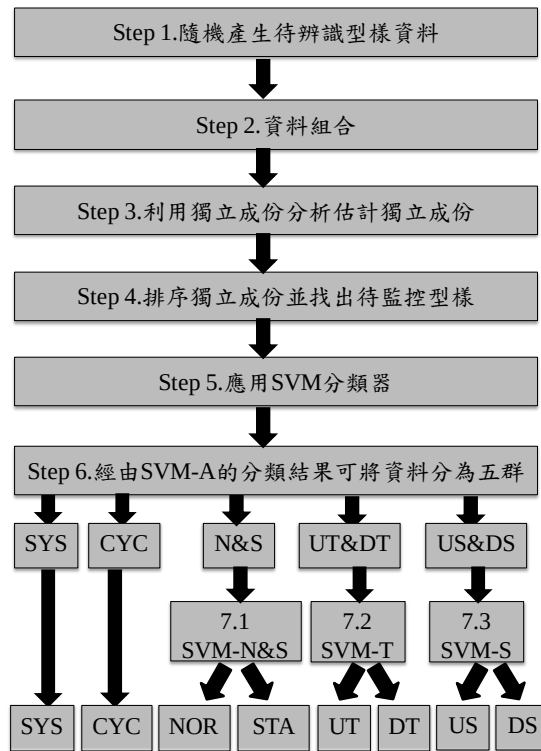


圖 4 測試階段流程圖

Step 1.隨機產生待辨識型樣資料：

由於本研究針對兩種問題，所以本實驗透過表 1 的計算公式分別產生 16000 筆的混合型樣資料和 16000 筆的權重變化型樣資料作為測試之用。混合型樣主要透過公式 14 將兩個單一型樣資料同時發生的情形模擬出來，而權重變化型樣則利用常態資料與干擾項間的權重 θ_1 和 θ_2 的變化來模擬其資料。所有的測試資料之參數變化、時間窗口和取樣方式接同訓練階段。

Step 2.資料組合：

獨立成份分析 ICA 在估計獨立成份時，必須有兩組觀測值，才能估計出兩個獨立成分，因此，本研究在針對權重變化型樣時，自行隨機產生一組常態資料 x_1^T ，分別搭配各種權重變化型樣資料 x_2^T ，形成多變量的測試資料 $X^T = (x_1^T, x_2^T)^T$ ，方能丟入獨立成份分析。另外，本研究針對混合型樣時，則不需產生一組常態資料來組合，因為混合型樣在混合過程中會產生兩組觀測值，為多變量製程。

Step 3.利用獨立成份分析估計獨立成份：

透過獨立成份分析，可將測試樣本 $X^T = (x_1^T, x_2^T)^T$ 估計出兩個獨立的成份 $Y = (y_1, y_2)^T$ ，分別為獨立成份 y_1 和 y_2 。

Step 4.排序獨立成份並找出待監控型樣：

根據文獻可知，峰態係數在標準常態資料下的峰態係數值會等於零；反之，非常態資料的峰態係數值則會大於或小於零(Hyvärinen and Oja, 2000; Hassan et al., 2003; Lee et al, 2004; Wang et al, 2009)，峰態係數公式為

$$Kurtosis = E\{y^4\} - 3(E\{y^2\})^2 \quad (11)$$

因此，本研究透過 Y_1 和 Y_2 的峰態係數來排序資料，可得知峰態係數不為零者為非常態資料，也就是本研究所要辨識的型樣資料。

Step 5.應用 SVM 分類器：

本研究利用支援向量機分類器來辨識不尋常混合型樣，將待監控制程資料丟入所建構的最佳參數 SVM 模型，來辨識管制圖型樣。

Step 6.經由 SVM-A 模型的分類結果可將資料分為五群：

首先將所有待監控型樣的獨立成份資料丟入訓練模型 SVM-A 中，可分類出五群不尋常型樣，分別為 NOR&STA 為一群，稱為 N&S 型樣；SYS 型樣；CYC 型樣；UT&DT 為一群，稱為 Trend 型樣；US&DS 為一群，稱為 Shift 型樣。

Step 7.將未分類之型樣資料分別丟入其餘三個分類器：

透過 Step 6 後，我們已成功的將 SYS 和 CYC 分類出來，尚未分類完成的有：N&S 型樣、Trend 型樣和 Shift 型樣。因此針對這些相似的型樣，本研究利用訓練好的分類模型分別對其分類，如以下步驟所示：

Step 7.1: 將 N&S 型樣利用訓練模型 SVM-N&S 來辨識，可分類出 NOR、STA。

Step 7.2: 將 Trend 型樣利用訓練模型 SVM-T 來辨識，可分類出 UT 和 DT。

Step 7.3: 將 Shift 型樣利用訓練模型 SVM-S 來辨識，可分類出 US 和 DS。

最後經由支援向量機分類器可清楚的辨識出八種管制圖型樣，分別為：
(1)normal；(2)stratification；(3)systematic；(4)cyclic；(5)up trend；(6)down trend；
(7)up shift；(8)down shift。

以下利用兩個例子來說明本研究所提出之 ICA-SVM 模型的辨識過程。首先針對混合型樣，本研究以 NOR+SYS 為例，圖 5(a1)和(a2)為製程中所觀測到的混合型樣資料所構成的管制圖，在管制圖中，很難用肉眼觀察出其所包含的基本不尋常型樣，因此在本研究的辨識過程中，首先透過本研究所提出之前處理程序 ICA，可將所觀察到之混合型樣分解成兩個獨立成份，分別為(b1)和(b2)，再經由峰態係

數之大小排序後，可明顯看出峰態係數較小者為 NOR 型樣，如圖 5(b1)所示，而峰態係數較大者為 SYS 型樣，也就是本研究所欲辨識之待監控型樣資料，如圖 5(b2)所示，最後再將待監控資料(b2)應用於所訓練好之辨識模型中，即可輕易地辨識其所歸屬的型樣。

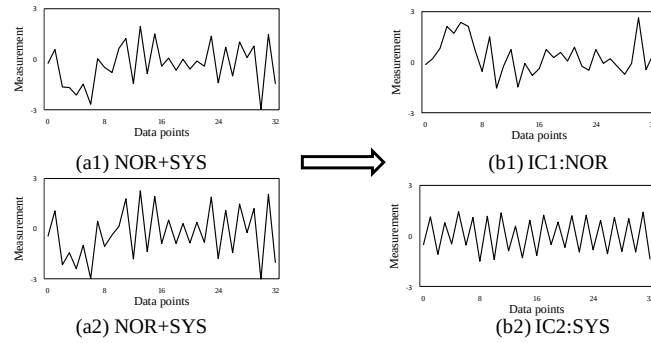


圖 5 混合型樣 ICA 分解範例

第二個範例則針對權重變化型樣，同樣地以 NOR+SYS 為例，分別舉例兩種權重變化相互比較。其中一種型樣為 $0.5 \cdot \text{NOR} + 1.5 \cdot \text{SYS}$ ，可由圖形表示如圖 6 中的(a1)，而另一種型樣則假設為 $1.5 \cdot \text{NOR} + 0.5 \cdot \text{SYS}$ ，可由圖形表示如圖 6 中的(c1)，比較兩種管制圖型樣(a1)和(c1)可發現到，當 NOR 權重較大時，很難察覺出其所隱藏的 SYS 型樣，也就是說，當面對 NOR 權重愈大的型樣時，愈不容易辨識。

因此在本研究的辨識過程中，首先透過本研究所提出之前處理程序 ICA，可分別將所觀察到之權重變化型樣分解成兩個獨立成份，但由於 ICA 的運作過程中，需要兩組觀測值才能估計兩個獨立成份，因此本研究個別在 $0.5 \cdot \text{NOR} + 1.5 \cdot \text{SYS}$ 型樣和 $1.5 \cdot \text{NOR} + 0.5 \cdot \text{SYS}$ 型樣中，各加入一組事先收集好之 NOR 觀測值，如圖 6(a2)和(c2)所示，然後分別利用 ICA 估計其獨立成份。其中 $0.5 \cdot \text{NOR} + 1.5 \cdot \text{SYS}$ 型樣估計出的獨立成份為圖 6(b1)和(b2)，而 $1.5 \cdot \text{NOR} + 0.5 \cdot \text{SYS}$ 型樣估計出的獨立成份為圖 6(d1)和(d2)，再經由峰態係數之大小排序後，可明顯看出峰態係數較小者為 NOR 型樣，如圖 6(b1)和(d1)所示，而峰態係數較大者為 SYS 型樣，也就是本研究所欲辨識之待監控型樣資料，如圖 6(b2)和(d2)所示，最後再將待監控資料(b2)和(d2)應用於所訓練好之辨識模型中，即可輕易地辨識其所歸屬的型樣。

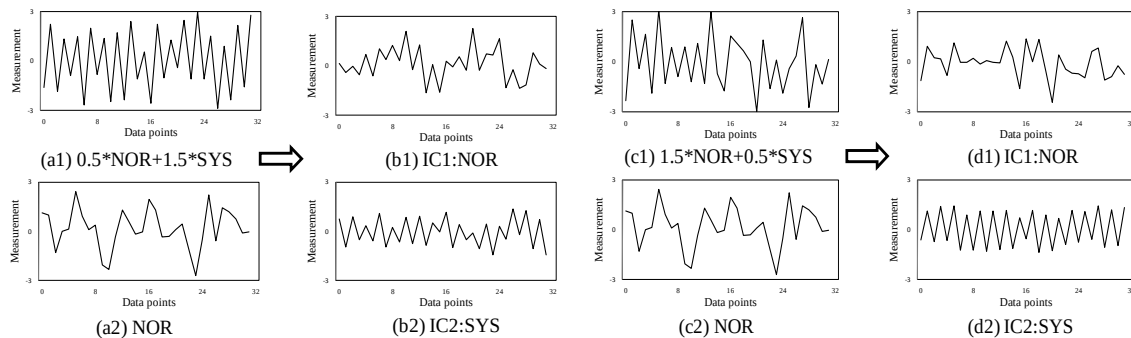


圖 6 權重變化型樣 ICA 分解範例

由以上範例可知，不論針對較難辨識之混合型樣或具權重變化型樣，透過本研究提出之 ICA 方法，均能將其拆解成較易辨識之 CCPs，再利用本文所建構的 SVM 分類器，既能提高其效率，也能提高其正確分類率。

參、實證分析

本研究提出 ICA-SVM 的建構模式來建立辨識系統，此外，為了證實本研究提出的 ICA-SVM 模式效果較佳，本實驗建立兩種沒有利用 ICA 的 SVM 模型與本研究提之方法作比較，其中一個 SVM 模型直接利用所觀察到的樣本點來作為輸入變數，建立管制圖型樣辨識模型，稱其為 SVM-D 模型；而另一個 SVM 模型的輸入變數則丟入觀察到的樣本點所算出的六個特徵值，稱其為 SVM-F 模型，其中特徵值包含 SB, ADIST, AASBP, AASL, SRANGE 和峰態係數，特徵值的選取參考文獻(Gauri and Charkaborty, 2006, 2009)，而此六個特徵值的定義可參考文獻(Gauri and Charkaborty, 2006, 2009)。本模擬實驗主要針對混合型樣和具權重變化型樣進行辨識，以下為兩種不尋常型樣之測試結果：

一、混合型樣

當兩個單一的基本不尋常管制圖型樣同時發生(concurrent)時，結合所產生混和管制圖型樣，其混合方式為：

$$\begin{bmatrix} X_1 \\ X_2 \end{bmatrix} = A \times \begin{bmatrix} S_1 \\ S_2 \end{bmatrix} \quad (12)$$

其中， X_1 、 X_2 為混合管制圖型樣， S_1 、 S_2 為兩兩相異的不尋常管制圖型樣， A 為混合矩陣，其設定參考 Wang et al. (2009)， $A = \begin{bmatrix} 0.6 & 0.4 \\ 0.4 & 0.6 \end{bmatrix}$ 。本研究利用 Grid 搜尋方式找出三個模型（SVM-F、SVM-D 和 ICA-SVM）中的最佳參數，表 2~5 中分別為 SVM-F、SVM-D 和 ICA-SVM 的測試結果，其中 SVM-F 模式的參數設定為 $C = 2^5$ ， $\gamma = 2^5$ ，十次平均分類正確率 36.04%，標準差為 0.037；SVM-D 模式的參數設定為 $C = 2^0$ ， $\gamma = 2^0$ ，十次平均分類正確率為 73.7%，標準差為 0.009；ICA-SVM 模式建構四個分類模型，其模型參數 C 和 γ 的設定分別為 SVM-A： $C = 2^5$ ， $\gamma = 2^5$ ；SVM-N&S： $C = 2^{-5}$ ， $\gamma = 2^{-5}$ ；SVM-T： $C = 2^0$ ， $\gamma = 2^0$ ；SVM-S： $C = 2^0$ ， $\gamma = 2^0$ ，十次平均分類正確率為 99.19%，標準差為 0.006。由表 2~5 的平均正確率列表比較得知，對於混合型樣的辨識結果，本研究提出的方法能有效的辨識混合型樣，相較於其他模式的辨識結果，ICA-SVM 模式較能獲得良好的辨識結果，且其標準差最低，分類結果最穩定。

表 2、SVM-index 模式之分類正確率矩陣列表(混合型樣)

真實型樣	辨識後型樣類別							
	NOR	STA	SYS	CYC	UT	DT	US	DS
NOR	77.5%	0.0%	0.0%	0.0%	8.0%	14.5%	0.0%	0.0%
STA	64.6%	35.4%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
SYS	11.2%	0.0%	70.5%	0.0%	9.8%	2.3%	0.0%	6.3%
CYC	4.6%	0.0%	0.0%	9.0%	36.6%	22.3%	19.0%	8.6%
UT	52.3%	0.0%	0.0%	0.1%	27.6%	18.3%	1.2%	0.7%
DT	48.7%	0.0%	0.0%	0.1%	27.3%	20.9%	2.4%	0.7%
US	21.8%	0.0%	0.0%	1.4%	41.8%	24.8%	7.0%	3.3%
DS	23.7%	0.0%	0.0%	0.5%	38.9%	23.0%	8.6%	5.5%
整體正確率	31.67% (5067/16000)							

表 3、SVM-data 模式之分類正確率矩陣列表(混合型樣)

真實型樣	辨識後型樣類別							
	NOR	STA	SYS	CYC	UT	DT	US	DS
NOR	99.9%	0.2%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
STA	1.3%	98.7%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
SYS	6.9%	0.0%	93.2%	0.0%	0.0%	0.0%	0.0%	0.0%
CYC	24.1%	0.0%	0.0%	75.9%	0.0%	0.0%	0.0%	0.0%
UT	27.0%	0.0%	0.0%	0.0%	56.0%	0.0%	17.0%	0.0%
DT	13.6%	0.0%	0.0%	0.0%	0.0%	74.5%	0.0%	11.9%
US	32.0%	0.0%	0.0%	0.0%	15.3%	0.0%	52.8%	0.0%
DS	20.2%	0.0%	0.0%	0.0%	0.0%	26.5%	0.1%	53.4%
整體正確率	75.54% (12086/16000)							

表 4、ICA-SVM 模式之分類正確率矩陣列表(混合型樣)

真實型樣	辨識後型樣類別							
	NOR	STA	SYS	CYC	UT	DT	US	DS
NOR	98.6%	1.2%	0.0%	0.0%	0.1%	0.1%	0.0%	0.0%
STA	1.0%	99.0%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
SYS	0.0%	0.0%	100%	0.0%	0.0%	0.0%	0.0%	0.0%
CYC	0.0%	0.0%	0.0%	100%	0.0%	0.0%	0.0%	0.0%
UT	0.0%	0.0%	0.0%	0.0%	100%	0.0%	0.0%	0.0%
DT	0.0%	0.0%	0.0%	0.0%	0.0%	100%	0.0%	0.0%
US	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%	100%	0.0%
DS	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%	100%
整體正確率	99.71% (15953/16000)							

表 5、各種模式之十次分類正確率與標準差(混合型樣)

實驗次數	SVM-index	SVM-data	ICA-SVM
1	31.67%	75.54%	98.81%
2	35.21%	73.43%	98.89%
3	38.49%	73.83%	98.59%
4	31.32%	72.64%	99.81%
5	32.86%	74.14%	99.37%
6	38.98%	72.23%	97.94%
7	41.24%	73.53%	99.78%
8	37.19%	74.11%	99.59%
9	32.94%	74.33%	99.39%
10	40.46%	73.26%	99.71%
平均數	36.04%	73.7%	99.19%
標準差	0.037	0.009	0.006

二、權重變化型樣

當單一不尋常管制圖型樣中自身的常態資料與干擾項變數間有權重變化之情形時，此型樣即為權重變化型樣，其模擬方式為常態資料 $(u + r_i\sigma)$ 與干擾項 $(d(-1)^i)$ 間加上權重比例 θ_1 和 θ_2 後，權重變化型樣的模擬公式便成為 $X_1 = \theta_1(u + r_i\sigma) + \theta_2(d(-1)^i)$ 。本研究針對隨機權重變化型樣，利用 SVM-F、SVM-D 和 ICA-SVM 三種辨識模式來進行測試並相互比較，表 6 中分別為 SVM-F、SVM-D 和 ICA-SVM 的測試結果，其中 SVM-F 的參數設定為 $C = 2^{-5}$ ， $\gamma = 2^5$ ，重複十次實驗後，平均分類正確率為 50.4%，標準差為 0.083；SVM-D 的參數設定為 $C = 2^{-5}$ ， $\gamma = 2^1$ ，十次平均分類正確率 68.29%，標準差為 0.211；ICA-SVM 模式建構四個分類模型，其模型參數 C 和 γ 的設定分別為 SVM-A： $C = 2^5$ ， $\gamma = 2^5$ ；SVM-N&S： $C = 2^1$ ， $\gamma = 2^{-5}$ ；SVM-T： $C = 2^{-2}$ ， $\gamma = 2^{-5}$ ；SVM-S： $C = 2^1$ ， $\gamma = 2^1$ ，十次平均分類正確率為 98.54%，標準差為 0.037。由表 6 的平均正確率列表比較得知，針對隨機權重變化型樣，本文所提出的辨識模式能有效的辨識任何權重變化的型樣，且由標準差也可看出 ICA-SVM 模式較其他模式穩定，變異較小且有較高的分類效果。

表 6、各種模式之十次分類正確率與標準差(隨機)

實驗次數	SVM-index	SVM-data	ICA-SVM
1	43.57%	70.52%	99.94%
2	57.53%	96.87%	99.69%
3	35.84%	47.15%	99.63%
4	50.38%	71.16%	99.46%
5	62.25%	72.19%	87.98%
6	43.88%	83.83%	99.57%
7	56.91%	82.47%	99.91%
8	50.99%	24.89%	99.81%
9	59.91%	87.31%	99.49%
10	42.73%	46.53%	99.88%
平均數	50.4%	68.29%	98.54%
標準差	0.083	0.211	0.037

三、整體正確率比較

將三種模式與辨識型樣種類的十次實驗平均正確率合併，如表 7 所示，由表可得知 ICA-SVM 不論針對混合型樣、權重變化型樣或基本不尋常型樣都具有高度的正確辨識效果，且透過標準差的比較得知，對於參數的隨機設定也有穩定的表現。因此由本實驗結果可知，透過本研究提出的 CCPR 模式，可以精準的辨識管制圖型樣，且能有效的分類出所屬型樣，找出製程出錯之根本原因，提高生產品質。

表 7、整體正確率對照表

建構模式 針對型樣	SVM-index		SVM-data		ICA-SVM	
	μ	σ	μ	σ	μ	σ
混合型樣	36%	0.037	73.7%	0.009	99.2%	0.006
權重變化型樣	50.4%	0.083	68.3%	0.211	98.5%	0.037

肆、結論與建議

一、研究結論

管制圖型樣辨識在現今的工業製程中已是不可或缺的監控技術之一，其主要功能為加強監控管制圖所不易偵測之不尋常型樣，透過辨識不尋常型樣之種類，可探討出其製程失控之可歸屬原因，藉以改善製程，並減少不良品的產生以提高生產品質；然而，目前大多數的研究只考慮了單一不尋常管制圖型態來進行辨識，經由兩個基本 CCPs 同時產生所形成的混合 CCPs，或型樣本身存在權重變化的

CCPs 並沒有討論到；因此，本研究提出一種整合 ICA 和 SVM 的辨識模式，針對混合型態 CCPs 與權重變化型態 CCPs 進行辨識。首先利用基本 CCPs 建構 SVM 分類模型，接著透過 ICA 將製程中所觀察到的混合型態分解，估計出 IC 並從中找出代表基本型樣的 IC 後，將這些 IC 應用於建立好之 SVM 模型進行 CCPR。經由本研究之模擬實驗可得到以下三個實證結果：

(一) 本研究之實證結果顯示，當待監控資料為混合型樣或權重變化型樣時，ICA-SVM 模式均能擁有平均高達九成九的正確分類率，由此可知，不論哪種型態的混合，本研究所提出之 ICA-SVM 模式皆可達到良好的辨識效果。

(二) 比起傳統其他辨識方法，例如：單純只用指標或直接利用觀測值結合 SVM 分類器所建構的辨識系統，本研究所提出的模式皆優於傳統辨識方法，不僅可以產生較高的正確分類率，且辨識結果較穩定。

(三) 本研究也將基本不尋常型樣應用於 ICA-SVM 模式中，與大多數文獻所針對的問題相互比較，由本實驗結果可知，ICA-SVM 模式針對一般型樣也能有效地辨識。

二、研究貢獻

(一) 本研究延伸 ICA 與 SVM 的應用範圍，針對 CCPR 中較少運用之 ICA 與 SVM，創新地將 ICA 與 SVM 整合應用於 CCPR 中。

(二) 本研究主要在於解決兩個型樣同時發生在同一製程中，導致製程監控人員難以辨識的狀況。利用本文所提出之 ICA 前處理方法，可有效的分解混合型樣，並拆解出容易辨識之基本組成型樣，最後在結合 SVM 分類器，可準確的辨識管制圖型樣。

(三) 本研究所提之方法也適用於文獻中較少討論到的具權重變化的管制圖型樣。

(四) 本研究所提之架構能有效處理較多樣化之基本型樣，包含了與 Normal 相類似之型樣：Stratification，以及 Trend 和 Shift 之正負趨勢問題。

三、未來研究方向與建議

本研究所處理之問題尚有不足之處，且此議題之研究也還有發展的空間，因此本研究列出以下幾點未來可研究方向與建議：

(一) 本研究處理之混合型樣與權重變化型樣僅針對 NOR 與各不尋常型樣間的混合，未來可朝向多個基本型樣之混合型樣邁進，多個混合型樣例如：NOR+SYC+CYC 所搭配合成的混合型樣。

(二) 本研究流程分為兩階段，首先將 NOR 與 STA、UT 與 DT、US 與 DS 歸為一類，經由第一階段分類後，第二階段再分別處理尚未分類之型樣，過程稍嫌複雜，因此建議未來可找出高相關性指標作為 SVM 之分類變數，直接將相似之型樣獨立辨識，取代兩階段分類。

伍、參考文獻

- (1) Western Electric, 1958. *Statistical quality control handbook*, Indianapolis: Western Electric Company.
- (2) Wang, C. H., and Kuo, W., 2007. Identification of control chart patterns using wavelet filtering and robust fuzzy clustering, *Journal of Intelligent Manufacturing*, 18, 343-350.
- (3) Gauri, S. K., and Chakraborty, S., 2006. Feature-based recognition of control chart patterns, *Computer and Industrial Engineering*, 51, 726-742.
- (4) Gauri, S. K., and Chakraborty, S., 2009. Recognition of control chart patterns using improved selection of features, *Computers and Industrial Engineering*, 56, 1577-1588.
- (5) Guh, R. S., Zorriassatine, F., Tannock, J. D. T., and O'Brien, C., 1999. On-line control chart pattern detection and discrimination - a neural network approach, *Artificial Intelligence in Engineering*, 13(4), 413-425.
- (6) Perry, M. B., Sporre, J. K., and Velasco, T., 2001. Control chart pattern recognition using back propagation artificial neural networks, *International Journal of Production Research*, 39(15), 3399-3418.
- (7) Guh, R. S., and Shiue, Y. R., 2005. On-line identification of control chart patterns using self-organizing approaches, *International Journal of Production Research*, 43(6), 1225-1254.
- (8) Lee, J. M., Yoo, C., and Lee, I. B., 2004. Statistical process monitoring with independent component analysis, *Journal of Process Control*, 14(5), 467-485.
- (9) Hyvärinen, A., and Oja, E., 2000. Independent component analysis: Algorithms and applications, *Neural Networks*, 13, 411-430.
- (10) Hyvärinen, A., Karhunen, J., and Oja, E., 2001. *Independent component analysis*, New York: Wiley press.
- (11) Vapnik, V., 1995. *The nature of statistical learning theory*, Springer Verlag.
- (12) Vapnik, V. N., 2000. *The nature of statistical learning theory*. Springer, Berlin.

- (13) Hsu, C. W., Lin, C. C., and Lin, C. J., 2003. A practical guide to support vector classification. Technical report, *Department of Computer Science and Information Engineering*, National Taiwan University, Taipei, available <http://www.csie.ntu.edu.tw/~cjlin/libsvm/>..
- (14) Cherkassky, V., and Ma, Y., 2004. Practical selection of SVM parameters and noise estimation for SVM regression, *Neural Network*, 17, 113-126.
- (15) Yang, J. H., and Yang, M. S., 2002. A fuzzy-soft learning vector quantization for control chart pattern recognition, *International Journal of Production Research*, 40(12), 2721-2731.
- (16) Yang, J. H., and Yang, M. S., 2005. A control chart pattern recognition scheme using a statistical correlation coefficient method, *Computers and Industrial Engineering*, 48, 205-221.
- (17) Assaleh, K., and Al-assaf, Y., 2005. Feature extraction and analysis for classifying causable patterns in control charts, *Computer and Industrial Engineering*, 49, 168-181.
- (18) El-Midany, T. T., El-Baz, M. A., and Abd-Elwahed, M. S., 2010. A proposed framework for control chart pattern recognition in multivariate process using artificial neural networks, *Expert Systems with Applications*, 37(2), 1035-1042.
- (19) Guh, R. S., 2005. A hybrid learning-based model for on-line detection and analysis of control chart patterns, *Computer and Industrial Engineering*, 49, 35-62.
- (20) Wang, C. H., Kuo, W., and Dong, T. P., 2009. A hybrid approach for identification of concurrent control chart patterns, *Journal of Intelligent Manufacturing*, 20, 409-419.