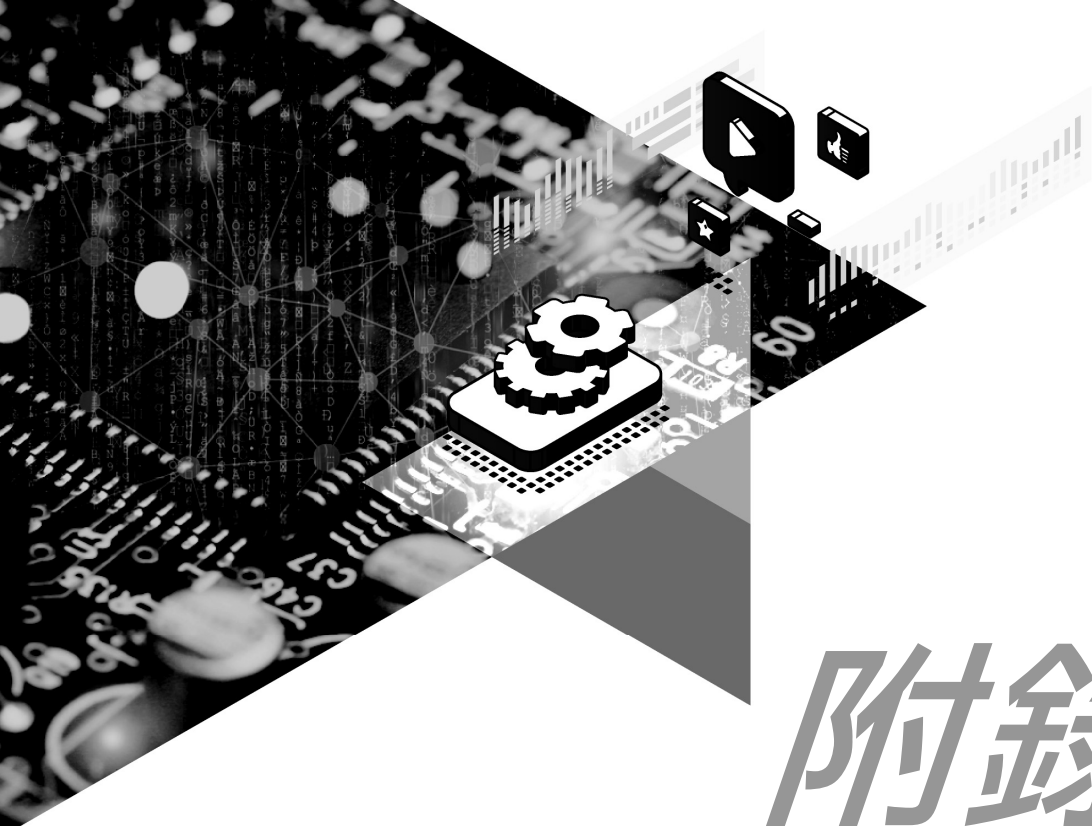




附錄E

統計製程管制與 先進製程控制

Statistical Process Control and
Advanced Process Control

- 
- E.1 統計製程管制
 - E.2 錯誤偵測與分類
 - E.3 先進製程控制
 - E.4 結語

在製造業中，產能與良率無疑是重要的命脈，有效地提高產能與良率是各個企業所關注的課題。本章節我們將專注於產品良率 / 品質，一方面，在將產品交付客戶前對品質做最後一道把關，另一方面，良率反應了實際的製程能力，透過對製程的監控、改善、與控制來提升製程能力，可大幅減少不良品或重工的成本。因此，本章節將依序介紹「統計製程管制」(statistical process control, SPC)、「故障偵測與分類」(fault detection and classification, FDC)以及「先進製程控制」(advance process control, APC)的思維與方法。這三方法分別對應到上述說明的製程監控（偵測異常情形發生）、製程改善（瞭解不良品的類型與可能的發生原因）以及製程控制（製程參數該如何調整與優化以縮小製程變異）。

E.1 統計製程管制

品質的定義包含了產品設計的品質與生產的一致性（也就是良率），前者於產品初期的「研發」(research and development, R&D)中進行改善（以實驗設計與模擬等方法來分析並調整產品設計），後者則於量產的製程中進行改善，本章節著重於後者。

統計製程管制是應用統計方法與工具對生產製程進行監控也就是，一方面可對機台或環境的製程參數（例如溫度、壓力、轉速等）進行監控，以確保機台加工過程如預期的設定；另一方面可透過品質檢測對產品的規格或量測結果（例如長、寬、高、膜厚等）進行監控，以確保產品品質是否合乎顧客規格。當系統性異常（變異）發生時發出警報，從而找出異常的原因，採取調整製程參數之行動，並以不斷地降低製程與品質的變異為目標，如圖 E.1 所示。實際上，產品的製程與其良率 / 品質為因果關係，製程變異 (process variation) 的存在造就了品質變異 (quality variation)。統計製程管制在於製程變異的控制；「統計品質管制」(statistical quality control, SQC) 則是在於品質變異的監控。因此為提升品質，對於製程（因）的監控、改善、控制更是根本之道。

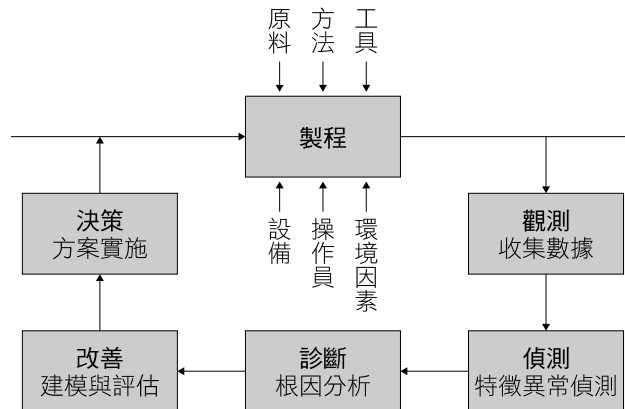


圖 E.1 統計製程管制

若我們進一步對製程變異剖析（利用資源 8M1I 等因素），主要可分成一般原因（common causes）與特殊原因（special causes），也就是隨機性變異與系統性變異的探討。一般原因像是隨機的，溫度與環境因子的變化、加工時的參數細微浮動、或量測誤差等，通常影響較小且難以避免；特殊原因像是錯誤工具、不當參數調整、原料品質變異、設備異常等，通常影響較大。因此，在統計製程管制中我們期望的是能減低這些系統性變異。

「統計製程管制」主要包含了品管（QC）七大手法：

- 「直方圖」(histogram)：用於呈現連續變數的分布。
- 「散佈圖」(scatter plot)：用於呈現兩變數間的相關程度。
- 「層別法」(stratification)：以階層式或分類別呈現各變數的關係。
- 「魚骨圖」(fishbones chart)：用於找出問題的根因。
- 「查檢圖」(check sheet)：用於記錄問題發生的類型、時間、次數等。
- 「帕瑞圖」(Pareto chart)：將問題的依據頻度排序。
- 「管制圖」(control chart)：以樣本平均數為中心，設置上下的管制界線作為監控。

如圖 E.2 所示，其中，直方圖、散佈圖、層別法是為了呈現變數本身與之間的關係與分布；魚骨圖、查檢圖、帕瑞圖、管制圖是釐清問題的原因、型態、重要性、監控。上述手法對數據與問題的剖析與釐清，與系統

化運算決策 (SYCON) 中的系統思考概念極為相似，強化點到系統面的思維。

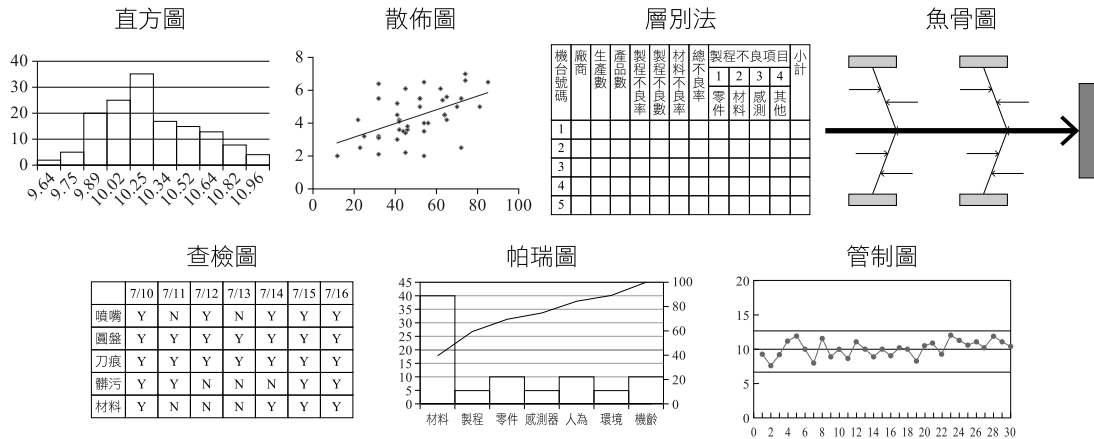


圖 E.2 統計製程管制

本小節將詳細說明統計製程管制中重要的管制圖基礎，以及不同類型的管制圖應用，包含了「單變量管制圖」(univariate control chart)、「時間加權管制圖」(time-weighted control chart) 以及「多變量管制圖」(multivariate control chart) (包含有母數與無母數方法)。

E.1.1 管制圖基礎

管制圖是一種被廣泛用於線上製程監控的視化工具，用以監測製程參數或品質特性隨時間變化的情形。其思維是期望穩定的製程參數或品質特性所服從的分配（或其統計量）維持不變，因此藉由統計的假設檢定來判定新收進來的數據分布特性是否發生改變。

典型的製程管制圖如圖 E.3 所示。包含一條「中心線」(centerline, CL)，代表穩定的製程參數或品質特性應有的期望值或目標值，同時也包含兩條水平線，分別為「上管制界限」(upper control limit, UCL) 與「下管制界限」(lower control limit, LCL)，是由中心線加減而其統計量的變異所構成的。管制圖的 x 軸為時間，y 軸為被監控的統計量（此圖範例為一組樣本的平均數 $\bar{X}$ ）。若僅有隨機性變異時，其統計量僅會落於上下管制界限間，被稱為製程或品質於「管制中」(in-control)；若因特殊原因而存在系統性的異常時，則會有統計量落於上下管制界限外，稱為製程或品質

「失控」(out-of-control)。這時線上製程監控系統會發出警報通知，讓工程師（決策者）或系統本身來檢視與判別異常，即時修正或控制。

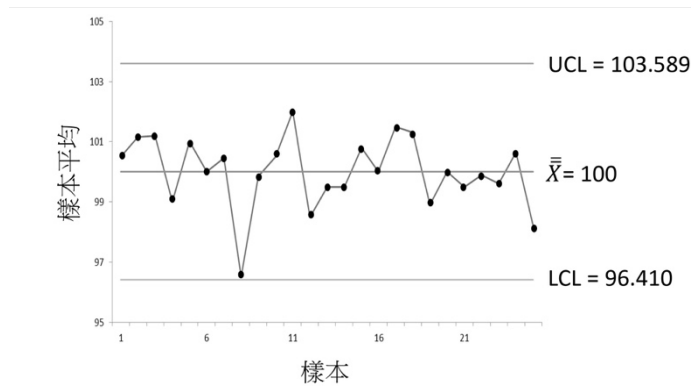


圖 E.3 典型的平均數 ($\bar{X}$) 管制圖

E.1.1.1 管制圖建構與應用

一般來說，管制圖的建構與應用可分為兩個階段：

- 「階段一」(phase I)：依據選定的製程參數或品質特性（連續或離散變數）決定管制圖的種類以及抽樣的樣本大小、頻率、方式，經初步的數據收集後，建構試用（trail）的管制界限來驗證收集的正常數據是否落於管制內，若有部分數據落於管制外，應檢討其正當性（是否應刪除）再重新建構，確保管制界限是由穩定製程的數據所建構的。
- 「階段二」(phase II)：利用階段一所建構的管制圖監控製程，比較每一組樣本的統計量是否落於管制界限內，以判定製程是否穩定，通常 UCL 與 LCL 會持續因製程狀況調整，以權衡偽陽性（假警報）與偽陰性（漏報）。

經由兩階段的管制圖建構與應用可逐步地改善造成異常的特殊原因，因而製程變異（也就是樣本統計量的變異）也將會被逐步縮小，形成更穩定的製程與高良率 / 品質的產出。

E.1.1.2 數據抽樣

管制圖是根據「合理樣本組」(rational subgrouping) 的概念來對數據抽樣，其思維是為了區分製程的「組內變異」(within subgroup variation)

與「組間變異」(between subgroup variation)，也就是我們期望每次抽樣時，一組樣本內的分配特性越為相似越好（組內變異越小），而不同組樣本間的差異才是我們想要找出的異常。一般而言，主要有兩種抽樣方式（如圖 E.4 所示）：

- 「瞬時法」(instant time method)：組內樣本盡可能地在一個很短的時間間距收集。此方法可以更精準地估計短時間內的樣本統計量，較容易偵測出瞬間劇變的趨勢，但另一方面，組間樣本之間的統計量無法緊密地被估計，因此，此方法的空間解析度高而時間解析度低。
- 「定時法」(period of time method)：將時間切割成固定的間距，並於間距內隨機抽樣。此方法可以估計連續時間間隔內的統計量，但另一方面，短時間的劇變可能會被組內隨機樣本稀釋，因此，此方法的時間解析度高而空間解析度低。

依據數據特性隨時間改變的抽樣方式與偵測機制將在後續章節「概念漂移與在線學習」有更詳細的介紹。

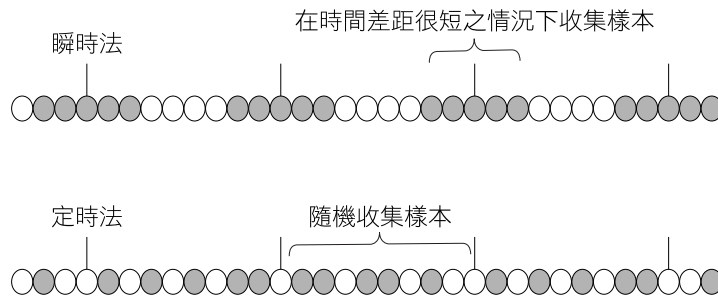


圖 E.4 數據抽樣方式比較

E.1.1.3 管制圖研判

儘管管制圖能偵測每組樣本統計量是否超出上下管制界限，但當所有統計量均落於管制界限內，也不能保證製程一定是穩定的。不穩定的例外可能是(1)可能很多統計量落於接近管制界限附近；(2)隨著時間統計量有特定的趨勢或樣型 (pattern)。因此我們可制定一些測試法則來更精細地研判出管制圖內的異常，例如有兩種常見的測試方法：

- 「區間測試」(zone tests) (又稱 Western Electric Handbook

Rules)：其包含以下規則：(1)一點落於管制界限外（例如落在樣本平均數抽樣分布的三倍標準差外）；(2)連續三點有兩點落於兩倍標準差外；(3)連續五點有四點落於一倍標準差外；(4)連續八點落於中心線同一側。

- 「連串測試」(run tests)：其包含以下規則：(1)連續七點落於中心線同一側；(2)連續十一點有十點落於中心線同一側；(3)連續十四點有十二點落於中心線同一側；(4)連續十七點有十四點落於中心線同一側；(5)連續二十點有十六點落於中心線同一側。

E.1.1.4 管制圖評估與製程能力

由於管制圖是基於假設檢定的方法，因此存在著「型一型二錯誤」(type I and type II error) 的風險，如圖 E.5 所示，左圖型一錯誤 α 是將穩定的製程平均誤判為失控的可能，為假警報；右圖型二錯誤 β 是無法偵測偏移的製程平均（為圖中平均數為 B 的情形）的可能，為未能觸發警報。因此當上下管制界限被調整時型一型二錯誤也會受之影響（例如更窄的管制界限其型二錯誤變小但型一錯誤變大）。

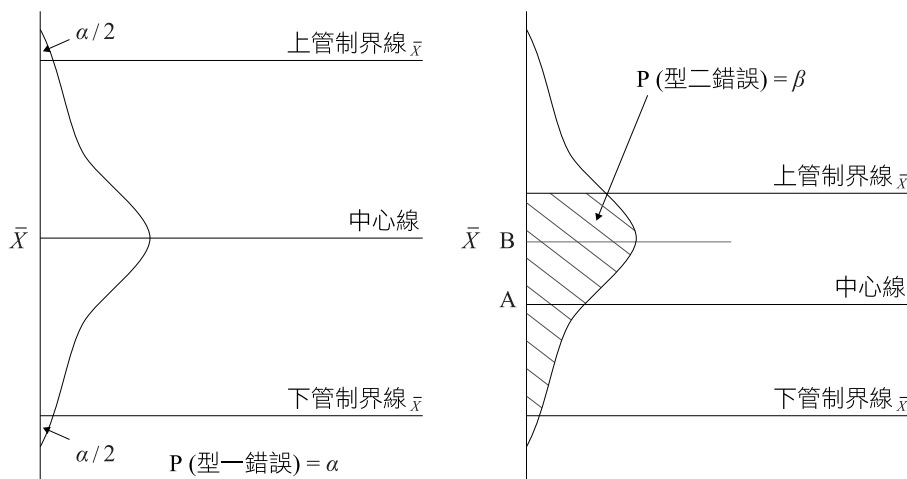


圖 E.5 管制圖的型一與型二錯誤

「平均連串長度」(average run length, ARL) 是一項連結型一型二錯誤與抽樣大小、頻率的重要指標，因此我們也會依據它來設定後兩者，可用以評估管制圖 (Montgomery, 2019)。平均連串長度的定義是在某個製程

狀態下超出管制界限的期望樣本數，如公式(E.1)所示，

$$ARL = \frac{1}{p} \quad (E.1)$$

其中 p 是任意樣本統計量超出管制界線的機率。舉例來說，在穩定製程的狀態下（平均數趨近於中心線），其根據常態分配超出三倍標準差管制界限的機率 p 等於 0.0027，因此其平均連串長度近似於 370。換句話說，在穩定製程下平均 370 個樣本才會發生一次假警報。在不穩定製程的狀態下（假設製程平均落在三倍標準差界限上且樣本大小為 1，也就是常態分配的平均值剛好落於管制界限上），其超出三倍標準差管制界限的機率 p 等於 0.5，因此其平均連串長度近似於 2。換句話說，在不穩定製程下平均 2 個樣本才會偵測出統計量的偏移。因此，當我們想以抽樣的角度來提升管制圖表現時，可以(1)提高抽樣頻率來縮減異常偵測所需的時間；(2)增加樣本數（sample size）來提高偵測能力（也就是也就是縮小估計的變異以提高檢定力 $1 - \beta$ ）。

另一方面，管制圖也能作為審視製程能力表現的工具。一般來說，製程參數與品質特性會事先給定一個特定的規格，也就是「上規格界限」（upper specification limit, USL）與「下規格界限」（lower specification limit, LSL）。因此，「製程能力指數」（process capability index, C_p 與 C_{pk} ）是呈現規格界限的能力評估指標，如公式(E.2)所示（在正負三倍標準差下），

$$\begin{aligned} PCR &= C_p = \frac{USL - LSL}{6\sigma} \\ PCR_k &= C_{pk} = \min \left[\frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma} \right] \end{aligned} \quad (E.2)$$

其中製程能力 C_p 為精確能力（capability of precision）（又稱製程能力比率 process capability ratio, PCR）是僅考慮製程變異並假設製程是對準中心（centerline）；而製程能力 C_{pk} （又稱 PCR_k ）則是同時考慮製程變異與偏離。因此在「六標準差製程」（six-sigma process）中，在正負六倍標準差下，完美符合規格的製程能力 C_p 為 1 而 C_{pk} 為 2，換句話說，也就是該製程僅有百萬分之 3.4 的可能會超出製程規格的上下界限。一般原則， C_p 或 C_{pk} 在 1.33 以下代表製程表現不佳；介於 1.33 到 1.67 之間代表製程表現

佳；1.67 以上代表製程能力表現傑出。切記，在管制圖上繪製規格界限或使用規格來決定管制界限是不對的作法，因為規格本身跟製程管制界限並無相關性。此外，在品管圖的實務應用中，其三倍標準差的管制界限也不會固定不變，應隨著產品組合、製程能力好壞、良率狀況、平均連串長度等，而提出動態調整的過程與機制。

E.1.2 單變量管制圖

若僅專注於某一製程參數或品質特性時，可使用「單變量管制圖」〔也就是修華特管制圖 (Shewhart control chart)〕進行監控。依據變數的特性，可分為連續變數的「計量管制圖」(variables control chart) 以及離散變數的「計數管制圖」(attributes control chart)。其中計量管制圖又可以依據每組樣本數的大小使用「平均數管制圖」(Xbar chart)、「標準差管制圖」(s chart)、「全距管制圖」(R chart)、以及「個別值與移動全距管制圖」(X-MR chart)；計數管制圖則包含了針對不良品的「不良率管制圖」(P chart) 與「不良數管制圖」(np chart)，以及包含了針對單一產品缺點的「缺陷數管制圖」(C chart) 與「平均缺陷數管制圖」(U chart)，如表 E.1 所示。此部分我們僅介紹平均數與標準差管制圖、不良率管制圖以及平均缺陷數管制圖。

表 E.1 計量與計數管制圖的種類

數據類型	缺陷定義	(組內) 樣本大小	管制圖類型
連續型變數 (計量) (variable data)		樣本數大於 10	平均數 ($\bar{X}$) 與標準差 (s)
		樣本數小於 10 且大於 1	平均數 ($\bar{X}$) 與全距 (R)
		樣本數等於 1	個別值 (X) 與移動全距 (MR)
離散型變數 (計數) (attribute data)	不良品數 (defective unit)	固定樣本數，且大於 50	不良數管制圖 (np chart)
		不固定樣本數，且大於 50	不良率管制圖 (P chart)
	缺陷數 (defect unit)	固定樣本數	缺陷數管制圖 (C chart)
		不固定樣本數	平均缺陷數管制圖 (U chart)

E.1.2.1 平均數與標準差管制圖 (Xbar-s chart)

基於「中央極限定理」(central limit theorem, CLT)，不論母體服從何種分配，當抽樣樣本數足夠大時，其樣本平均數的抽樣分配會服從常態分配。因此，我們期望監控的製程參數或品質特性每組抽樣的樣本平均數

($\bar{x}$) 會服從常態分配，其平均數與標準差如公式(E.3)所示，其中 n 與 σ 分別為樣本數與標準差。

$$\bar{x} \sim N\left(\bar{\bar{x}}, \frac{\sigma}{\sqrt{n}}\right) \quad (\text{E.3})$$

以正負三倍標準差為例的平均數管制圖 ($\bar{X}$ chart from s) 依據上式可計算出其中心線 (CL) 與上下管制界限 (UCL, LCL) 如公式(E.4)所示，

$$\begin{aligned} CL &= \hat{\mu} \\ UCL &= \hat{\mu} + 3 \frac{\hat{\sigma}}{\sqrt{n}} \\ LCL &= \hat{\mu} - 3 \frac{\hat{\sigma}}{\sqrt{n}} \end{aligned} \quad (\text{E.4})$$

$$\text{其中 } \hat{\mu} = \bar{\bar{x}} = \frac{1}{mn} \sum_{i=1}^m \sum_{j=1}^n x_{ij} \text{ 與 } \hat{\sigma} = \bar{s}/c = \sum_{i=1}^m \sqrt{\frac{\sum_{j=1}^n (x_{ij} - \bar{x}_i)^2}{n-1}} / cm$$

其中 n 為組內樣本數，且 m 為組間數（批次數）。 c 是當小樣本時依據樣本數修正標準差的常數，如表 E.2。這是由於樣本標準差 s 是母體標準差 σ 的偏估計量 (bias estimator)，當小樣本數時，需要做估計上的修正，也就是 $\hat{\sigma} = \bar{s}/c$ 。標準差管制圖 (s chart) 則可表示如公式(E.5)所示。

$$\begin{aligned} CL &= \bar{s} = c\hat{\sigma} \\ UCL &= \bar{s} + 3 \frac{\bar{s}}{c} \sqrt{1-c^2} \\ LCL &= \bar{s} - 3 \frac{\bar{s}}{c} \sqrt{1-c^2} \end{aligned} \quad (\text{E.5})$$

因此，我們可由「平均數管制圖」監控平均數的改變以及「標準差管制圖」監控標準差的變動。

表 E.2 依樣本數修正標準差的常數 c

n	2	3	4	5	6	7	8	9	10	15	20	25
c	0.7979	0.8862	0.9213	0.9400	0.9515	0.9594	0.9650	0.9693	0.9727	0.9823	0.9869	0.9896

E.1.2.2 不良率管制圖 (P chart)

不良率管制圖是計數且為分數 (fraction) 形式的管制圖，主要適用於

管制比率，常用於例如製造商每生產批量下的不良率、化學工廠連續生產下轉換牌號所產生過渡料（或廢料）的重量比率等。 D 為表示在樣本中不良品的數目，其為隨機變數且服從二項分配（binomial distribution）含未知參數 p ，因此每個樣本不良率隨機變數為 $p = D/n$ 且變異數為 $\sigma_p^2 = \frac{p(1-p)}{n}$ 。若 D_i 為第 i 個樣本的不良品數，則可觀察到的平均不良率為 $\bar{p} = \frac{1}{mn} \sum_{i=1}^m D_i$ ，用以估計未知參數 p 。以正負三倍標準差為例的不良率管制圖則可表示如公式(E.6)所示。

$$\begin{aligned} CL &= \bar{p} \\ UCL &= \bar{p} + 3 \sqrt{\frac{\bar{p}(1-\bar{p})}{n}} \\ LCL &= \bar{p} - 3 \sqrt{\frac{\bar{p}(1-\bar{p})}{n}} \end{aligned} \quad (E.6)$$

E.1.2.3 平均缺陷數管制圖（U chart）

單位缺陷數圖是計數且呈現每一單位空間或時間下的次數管制圖，主要適用於管制個數，常用於例如半導體製造商記錄每個晶片（粒）的缺陷顆數、醫院記錄的每個月感染病例個案數目等。假如每個樣本批量包含 n 個單位（樣本大小）且該批全部有 C 個缺陷（總缺陷數），如果 C 為一隨機變數且服從卜瓦松分配（Poisson distribution）含未知參數 λ （期望值與變異數）。令 $u = C/n$ 為某一樣本的平均缺陷數，因此 u 為隨機變數，其期望值為 λ 且變異數為 λ/n 。若 C_i 為第 i 個樣本的缺陷數，則可觀察到的平均缺陷數為 $\bar{u} = \frac{1}{mn} \sum_{i=1}^m C_i$ ，用以估計未知參數 λ 。以正負三倍標準差為例的平均缺陷數管制圖則可表示如公式(E.7)所示。

$$\begin{aligned} CL &= \bar{u} \\ UCL &= \bar{u} + 3 \sqrt{\frac{\bar{u}}{n}} \\ LCL &= \bar{u} - 3 \sqrt{\frac{\bar{u}}{n}} \end{aligned} \quad (E.7)$$

E.1.3 時間加權管制圖

然而，單變量管制圖於階段二的應用上存在著一個缺點，就是對於微小統計量的偏移難以偵測，如圖 E.6 所示，平均數僅偏移一倍標準差也會都落於管制界限內。雖我們可藉由前述管制圖研判提及的區間測試與連串測試來輔助，然而這些規則偵測實際異常所需的平均連串長度可能過多（也就是不易權衡型一型二錯誤）。因此，時間加權管制圖是將階段二的樣本序列資訊一併納入統計量的計算中，來增加偵測微小偏移的敏感度，其中包含的兩個主要方法為「累積和管制圖」（cumulative sum control chart, CUSUM chart）與「指數加權移動平均管制圖」（exponentially weighted moving average control chart, EWMA chart）。

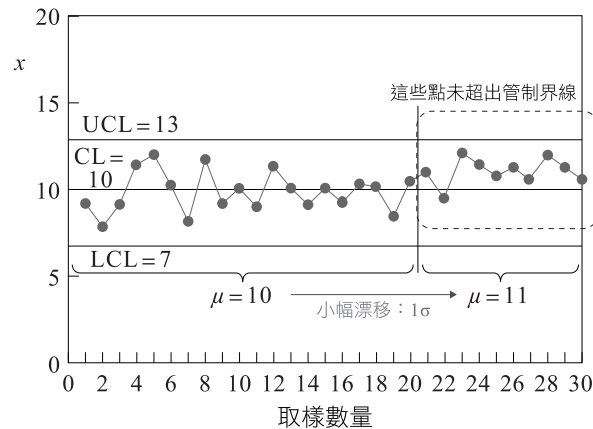
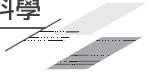


圖 E.6 無法偵測的平均數微小偏移

E.1.3.1 累積和管制圖（cumulative sum control chart, CUSUM chart）

累積和管制圖利用了階段二每一個收集的新數據去計算統計量，來增強偵測的敏感度。此方法是不斷地累加樣本平均數與中心線的偏移量，將正向與負向的偏移各別建置一個單邊的管制圖，對觀測值 i 來說，如公式 (E.8) 所示。

$$\begin{aligned} C_i^+ &= \max[0, x_i - (\mu_0 + K) + C_{i-1}^+] \\ C_i^- &= \max[0, (\mu_0 + K) - x_i + C_{i-1}^-] \end{aligned} \quad (E.8)$$



其中 C_i^+ 為大於中心線的正值累加， C_i^- 為小於中心線的負值累加，初始值 $C_0^+ = C_0^- = 0$ 。 K 為參考值是一寬鬆常數用於調整敏感度，通常設為製程偏移量的一半，也就是如果製程樣本平均數從中心線 μ_0 偏移至 μ_1 ，則 $K = (\mu_1 - \mu_0)/2$ 。管制界線 H （又稱為決策區間，decision interval）是設置的常數。若與目標值的偏差這些累積和大於 K ，兩個量 C_i^+ 或 C_i^- 會在變為負數時重置為零（也就是公式(E.8)）。如果 C_i^+ 或 C_i^- 有一者超過常數 H ，則稱製程失控。此二常數 K 與 H 需要藉由平均連串長度來輔助設定（可透過查表方式），用以較更好地權衡微小偏移的型一型二錯誤。而當累積和管制圖進行監控時偵測到微小偏移後，經異常診斷後需將累積和設為零，再進行後續監控。

E.1.3.2 指數加權移動平均管制圖（exponentially weighted moving average control chart, EWMA chart）

指數加權移動平均管制圖同樣地利用了階段二的新數據來增強偵測的敏感度，但也盡可能地強化當下時間的資訊，使其可以達到典型的平均數管制圖對大幅偏移以及累積和管制圖對微幅偏移的偵測。此方法中每個時間點 t 的統計量（例如樣本平均數 $\bar{x}_t$ ）是以指數加權平均將樣本序列進行加總，如公式(E.9)所示。

$$\begin{aligned}
 z_t &= \lambda \bar{x}_t + \lambda(1-\lambda)\bar{x}_{t-1} + \lambda(1-\lambda)^2\bar{x}_{t-2} + \cdots + \lambda(1-\lambda)^{t-1}\bar{x}_1 \\
 &\quad + (1-\lambda)^t\mu_0 \\
 &= \lambda \bar{x}_t + (1-\lambda)z_{t-1} \\
 &= \sum_{k=0}^{t-1} \lambda(1-\lambda)^k \bar{x}_{t-k} + (1-\lambda)^t\mu_0
 \end{aligned} \tag{E.9}$$

其中 μ_0 為管制中的製程目標（target）或歷史平均數，初始值 $z_0 = \mu_0$ ，而參數 λ 為一個介於 0 到 1 的常數，為當下樣本的權重大小，而第 k 期的權重則是指數加權的權重 $\lambda(1-\lambda)^k$ 。實際上，當我們將參數 λ 設定為 1 時，等於只看當下的樣本，此時管制圖會近似於平均數管制圖；當我們將參數 λ 設定為一個很小的數值時，近似於賦予所有樣本相同的權重 λ ，此時管制圖會近似於累積和管制圖。因而，計算上述統計量 z_t 的標準差後，可得出以正負三倍標準差為例的指數加權移動平均管制圖如公式(E.10)所示。

$$\begin{aligned}
 UCL &= \mu_0 + 3 \frac{\sigma}{\sqrt{n}} \sqrt{\frac{\lambda}{2-\lambda} [1 - (1-\lambda)^{2t}]} \\
 LCL &= \mu_0 - 3 \frac{\sigma}{\sqrt{n}} \sqrt{\frac{\lambda}{2-\lambda} [1 - (1-\lambda)^{2t}]}
 \end{aligned} \tag{E.10}$$

其中 $\sigma/\sqrt{n}$ 為管制中的製程標準差，此常數 λ 同樣須藉由平均連串長度來輔助設定（查表），用以更好地權衡偏移的型一型二錯誤。注意，上下管制界線並非隨時間等寬於中心線，然而標準差最後會收斂於 $\frac{\sigma}{\sqrt{n}} \sqrt{\frac{\lambda}{2-\lambda}}$ ，其中 σ 估計同 $\bar{X}$ 管制圖中依樣本數修正為 $\hat{\sigma} = \bar{s}/c$ 。

E.1.4 多變量管制圖

現今製程日益複雜，一道製程通常具有多個製程參數與品質特性，倘若我們獨立地以單變量管制圖對每一製程參數或品質特性進行監控，可能因管制界限過於寬鬆而無法偵測出離目標值過遠的樣本統計量，以及缺少考慮變數間相關性的資訊，如圖 E.7，單變量管制圖皆於管制中，但考慮變量間相關性後，多變量管制發生失控。因此，多變量管制圖的發展逐步地從獨立的單變量管制圖，到考慮變數間相關性的「霍特林 T 平方多變量管制圖」（Hotelling's T-squared multivariate control chart），甚至是不需要假設服從分配的無母數「基於核距離多變量管制圖」（kernel-distance based multivariate control chart, K chart）。如圖 E.8 所示，(a)表示傳統修華特單變量管制圖，它僅能管制一個變數並以常態假設為前提，聯立分別監控兩個以上的變數；(b)與(c)表示以多變量是否相關的情況下建構 Hotelling's T^2 並服從多元常態分配，二維空間以橢圓的輪廓包覆資料；(d)是適應真實數據，不藉由任何分配情況下建構的 K 管制圖，並以這群數據中的支持向量為基底，建構不規則的輪廓將數據包覆。

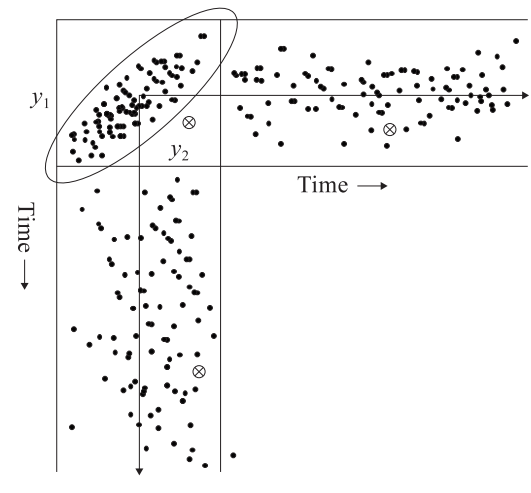


圖 E.7 單變量皆於管制中但考慮變量間相關性後多變量管制失控

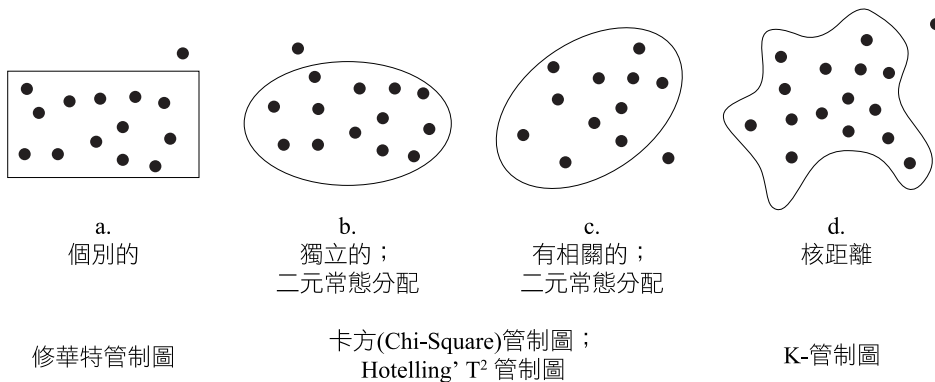


圖 E.8 多變量管制圖的發展 (Sun and Tsung, 2003)

E.1.4.1 霍特林 T 平方多變量管制圖 (Hotelling t-squared multivariate control chart)

我們將管制圖由單變量常態分配拓展至多變量常態分配，實際上是同時對多個變量的母體平均數與其相關性做假設檢定，如公式(E.11)所示，

$$\begin{aligned} H_0: \boldsymbol{\mu} &= \boldsymbol{\mu}_0 \\ H_a: \boldsymbol{\mu} &\neq \boldsymbol{\mu}_0 \end{aligned} \quad (\text{E.11})$$

其中 $\boldsymbol{\mu}_0$ 是 p 個維度的平均數向量。由此可計算 T 平方統計量，如公式(E.12)所示，

$$t^2 = (\bar{\mathbf{x}} - \boldsymbol{\mu}_0)' \boldsymbol{\Sigma}_{\bar{\mathbf{x}}}^{-1} (\bar{\mathbf{x}} - \boldsymbol{\mu}_0) \quad (\text{E.12})$$

其中 $\bar{\mathbf{x}}$ 是 p 個維度的樣本平均數， $\boldsymbol{\Sigma}_{\bar{\mathbf{x}}}$ 是平均數的共變異矩陣， $\boldsymbol{\Sigma}$ 是共變異矩陣，則 $\boldsymbol{\Sigma}_{\bar{\mathbf{x}}} = \boldsymbol{\Sigma}/n$ 。並且此 T 平方統計量服從 p 個維度的卡方（chi-squared）分配。當我們應用至多組樣本的多變量管制圖時，每組樣本 i 的 T 平方統計量如公式(E.13)所示，

$$t_i^2 = (\bar{\mathbf{x}}_i - \bar{\bar{\mathbf{x}}})' \mathbf{S}_{\bar{\mathbf{x}}}^{-1} (\bar{\mathbf{x}}_i - \bar{\bar{\mathbf{x}}}) \quad (\text{E.13})$$

其中 $\mathbf{S}_{\bar{\mathbf{x}}}$ 為平均數的樣本共變異矩陣。此時 T 平方統計量將服從以下分配，其中 $F_{p, mn-m-p+1}$ 為 F 分配包含兩個參數 p 與 $mn - m - p + 1$ 。

$$t^2 \sim T_{p, (m-1)(n-1)}^2 = \frac{p(m-1)(n-1)}{mn - m - p + 1} F_{p, mn-m-p+1} \quad (\text{E.14})$$

因此「霍特林 T 平方多變量管制圖」可表示如公式(E.15)所示，

$$\begin{aligned} UCL &= \frac{p(m-1)(n-1)}{mn - m - p + 1} F_{\alpha, p, mn-m-p+1} \\ LCL &= 0 \end{aligned} \quad (\text{E.15})$$

其中上管制界限是由給定信心水準 α 下的 F 分配（實際上是為了計算對應的 p 值，將 T 平方統計量轉換至 F 分配中計算），而下管制界限是平均數的中心因此為零。霍特林 T 平方多變量管制圖是一個單邊的管制圖（僅有 UCL ），我們可由其 T 平方統計量的幾何意義來理解。 T 平方統計量實際上是計算在多變量高維度的空間中，樣本平均數至整體平均數在考慮變數間相關性下的加權距離（概念上可參考馬氏距離，Mahalanobis distance），因此離中心越近時， T 平方統計量會趨近於零。此外，我們也可將多變量管制圖拓展至時間加權管制圖上，形成多變量累積和管制圖與多變量指數加權移動平均管制圖（Multivariate EWMA chart）。

E.1.4.2 基於核距離多變量管制圖（kernel-distance based multivariate control chart, K chart）

典型的多變量管制圖大多是基於有母數的分配，通常假設統計量服從常態分配（例如霍特林 T 平方多變量管制圖），但實務上要符合假設的分配是不容易（尤其是多變量的分配）。基於核距離多變量管制圖（Sun and

Fugee, 2003) 則是一個無母數的管制圖，其常用方法「支持向量數據描述」(support vector data description, SVDD) 的方法將管制中的樣本以一個「超球體」(hypersphere) 包覆(完整方法介紹請參閱章節「支持向量機」)，此超球體不需服從任何分配而是直接的基於數據特性所建構的，因而能更好地適應變數關係複雜的數據。相較於典型管制圖需求出樣本統計量，基於核距離多變量管制圖則是以支持向量數據描述所計算的「核距離」(kernel distance) 作為管制圖監控的對象，而其管制界限則是由支持向量所建構而成的，然而我們必須確保支持向量不包含異常值(可透過調整超參數來觀察)。以圖 E.9 說明，左圖是將數據以二維方式呈現，可以發現共有 7 個支持向量包覆在群體數據的外側，同時也可以發現有 7 個點明顯落於群體數據較遠處；右圖將這些結果以管制圖的方式呈現，它的管制界限是用支持向量計算出，因此共有 7 個點落於界限上，並且會有 7 個屬於群體數據外側超過管制界限。

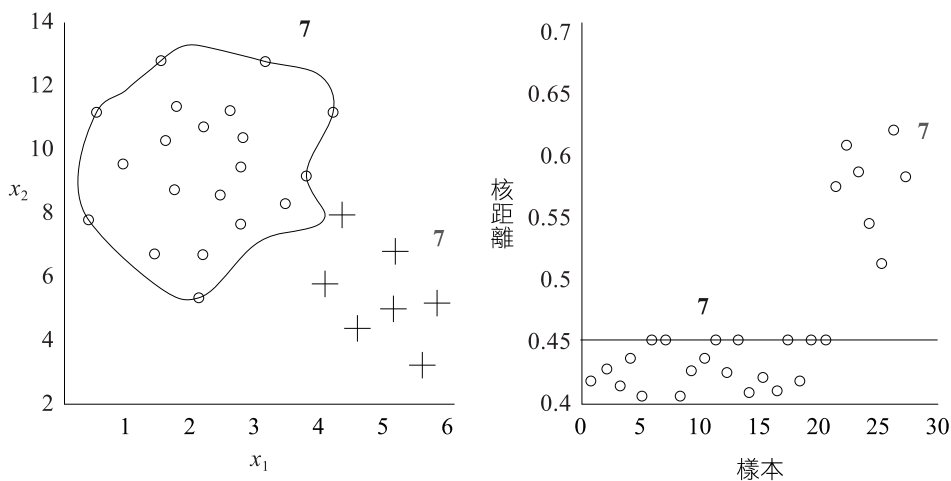


圖 E.9 K 管制圖的建構

在產線上，管制圖是線上製程監控不可或缺的一項利器，也是維持與降低製程變異的一道重要防線，在實務上，單變量管制圖可用以針對單一變數的監控；時間加權管制圖用於更精準地偵測製程微小偏移；多變量管制圖可用於同時對多個變數進行監控；更值得注意的是基於核距離多變量管制圖不需任何分配的假設，可用以對變數關係複雜的數據進行監控。

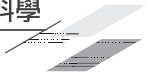
然而，SPC 仍存在著某些缺點。首先 SPC 透過數據抽樣進而繪製管

制圖，由於計算一批樣本的統計量有時間延遲的狀況，因此當人員採取警報或異常處理措施通常都稍嫌太晚。第二，大多數管制圖有分配的假設，然高複雜度的製程（例如半導體製造）其設備運轉狀況通常具有非平穩性（non-stationary）、自相關性（auto-correlated）及交互關聯性（cross-correlated）。以半導體製造為例，在微影製程中，曝光前烘烤的時間及溫度、曝光光源強度、焦距及顯影時間和溫度等皆存在高度相關，而非個別獨立的製程參數，若同時考慮這些相互關聯的變數，落在管制界線內的機率不等於個別使用管制圖時的結果，因此管制圖可能會發生誤導，即使使用多變量管制圖，在變數相當多以及變數間高階（high-order）交互作用下，依然會有許多困難與挑戰，這也是我們要努力的方向。

E.2 錯誤偵測與分類

所謂的異常（anomaly），就是生產過程中的觀測值、監控值或其參數偏離可接受範圍。故障（fault）可定義為過程異常，例如發現反應室中溫度異常（果），而造成溫度異常的原因為冷卻泵或其控制器故障（因），這我們稱為根本原因（root cause）。錯誤偵測與分類（fault detection and classification, FDC）主要用於異常事件的偵測與故障排除，協助自動化工程分析系統的發展。FDC 分為兩個部份，分別是錯誤偵測（fault detection, FD）以及錯誤識別（fault classification, FC）。FD 是藉由建構特徵或相關指標監控系統，使工廠運作發生異常時可以提早警報（pre-alarm）處理，以避免異常事件造成重大損失，達到預防的效果，同時能降低停機時間（減少產能損失）及提高生產品質以提升產業競爭力，前述章節 SPC 多著墨於 FD。FC 是當異常發生時工程人員將需要從異常品生產的紀錄或是機台的參數找出蛛絲馬跡，最後依照分析方法與專業知識經驗將問題分門別類，建立異常管理資料庫，以縮短未來相似故障發生時的異常解析時間。換言之，錯誤診斷任務可視為分類問題，以建構診斷系統（fault diagnosis system）或診斷分類器（diagnostic classifier）。在 FDC 分析中，對一個回饋系統來說（例如機台、設施、製造執行系統等），常見之故障類別如下，每種故障的分析手法也不盡相同。

- 系統參數變異：當外在環境發生變化，干擾將透過系統邊界 / 介



面傳遞到系統本身，就會影響到參數而導致錯誤與故障，例如機台中的粒子粉塵（particles）導致產品電路短路、熱交換器中的污垢引起的傳熱係數的變化。

- **機構故障**：由設備故障所產生的異常，該故障將造成變量間之數據發生變化，例如傳動系統中頻繁正反向運動、進給運動的負載過大、潤滑不良、螺釘鬆動等所造成的機構故障或損毀。
- **感測器或儀器故障**：可能是儀器本體故障，有恆定偏差（正或負）或超過範圍設定之故障。某些關鍵儀器或量測機台對於製程相當重要，若無法及時檢測到故障並調整，則儀器故障可能導致生產狀態變量偏離可接受的管制界限。

E.2.1 錯誤偵測與分類方法

隨著資訊科技的進步，由人員手動控制的動作（例如調節控制中的開關閥門），漸漸被自動控制所取代。在大數據時代下，現代工廠的設備中裝設大量的感測器以收集海量資料進行製程行為的分析。以半導體廠為例，晶圓製造的製程複雜且冗長，其中包含大量製程、配方以及設備，必須憑藉著先進的資訊傳遞技術與感測器，自動記錄生產中的數據。在此情形下，系統須能負荷每秒鐘所產生大量的資料，同時也要求系統能快速響應，反應即時狀況且降低因異常造成的產能損失。FDC 在實務上的運用，多年來已經開發了不同類型之方法，從建模的角度來看，大致上分為三大類別（Venkatasubramanian et al., 2003a; 2003b; 2003c）：

E.2.1.1 基於模型的方法（model-based method）

藉由物理模型或數學函數表示特定製程的輸入與輸出關係式。此類型的優點在於特定情境下準確度高，參數確定後計算速度快，適用於系統變異小的情況。其缺點為若隨環境變化大其參數估計不易且需耗費計算資源，因此不易於全面性實施。以下介紹三種基於模型方法的模型：

1. **觀測器（observer）**：此方法開發了一組觀測器，每個觀測器會生成一個殘差，透過這個殘差能檢測及識別特定的故障類型（failure mode），可藉由監控輸入值找出與之產生顯著殘差的觀測器以此識別故障，其優點為決策較少受到噪聲及各類不確定性影響。
2. **奇偶關係（parity relations）**：此模型是檢查傳感器的輸出值（測量值）與過程輸入的奇偶性（一致性），在理想的穩態運行條件下，

奇偶方程的值（殘差項）為零。在實際情況中，由與各式各樣的原因會使得殘差非零（例如傳感器和執行過程中的錯誤或故障）。

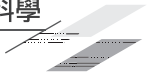
3. 卡爾曼濾波器（Kalman filters）：製程運作時伴隨干擾，通常只有統計參數是已知，在此狀況下，估計值將伴隨的誤差使得診斷系統表現不佳。卡爾曼濾波器是一種最佳化狀態估計的遞迴算法，目標是設計一個估計誤差最小的狀態估計器。

E.2.1.2 基於知識的方法（knowledge-based method）

仰賴於專業知識為主體，藉由模擬人類專家在實務中解決問題的邏輯行為，建立 if-then 規則和推理引擎，來執行診斷。但其模型對於人的依賴性高，在實務上特徵間的關係錯綜複雜，因此較為主觀，不易快速確認異常原因。以下介紹基於知識方法的模型：

1. 因果有向圖（directed graph based causal models）：此方法是將因果模型中的變量集都視為頂點集合，其中互相關連的變量使用有方向的邊進行連結，被指向的那一方為某變量的影響對象，並且利用+或-表示正或負影響的因果關係，知識圖譜（knowledge graph）是常用的方法之一。
2. 故障樹（fault trees）：故障樹是一種由上往下的邏輯推論，樹的頂層代表著某故障或狀態，藉由對這狀態抽絲剝繭的分析（why/what/how），每一層都建立節點（代表不同的狀態），不斷地往下延伸直到為獨立狀態為止，其中每一層的節點使用 and 或 or 的邏輯符號進行連結。優點為相較於因果關係圖簡潔易懂，缺點是不易開發且無法解釋所有因果關係、模型準確性評估困難。
3. 抽象階層（abstraction hierarchy）：抽象架構是一個多層的表示架構，由物理和功能單元組成，本方法在彙集各項子系統行為規則來推論出整個系統狀態。系統分解的一個重要原則是局部性原則：某一部分的規則不能具體引用任何其他部分。
4. 德爾菲法（Delphi method）：是一種透過平台協同多位專家各自獨立發表意見，平台反覆彙集討論資訊並客觀摘要說明，以收斂達成最後共識提出結論。

在以基於知識的方法架構完成因果關係後，診斷搜索策略的類型在故障診斷中基本上有兩種不同的搜索方法：



1. **拓撲搜索 (typology search strategies)**：拓撲搜索以正常運作的模板為基礎進行故障搜索，當故障發生時，模板將在系統某個位置發生不匹配之情形，此時將依據系統發生不匹配位置識別出異常。
2. **症狀搜索 (symptomatic search)**：使用一組系統異常狀態的觀察結果作為搜索模板，在不同的異常狀態下從已知的故障資料庫中進行比對找出故障，這種類型的搜索稱為症狀搜索。該方法的特點是它們的決策是從數據集的結構及其內部關係中得出的，而不是從系統屬性的拓撲結構中得出的。優點為從資訊運算的角度來看，症狀搜索是有利的；缺點是必須有實際的異常狀態存在資料集內，因此不易識別多個故障和新型態的干擾。

E.2.1.3 基於數據驅動的方法 (data-driven method)

該模型並未有預先假設或先驗知識 (priori knowledge)，僅依賴歷史數據進行分析，分析的方法包含統計分析、數據挖掘和機器學習。此類方法從大量的歷史傳感器數據資料中學習製程的設備狀況與行為，並將轉換後的特徵變為先驗知識提供人員進行工程驗證，也就是特徵工程 (feature engineering)，特徵可為離散型或連續型。建模上主要有定量與定性方法，其中定量方法包含機器學習、數據科學與統計方法等，可參閱本書相關章節，此處介紹定性方法敘述如下：

- **專家系統 (expert system)**：根據領域知識中的規則 (rule) 與特徵進行提取。在實務上，專家系統可以解決特定專業領域中的問題，該系統包含了四個部分，包括知識獲取、知識表示的選擇、知識庫中的編碼、診斷推理程序的開發。其優點為開發容易且能為結果提供物理解釋，然由於針對特定領域知識所建構的專家系統僅能處理特定問題，所以應用廣度受限。
- **定性趨勢分析 (qualitative trend analysis, QTA)**：趨勢建模可用於解釋過程中發生的各種重要事件，進行故障診斷和預測未來狀態，為了不受干擾因素的影響，通常搭配濾波器進行分析。假設 x 軸為時間與 y 軸為響應值，當故障發生時，會產生明顯的變化趨勢，使用者可以藉由趨勢對過程進行推理也可以識別過程中的潛在異常，具有預警式的效果。

本節針對 FDC 系統中常見的三種模型：基於模型、基於知識、基於

數據驅動，進行概括性的介紹，圖 E.10 描述基於不同模型所搭配的技術。

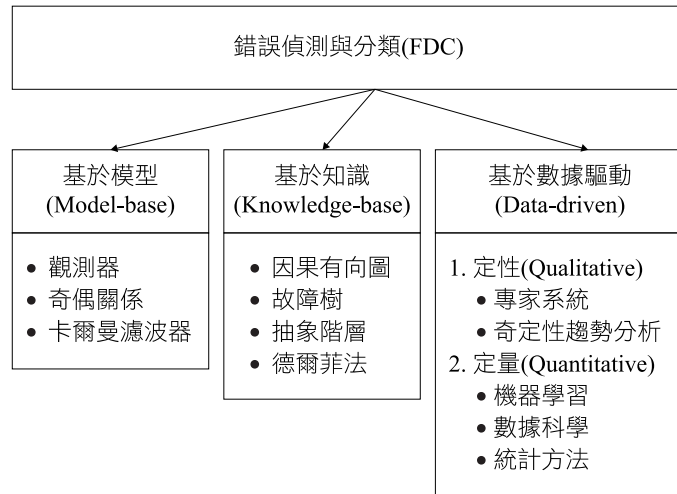


圖 E.10 FDC 模型種類及所搭配之技術

在製造業大數據的時代，製程複雜度與生產規模以指數爆增，各式各樣的故障類型參雜其中。單純倚靠專業領域的經驗及先驗知識為主的方法（例如基於模型與基於知識），在製程上下游與變數之間交互作用的情況下，難以將數學函數及物理關係完整地呈現。因此，透過數據科學的手法將歷史資料中潛在的先驗知識提供給使用者進行參考的數據驅動方法已然成為趨勢，與前兩者相比，更適合用在複雜製程進行故障檢測與診斷。

E.2.2 錯誤偵測與分類系統流程、架構與績效指標

FDC 系統的分析流程，主要從傳感器裡的資料輸入系統後，將經由一系列的計算最終轉換為診斷結果供人員進行判斷。學理上來說，分析過程相當於是進行函數的映射（mapping），從某個空間轉換到另一個空間的過程。如圖 E.11 所示，從測量空間（measurement space）、特徵空間（feature space）、決策空間（decision space）、到類別空間（class space）的轉換過程，其中轉換可依據基於模型、基於知識、或數據驅動方法，以先驗知識搭配搜尋方法進行。

- **測量空間**：從感測器所直接得到的原始資料且不包含先驗知識，通常數據的解析度跟感測器所收集的變數與抽樣頻率有關。

- **特徵空間**：可利用先驗知識或數據科學的方法進行相關特徵工程，目的為了將資料轉換為資訊以利診斷。
- **決策空間**：當資訊映射在決策空間時會依所設定之目標（損失）函數（例如最小化錯誤分類）建立判別函數，其決策邊界（decision boundary）可使用學習演算法建構，經常權衡於多目標（multiple objectives）之間，其中最簡單的方法可為設定一個閾值（threshold）即可實現，例如判別正常與故障。
- **類別空間**：此空間可以是一組整數集合 $c = [c_1, \dots, c_M]$ ，以呈現 M 個故障類別（水準），根據情境可使用不同的方法，例如閾值函數、樣型匹配或符號邏輯，將決策空間映射到類別空間進行檢索，最後明確指示診斷模型所判斷的故障類別供使用者參考。

一般來說，決策空間與類別空間在大多數情況下具有相同的維度，以方便歸類。但某些情況下分類器無法給出明確的判斷，因此決策空間與類別空間仍可保持其獨立性。例如神經網路中，我們需要從決策空間（output layer）基於某函數（例 softmax）作出分類解釋結果。

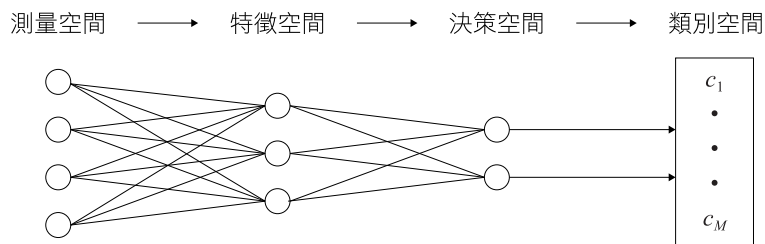


圖 E.11 資料轉換流程與神經網路的運作模式

接下來介紹基於數據驅動的 FDC 建模流程，可依照使用者需求或資料狀態找出適合的模型，如圖 E.12。共區分三個區塊，類別空間需要先驗知識或資料分析後的結果進行設置，設置依據可依照需求以及製程狀況而定。其中最重要的程序為測量空間至特徵空間，決定 FDC 的績效結果主要在這個階段，包含資料預處理以改善數據品質，以及特徵工程以轉換出重要特徵。因此針對感測器的資料需進行資料品質評估，再利用有限的先驗知識或數據分析手法進行特徵挑選（去除與響應無關之變數）或特徵萃取的動作（降低維度），用以提升模型對於診斷的精確性以及運算速度

(降低模型複雜度)。最後再依照資料類型與診斷目的，選擇適合的模型以權衡預測準確性與可解釋性（各方法細節請參閱本書相關章節）。

為了輔助使用者依照系統需求選擇合適的模型，文獻中提出診斷系統評估的十個面向，以瞭解診斷系統的整體表現。學理上，沒有一個單獨的模型可以完全滿足這十項特性，某些特性存在著相互權衡的關係，需要使用者依照系統需求進行取舍（Venkatasubramanian et al., 2003c），如表 E.3 所示。

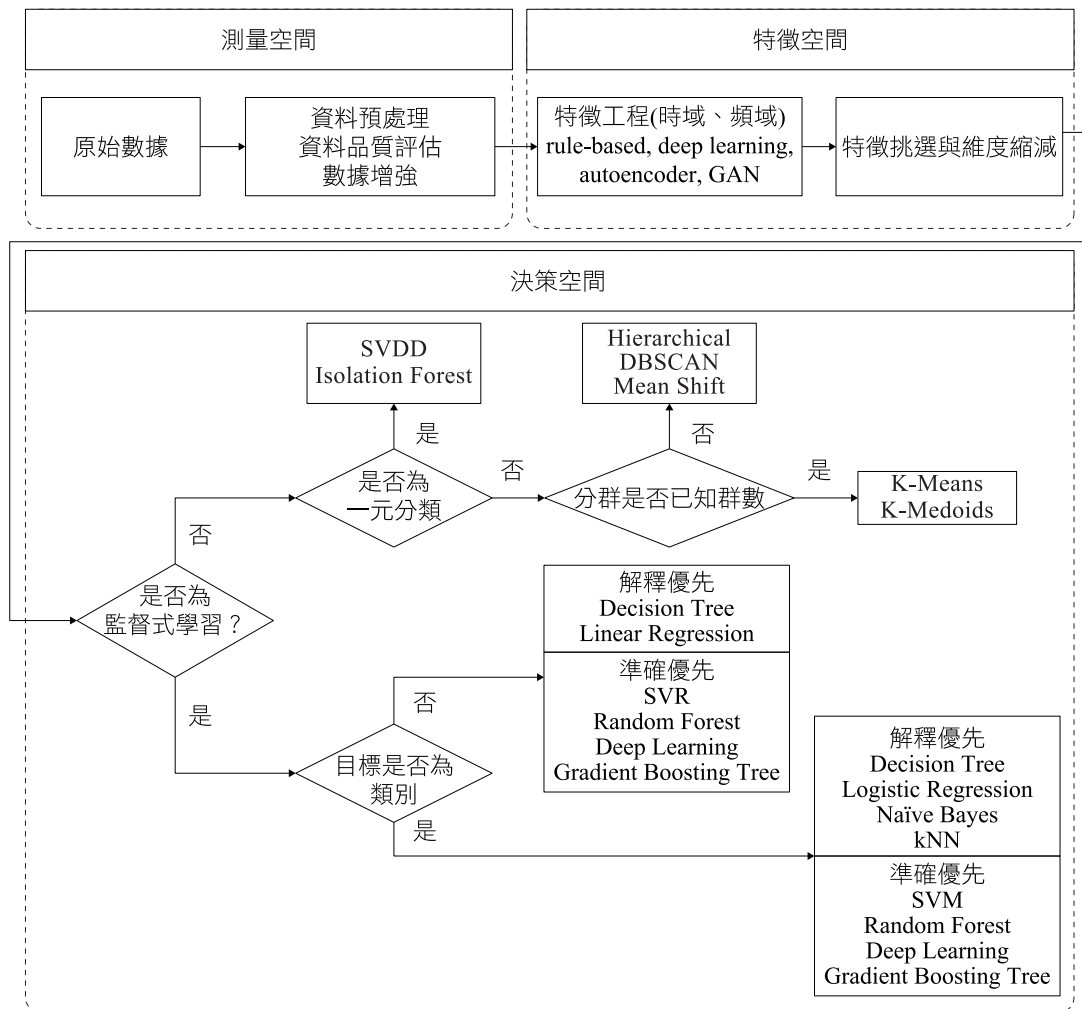


圖 E.12 數據驅動的 FDC 建模流程

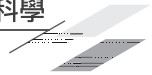


表 E.3 FDC 系統的 10 項評估特性

編號	特性	解釋
1	快速偵測和診斷 (quick detection and diagnosis)	診斷系統對於異常應快速反應並立即進行故障排除建議根本原因，但這同時代表著系統對於噪音 (noise) 敏感，可能導致誤報的情況。
2	可隔離性 (isolability)	診斷系統能有效區分不同故障之能力。
3	穩健性 (robustness)	診斷系統對各種噪音和不確定性具有穩健性，與第一點相互權衡。
4	新穎可識別性 (novelty identifiability)	系統能確定當下運作是否正常，並且當異常發生時可以區別是已知故障類型還是未知的新故障。
5	分類錯誤 (classification error)	診斷系統要能計算或估計分類錯誤，並應減少分類錯誤情況或權衡型一與型二錯誤成本，並能建立模型可靠的信心指標。
6	適應性 (adaptability)	系統要能夠隨者環境外部輸入變化或新串流資料不斷地收集，並透過訓練機制 (或重新訓練) 進行自我適應與調整。
7	解釋性 (explanation facility)	對於診斷結果可解釋其因果關係或物理上的意涵。
8	建模需求 (modelling requirements)	對於開發診斷系統所需的建模工作應盡可能地減少，即減少多次開發或多次建模的情形。
9	存儲和計算需求 (storage and computational requirements)	診斷系統應具備數據儲存、計算與分析的功能。
10	多重故障識別能力 (multiple fault identifiability)	在實務上缺陷有時不會單一出現，可能同時伴隨著多種類型的故障，診斷系統必須具備識別多重故障的能力及其多重故障間的關聯。

E.2.3 半導體製造廠案例

此處以舉半導體製造廠導入數據驅動 FDC 系統的進行介紹，為了提升良率並快速故障排除 (troubleshooting)，針對良品晶圓與缺陷晶圓的數據收集，透過數據科學方法鑑別關鍵的狀態變數 (state variable identification, SVID) (例如溫度、壓力、濃度等)，並找到在加工過程中關鍵的加工時段 (processing time step)，以協助製程診斷 (Fan et al., 2020)。

例如針對化學氣相沉積 (CVD) 設備中的腔體 (chamber)，FDC 相關數據收集如圖 E.13 所示，其數據立方體呈現數據收集的三個維度 (SVID 變數、晶圓樣本、紀錄時間)，標籤為二元變數 (正常與異常) 並可以不同的視角來進行達成不一樣的分析目的。例如欲探討不同加工時間點下，影響良率關鍵的 SVID 變數是否有所不同，此時我們可使用數據集

視角為圖 E.13 右上圖，此為針對某一加工時段下的樣本與變數的表單；另一方面，若欲探討不同 SVID 變數下，影響良率關鍵的加工時段是否有所不同，此時我們可使用數據集視角為圖 E.13 右下圖，此為針對某一變數下的樣本與加工時段的表單。

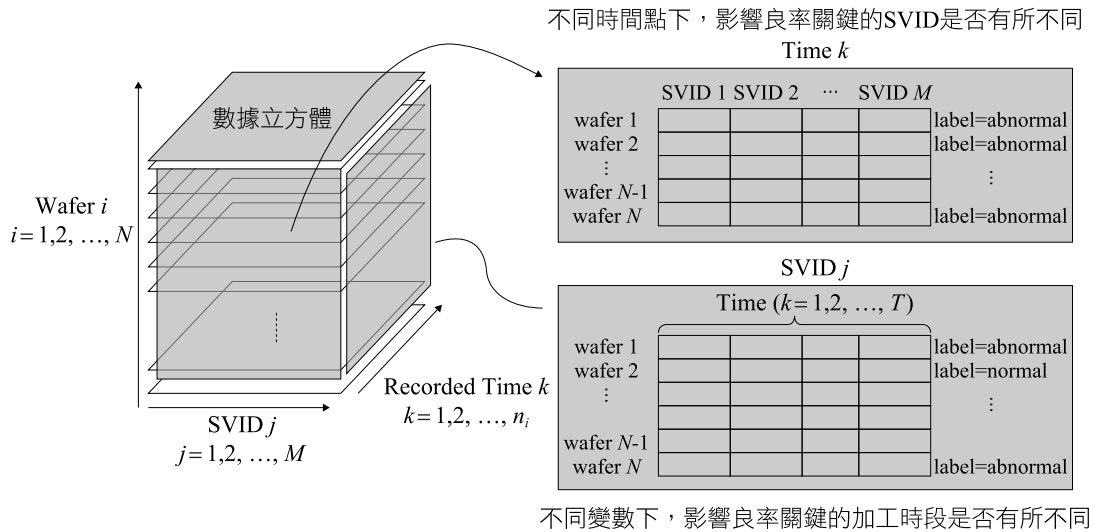


圖 E.13 FDC 數據與不同視角的分析目的 (Fan et al., 2020)

圖 E.14 半導體製造 FDC 建模流程，以流程圖的方式呈現。首先，在測量空間需要對原始數據進行預處理與資料品質評估，並改善數據不平衡、去除無意義變數、以及將傳感器數據與測量時間點不同步的問題進行修正。在特徵空間，透過特徵挑選方法對各 SVID 變數進行重要性評估，同時以分群技術，分出重要性高的 SVID 與重要性低的 SVID，並自動化產生閾值 (threshold)。由於在不同機台或腔體 (chamber) 可能會有所差異 (也就是「機差」問題)，因此可對平行機台分別找出重要 SVID 後再取交集。在決策空間中，可對每個單一 SVID 各別建構獨立診斷器，最後以集成學習進行分類。在集成學習中通常選擇弱分類器 (weak classifier) 例如 K 最近鄰法 (kNN) 或樸素貝式分類器 (Naïve Bayes)，可根據分類器結果評估的靈敏度 (sensitivity) 與特異度 (specificity) 合成評價指標，找出最後關鍵的 SVID 變數。接著可進行集成學習以投票方式決定分類，或是以串接方式將所有重要 SVID 併在一起建構單一分類器後分類，

根據最後結果可選擇適當的分類器作為決策空間的主架構。

基於數據驅動的 FDC 模型可以從上述的說明與案例，瞭解其建模彈性大，方法模型可互相搭配或使用集成的技術使分類效果更佳，並且容易權衡十個 FDC 系統評估特性。在先驗知識方面，專家的經驗可以輔助模型學習，另外模型可以提供不易直接觀察到的因果關係或進階交互作用，供從業人員進行參考，創造人機互相成長的學習環境。

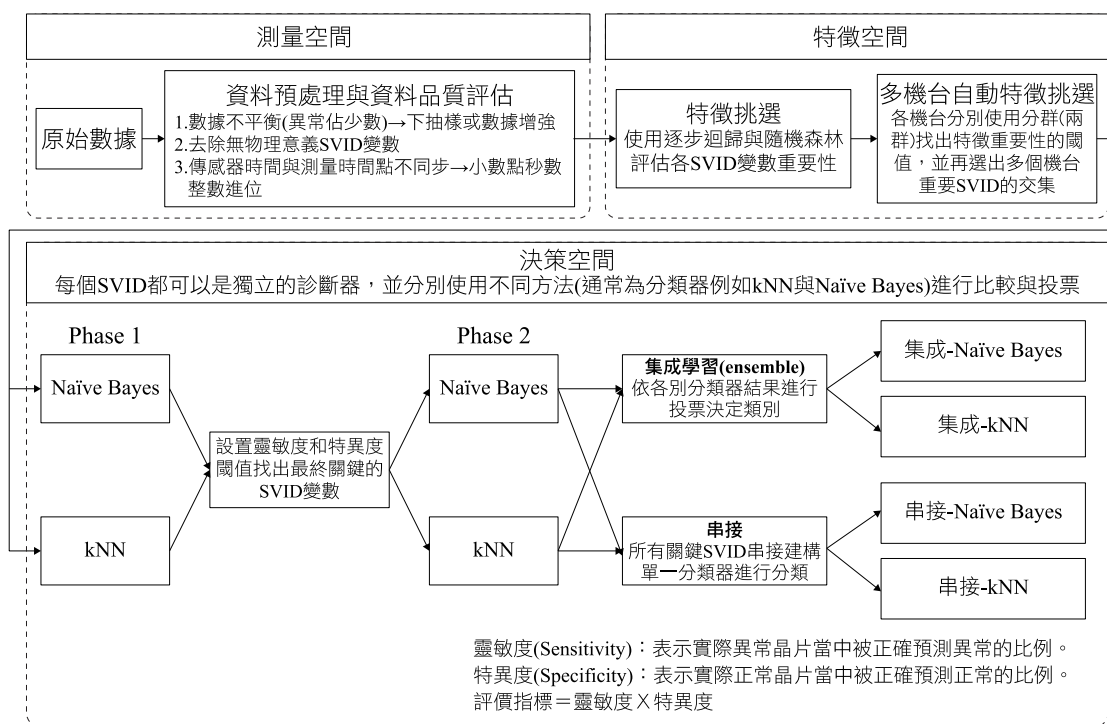


圖 E.14 半導體製造 FDC 建模流程 (Fan et al., 2020)

E.3 先進製程控制

在瞭解偵測異常的統計製程管制 (SPC)，以及故障排除根因判別的錯誤偵測與分類 (FDC) 後，會發現兩者都屬於被動的分析且通常會有時間延遲的情況。對於現場的即時參數調整，先進製程控制 (advanced process control, APC) 是常用的方法之一。實務上，一個 APC 系統兼具多項功能，包含即時機台狀態監控、即時錯誤診斷及分類、預測保養、以及

前饋 / 反饋批次控制 (feedback/feedforward run-to-run control, R2R) 等，此章節我們著重討論**批次控制**。批次控制為一種在批次與批次之間調整機台配方 (recipe) 參數或其參數組合的控制方法，通常用於消除**初始參數偏差** (bias)、**製程偏移** (shift) 或**環境干擾**等，以確保每個批次的運行保持最佳狀態。圖 E.15 左圖呈現根據每批 SPC 量測值結果，會將前量值與後量值回傳到 APC 系統，以進行批量控制的參數補償。圖 E.15 右圖呈現控制方塊圖，以說明典型的批次控制架構 (Liu et al., 2018)。一般來說，機台的動作控制會利用內部回授機制使機台的輸出動作達到目標穩態，也就是右圖中的**內循環** (inner loop)。然而為避免因環境變化、前製程變異、刀具磨損、或零組件老化等因素使得輸出動作有偏差或錯誤導致產品不良，我們可利用加裝感測器獲得當下機台的相關輸出參數或量測值 y_j (如產品尺寸、位置、壓力、溫度等)，其中 $j = 1, \dots, m$ 代表批次數。再根據批次控制與應達到的目標值 T 調整下一批次參數 u_{j+1} (需補償的參數)，也就是右圖中的**外循環** (outer loop)。

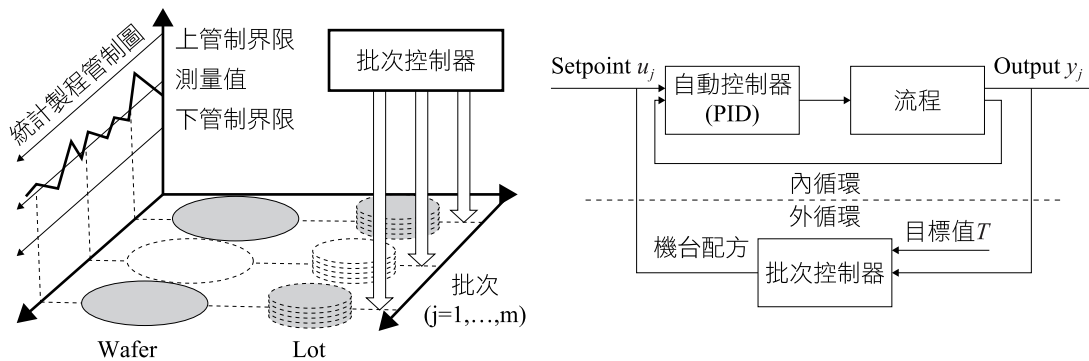


圖 E.15 半導體製造批次控制與其流程架構 (Liu et al., 2018)

E.3.1 先進製程控制器

在實務上，有許多種不同類型的批次控制器，以下我們介紹幾種常用的典型控制器，包含「指數加權移動平均控制器」(exponentially weighted moving average controller, EWMA controller)、「雙重指數加權移動平均控制器」(double exponentially weighted moving average controller, D-EWMA controller)、「卡爾曼最小變異控制器」(Kalman minimum variance controller)、「最佳自適應品質控制器」(optimizing adaptive quality

controller, OAQC) (Castillo and Yeh, 1998; Liu et al., 2018)。

E.3.1.1 指數加權移動平均控制器

指數加權移動平均控制器 (EWMA controller) 是現場常被推薦的控制器之一，其主要理由為簡單易懂且直觀的線性補償 (compensate)，適用於製程沒有遭遇嚴重的漂移 (drift) 或較大的偏移 (shift)，但其無法處理二階 (second order) 系統的補償。EWMA 控制器的方塊圖如圖 E.16，其中 B 為「倒偏移運算元」(backshift operator) $By_j = y_{j-1}$ 。

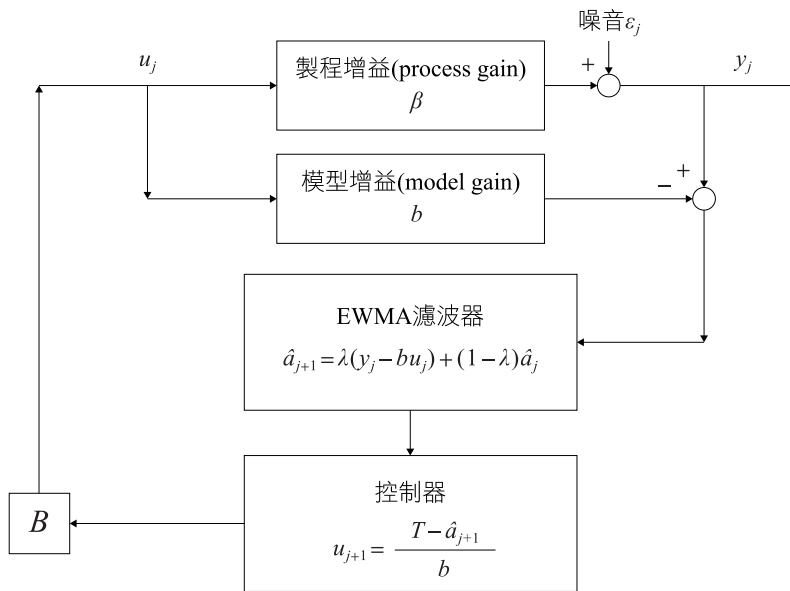


圖 E.16 指數加權移動平均控制器

假設在第 j 批次的輸出值 y_j 與輸入值 u_j 為線性關係，則可表示如公式 (E.16) 所示，

$$y_j = \alpha_j + \beta u_j + \varepsilon_j, \quad \varepsilon_j \stackrel{\text{i.i.d.}}{\sim} N(0, \sigma_\varepsilon^2), \quad j = 1, \dots, m \quad (\text{E.16})$$

其中 α_j 為截距項 (shift)， β 為斜率項表示製程增益 (process gain)， ε_j 為批次 j 的製程干擾。由於 α_j 和 β 未知，因此可由實驗設計或歷史資料迴歸估計進而建立預測模型為： $\hat{y}_j = \alpha_j + \beta u_j$ ， β 為模型增益 (model gain)。假設輸出的目標值 (target) T ，若依預測模型的調整輸入值為 $u_j = (T - \alpha_j)/\beta$ ，則會造成與目標值之間有個初始偏差 κ_0 ，如公式 (E.17) 所示，

$$\kappa_0 = \alpha_j + \beta \frac{(T - a_j)}{b} - T \quad (\text{E.17})$$

為了消弭此偏差，我們利用遺忘因子 λ （即採用部分上一次的比例， $0 < \lambda \leq 1$ ），不斷地調整預測模型的截距項 a_j ，進而更新製程的輸入值，其方法如公式(E.18)所示，

$$\hat{a}_{j+1} = \lambda(y_j - bu_j) + (1 - \lambda)\hat{a}_j \quad (\text{E.18})$$

因此，控制器的輸入建議將修正如公式(E.19)所示，

$$u_{j+1} = \frac{T - \hat{a}_{j+1}}{b} \quad (\text{E.19})$$

直到調整至可使期望輸出值達到目標值 T 的截距項（此時 $\kappa = 0$ ），如公式(E.20)所示，

$$a' = T - \left(\frac{b}{\beta}\right)(T - a) \quad (\text{E.20})$$

上述參數調整過程如圖 E.17 所示。首先， y_j 為在第 j 批次時的真實輸出， $\hat{y}_j$ 為模型的預測輸出，在一開始設定輸入值 u_0 後可得到首批的輸出 y_0 ，其與目標值 T 之間的差異為初始偏差 κ_0 。為了使預測模型達到 T ，我們藉由移動截距項，從截距 a 修正至 a' （由實線移至虛線），同時得到模型推薦的輸入值 u^* 。

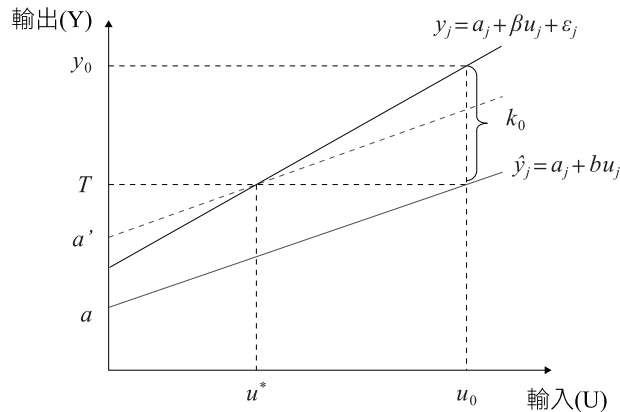


圖 E.17 EWMA 控制器參數調整示意圖

E.3.1.2 雙重指數加權移動平均控制器

雙重指數加權移動平均控制器（D-EWMA controller）修正了指數加權移動平均控制器無法有效地處理製程發生確定性的漂移趨勢（例如逐漸向上或逐漸向下）。此控制器建構兩個 EWMA 濾波器，其中第一個濾波器用於估計真實的輸出，如同典型單一的 EWMA；第二個濾波器用於估計趨勢（trend） d ，因此可以補償目標追蹤的延遲（lag），當製程正在經歷漂移的過程。D-EWMA 控制器的方塊圖如圖 E.18 所示。

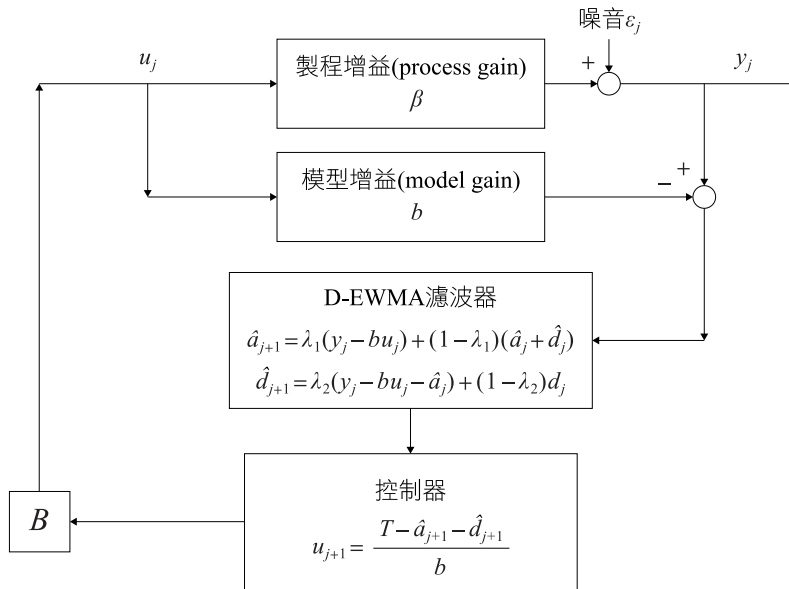


圖 E.18 雙重指數加權移動平均控制器

隨著機台運轉的時間與次數增加，機台內部由於磨損老化、沉積物增加等，造成了製程漂移的現象。為了解決這個情況，我們將公式(E.16)增加了一個截距漂移項 d_j ，而關係式改為如公式(E.21)所示，

$$y_j = \alpha_j + \beta u_j + d_j + \varepsilon_j, \quad \varepsilon_j \stackrel{\text{i.i.d.}}{\sim} N(0, \sigma_\varepsilon^2), \quad j = 1, \dots, m \quad (\text{E.21})$$

由於截距漂移項 d_j 會影響到截距項 α_j ，因此我們需要分別控制這兩項。D-EWMA 控制器由兩個 EWMA 控制器組成，主要控制估計式的截距項 α_j 及截距漂移量 d_j ，截距項與截距漂移量的控制分別如公式(E.22)與公式(E.23)所示，

$$\begin{aligned}\hat{a}_{j+1} &= \lambda_1(y_j - bu_j) + (1 - \lambda_1)(\hat{a}_j + \hat{d}_j), \quad j = 1, \dots, m, \\ 0 &< \lambda_1 \leq 1\end{aligned}\quad (\text{E.22})$$

$$\begin{aligned}\hat{d}_{j+1} &= \lambda_2(y_j - bu_j - \hat{a}_j) + (1 - \lambda_2)\hat{d}_j, \quad j = 1, \dots, m, \\ 0 &< \lambda_2 \leq 1\end{aligned}\quad (\text{E.23})$$

其中，初始截距項 a_0 由實驗設計或歷史資料迴歸估計獲得，初始截距漂移量 $d_0 = 0$ 。因此，控制器的輸入建議將修正如公式(E.24)所示。

$$u_{j+1} = \frac{T - \hat{a}_{j+1} - \hat{d}_{j+1}}{b} \quad (\text{E.24})$$

E.3.1.3 卡爾曼濾波器

本章節介紹卡爾曼濾波器 (Kalman filter)，由於意將輸出與目標值之間的誤差雜訊最小化，相當於將訊號中的「高頻誤差項過濾掉」，故稱之為「濾波器」。在完成一個批次之後，我們可獲得系統狀態的資訊。然而，但若要事先估計系統狀態以提供下一批次的輸入參數就有相當困難與挑戰，而卡爾曼濾波器提供了有效的方法，利用遞迴的方式進行準確的估計。卡爾曼濾波器由狀態方程式（描述系統中的變化）、空間狀態干擾模型、及輸出方程式（描述系統變化後的輸出）組成，如公式(E.25)、公式(E.26)與公式(E.27)所示，

$$\text{狀態方程式：} x_{j+1} = Cx_j + Du_j + w_j \quad (\text{E.25})$$

$$\text{輸出干擾模型：} a_{j+1} = a_j + \varepsilon_j \quad (\text{E.26})$$

$$\text{輸出方程式：} y_j = bx_j + v_j \quad (\text{E.27})$$

其中在狀態方程式中， x 為系統狀態向量、 u 為系統輸入、 C 為批次間的狀態轉移矩陣、 D 為控制輸入矩陣、 w 為系統干擾向量（高斯噪音）。 a 為輸出截距干擾模型。在輸出方程式中， y 為系統輸出向量、 b 為觀測矩陣、 v 為量測干擾向量（高斯噪音）。此外，由於系統干擾 w 、量測干擾 v 、與系統狀態 x 三者互相獨立，且是多維度，因此其性質如公式(E.28)所示，

$$E \left\{ \begin{bmatrix} w_j \\ v_j \end{bmatrix} \begin{bmatrix} w_j & v_j \end{bmatrix} \right\} = \begin{bmatrix} Q_w & 0 \\ 0 & R_v \end{bmatrix} \quad (\text{E.28})$$



其中 Q_w 為 w 的共變異數，可利用歷史參數資料與輸入值的基線設定求得；而 R_v 為 v 的共變異數，可參考機台操作手冊提供的量測誤差值得出。初始值設定 $E[x_0 w_j^T] = 0$ 且 $E[x_0 v_j^T] = 0, j = 1, \dots, m$ 。為了運算方便，我們將狀態方程式與輸出干擾模型結合成一個系統，如公式(E.29)與公式(E.30)所示，

$$E\mathbf{X}_{j+1} = \mathbf{C}\mathbf{X}_j + \mathbf{D}\mathbf{U}_j + \mathbf{w}_j \quad (\text{E.29})$$

$$\mathbf{Y}_j = \mathbf{b}\mathbf{X}_j + \mathbf{v}_j \quad (\text{E.30})$$

其中 $\mathbf{X}_j = \begin{bmatrix} x_j \\ a_j \end{bmatrix}$ 、 $\mathbf{C} = \begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix}$ 、 $\mathbf{D} = \begin{bmatrix} D \\ 0 \end{bmatrix}$ 、 $\mathbf{U}_j = u_j$ 、 $\mathbf{w}_j = \begin{bmatrix} w_j \\ \varepsilon_j \end{bmatrix}$ 、 $\mathbf{Y}_j = y_j$ 、與 $\mathbf{b} = [b \quad I]$ 。因此卡爾曼濾波器的形式如公式(E.31)所示，

$$\begin{bmatrix} \hat{x}_{j+1} \\ \hat{a}_{j+1} \end{bmatrix} = \begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} \hat{x}_j \\ \hat{a}_j \end{bmatrix} + \begin{bmatrix} D \\ 0 \end{bmatrix} u_j + L_{j+1}(y_j - b\hat{x}_j - \hat{a}_j) \quad (\text{E.31})$$

其中 L 稱為卡爾曼增益，提供願意相信是此次製程狀態轉移誤差較小或是量測器誤差較小的選擇，如公式(E.32)所示，令 $\Sigma_{j+1|j}$ 為利用已給定的初始值從批次 j 預測下一個批次 $j+1$ 的狀態變異矩陣，如公式(E.33)所示，

$$L_{j+1} = C\Sigma_{j+1|j}b^T(b\Sigma_{j+1|j}b^T + R_v)^{-1} \quad (\text{E.32})$$

$$\Sigma_{j+1|j} = C\Sigma_j C^T + Q_w - C\Sigma_j b^T(b\Sigma_j b^T + R_v)^{-1} b\Sigma_j C^T \quad (\text{E.33})$$

其中初始值設定為 $\hat{X}_0 = E[X_0] = \bar{X}_0$ 且 $\Sigma_0 = \text{Var}(X_0)$ 。由此反覆迭代得出系統狀態 $\mathbf{X}$ ，並可進而與其他控制器結合進行拓展，例如用於非線性系統（例如導航系統、批次進料化學反應爐等）的擴展卡爾曼濾波器（extended Kalman filter, EKF）（Hoshiya and Saito, 1984; Wilson et al., 1998）；抑或因製程干擾具有時間序列特色（例如金屬沉澱物隨時間越積越厚），因而與 d-EWMA 進行結合（Chen et al., 2007）。

E.3.1.4 最佳自適應品質控制器（OAQC）

最佳自適應品質控制器（OAQC）是一種可進行動態調整模型，結合「遞迴最小平方法」（recursive least squares, RLS）及非線性優化器（nonlinear optimizer）的方法，其具有以下特色：

- 在確認控制參數的過程中，其可以進行 on-line 的多變量 RLS 模型

配適 (model-fitting)。

- 基於最佳化準則的實驗設計。
- 可結合 SPC 管制圖的上下界限以減少雜訊的影響，同時監測製程是否失控。

以下先說明 RLS 的運作方式。OAQC 採用了「二階多輸入多輸出漢默斯轉移模型」(second-order MIMO Hammerstein transfer model) 進行製程的模型配適，如公式(E.34)所示，

$$y_{j+1} = y_0 + Nz_j + Mt + (I_p - CB)\varepsilon_{j+1} \quad (\text{E.34})$$

其中 y_0 向量為有 p 個參數截距或偏移、 z_j 為控制參數向量包含其二次展開項 $(u_j, u_j^2, u_j u_{j'})$ 、 t 為在 p 個分量裡的批次指標 j 、 B 為倒偏移運算元、 ε_j 為製程干擾、 N, M, C 皆為迴歸係數。而藉由「向前一步最小均方差預測」(one-step ahead predictor minimum mean square error prediction) 推得如公式(E.35)所示，

$$\hat{y}_{j+1} = Ly_{j+1} + Nz_{j+1} + M(t+1) \quad (\text{E.35})$$

同時我們可進一步改寫為 $\hat{y}_{j+1} = \hat{\theta}'_{j+1}\phi_{j+1}$ ，其中需要估計的未知參數（迴歸係數） $\hat{\theta}' = [\hat{L}, \hat{N}, \hat{M}]$ ，且 $\phi_{j+1} = [y_j, z_j, t+1]$ ，而我們希望在滿足公式的條件下，找出 $\hat{\theta}'$ 的最佳估計量。我們透過下列三個步驟來協助估計：

- 步驟 1：計算過程增益向量，其中 P_j 為 ϕ_j 的共變異矩陣之逆矩陣，且 λ 為折損因子 (discounting factor) 改善暫態 (transient) 績效，若 $\lambda = 1$ 則呈現標準的 RLS 演算法（調整不同的 λ 可以得到不同的狀態，在此我們僅討論 λ 為 1 時的標準 RLS 方法）。

$$K_{j+1} = \frac{P_j \phi_{j+1}}{\lambda + \phi_{j+1}^T P_j \phi_{j+1}} \quad (\text{E.36})$$

- 步驟 2：計算向前一步預測誤差 (the one step ahead forecast error)。

$$e_{j+1} = y_{j+1} - \hat{\theta}'_{j+1}\phi_{j+1} \quad (\text{E.37})$$

$$\hat{\theta}_{j+1} = \hat{\theta}_j + K_{j+1}e_{j+1} \quad (\text{E.38})$$

- $$P_{j+1} = \left[P_j - \frac{P_j \phi_{j+1} \phi_{j+1}^T P_j}{\lambda + \phi_{j+1}^T P_j \phi_{j+1}} \right] / \lambda \quad (\text{E.39})$$

- 接下來，利用非線性優化器 J 使得輸出 $\hat{y}_j$ 盡可能地靠近目標 T ，因此各別賦予輸入及輸出有其限制範圍，如下公式，

$$J = \min \{ (\hat{y}_j - T)' W (\hat{y}_j - T) + (u_j - u_{j-1})' \Gamma (u_j - u_{j-1}) \} \quad (\text{E.40})$$

$$u_j = (\hat{N}'W\hat{N} + \Gamma)^{-1}\hat{N}W[T - \hat{M}(t+1) - \hat{L}y_{j||i-1}] + \Gamma u_{j-1} \quad (\text{E.41})$$

此即為 OAQC 的方法，其方塊圖如圖 E.19 所示。倘若輸出落在 SPC 管制圖的目標界限範圍內（例如 LCL 或 UCL），則不調整模型進行下一批次的作動；反之，若脫離了目標管制，則此批次的輸入與輸出值將會進入 OAQC 模型進行調整。

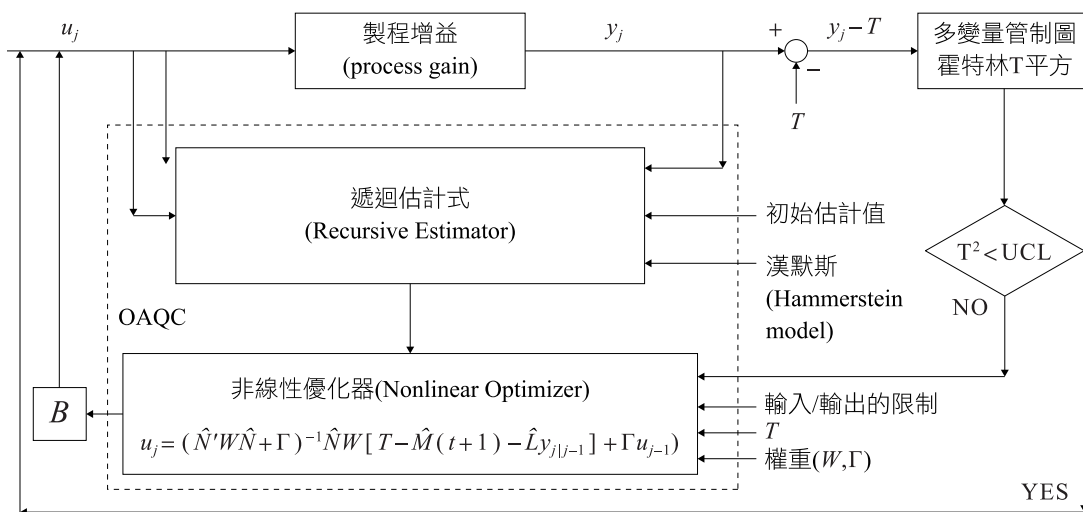


圖 E.19 最佳自適應品質控制器

批次控制多年以來已發展出多種控制器，其以統計模型為基礎並以數據科學進行驅動的概念，讓我們能藉由量測一批成品的輸出數據（如產品尺寸、位置、壓力及溫度等）進行下一批次輸入參數的修正，同時整合傳統物理公式的自動控制想法，甚至能更加即時、精確地修改輸入參數。實務上，批次控制發展至今已被許多製造業者導入進行製程優化，也獲得了不錯的成果。然而，目前大多使用在單機台的輸入輸出上，應用場合受限，若能發展多機台或是跨上下游雙機台的批次控制，如圖 E.20，則能為產業帶來更全面的應用與經濟效應。事實上，如圖虛線方框處，由於部分中間機台無法在輸出端裝置感測器，因此難以獲得輸出資訊，再加上機台 A 的輸出 y_{Aj} 是受到干擾後的結果，因此機台 B 從機台 A 所獲得的輸入並不完美，在此架構下的批次控制必須面對雙重干擾。同時，如何僅從機台 B 所獲得的輸出 y_{Bj} 回推，並分別為兩台不同機台進行輸入參數（ $u_{A(j+1)}$, $u_{B(j+1)}$ ）的調整，其中的分配比例以及誤差溯源是困難且有挑戰性的地方。

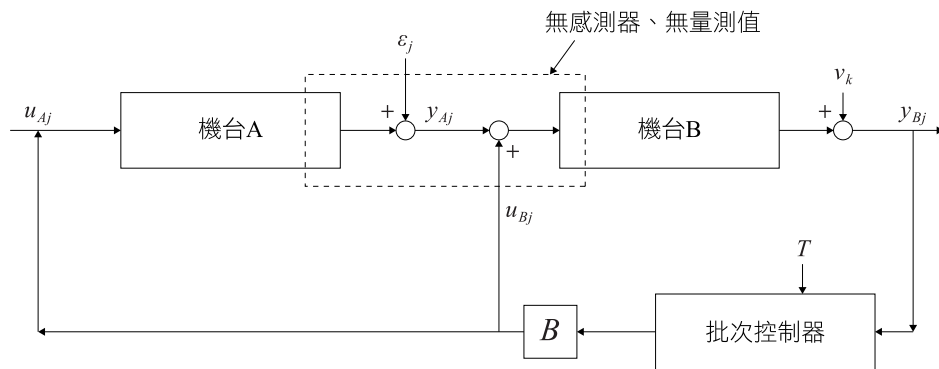
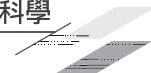


圖 E.20 跨上下游雙機台之控制器示意圖

E.4 結語

本章節介紹了統計製程管制（SPC）、故障偵測與分類（FDC）以及先進製程控制（APC）中的 R2R 控制器，這些方法論在製造（或服務）現場的重要性與日俱增，尤其是在製造業陸續開發了許多相關的資訊系統，並與製造執行系統（MES）或其他工程自動化系統連結，以達到提早



警報 (pre-alarm) 與精度控制品質提升。在同步開發眾多資訊系統的同時，也應同時留意以下幾點延伸議題，包含(1)IT 基礎建設；(2)大數據分析引擎；(3)系統整合；(4)資料庫設計；(5)資訊安全；(6)天災突發事件，並適當地建立對策與管理機制，才能有穩定的成長與改善。

參考文獻

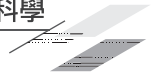
- [1] Castillo, E. D., and Yeh, J. Y. (1998). An adaptive run-to-run optimizing controller for linear and nonlinear semiconductor processes. *IEEE Transactions on Semiconductor Manufacturing*, 11(2), 285-295.
- [2] Chen, J.-H., Kuo, T.-W., and Lee, A.-C. (2007). Run-by-run process control of metal sputter deposition: combining time series and extended Kalman filter. *IEEE Transactions on Semiconductor Manufacturing*, 20(3), 278-285.
- [3] Fan, S.-K., Hsu, C.-Y., Tsai, D.-M., He, F., and Cheng, C.-C. (2020). Data-driven approach for fault detection and diagnostic in semiconductor manufacturing. *IEEE Transactions on Automation Science and Engineering*, 17(4), 1925-1936.
- [4] Liu, K., Chen, Y., Zhang, T., Tian, S., and Zhang, X. (2018). A survey of run-to-run control for batch processes. *ISA Transactions*, 83, 107-125.
- [5] Montgomery, D. C. (2019). *Introduction to Statistical Quality Control*. 8th edition, John Wiley & Sons.
- [6] Sun, R., and Tsung, F. (2003). A kernel-distance-based multivariate control chart using support vector methods. *International Journal of Production Research*, 41(13), 2975-2989.
- [7] Venkatasubramanian, V., Rengaswamy, R., Yin, K., and Kavuri, S. N. (2003a). A review of process fault detection and diagnosis: Part I: Quantitative model-based methods,” *Computers & Chemical Engineering*, 27(3), 293-311.
- [8] Venkatasubramanian, V., Rengaswamy, R., and Kavuri, S. (2003b). A review of process fault detection and diagnosis: Part II: Qualitative models and search strategies. *Computers & Chemical Engineering*, 27(3), 313-326.
- [9] Venkatasubramanian, V., Rengaswamy, R., and Kavuri, S. (2003c). A review of process fault detection and diagnosis: Part III: Process history based methods. *Computers & Chemical Engineering*, 27(3), 327-346.
- [10] Hoshiya, M., and Saito, E. (1984). Structural identification by extended Kalman filter. *Journal of Engineering Mechanics*, 110(12), 1757-1770.

- [11] Wilson, D. I., Agarwal, M., and Rippin, D. W. T. (1998). Experiences implementing the extended Kalman filter on an industrial batch reactor. *Computers & Chemical Engineering*, 22(11), 1653-1672.

問題與討論



1. (a)試說明品質的定義；(b)承接(a)答案，就您的觀點從這定義中還有什麼需要注意或延伸的補充？(c)根據您的經歷，製程變異包含哪些問題？其潛在的根本原因可能為何？
2. 試說明(a)品管七大手法；(b)承接(a)，各方法使用的時機與優劣為何？
3. 試說明(a)管制圖兩階段的建構與應用；(b)如何評估品管圖優劣？常用方法為何？(c)何謂製程能力？其常用指標為何？(d)品管圖的上下管制界限與製程能力之間有何關聯？
4. 在 UCI Machine Learning Repository 開放數據中的包含時間序列數據（Synthetic control chart time series data set, <https://archive.ics.uci.edu/ml/datasets/synthetic+control+chart+time+series>），一共包含了 600 個觀測值並每個觀測值有 60 個樣本點，其呈現六種時間序列樣型（pattern），包含正常（normal）、週期（cyclic）、上升趨勢（increasing trend）、下降趨勢（decreasing trend）、向上漂移（upward shift）、以及向下漂移（downward shift）。試著參考網路資源學習並撰寫程式，使用此數據回答下列問題：
 - (a) 試針對「正常」的樣型（觀測值 1-100），任取一筆觀測值，建構 $\bar{X}$ 管制圖，包含計算 UCL、CL、與 LCL。
 - (b) 承上題(a)，若今任取一筆筆「向上漂移」的數據（觀測值 401-500），並放入(a)所建構的管制圖中，請問是否能有效地偵測管制失控？為什麼？
 - (c) 若假設製程規格為 30.00 ± 5.00 ，計算製程能力 PCR 與 PCR_k ，該製程表現如何？是否滿足 $PCR \geq 1.33$ ？
 - (d) 若現今欲達到「六標準差製程」，製程變異數 σ^2 應要降低至多少，才能達到 $PCR_k = 2.0$ ？
 - (e) 若現今製程平均值突然向上漂移至 30.5，試問下一個樣本能偵測到這個漂移的機率為何？漂移後的平均連串長度（ARL）是多少？
5. 在 UCI Machine Learning Repository 開放數據中的包含時間序列數據（Synthetic



control chart time series data set, <https://archive.ics.uci.edu/ml/datasets/synthetic+control+chart+time+series>), 一共包含了 600 個觀測值並每個觀測值有 60 個樣本點, 其呈現六種時間序列樣型 (pattern), 包含正常 (normal)、週期 (cyclic)、上升趨勢 (increasing trend)、下降趨勢 (decreasing trend)、向上漂移 (upward shift)、以及向下漂移 (downward shift)。試著參考網路資源學習並撰寫程式, 使用此數據回答下列問題:

- (a) 試針對「正常」的樣型 (觀測值 1-100), 任取一筆觀測值, 建構累積和 CUSUM 管制圖。
 - (b) 承上題(a), 若將六種樣型各任取一筆, 放入(a)所建構的管制圖中, 請問是否能有效地偵測管制? 試比較其結果。
 - (c) 試針對「正常」的樣型 (觀測值 1-100), 任取一筆觀測值, 建構指數加權移動平均 EWMA 管制圖。
 - (d) 承上題(c), 若調升指數加權的權重 λ , 管制界限 UCL 與 LCL 如何變化? 若減少權重 λ , 其管制界限如何變化?
 - (e) 若從「正常」與「週期」的樣型, 各任取一筆觀測值, 當作製程參數 x_1 與 x_2 , 試建構霍特林 T 平方多變量管制圖。
 - (f) 承上題(e), 若從「上升趨勢」與「向下漂移」的樣型, 各任取一筆觀測值, 當作製程參數 x_1 與 x_2 所發生的變化, 放入(e)所建構的管制圖中, 請問是否能有效地偵測管制? 為什麼?
6. (a)試簡要說明錯誤偵測與分類 (FDC) 的三大方法; (b)就您觀點, 這三大方法的使用時機與優劣為何?
7. (a)FDC 的分析流程相當於是轉換函數的映射, 試解釋其轉換的空間依序為何? (b)承接(a)答案, 就您觀點, 哪一空間的轉換最顯著影響預測績效? 為什麼? 如何改善? (c)FDC 診斷系統評估十項特性為何? (d)承接(c)答案, 就您觀點來說哪一項最重要? 為什麼? 如何改善?
8. 在 Kaggle 開放數據中的包含三相電力系統的數據 (Electrical fault detection and classification, <https://www.kaggle.com/esathyaprakash/electrical-fault-detection-and-classification>), 在資料集 classData.csv 中, 一共包含了 7,861 個觀測值, 6 個有關電壓與電流的欄位, 以及複合標籤 (含 G、C、B、A 四個欄位) 呈現正常與五種錯誤類型。試著參考網路資源學習並撰寫程式, 使用此數據回答下列問題:
- (a) 試建構 FDC 數據科學模型來進行分類與診斷, 請繪製數據科學架構圖或建模流程圖。

- (b) 試說明資料預處理，是否有類別不平衡？如何處理？
 - (c) 試說明是否需要特徵工程？為什麼？如需要，該如何進行？
 - (d) 試挑選至少兩種分類器，其各別交叉驗證後的分類效果如何？
 - (e) 針對五種錯誤類型，分別分析其重要影響因子排序為何？
9. (a)試說明先進製程控制（APC）的目的為何？(b)APC 系統如何與 SPC 系統跟 FDC 系統有所關聯？這三系統彼此互補或各自獨立？(c)就您觀點，在您日常處理的業務範疇中，試就一項業務繪製流程圖，試說明流程中哪裡可以透過 SPC、FDC 或 APC 的觀點進行改善？
- 10.(a)常見的批次控制方法有哪些（不限於課本內容）？列舉三種並簡要說明其運作原理。(b)承接(a)答案，試說明此三種方法的使用時機與其優劣。