

行政院所屬各機關因公出國人員出國報告書
(出國類別:實習)

參加東南亞國家中央銀行研訓中心

「計量模型建構與預測」

基礎課程出國報告

服務機關：中央銀行

姓名職稱：曾虹瑋/四等專員

派赴國家：馬來西亞吉隆坡

出國期間：106 年 4 月 15 日至 4 月 22 日

報告日期：106 年 7 月

目錄

壹、 前言	2
貳、 建立計量模型之功用與挑戰	3
參、 時間序列模型及其估計方法	7
肆、 建立預測模型之準則	14
伍、 課程重要範例	18
陸、 課程心得及建議事項	21

壹、前言

職奉准於民國 106 年 4 月 15 日至 4 月 22 日參加東南亞國家中央銀行研訓中心(SEACEN Centre)於馬來西亞吉隆坡舉辦為期 6 日之「計量模型建構與預測」基礎課程(SEACEN Foundational Course on Economic Modelling and Forecasting)。本次學員來自 11 國共 23 位¹。3 位課程講師均來自 SEACEN Centre 總體經濟及貨幣政策管理部門；其中，主講者 Dr. Ole Rummel 講述內容更是融合其過去任職英格蘭銀行參與貨幣政策分析的實務經驗，為課程帶來更高的實用性。

本次課程主要目的為訓練學員如何建構實證計量模型，培養量化分析總體經濟及貨幣政策的能力。除講述計量基礎觀念及方法外，全程更搭配計量軟體 Eviews 進行案例演練。課程內容編排上，先介紹計量模型之功用及實務挑戰，再介紹時間序列模型（time series model）、資料特性處理、一般最小平方法（ordinary least squares, OLS）、資料平穩性、季節性處理、濾波法及向量自我迴歸模型等常用計量方法，並以此應用於計算產出缺口、估計新凱因斯模型、預測通貨膨脹率等。

本報告第一節為前言，簡介課程內容；第二節為建構實證計量模型之功用與運用限制，以瞭解量化分析雖能提供客觀數據，惟仍牽涉許多主觀判斷；第三節介紹時間序列模型、資料處理、估計及模型選擇方法；建立模型除瞭解現況外，亦提供預測價值，第四節說明預測須注意之通則；第五節挑選課程中數個與央行應用有關之重要案例進行說明；第六節則為心得與建議。

¹ 馬來西亞(5)、印尼(4)、泰國(3)、柬埔寨(2)、斯里蘭卡(2)、菲律賓(2)、巴布紐新幾內亞(1)、印度(1)、斐濟(1)、孟加拉(1)、台灣(1)。

貳、建立計量模型之功用與挑戰

建立計量模型進行實證，有助瞭解影響經濟金融體系運作的重要因素，藉由模型預測及政策模擬，可提供決策者較科學及客觀的分析結果，有利提升決策品質。然而，實務上面臨相當多的挑戰。

一、多模型之必然性

基於經濟體運作相當複雜，其真實樣貌及可能面臨的衝擊，均面臨相當多的不確定性，以不同的模型估計及預測，即可能產生不同的結果。此外，單一模型通常僅能解決特定的問題，無法完整回應決策者的需要。因此，先進國家央行常使用多種模型（甚或加權各項模型結果），以一或兩個模型作為核心模型，並建立衛星模型，提供核心模型所需之資料。除可考量更多面向外，亦以降低使用單一模型可能衍生判斷錯誤的風險。

常見的模型類型如下：

- （一）活頁簿（Spreadsheets）：如 Excel 的活頁簿
- （二）從組模型（Suite models）：如簡單的單變量及多變量模型
- （三）小型結構型模型（small-scale structural models）
- （四）大型動態隨機一般均衡（Dynamic Stochastic General Equilibrium, DSGE）模型

二、原始資料品質影響結果

良好的模型需配合高品質的原始資料，始能產生較佳的實證結果，因此，須積極解決資料可取得性、資料品質、資料衡量等問題（表1）。若資料品質差，要儘量採用符合理論基礎的模型，且模型結構要儘量簡單，如勞動市場資料品質不佳，則考慮不要納入勞動市場，抑或找尋其他可以取代的變數。

表1 資料問題類型

類型	問題或解決方法
資料可取得性 (如：沒有足夠資料，樣本不夠)	解決方法 ● 利用央行的可信度蒐集資料 ● 插補資料將低頻資料轉為高頻資料 ● 自行編製資料 ● 尋找連動性高的替代指標 ● 使用小樣本方法估計
資料品質	解決方法 ● 自行編製資料 ● 資料品質控制，如使用IMF、OECD等國際組織的資料。
資料衡量	問題類型 ● 不同的定義 ● 統計方法的改變 ● 重新設定基期 ● 資料和理論無法連結

資料來源：課程講述內容

三、符合經濟理論及通過實證測試間之權衡

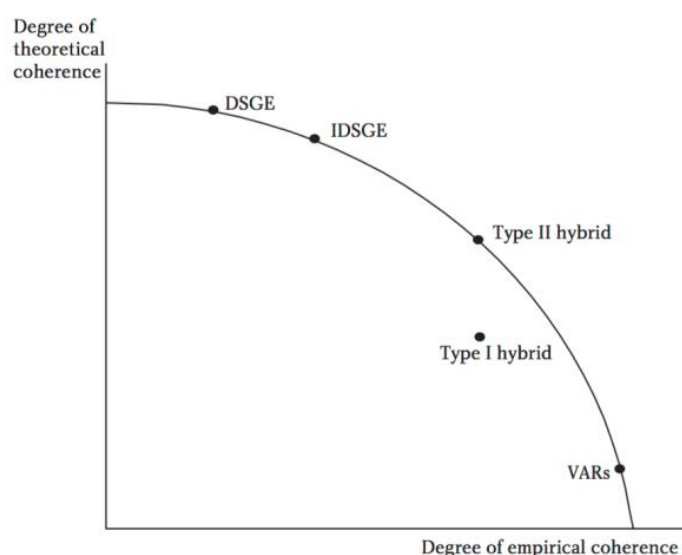
良好的模型須具備經濟理論基礎，要能清楚表現家計部門、廠商及政府行為，並進行福利分析（welfare analysis）。此外，模型要能通過實證測試，包括能與資料配適，並在政策架構沒有太大變化的狀況下，維持模型穩定性。

然而，在實務上，較吻合經濟理論之模型，實證結果未必較佳，實證結果較佳的模型，有時缺乏經濟理論，致決策者難以向大眾說明經濟現象。

央行採用的模型主要可分為兩種類型。第一類，著重在結構性基礎（structural foundation），模型內包含代表一般均衡的方程式。第二類，則著重在資料的配適程度，此主要是時間序列模型。Pagan（2003）根據理論基礎及資料配適程度兩大面向，將央行常用的模型排序（圖1），其中，以 DSGE 模型最具理論基礎，惟實證結果最差，而 VAR 模型資料配適程度最高，惟缺乏理論基礎。

前述兩類模型均為央行採用，而貨幣政策架構是影響模型選擇的重要因素之一。若央行採通膨目標，因對群眾溝通非常重要，故須仰賴 DSGE 或小型結構模型，以此做為闡述政策的基礎。若採管理浮動匯率制度或貨幣數量目標，則因較不需對群眾溝通，可能以 Excel 活頁簿計算，就足以瞭解通膨發生的原因。此外，在選擇模型時，也須考量金融市場發展程度，如一國金融及價格管制相當嚴格，就不宜採取 DSGE 模型（其前提假設為完全競爭市場），且做決策時，需加入更多主觀判斷。

圖1 央行常用之各類模型符合理論及實證之程度



資料來源：Pagan（2003）

四、過去資料配適度佳，未必有較佳的預測能力

過去資料與模型配適度較高，正常情況下，可推論其具較佳的預測能力，惟若經濟體發生巨大衝擊，產生結構性變化，模型可能不再適用，須重新進行估計，以修正模型。

五、模型無法取代主觀判斷

決策者為避免決策錯誤，擬定貨幣政策時，不宜僅靠模型擬訂決策，尚需具備主觀判斷能力，將更多未能涵括在模型中卻有價值的資訊列入考量，以彌補模型未能考慮所有變數的缺點。

六、同時肩負總體經濟及金融穩定職責增加模型複雜度

2008 年美國次貸危機後，央行權責大幅增加，不僅須達成總體經濟穩定，更需兼顧金融穩定，政策複雜度大幅提升，致建構模型的複雜度及困難度也隨之增加。

七、其餘須考量之重要因素

(一) 資料長度影響估計結果

用以估計的樣本資料涵蓋期間長度會影響估計結果，如樣本資料過少，可能產生偏誤，有高估或低估之風險。

(二) 貨幣政策傳導機制及傳導速度

對於小型開放經濟體而言，匯率管道非常重要，其影響貨幣政策傳導速度，例如在浮動匯率制度下，匯率與利率的關係較為緊密，傳導速度較快。

(三) 足夠的人力

相較其他模型，DSGE 模型非常複雜，不僅軟體成本較高，也需要足夠且研究經驗豐富的高階研究人員從事研究。

參、時間序列模型及其估計方法

一、經濟金融分析常採時間序列模型

時間序列模型是指變數由自身及（或）其他變數之落後項決定，此隱含歷史資料將影響未來。由於經濟金融現象常有持續性，如本期發生通膨，下期發生通膨的可能性高，故常採時間序列模型分析。

二、時間序列資料來自未知的資料產生過程

大多數的經濟金融資料得依發生時點之先後排序，可表示為 $\{y_t\}$, $t = 1, 2, 3, \dots, T$, $\{y_t\}$ 即為一時間序列資料。

y_t 為隨機變數 (random variable)，該隨機變數來自於不可觀察的母體 (unobserved population)，而每一個在 t 時點的觀察值，只是其中一個隨機資料產生過程 (data-generating process, DGP) 的結果。因此，在 t 時點觀察到的值，可視為所有可能 DGP 結果的平均值。

三、時間序列資料之統計量

假設有一時間序列為 $\{y_t\}$ ，則其統計量如下：

- (一) 期望值為 $E[y_t] = \mu_t$
- (二) 共變數為 $Cov(y_t, y_{t-j}) = E[(y_t - \mu_t)(y_{t-j} - \mu_{t-j})]$, j 表時間間隔
- (三) 若 $j = 0$ ，則代表變異數 $Var(y_t) = E[(y_t - \mu_t)^2]$
- (四) 若 $\mu_t = 0$ ，則變異數更可簡化為 $Var(y_t) = E[y_t^2]$

四、時間序列若符合弱穩定使分析更為便利且準確

若一時間序列資料符合下列性質，稱為弱穩定 (weak stationary) 或共變異數穩定 (covariance stationary)：

- (一) 期望值固定，代表時間序列不會出現時間趨勢，即 $E[y_t] = \mu$ 。
- (二) 共變數只由時間間隔 (j) 決定，意即 $Cov(y_t, y_{t-j}) = \gamma(j)$ 。

(三) 符合第 2 條件代表所有變異數($j = 0$)是固定的。

在弱穩定情況下進行時間序列分析，可使用較簡單的模型或估計方法。反之，若數列不穩定且無共整合，誤用模型容易出現假性迴歸 (spurious regression)，即便自變數及應變數間無因果關係或經濟意涵，統計結果卻顯現關係顯著，此將使研究結論發生錯誤；又或使用錯誤的估計方式，使參數估計結果產生偏誤。

因此，採用時間序列模型，應先進行單根檢定 (unit root test)，確定其為定態或非定態。一旦完成模型估計，由於忽略重要自變數 (omitted variable)、模型設定錯誤、衡量誤差 (measurement error) 等錯誤會藏在殘差中，須再次檢測殘差是否符合白噪音²特質，以確認模型正確性。

五、單變量單因子時間序列模型

(一) AR 模型

變數由自身落後項決定，代表資料本身具有延續性。模型表示為 AR(p)，p 代表變數落後期數，如以 AR(1)為例：

$$y_t = \rho y_{t-1} + \varepsilon_t$$

(二) MA 模型

變數由殘差落後項決定，殘差為預測誤差，將前期殘差列入考量，隱含修正誤差特質。模型表示為 MA(q)，q 為殘差落後期數，如以 MA(1)為例：

$$y_t = \theta \varepsilon_{t-1} + \varepsilon_t$$

² 白噪音定義為時間序列變數符合 (條件) 期望值為 0、(條件) 變異數為固定常數、(條件) 共變異數為 0。

(三) ARMA 模型

又稱 Box-Jenkins 模型，變數由自身落後項及殘差落後項決定。模型表示為 ARMA(p, q)，如以 ARMA(1,1)為例：

$$y_t = \rho y_{t-1} + \theta \varepsilon_{t-1} + \varepsilon_t$$

模型不僅包含資料延續性，也涵蓋誤差修正特質，實務上最具短期預測能力。此外，Box and Jenkins (1976)更是強調 ARMA 有助精簡參數，相較單純的 AR 或 MA 模型，ARMA 模型使用較少的落後期數即可達到相同模型表現。

(四) ARIMA 模型

若一數列經過 d 階差分，可成為一穩定數列，再以 ARMA(p, q) 模型建模，稱為 ARIMA (p, d, q) 模型。

(五) 單變量模型落後期數之決定

有關單變量模型的類型及自變數落後期數，可以自我相關函數 (autocorrelation function, ACF) 及偏自我相關函數 (partial autocorrelation function, PACF) 判定。

1. 自我相關函數

自我相關函數 $\rho(j)$ 是在探討變數當期及其落後項間的關係，其定義如下，數值介於 ± 1 ，j 表落後期數。

$$\rho(j) = \text{Cov}(y_t, y_{t-j}) / \text{Var}(y_t) , \quad -1 \leq \rho(j) \leq 1$$

2. 偏自我相關函數

偏自我相關函數則係探討去除該變數所有落後期數小於 j 的影響後，探討當期變數與落後 j 期的變數間之關係，即為線性迴歸式中 y_{t-j} 的係數 β_j ，方程式如下：

$$y_t = c + \beta_1 y_{t-1} + \beta_2 y_{t-2} + \cdots + \beta_j y_{t-j}$$

如以 Eviews 觀察 ACF 及 PACF 圖形，AR(p)模型特性為 ACF 無限，而 PACF 在超過 p 期後，就會趨近於 0；MA(q)模型特性為 ACF 在超出 q 期後，將趨近於 0，而 PACF 無限。若缺乏前述情形，則可判定為 ARMA 模型。

六、多變量時間序列模型

(一) 向量自我迴歸 (Vector Autoregression, VAR)

1980 年代以前，經濟實證模型以大型凱因斯模型為主，該模型建構在經濟理論上，以方程式描寫變數間的理論關係，變數包括內生變數（系統內決定）及外生變數（來自系統外），其中外生變數不受內生變數的當期或落後期影響。然而，因經濟體錯綜複雜，內生或外生變數之設定並不容易，且可能出現設定錯誤。

VAR 模型之優點係將所有變數視為內生變數，不主觀假設變數間的關係，以避免任意設定外生變數及內生變數衍生的問題。VAR 模型係以變數自身及其他變數的落後項所組成，以兩個變數及落後項為 1，說明 VAR(1)模型：

$$x_t = b_{10} + b_{11}x_{t-1} + b_{12}y_{t-1} + \varepsilon_{xt}$$

$$y_t = b_{20} + b_{21}x_{t-1} + b_{22}y_{t-1} + \varepsilon_{yt}$$

VAR 落後期數得由 LR 統計量決定，在決定落後期數後，若採一般化最小平方法 (OLS) 估計，須檢測殘差是否穩定，否則將採似不相關迴歸模型 (Seemingly unrelated regression model, SUR) 等其他方法估計。此外，若進行 Granger-Causality 檢定，得檢定變數的落後項是否對另一變數有所影響。

實證結果顯示，在短期預測能力表現上，VAR 相較傳統大型結

構模型有較佳的表現。

(二) 結構式向量自我迴歸 (Structural Vector Autoregression, SVAR)

VAR 模型之短期預測能力佳且被廣泛地用於經濟金融實證，惟該模型無法表現變數同期間之交互影響，與經濟金融實況常不符。

SVAR 除加入同期變數間的關係外，亦可依循經濟理論，對變數間的關係給予條件限制，除符合實際狀況外，亦有助利用參數估計結果，分析經濟狀況及變數間的互動關係。

七、資料特性及處理方式

由於資料會影響模型估計結果，為避免估計產生偏誤，須根據資料特性，做進一步的處理，常見處理方式如表 2。

表2 資料特性及處理方式

特性	處理方法
趨勢	● 得取對數 (log) 平滑資料，再取一階差分分析。 ● 加入時間趨勢作為自變數，以捕捉趨勢因素。
季節性	● 基於季節性因素多為暫時，難提供太多有用的資訊，故須將歸於季節性的波動移除。
極端值	● 加入虛擬變數，以排除該極端值對模型的影響。 ● 須留意可能是經濟體產生結構性變化。
異質性	● 係指資料波動不齊一的現象，此時資料將呈群聚 (cluster) 現象，宜用 GARCH 模型處理該現象。

資料來源：Ole Rummel (2017)，"Key Features of Data," SEACEN 上課講義

八、一般化最小平方法估計

在估計迴歸式的係數時，一般最小平方法（Ordinary least squares, OLS）係常見且基本的方法，其概念是尋求能極小化所有資料點殘差平方和的係數值。

根據高斯－馬爾可夫定理(Gauss-Markov Theorem)，如迴歸式滿足殘差期望值為零、自變數與殘差無相關、殘差具均值變異、殘差無序列相關、自變數間無共線性及殘差為常態分配之條件，則以 OLS 估計，係數將具有不偏及波動度最小的性質（BLUE）。因此，如不符合條件，即可能會產生錯誤，須進行修正，修正方法如表 3。

表 3 違反基本假設的檢查方式、估計錯誤的結果及可修正的方法

符合 OLS 估計之條件	檢查方式	錯誤結果	可修正的方法
殘差期望值為零	觀察殘差圖形	OLS 估計偏誤	● 定義新的截距
自變數與殘差無相關	兩者相關係數高	OLS 估計偏誤	● 拿掉自變數
殘差具均值變異	White 檢定	OLS 估計雖不偏，但不效率	● 重新調整模型 ● 取對數 ● 改用 HAC 估計
殘差無序列相關	● Durbin-Watson 檢定 ● Brusch-GoefferyLM 檢定	OLS 估計雖不偏，但不效率	● 採 AR 模型 ● 改用 HAC 估計
自變數間無共線性	● 自變數間相關性高 ● 判定係數（ R^2 ）很高，但係數均不顯著	● 估計係數不正確且不穩定	● 拿掉自變數
殘差為常態分配	● Jacque-Bera 檢定	● 在大樣本中，因將趨近常態分配，此可忽略。	● 看是否有極端值 ● 看是否出現結構性轉變 ● 增加樣本數

資料來源：Vincent Lim Choon Seng (2017)，"Connecting the Dots - A Reference to Ordinary Least Squares," SEACEN 上課講義

九、模型選擇準則

由於真實模型未知，根據估計結果，可能產生許多預選模型，可採模型選擇準則挑出最適模型。常用的模型選擇準則包括：

1. AIC (Akaike' s information criterion)

單變量： $AIC = \ln(\sigma^2) + 2(k/T)$

多變量： $MAIC = \ln|\Sigma| + 2(k/T)$

2. SBIC (Schwarz' s Bayesian information criterion)

單變量： $SBIC = \ln(\sigma^2) + \ln T(k)$

多變量： $MSBIC = \ln|\Sigma| + \ln T(k/T)$

3. HQIC (Hannan-Quinn criterion)

單變量： $HQIC = \ln(\sigma^2) + 2\ln(\ln T)(k/T)$

多變量： $MHQIC = \ln|\Sigma| + 2\ln(\ln T)(k/T)$

σ^2 及 Σ 表殘差平方和， k 表待估參數總數， T 表樣本總數。

模型選擇準則大體上考量兩個因素，第一為模型配適程度，此由殘差平方和所決定，值越低越佳，如單變量 AIC 準則中的 $\ln(\sigma^2)$ ；另一為增加參數帶來的懲罰 (penalty)，隨參數增加，殘差平方和將持續降低，惟參數過多，將造成估測準確度降低，因此若參數較多，此項值會較高，如單變量 AIC 準則中的 $2(k/T)$ 。不論是模型選擇準則為何，其整體值越低，如單變量 AIC 之值，代表模型越佳。

十、衝擊反應函數及變異數分解

觀察模型係數值可瞭解變數間的靜態關係，而衝擊反應分析 (Impulse Response) 及變異數分解 (Variance Decomposition) 可觀察變數間長期的交互作用。

衝擊反應分析，係指在其他變數不變的情況下，某個變數出現非預期性的衝擊，對自身及其他變數長期的影響。以數學表達，即是計算每個變數對落後 1~N 期殘差的偏微分。至於 N 應取到第幾期，可視衝擊反應的收斂情況來決定。

變異數分解則是在分析 1~N 期間，每個時點上，各變數的預測變異佔總變異的比率，變數預測變異佔總變異比率越高，代表變數影響力越深。

肆、建立預測模型之準則

一、預測的價值

建立模型並進行預測的目的，係在提供政策建議、改善決策品質，或作為向大眾說明政策之佐證。

二、預測應考慮的面向

即便預測有相當多的益處，然預測並不直覺，且牽涉許多主觀判斷。在進行預測時，須考量下列面向：

(一) 預測錯誤造成之損失

由於預測正確幾無可能，且各模型預測結果可能不盡相同，可採損失函數，進行模型預測能力優劣的比較。

損失函數係在衡量實際結果不如模型預期產生的成本（損失），各種不同的誤差可能產生不同的成本，如究竟是高估或低估通膨較為嚴重。損失函數係由決策者考量其環境及偏好決定，決定函數後，再將各資料點的預測誤差帶入計算。

損失函數必須符合三大特性：

1. 完美預測時，即預測誤差為零時，損失為零。

2. 類似的預測誤差，其損失應該雷同。
3. 誤差越大，損失越大。

(二) 預測的目的

預測目的主要可分為預測時間數列（如貨幣政策時間落後效果）、預測情境發生的時點（如道瓊指數何時漲到特定值）或模擬各種情境的結果（如執行不同政策的結果，以決定採何種策略）。

(三) 預測呈現方式

預測可依其呈現的方式，分為點預測（僅呈現單一數值）、區間預測（呈現一範圍的數值，如下滑 5% 內）、密度預測（呈現未來可能出現的值及其機率）。

由於預測結果會隨各項變數及外在環境的改變，而產生變化，因此，點預測發生的機率趨近於零，提供的資訊最少，其次是區間預測，密度預測較為周全，決策者可瞭解所面對的環境不確定性有多高。

(四) 良好模型須達成的條件

透過建立模型，得以紀律且一致性的方式，建立預測架構。政策當局通常依據貨幣傳導機制以及衝擊（shock）如何影響經濟及通貨膨脹建立模型，並權衡實證結果及經濟理論採用模型。一個好的模型，必須達成下列條件：

1. 提供一致性（consistency），此有助解釋為何預測錯誤。
2. 提供長期預測。
3. 能夠解釋情境分析及風險。
4. 能夠分解過去的誤差。
5. 能夠提供央行內部決策者一致的討論點。

(五) 預測時間長度

分為事前預測 (ex ante forecast) 及事後預測 (ex post forecast) (one-step-ahead)。事前預測係使用所有的樣本資料，建立模型進行預測，惟須待其未來真實值發生時，方能檢測模型預測能力；事後預測則係將可使用的樣本資料分成兩部分，先使用一部分樣本，建立模型進行預測，再將預測值與其餘樣本比較，此法可立即比較模型預測能力，實務上較常使用。

(六) 模型預測績效的評估

建立模型後，須持續監控及評估其預測績效，可逕觀察誤差圖形以瞭解誤差分布狀況（如總是高估）或與其他模型及外部預測比較。

此外，得計算誤差均方根 (Root Mean Squared Error, RMSE) 或平均誤差絕對值 (Mean Absolute Error, MAE) 衡量預測績效，指標值越小，代表預測力越佳。雖可根據此排序出模型優劣，惟因渠等計算方式隱含所有誤差均為同等嚴重，然事實可能非如此，如將誤差帶入損失函數，再依此進行模型排序，即可能導出不同的排序結果。

三、結合各種預測結果

決策者通常會有各種不同模型做出的預測結果，可結合各種不同模型計算出的結果，用變異數及共變異數方法 (variance-covariance method) 及迴歸方法 (regression method) 等，計算各模型預測結果之權數，再加權後得到預測結果。

假設今有a及b兩種模型預測結果，模型配置權重分別為 ω 及 $(1 - \omega)$ ，其加權結果為c，h為預測期間。以兩種方法說明權重計算方式：

(一) 變異數及共變異數方法

兩種模型加權後的誤差值為：

$$e_{t+h|t}^c = \omega e_{t+h|t}^a + (1 - \omega) e_{t+h|t}^b$$

其中， e 表預測誤差

加權後的誤差變異數為：

$$\sigma_c^2 = \omega^2 \sigma_{aa}^2 + (1 - \omega)^2 \sigma_{bb}^2 + 2\omega(1 - \omega)\sigma_{ab}^2$$

其中， σ_{aa}^2 表模型 a 之誤差變異數；

σ_{bb}^2 表模型 b 之誤差變異數；

σ_{ab}^2 表模型 a 及 b 誤差之共變異數；

求使加權後誤差變異數極小化的權重 ω^* ，可得：

$$\omega^* = (\sigma_{bb}^2 - \sigma_{ab}^2) / (\sigma_{aa}^2 + \sigma_{bb}^2 - 2\sigma_{ab}^2)$$

(二) 迴歸方法

根據 Bate and Granger(1969)，加權後模型的不偏預測值為：

$$y_{t+h|t}^c = \omega y_{t+h|t}^a + (1 - \omega) y_{t+h|t}^b$$

其中， y 表模型不偏預測值。

該方法係直接將兩種模型之不偏預測值作為自變數、實際發生值作為應變數，建立迴歸式， $y_{t+h|t}^a$ 的係數 β_a 即為 ω ， $y_{t+h|t}^b$ 的係數 β_b 即為 $(1 - \omega)$ 。

$$y_{t+h} = \beta_a y_{t+h|t}^a + \beta_b y_{t+h|t}^b + \varepsilon_{t+h}$$

伍、課程重要範例

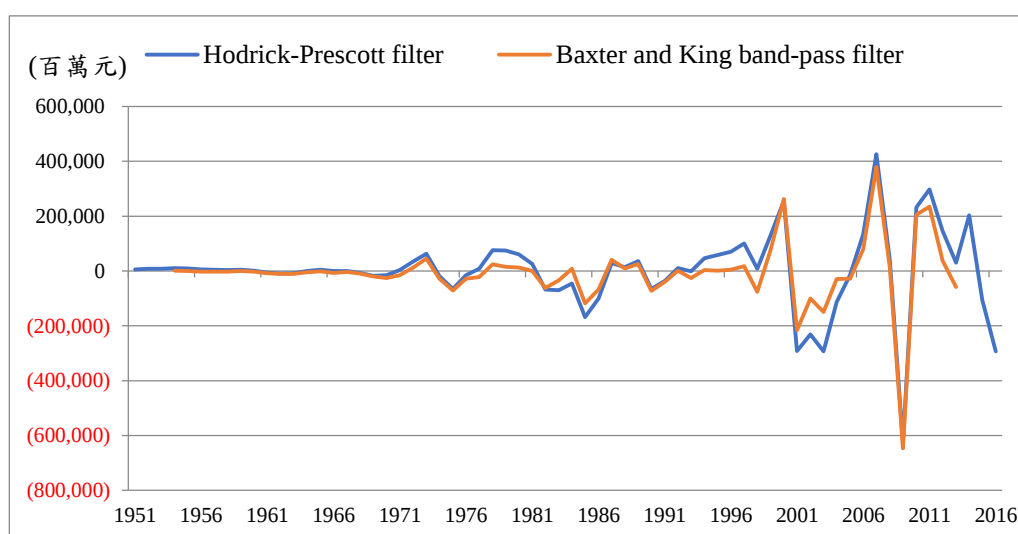
一、產出缺口之估計

產出缺口係指實際產出與潛在產出間之差距，可用以衡量通膨情勢，因此對於一國央行貨幣政策之制訂與執行具有重大涵義，央行可藉由觀察產出缺口之大小，瞭解景氣榮枯，進而預期未來通膨可能變化，採行合宜的貨幣政策因應。

產出缺口之衡量包括實際產出及潛在產出，除實際產出資料品質相當重要外，潛在產出因無法觀察，必須估計，而估計之方法很多，各有優缺點，且結果可能並非一致，文獻多採濾波法。

濾波法是將產出分為趨勢項（trend）及循環項（cyclical），而潛在產出係擷取其中的趨勢項。本課程以斯里蘭卡實質產出，演練設定時間趨勢、Hodrick-Prescott、Hamilton model 及 Baxter and King band-pass 濾波等方法。圖 2 係作者自行以 Hodrick-Prescott 及 Baxter and King band-pass 濾波法，計算臺灣實質產出缺口，可見兩者趨勢大致相同，但在部分時點會出現正負不一致的情形。

圖 2 台灣產出缺口



資料來源：作者自行計算

二、新凱因斯模型 IS 曲線估計

新凱因斯模型 (canonical new Keynesian model) 包涵三條重要的方程式，分別是衡量總合供給的菲利普 (Phillips) 曲線、衡量總合需求的 IS 曲線及衡量央行貨幣政策反應函數的泰勒法則 (Taylor rule)。其中，由於價格具僵固性，央行可藉調整名目利率，影響實質利率及產出 (即總合供給及總合需求決定之產出水準)。

課程中將 IMF (2010) “A Monetary Policy Model Without Money for India.”一文中的 IS 曲線，以 Eviews 計量分析軟體重新做實證。IMF 報告做出的 IS 曲線中，第 10 條方程式係採 OLS 估計 (圖 3)，而本課程先將圖 3 中所有參數放入方程式，同樣以 OLS 估計，再逐項刪除不顯著的變數，所得之估計結果 (圖 4)。將圖 3 及圖 4 之結果比較，兩者並不相同，顯示即便採同樣的資料及同樣的估計方法，可能因實證步驟不同，而產生不同的結果。

圖 3 IMF (2010) 對印度 IS 曲線的預估

Variable	Alternative Specifications										
	2	3	4	5	6	7	8	9	10	11	12
Dependent Variable: YGAPSA (Sample Period: 1997:2-2009:3)											
	Using GDP deflator									Using WPI	
Constant	0.32 (3.3)	0.37 (4.8)	0.37 (5.6)	0.25 (2.4)	0.08 (1.0)	0.12 (1.4)	0.11 (1.8)	0.05 (0.6)	0.16 (1.3)	0.04 (0.4)	-0.03 (0.6)
RPR{-3}	-0.10 (2.8)		-0.14 (5.2)								
RPR{-5}						-0.06 (1.8)		-0.03 (0.8)	-0.08 (1.9)		
RPR{-6}				-0.13 (3.4)			-0.06 (2.2)				
RPR{-7}					-0.05 (1.1)						
RPR{-8}		-0.16 (4.4)									
RPRWPI{-5}										-0.01 (0.2)	-0.01 (0.3)
YGAPSA{-1}	0.59 (8.0)	0.40 (7.5)	0.28 (4.2)			0.39 (8.5)	0.39 (8.5)	0.24 (3.6)	0.26 (2.6)	0.27 (2.8)	0.30 (5.9)
YGAPSA{+1}				0.55 (9.0)	0.35 (6.3)	0.44 (8.0)	0.43 (9.1)	0.33 (4.5)			0.49 (6.1)
YGAPAGRSA}	0.13 (3.0)	0.23 (7.2)	0.21 (5.7)	0.24 (6.2)	0.27 (7.7)	0.12 (2.8)	0.11 (2.8)	0.20 (3.9)	0.25 (8.9)	0.25 (4.7)	0.12 (2.7)
WEXPRGAPSA}			0.12 (8.0)		0.09 (7.5)			0.06 (4.2)	0.09 (3.3)	0.09 (5.5)	0.04 (2.3)
REER36GAPSA{-2}			-0.06 (3.0)		-0.02 (0.8)			-0.02 (1.2)	-0.08 (2.9)	-0.08 (3.7)	-0.00 (0.2)

資料來源：IMF (2010) “A Monetary Policy Model Without Money for India.”

圖 4 課程對印度 IS 曲線的預估

Dependent Variable: YGAPSA
Method: Generalized Method of Moments
Date: 04/04/11 Time: 16:08
Sample (adjusted): 1997Q3 2009Q2
Included observations: 48 after adjustments
Kernel: Bartlett, Bandwidth: Fixed (2), No prewhitening
Simultaneous weighting matrix & coefficient iteration
Instrument specification: C RPRWPI(-1 TO -2) YGAPSA(-1 TO -2)
YGAPAGRSA(-1 TO -2) WEXPRGAPSA(-1 TO -2) REER36GAPSA(-1 TO -2) DLNFC(-1 TO -2) DLBSESA(-1 TO -2) CRR(-1 TO -2)

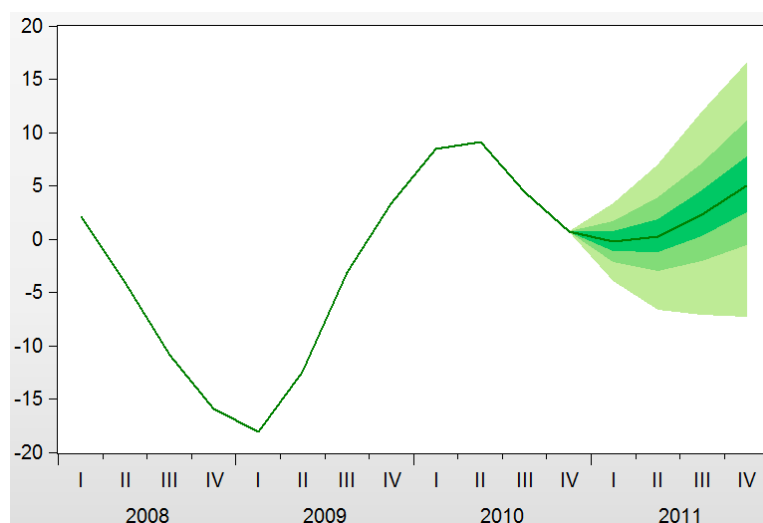
Variable	Coefficient	Std. Error	t-Statistic	Prob.
C	-0.011536	0.040159	-0.287256	0.7754
RPRWPI(-5)	-0.019096	0.017839	-1.070506	0.2907
YGAPSA(-1)	0.278050	0.050768	5.476841	0.0000
YGAPSA(1)	0.524015	0.082225	6.372935	0.0000
YGAPAGRSA	0.141671	0.041096	3.447302	0.0013
WEXPRGAPSA	0.028701	0.016352	1.755165	0.0867
REER36GAPSA(-2)	-0.004862	0.019489	-0.249492	0.8042
R-squared	0.778790	Mean dependent var	-0.010126	
Adjusted R-squared	0.746418	S.D. dependent var	1.199258	
S.E. of regression	0.603909	Sum squared resid	14.95294	
Durbin-Watson stat	2.547380	J-statistic	0.154246	

資料來源：課程實作案例

三、預測英國房價並繪製扇形圖

英格蘭央行通膨報告常以扇形圖，呈現其對經濟變數的預測。課程中以 1962Q3 至 2010Q4 之英國經濟成長率、通貨膨脹率、英格蘭銀行政策利率（Bank Rate）及房價成長率資料建立向量自我迴歸模型（VAR），再以此預測 2011Q1 至 2011Q4 之房價成長率，並將結果繪製成扇形圖（圖 5）。

圖 5 課程對英國房價成長率預估的扇形圖



資料來源：課程實作案例

陸、課程心得及建議事項

一、課程心得

本次課程安排大量實作演練，由於經濟體複雜且各國國情不同，不論是變數或模型的選擇，常常涉及許多主觀的判斷，講師亦保持開放態度，並無一定的答案。此呼應課程中的重要精神，雖然計量模型可提供一客觀標準，供做量化分析及比較，惟實務上仍需要主觀判斷能力。此外，不同的假設或模型即可能導出不同的結果，凸顯使用單一模型的風險。

對於計量分析而言，比模型更重要的是資料品質。若資料品質不佳，即便使用再精密的模型，亦容易導出錯誤結論。因此，央行能蒐集大量金融數據，更顯珍貴，如能持續提升資料品質，必能在現有研究基礎上，提升計量分析結果之參考價值。

二、建議事項

在資訊爆炸時代，如何萃取有用資訊非常重要。對個人而言，量化分析將成為必備之基礎能力。對組織而言，資料應該成為一種戰略，有系統的蒐集及歸類，供組織內之個人進行量化分析。謹提出建議如下：

（一）業務上更積極運用量化分析

以適當的計量模型實證量化資訊，佐以質化資訊進行判斷，除有客觀資訊可比較之優點外，亦能擁有質化資訊具備之彈性空間，進而提升本行研究品質。

本次訓練課程提供豐富的基礎計量觀念及軟體演練，收穫甚豐，期能將所學量化技巧用於研究（如嘗試推估民營企業部分金融性資產負債科目，以提高資料頻率），以提高統計品質及效用。

(二) 讓資料成為一種戰略

資料正確性對計量分析至關重要。就單一統計資料而言，宜依據使用目的，訂定明確的資料範圍及定義，以效率性及正確性蒐集資料。就組織而言，應檢視各統計資料蒐集方式及定義，進行內部跨單位的整合（如銀行申報一張含所有單位所需資料的綜合報表，各使用單位則依據不同的權責，讀取所需內容），以達資料的一致性，亦可降低作業成本，避免重複工作。

(三) 增加相關人力培訓，提升研究品質

為加強行員量化分析能力，建議有計畫地在行內安排一系列的計量分析或資料處理等理論課程，並搭配 EViews、R 等軟體進行實際操作，以提升研究品質。

參考文獻

SEACEN Foundational Course on Econometric Modelling and Forecasting Programme Materials, April 2016.

- Ole Rummel, “Overview of Econometric Modelling and Forecasting in Central Banks.”
- Ole Rummel, “Key Statistical Time Series Concepts for Model Building.”
- Ole Rummel, “Key Features of Data.”
- Ole Rummel, “Concepts, Methodologies and Practical Challenges in Estimating Output Gaps.”
- Ole Rummel, “Estimating the Sri Lankan output gap.”
- Ole Rummel, “Estimating the Sri Lankan output gap with the Hamilton(2016) model.”
- Ole Rummel, “The New Keynesian (NK) Model.”
- Ole Rummel, “Estimating a monetary policy model for India.”
- Ole Rummel, “Principles of Forecasting.”
- Ole Rummel, “(Conditional) point, range and density VAR forecasts in EViews.”
- Victor Pontines, “An introduction to VARs.”
- Vincent Lim Choon Seng, “Connecting the Dots - A Reference to Ordinary Least Squares.”