

實時預測澳門經濟增長： 動態因子模型的應用

梁展翹*

澳門金融管理局研究及統計廳

內容提要

實時的經濟決策是基於對當前和未來經濟狀況的評估，惟這些評估所使用的數據並不一定完整。鑒於大部分統計數據的發佈存在滯後性，且經常進行後續修訂，因此，獲取對當前季度狀況的可靠估算和預測，對貨幣當局來說是至關重要的。最新的統計和“大數據”技術使構建實時預測（Nowcasting）——實時監察經濟狀況變化的半自動化平台變得可行。本文介紹動態因子模型的統計框架，以及將其應用於實時預測澳門經濟。實證結果顯示，動態因子模型在預測準確度方面普遍優於傳統的時間序列方法。我們展示了實時預測不僅是數字上的估算；該技術能夠提供一個適用於分析實時數據流的科學框架，以了解經濟基本狀況的變化和識別市場轉折點。

* 作者非常感謝參與澳門大學經濟學工作坊（2023 年 6 月 12 日）的學者們對本研究提供寶貴的意見。文中如有任何錯漏，概由作者負責。

1. 導言

“實時預測”（Nowcasting）一詞源自 *現在*（Now）和 *預測*（Forecasting）的英文合成詞，泛指使用實時信息流估算目標變量當前或未來的狀態（Hopp 2022）。實時預測在經濟學上具重要意義，因為反映經濟現狀的重點統計數據往往存在一定滯後。對於按季度收集的數據來說更是如此，本地生產總值（GDP）正是當中最具代表性的例子。以澳門為例，一般在當季結束約兩個月後，方會發佈 GDP 的官方初步估算數據。

雖然在正常時期，傳統的季度 GDP 數據或足以充分反映經濟活動的變化，惟新冠疫情突顯了及時和系統分析的重要性，尤其在危機時期能為政策決定提供重要信息。在這方面，實時預測是能夠及時且科學地測量經濟現狀的量化技術，並可用於生成 GDP 和其他宏觀經濟指標的初步估算。

實時預測的關鍵在於盡可能使用一些比目標變量更高頻率且更及時的數據，從而在官方統計發佈前進行預測或初步估算（Bańbura *et al.* 2013）。以追蹤澳門的 GDP 為例，我們可以使用入境旅客人次和幸運博彩毛收入等按月公佈的數據，上述兩者均是澳門經濟重要的推動力。我們亦可以納入如金融變量等具有前瞻性的指標；金融變量的頻率較高，且原則上包含一些關於市場對未來經濟發展預期的信息。無論頻率如何，任何數據均有可能影響宏觀經濟的預測和準確度。從這角度來看，我們並沒有捨棄任何信息的理由，但關鍵在於需要了解每項數據是否能忠實地作為反映經濟狀況的訊號。

實時監察經濟狀況在本質上是一個“大數據”問題（Bok *et al.* 2017），對於 GDP 波動性較高的微型開放經濟體更是別具挑戰性。在實踐中，經濟實時預測主要面臨三個涉及數據方面的難點。一是經濟數據的採樣頻率是混合的，如月度或季度等。二是數據公佈時間並不同步，最新可觀測時間因不同數列而異，以致數據出現“參差不齊”的問題。因此，任何實時預測方法均需要為不完整或部分完整的數據做好準備。三是數據量可能非常龐大，導致傳統的統計技術無法有效

應用於實時預測（Buono *et al.* 2017）。然而，近年時間序列計量經濟學的突破和計算力的提高，使經濟學家能迎接大數據的挑戰。

我們在文中描述一個專為實時預測澳門 GDP 而設的計量經濟框架，並說明如何將其應用於分析實時數據流。本文其餘部分的結構如下。第二節回顧與經濟實時預測相關的實證文獻，以及“大數據”帶來的潛在挑戰。第三節概述理論框架和處理複雜數據的方法。第四節展示實證應用和結果。第五節作總結，並提出一些未來研究的方向。

2. 文獻回顧

2.1 監察經濟狀況

一直以來，學術界和政策機構的經濟學家都試圖使用高頻數據來預測季度 GDP 和經濟周期的轉折點（Dauphin *et al.* 2022）。經濟學家恆常地從統計機構、私人和公共調查，或其他來源收集大量數據，並用於分析經濟的健康狀況。然而，從雜亂的數據中提取有意義的信息並非一個簡單的任務。首個關於經濟波動的系統性分析研究可以追溯至 A. F. Burns 和 W. C. Mitchell，他們是在 1930 年代後期於美國國家經濟研究局（NBER）開創經濟周期研究的經濟學家。他們仔細研究數百組經濟數列，務求從中尋找模式和規律；文獻發現數列之間存在一些系統性的聯動，而這些共同波動橫跨並擴散在不同經濟部門和不同類型的經濟活動之中。Burns 和 Mitchell 的研究之所以如此具有創新性，是因為他們想要尋找的規律是未知的——隱含模式識別正是在幾十年後，我們在機器學習領域中稱為“無監督分類”的課題。

Sargent 和 Sims（1977）使用基於因子模型的方法，為上述經濟數列聯動的概念賦予統計學的框架。文獻發現共同波動可以通過少量“潛在”因子來捕捉，而這些因子是能通過數據進行估算的。然而，早期的模型和計算力僅適用於有限

的變量和因子數量。Stock 和 Watson (1989) 是率先將因子模型應用於宏觀經濟建模的學者，他們從一些為數不多的經濟變量，例如就業、工業生產、銷售和收入等數列之中提取信息，並以此構建一個基於單因子模型的經濟周期指數。

隨着時代進步，因子模型得到了延伸和改進，可足以容納更多的變量和構建多項經濟指標。Watson (2004) 證明了 Sargent 和 Sims (1977) 的研究結果同樣適用於高維度的經濟數列。Giannone *et al.* (2004) 擴展了因子模型的計量經濟學框架，並提出使用“狀態空間表達式”為複雜且大型的經濟系統進行降維處理。Evans (2005) 設計了一種在得到新數據時，可實時更新 GDP 估算的算法機制。自此之後，因子模型被廣泛應用於從大量的經濟指標中提取和整合信息。它們現時已成為量化經濟學家用於應對大數據和實時監察經濟狀況的主流方法。特別是，Giannone、Reichlin 和 Small (2008) 提出的動態因子模型是首個被美國聯邦儲備理事會用於實時預測美國 GDP 的方法。隨後，針對不同的經濟體系，其他中央銀行和國際金融機構亦相繼構建和應用不同版本的動態因子模型，其中包括歐洲中央銀行 (Bańbura *et al.* 2011) 和國際貨幣基金組織 (Matheson 2011)。

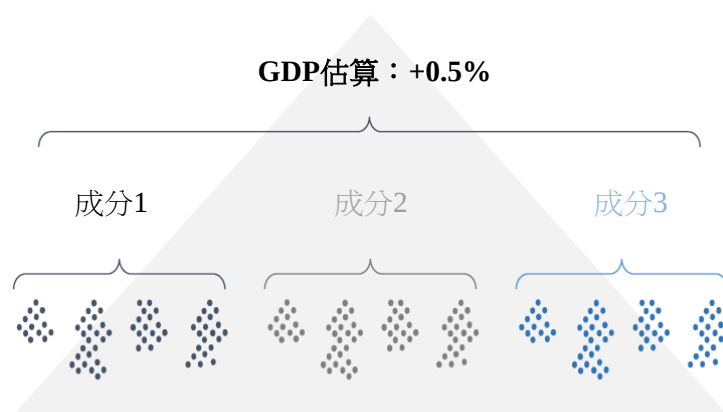
實時預測經濟狀況不僅是在取得新數據時更新 GDP 的估算，同時需要能夠解釋一系列修訂的原因 (Marcellino 和 Sivec 2021)。動態因子模型其中一個理想特徵在於，它允許我們在每次新數據發佈時，計算一個“基於模型的新信息”，以及隨着季度推移，得出每個信息對 GDP 估算的邊際影響。如此一來，我們可以系統地追蹤和評估不同類型或組別數據於傳達經濟活動訊號方面的作用。從本質來看，實時預測技術模仿了市場參與者和專業預測機構的行為，其通過及時監察所有相關的經濟指標以形成一些預期，並持續根據意外衝擊或新數據發佈對預測作出相應的調整。

2.2 實時信息流

除了用於監察經濟狀況的量化技術的進步外，數據可用性亦得以同步發展。傳統上，經濟學家主要關注幾個重點宏觀經濟指標和國民帳戶的組成部分。這些

時間序列數據是在一個協調一致的框架內，有系統地測量所有經濟活動所得出的結果。從概念上看，這種測量是基於核算法，而非統計學或計量經濟學的方法，並一般稱之為**自下而上**，或“數豆子”式的方法。

圖 1：自下而上式（“數豆子”）方法

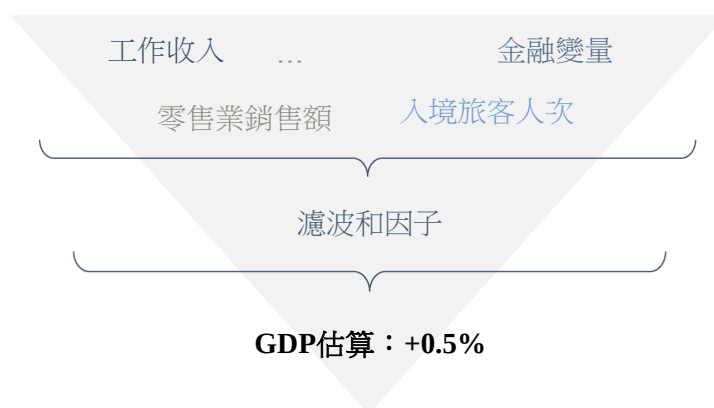


資料來源：原始資料來自 Sonntag 和 Little (2017)，經作者修訂。

然而，“數豆子”方法在大數據的背景和挑戰下顯得過於正規——該方法對細節的注重固然是一種優勢，但同時也是一種缺點，因為細分數據並不容易獲取或存在一定滯後。這對需要及時監察經濟動態的政策制定者來說是巨大的挑戰，因為經濟是複雜且持續演變的。此外，“數豆子”方法並未能完全捕捉或反映不同經濟部門之間的截面相關性（Sonntag 和 Little 2017）。

二十一世紀其中一個典型特徵無疑是數據儲存量、生成速度和種類的爆炸式增長（Hopp 2022）。有鑒於此，第二種方法，即**自上而下**式模型，採用更全面和更直接的估算方法。這種模型從大量數據中提取信息，並將其濃縮成單一指標。因此，實時預測在本質上是一種自上而下式的方法，其一方面盡可能收集更多的數據，另一方面將數列的聯動匯總成少數共同因子。這些因子的動態可通過一些**濾波技術**（附錄 2）追蹤，繼而用於評估經濟當前的狀況。

圖 2：自上而下式估算



資料來源：原始資料來自 Sonntag 和 Little (2017)，經作者修訂。

然而在實踐中，經濟數據的複雜性使自上而下式模型的應用變得困難，例如混合頻率、非同步發佈、高維度，或某些數據來源由於樣本收集時間較短而出現早期數據缺失等。雖然 Giannone *et al.* (2008) 的框架能部分應對數據“參差不齊”的問題，惟他們的方法並不能直接用於數列長度不一的混合頻率面板數據，更概括來說，其不適用於任何數據缺失的情況。此外，文獻使用的估算方法為普通最小平方法，於少量觀測樣本的情況下效率較低，在實時預測澳門經濟時可能導致預測不穩定。

為克服大數據和實時數據流帶來的技術挑戰，Mariano 和 Murasawa (2003) 描述了一種混合頻率因子模型的方法，通過將季度變量表達為可部分觀測及相應的月度變量，繼而將其納入至模型框架之中。Matheson (2011) 在狀態空間模型的框架下，使用卡爾曼濾波和平滑器以處理“參差不齊”的數據，在提取信息的同時，有效地利用不完整橫截面之間的動態關係。Bańbura 和 Modugno (2010) 改進了動態因子模型的估算技術，文獻提出以似然法替代普通最小平方法，並採用最大期望演算法，以應對任意模式的數據缺失。關於該估計量的詳細漸近分析可參考 Doz *et al.* (2012)；其研究證明了似然法在橫截面誤設、特異誤差存在自相關性，以及在數據不遵從正態分佈的情況下均具有穩健性。最重要的是，給定樣本數量 T ，似然法對時間序列數量 n 的限制較小，這確保該方法能適用於

“高維稀疏”，即橫截面維度 n 較高，而可觀測時間 T 相對較少的樣本（Bok *et al.* 2017）。

上述文獻均為似然法於因子模型的應用奠定了理論基礎，該方法對模型誤設具有穩健性，並適用於實時數據流的場境。Bańbura *et al.*（2013）提出使用最大期望演算法，以對動態因子模型進行最大似然估算，繼而將模型應用於實時預測。上述方法被紐約聯邦儲備銀行所使用，並對美國實質 GDP 增長進行高頻率監測（FRBNY 2016）。此外，Ankargren 和 Lindholm（2021）以實證證明，通過最大期望演算法訓練的動態因子模型，能夠在疫情期間做出有力的反應。

3. 數據和方法

在本節中，我們參考 Giannone *et al.*（2008）的計量經濟學框架，使用狀態空間表達式及參數化動態因子模型以應對高頻數據，為實時預測的課題提供一個解決方案。本文的目的是探究實時預測技術能如何應用於預測和及時追蹤澳門的 GDP、分析月度和高頻數據與宏觀經濟指標之間的聯繫，以及其是否能夠為精確估算季度累計指標帶來貢獻。在估算技術方面，我們採用 Bańbura *et al.*（2013）的方法，以最大期望演算法訓練並對模型進行估算。接下來將介紹數據集和模型的框架，以及相關的估算程序。

3.1 數據選取

在構建經濟實時預測模型時，我們需要選擇一些及時和有助於預測實質 GDP 增長的變量（Kunovac 和 Špalat 2014）。在本文中，我們收集的混合頻率數據集包括合計 25 組時間序列，其中大部分為月度數列，因此我們的實時預測模型採用月度頻率。對於如金融變量等頻率高於月度的數據，我們可以累計並計算其月度平均值。表 1 展示了本文採用的變量及頻率。數據集的時間跨度由 2002 年 1 月至 2023 年 4 月，而完整的數據集可自 2008 年 1 月起取得。大部分的數據均遵從可預示的統

計公佈時間表，為在澳門應用實時預測的方法提供了合理性。

值得注意的是，數據集需要足夠多樣化以具包容性。選取過於集中在特定經濟部門的數列將導致模型參數出現偏差，繼而降低模型在估算增長和潛在因子的有效性（Matheson 2011）。因此，我們仔細地從廣泛的經濟部門之中挑選適當的變量，並將其分類成八個組別（表1）。每組時間序列均以平穩型態輸入至模型之中：這普遍需要進行對數變換、差分，或兩者兼備；異常值則從數據集中移除。為實時預測澳門GDP，我們的數據集取自於統計暨普查局（統計局）、博彩監察協調局（博監局），以及澳門金融管理局（金管局）。

表 1：數據集

序號	組別	時間序列	頻率
1	產出和商品貿易	實質本地生產總值	Q
2		貨物出口	M
3		貨物進口	M
4		工業生產指數：總指數	Q
5	勞動力市場	失業率	M
6		就業不足率	M
7		勞動人口	M
8	價格和收入	綜合消費物價指數：總指數	M
9		進口單位價格指數：總指數	M
10		月工作收入中位數	Q
11	房地產和建築	住宅單位交易數目	M
12		住宅單位交易價值	M
13		獲發動工批示（新動工）住宅單位數目	M
14	利率和匯率	1 個月為期的金融票據收益率	D
15		3 個月為期的金融票據收益率	D
16		12 個月為期的金融票據收益率	D
17		按貿易額加權計算的名義有效匯率	D
18	貨幣和信貸	貨幣量 M1	M
19		貨幣量 M2	M
20		本地信貸	M
21	旅遊業和零售業	入境旅客人次	M
22		旅客人均消費	Q
23		酒店業平均入住率	M
24		零售業銷售額	Q
25	博彩業	幸運博彩毛收入	M

附註：頻率列註明變量的發佈頻率，取值為季度（Q）、月度（M）或每日（D）。

資料來源：統計局、博監局及金管局。

另一項選擇標準為，我們只採用每組經濟數列的總體指標。以綜合消費物價指數為例，我們僅採用其總體指數。Barhoumi *et al.* (2009) 通過檢驗證明了細分數據對實時預測準確度的邊際影響較小，此結論與市場普遍更為關注總體指標的行為一致。

3.2 計量經濟框架

3.2.1 應對大數據

高維度數據使統計模型的複雜性不斷增加。顯然地，使用過於簡化的模型來處理巨量數據可能會遺漏重要的數據特徵，從而導致模型出現誤設。相反，使用過多的變量將令模型的參數量過度增加，意味參數的估算存在較大不確定性，這將使傳統方法的估算結果變得不可靠和不穩定。此問題在統計科學中一般稱為“維度災難”。從計量經濟學的角度來看，在橫截面維度 n 較高，而可觀測時間樣本 T 較少的情況下，即“高維稀疏”的數據集，參數的估算是尤其困難的。

然而，Burns和Mitchell (1946) 的研究指出，大部分的經濟波動都是被一些共同因素所驅動的，在統計學中一般稱這些共同因素為因子。繼此開創性的發現後，Forni *et al.* (2000) 及Stock和Watson (2002) 的貢獻同樣重要，他們在應對巨量數據的時候，引進了主成分分析法為大型因子模型進行降維處理，我們認為這是該領域最前沿的技術之一。“大數據計量經濟學”旨在將維度缺憾轉為一個優勢，並以精簡的形式捕捉多組經濟時間序列的特徵和相互之間的關係 (Bok *et al.* 2017)。

3.2.2 混合頻率數據的時間聚合法

正如文獻中所強調，經濟實時預測的另一技術難點為應對混合頻率數據。鑒於本文的目標是構建一個適用於監測澳門GDP的月度指標，因此我們有必要為季度數據，以及與其對應的、僅部分可觀測的月度數值之間建立一個數學上的關係。

該關係取決於數列是流量變量或存量變量，以及其在輸入模型之前的變換方法。對於以GDP為代表的流量變量，關鍵是首先將季度數列表達為月度數列之和（Harvey 1990），然後使用幾何平均值作為算術平均值的近似表達式：

$$Y_t^Q = Y_t^M + Y_{t-1}^M + Y_{t-2}^M \\ \approx 3(Y_t^M Y_{t-1}^M Y_{t-2}^M)^{\frac{1}{3}},$$

當中的 t 代表月份，因此上述表達式在 $t = 3, 6, 9, \dots$ 的情況下是正確的。如果我們對GDP取對數差分變換，並將其作為變動率的近似表達式，那麼以下的算法將是正確的（Mariano和Murasawa 2003）：

$$\begin{aligned} \Delta^Q y_t^Q &= \log(Y_t^Q) - \log(Y_{t-3}^Q) \\ &= \frac{1}{3} \left(\Delta \log(Y_t^M) + 2\Delta \log(Y_{t-1}^M) + 3\Delta \log(Y_{t-2}^M) \right. \\ &\quad \left. + 2\Delta \log(Y_{t-3}^M) + \Delta \log(Y_{t-4}^M) \right) \\ &= \frac{1}{3} (\Delta y_t^M + 2\Delta y_{t-1}^M + 3\Delta y_{t-2}^M + 2\Delta y_{t-3}^M + \Delta y_{t-4}^M). \end{aligned}$$

如此一來，若我們進一步定義以下矩陣（Kunovac 和 Špalat 2014）：

$$A_Q = \begin{bmatrix} \ddots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \dots & 3 & 2 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ \dots & 0 & 1 & 2 & 3 & 2 & 1 & 0 & 0 & 0 \\ \dots & 0 & 0 & 0 & 0 & 1 & 2 & 3 & 2 & 1 \end{bmatrix},$$

那麼我們便可以在 GDP 的季度數列和月度數列兩組向量之間建立一個數學關係： $\Delta^Q Y^Q = \frac{1}{3} A_Q \Delta Y^M$ ，當中的 $\Delta^Q Y^Q = (\dots, \Delta^Q y_{T-6}^Q, \Delta^Q y_{T-3}^Q, \Delta^Q y_T^Q)'$ 代表季度數列，而 $\Delta Y^M = (\dots, \Delta y_{T-2}^M, \Delta y_{T-1}^M, \Delta y_T^M)'$ 則代表與之對應的月度數列。

3.2.3 動態因子模型

我們現為表1的經濟變量建立一個狀態空間表達式及參數化動態因子模型。

如果數列之間存在高度的聯動，那麼這些具擴散性的共同波動將能夠被少量的潛在因子捕捉，而這些因子可以通過數據估算。動態因子模型的前提在於，其假設多組觀測數據 Y_t 是由少量潛在因子 F_t 所驅動，個別數列特有的特徵則歸納在特異部分 E_t 之中，如個別數列的測量誤差等。參考此類模型的典型設定，假設因子的動態遵從向量自回歸 過程。我們的實證模型可以通過以下方程總結：

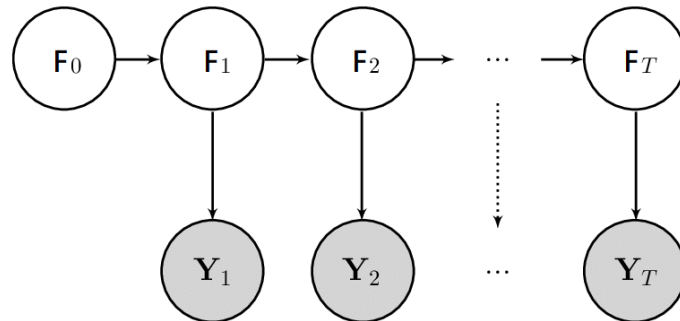
$$Y_t = \mu + \Lambda F_t + E_t, \quad E_t \sim i.i.d. \mathcal{N}(0, \Sigma_E), \quad (1)$$

$$F_t = \Phi(L)F_t + U_t, \quad U_t \sim i.i.d. \mathcal{N}(0, \Sigma_U), \quad (2)$$

當中的 $\Phi(L)$ 代表滯後算子， U_t 則是對共同因子造成衝擊的白噪聲 過程。向量 μ 和矩陣 Λ 分別代表無條件平均值和因子載荷；它們是除了潛在因子外，亦需要進行估算的模型參數。因子的向量自回歸結構對實時預測尤其重要，因為在數據“參差不齊”的情況下，橫截面和動態信息均是非常具有價值的。

方程（1）—（2）共同組成了一個狀態空間表達式，當中的共同因子和特異部分則是一些“潛在”的狀態。方程（1）是將觀測變量和不能觀測的狀態聯繫起來的測量方程，方程（2）則是描述整個系統動態的轉移方程（這正是該模型被稱為“動態”因子模型的原因）。如此一來，動態因子模型將目標變量和預測變量的聯合動態匯總在一個統一的系統之中。

圖 3：狀態空間表達式



資料來源：原始資料來自Murray et al. (2013)，經作者修訂。

為避免模型參數過多，我們假設共同衝擊 U_t 和特異衝擊 E_t 的橫截面均是正交的。我們進一步假設它們相互之間亦是正交的。然而，有別於在模型中假設為白噪聲的共同衝擊，特異衝擊允許存在較弱的自相關性（Cepni *et al.* 2018）。

3.2.4 估算程序

鑒於因子 F_t 是不可觀測的，動態因子模型的參數，即 $\theta = (\mu, \Lambda, \Phi, \Sigma_E, \Sigma_U)$ 的最大似然估計一般不存在閉式解。另一方面，直接以數值方法求最大似然需要龐大的計算力，對於截面維度高及模型參數多的情況更是如此（Jungbacker 和 Koopman 2008）。傳統上，小型因子模型能通過一些算法在時域中進行優化。如 Engle 和 Watson（1981）在狀態空間模型中，使用卡爾曼濾波計算似然函數，並使用一些優化方法求得參數的最大似然估計值。Watson 和 Engle（1983）提出另一種方法，他們將 Dempster *et al.*（1977）的最大期望演算法用於估算因子模型。

最大期望演算法的設計初衷在於，在面臨潛在或（部分）數據不完整等棘手情況時，為似然函數的估算提供一個通用的辦法。然而，在動態因子模型的框架下，早期的文獻普遍假設觀測變量並無數據缺失，該演算法的用途亦僅限於估算潛在因子。因此 Bańbura 和 Modugno（2010）對最大期望演算法進行補強，並將其延伸至適用於任意數據缺失的情況，為大型因子模型的估算提供了辦法；相關數學細節收錄在附錄 3。最大期望演算法的原理是，將似然函數表達為觀測數據和不可觀測數據的方程，並迭代以下兩個數學運算：(i) 根據數據和上一次迭代得出的參數作為假設條件，“填補”缺失的數據和計算對數似然函數的期望值；(ii) 求解和重估模型參數，從而將對數似然函數的期望值最大化。

在估算方程（1）—（2）時，最大期望演算法將迭代上述兩個步驟直至收斂，並在每一步因應共同因子的估算和不確定性，對模型參數作出相應校正（Watson 和 Engle 1983）。為了對演算法進行初始化，以及獲取模型參數的初始值 $\theta(0)$ ，我們可從標準正態分佈 $\mathcal{N}(0, 1)$ 中隨機抽取數值以替換 Y_t 內缺失的數據，並對原始數據集使用主成分分析法（附錄 1）。

演算法 1：最大期望演算法

1. 選擇和輸入模型參數的初始估計值 $\theta(0) = (\hat{\mu}, \hat{\Lambda}, \hat{\Phi}, \hat{\Sigma}_E, \hat{\Sigma}_U)$ ；
2. **重複**
3. 計算對數似然函數 $l(Y, F; \theta)$ ；
4. **求期望值**：根據數據集 Ω_v 和上一次迭代得出的模型參數 $\theta(j)$ 作為假設條件，計算對數似然的期望值：

$$L(\theta, \theta(j)) = \mathbb{E}_{\theta(j)} [l(Y, F; \theta) | \Omega_v]；$$

5. **求最大化**：求解模型參數 θ 以使對數似然函數的期望值最大化，重估所得的取值 $\theta(j+1)$ 即為新的模型參數值：

$$\theta(j+1) = \operatorname{argmax}_{\theta} \{L(\theta, \theta(j))\}；$$

6. **直至** 收斂至（局部）最優的 $l(Y, F; \theta)$ ；
 7. **輸出** 模型參數的最優解向量 $\theta = (\mu, \Lambda, \Phi, \Sigma_E, \Sigma_U)$ 。
-

資料來源：原始資料來自 Bańbura 和 Modugno（2010），經作者修訂。

Doz *et al.*（2012）證明了最大似然法在動態因子模型和方程（1）—（2）等大型系統的一致性。其研究檢驗了最大似然估計在不同誤設情況下的漸近性質，包括在特異部分遺漏橫截面相關性和自相關性，或數據為**非高斯分佈**等。文獻中的一個關鍵發現在於，在橫截面維度較高和觀測樣本足夠多的時候，模型誤設對於估算因子的影響是較小的。最大似然法的另一優點在於，其允許我們對模型參數施加限制。例如，Bańbura 和 Modugno（2010）在因子載荷上施加一些條件，以反映混合頻率數據的時間聚合。Bańbura *et al.*（2011）則針對不同變量組別，引入了某些組別特有的因子。

最後，在估算因子和模型參數後，季度變量 k 在目標時間點 t 的預測值或實時預測值可通過對現有的因子估算進行時間聚合，形成一個簡易預測表達式：

$$y_{k,t}^Q = \alpha + \beta F_{t|\Omega_v}^Q + e_{k,t}^Q, \quad t = 3, 6, 9, \dots \quad (3)$$

3.2.5 實時預測的更新和新信息

承接前節討論，在實時預測 GDP 的過程中，我們需面對非同步發佈的連續數據流，且發佈的滯後程度因不同數列而異。因此，對目標時段僅作出單次預測的情況並不常見，在現實中更常見的是多次的預測，這些一系列的測算將在接收新資訊時進行更新（Schumacher 和 Breitung 2008）。

直觀來看，只有數據中的*新信息*，或“意料之外”的部分方會修訂模型估算。因此，實時預測的關鍵在於，提取*新信息*和將其與後續修訂建立一個系統性的聯繫。接下來將介紹“基於模型的新信息”的概念，並說明動態因子模型如何提取這些訊號以進行實時預測。

我們以 Ω_v 表示在發佈日 v^1 可獲得的數據集。首先，我們應分析兩組連續的數據集 Ω_v 和 Ω_{v+1} 之間的差異。信息集 Ω_v 和 Ω_{v+1} 可能因以下兩個原因而不同。一是 Ω_{v+1} 包含一些最新發佈的數據 $\{y_{j,t_{j,v+1}}, j \in \mathbb{J}_{v+1}\}$ ，而這些數據並不存在於上一次的信息集 Ω_v 之中。二是可能某些舊數據出現修訂。為免使符號變得更複雜，對於以往的數據，我們採用其最新修訂的數字，並得出：

$$\Omega_v \subset \Omega_{v+1} \quad \text{和} \quad \Omega_{v+1} \setminus \Omega_v = \{y_{j,t_{j,v+1}}, j \in \mathbb{J}_{v+1}\},$$

因此信息集是具有“擴充性”的。鑒於不同種類的經濟數據的發佈滯後程度並不一致，故在 $j \neq l$ 的情況下，一般來說 $t_j \neq t_l$ 。我們現探討變量 k 在目標時間點 t_k 在連續兩次實時預測之間的更新機制。新數據發佈 $\{y_{j,t_{j,v+1}}, j \in \mathbb{J}_{v+1}\}$ 普遍包含一些關於目標變量 y_{k,t_k} 的新資訊，而這些新資訊可能引致模型對估算作出修訂。根據條件期望及其作為正交投影算子的特性，可得出：

$$\underbrace{\mathbb{E}[y_{k,t_k}^o \mid \Omega_{v+1}]}_{\text{新預測}} = \underbrace{\mathbb{E}[y_{k,t_k}^o \mid \Omega_v]}_{\text{舊預測}} + \underbrace{\mathbb{E}[y_{k,t_k}^o \mid I_{v+1}]}_{\text{修訂}}, \quad (4)$$

¹ 我們不以 j 對數據集進行記數，因為經濟數據可能會因不同步發佈而出現更高頻率的情況。

當中的 I_{v+1} 是信息集 Ω_{v+1} 中的一個子集，其元素與 Ω_v 中所有的元素均是正交的。根據上述列明的 Ω_v 和 Ω_{v+1} 之間的差異，可得出：

$$I_{v+1} = (I_{v+1,1}, \dots, I_{v+1,J_{v+1}})' \quad \text{和} \quad I_{v+1,j} = y_{j,t_{j,v+1}} - \mathbb{E}[y_{j,t_{j,v+1}} | \Omega_v],$$

當中的 J_{v+1} 代表 $\mathbb{J}_{v+1}$ 的元素數量。換句話說， I_{v+1} 是數據發佈中“意料之外”的部分（對於模型而言），我們將之稱為新信息。值得注意的是，對預測造成修訂的並不是數據發佈本身，而是這些新信息。縱使可能性不大，但假設於數據發佈與模型預測完全一致的情況下，即“沒有新信息”，模型預測並不會作出修訂。另一方面，我們可以直覺地預想，例如在旅客人次出現負面新信息的情況下，模型將會下調澳門 GDP 的預測值。

方程（4）可以進一步擴展以量化這些“基於模型的新信息”，從而使我們能夠評估每次數據發佈的影響。Bańbura *et al.*（2013）證明了在動態因子模型的框架下，假設所有數據都遵從高斯分佈，那麼我們可以找到一個滿足以下條件的向量 $\mathcal{D}_{v+1} = (\delta_{v+1,1}, \dots, \delta_{v+1,J_{v+1}})$ ：

$$\underbrace{\mathbb{E}[y_{k,t_k}^Q | \Omega_{v+1}] - \mathbb{E}[y_{k,t_k}^Q | \Omega_v]}_{\text{修訂}} = \mathcal{D}_{v+1} I_{v+1} = \sum_{j=1}^{J_{v+1}} \delta_{v+1,j} \underbrace{(y_{k,t_k} - \mathbb{E}[y_{k,t_k} | \Omega_v])}_{\text{新信息}} \quad (5)$$

換句話說，預測的修訂可以分解為新信息的加權平均，而這些新信息可以在最新的數據發佈中提取。因此，與我們早前的預想一致，修訂的幅度一方面取決於新信息的大小，另一方面取決於其與目標變量的關係，即權重的數值 $\delta_{v+1,j}$ 。事實上，方程（5）可以理解為卡爾曼濾波更新方程（Harvey 1990）在應對數據非同步發佈時的通用版本。方程（5）可以幫助追蹤預測修訂的源頭，其能夠將一系列的修訂分解為來自不同個別，或如表1裡不同組別數據所貢獻的新信息。此外，我們可依此提出諸如“鑒於最新發佈的旅客人次數據高於預期，因此澳門實質GDP的增長預測相應向上調整”等表述。

4. 實證應用

4.1 實時預測澳門實質 GDP 增長

我們將在本節展示上文的概念，並應用動態因子模型於實時預測澳門的實質 GDP 增長。本節的目的是闡明實時數據流於連續更新和實時預測的作用。在每一次更新時，我們都會檢視不同組別的數據發佈將如何修訂預測，以及其帶來的不確定性。實時預測模型會以*遞迴形式* 進行*樣本外的回測*，從而在每一樣本複現該指定時間點所能獲取的可用數據集，即一個模擬預測環境。為簡化流程，我們遵從一個既定和程式化的數據發佈時間表，而對於以往的數據，我們則採用其最新修訂的數字，即一個*虛擬仿真* 的實時平台設計。

在本文的實證應用中，我們將以月度頻率生成 GDP 的實時預測；我們亦可以將其延伸至更高的預測頻率。在每次進行月度更新時，我們使用當月月初可得的數據集。準確來說，每次月度更新的時機將緊接博監局發佈幸運博彩毛收入的月度數據之後。例如，在五月初，幸運博彩毛收入的最新可觀測數據為四月的數據，而入境旅客的數據則截至三月底。相應地，在每月月初的時間點，幸運博彩毛收入和入境旅客的數據將分別存在一個月及兩個月的“參差不齊”。因此，若我們在五月初評估動態因子模型，那麼我們可以得出四月的幸運博彩毛收入數據對模型預測的影響。相同的機制適用於表 1 裡所有的變量。

作為一個示範，我們對 2022 年第二季實質 GDP 的按年增長率進行一系列的預測²和實時預測。從 2022 年一月初至 2022 年八月底期間，即在統計局發佈官方 GDP 估算之前，我們每月都會納入新發佈的數據（表 2），並計算實時預測值。2022 年第二季是一段理想的回測時段，其特別適用於展示動態因子模型在應對市場環境變化的能力——於 2021 年第二季，澳門經濟從疫情中展現強勁的復甦，惟於 2022 年 6 月 18 日卻再次經歷了本地疫情的爆發。

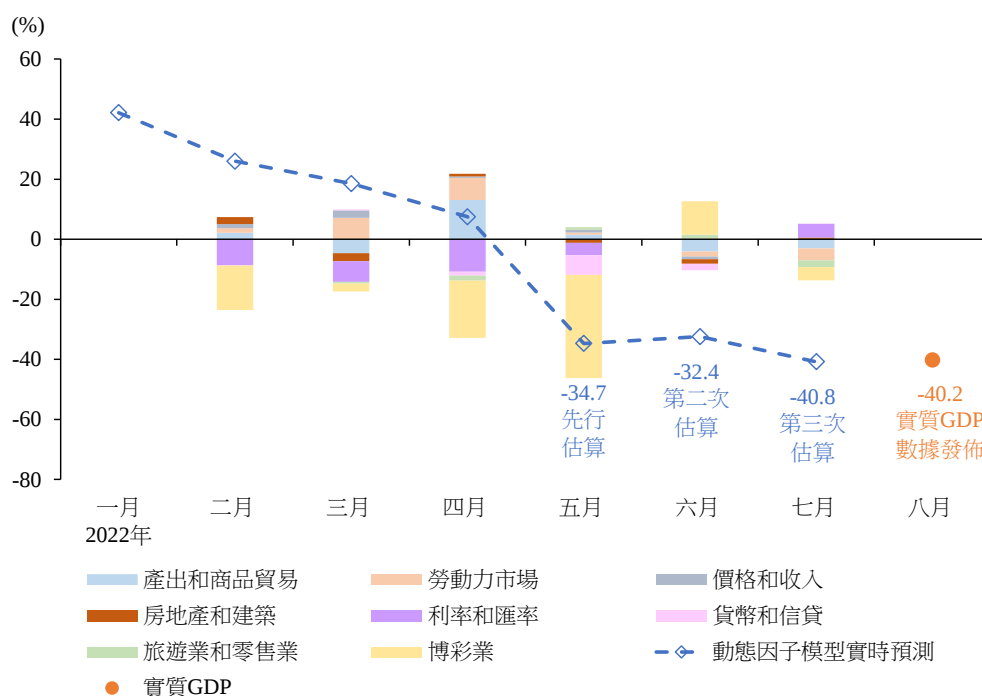
² 在本文的術語中，這泛指在上一季進行的測算。

表 2：實時預測澳門 GDP 的機制

時間點	實時預測序列
2022 年 5 月初	使用截至四月底的數據，即測算季度內第一個月的數據，對 2022 年第二季的 GDP 進行 先行估算 。
2022 年 6 月初	使用截至五月底的數據，即測算季度內第二個月的數據，對 2022 年第二季的 GDP 進行 第二次估算 。
2022 年 7 月初	使用截至六月底的數據，即測算季度內第三個月的數據，對 2022 年第二季的 GDP 進行 第三次估算 。

資料來源：作者。

圖 4：實時預測澳門 GDP



附註：橫軸表示進行預測或實時預測的時間點。

資料來源：原始資料來自統計局、博監局、金管局及作者估算。

我們現評論 GDP 實時預測的結果。在第二季前，模型的預測值維持在正數的區間。然而，在四月的幸運博彩毛收入按年下跌 68.1% 的數據發佈後，第一次的實時預測隨即預示了 2022 年第二季的實質 GDP 將呈現負增長。鑒於幸運博彩毛收入明顯低於動態因子模型的預期，因此該負面的新信息被反映在模型的

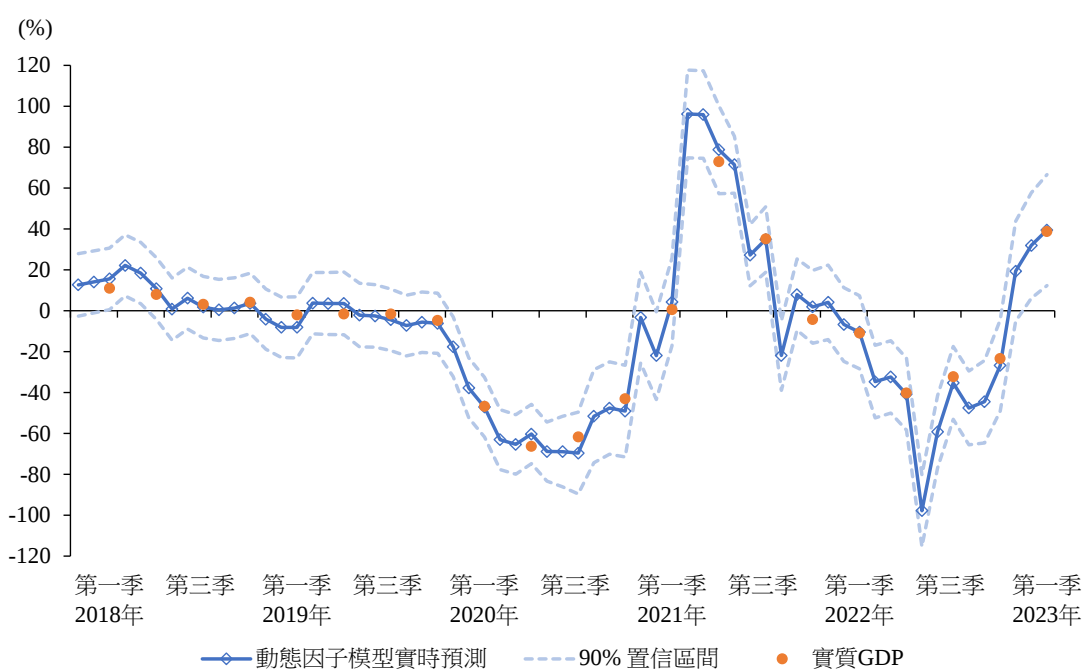
先行估算之中。事實上，GDP 實時預測在最新的數據發佈後同樣呈現下調，部分反映了 2022 年六月中旬在本地爆發的疫情。

根據我們觀察，動態因子模型能夠對持續變化的市場環境做出迅速的反應，且 GDP 的實時預測序列會隨着時間推移而向真實數值收斂。更重要的是，圖 4 描繪了實時預測的核心思想：該技術能夠從實時數據流中提取資訊，並將新信息分解和歸類為來自於不同數據組別的貢獻，即圖 4 中的彩色條形圖，使經濟學家能夠明白經濟基本狀況的變化，以及量化每個經濟部門的影響。

4.2 回測和預測評估

為進行穩健性測試，以及在更長的回測時段內評估模型的預測能力，我們在 2018 年第一季至 2023 年第一季期間以遞迴形式進行樣本外的回測。圖 5 展示了澳門實質 GDP 按年增長率的虛擬實時預測。圖中的橫軸現表示預測的時間點，從而使我們能更直觀地將預測結果與真實數值進行對比。

圖 5：樣本外的 GDP 預測



資料來源：原始資料來自統計局、博監局、金管局及作者估算。

如圖 5 所示，動態因子模型在實時追蹤澳門實質 GDP 增長的表現卓越。尤其是，該模型能有效地識別樣本的轉折點——從 2020 年疫情引致的經濟低迷、2021 年的反彈，2022 年的收縮，乃至 2023 年第一季的復甦。我們觀察到，信息的接收和積累對預測準確度的提升是非常明顯的，表明了該模型是由數據支撐的特性。此外，動態因子模型的預測結果與 GDP 的真實數值頗為一致。

在本文中，因子數量或滯後階數等模型選擇，是基於均方根預測誤差的最小化。我們亦可使用該指標評估動態因子模型，從而與一些基準模型作對比，例如自回歸模型和簡易隨機遊走等。

表 3 展示了動態因子模型和一階自回歸模型的預測結果，以及其與簡易隨機遊走的相對均方根預測誤差，我們以後者在兩季前進行的預測作為評估的基準單位。若模型在給定評估樣本內的相對均方根預測誤差數值小於一，則表示該模型的預測能力優於簡易隨機遊走的基準。

表 3：基於相對均方根預測誤差的評估

回測時段：2018 年第一季至 2023 年第一季			
	動態因子模型	一階自回歸模型	隨機遊走
兩季前進行的預測			
基準預測	0.81	0.89	1.00
上一季進行的預測			
第一次預測	0.82	0.89	1.00
第二次預測	0.77	0.76	0.83
第三次預測	0.57	0.76	0.83
當季進行的預測			
先行估算	0.39	0.76	0.83
第二次估算	0.22	0.52	0.55
第三次估算	0.08	0.52	0.55

附註：粗斜體表示該模型在給定評估樣本的均方根預測誤差是最優的。

資料來源：原始資料來自統計局、博監局、金管局及作者估算。

表 3 的結果證明了動態因子模型在預測澳門 GDP 方面比其他方法更準確。對於較長期的預測，我們觀察到均方根預測誤差的差異較小，而動態因子模型的表現普遍優於單變量的基準模型。對於較短期的預測，隨着數據逐漸豐富，動態因子模型展現了卓越的預測能力。從基準預測到第三次預測之間，模型的均方根預測誤差減少了約 30%，而到達第三次估算，即最後一次實時預測時更是縮減了超過 90%。此實證證明了使用高維度數據以提取實時信息對經濟預測的重要性（Bernanke 和 Boivin 2003）；這正是我們在現代語言所說的“大數據分析”。

5. 總結

我們在本文說明了現代量化技術能如何應用於構建一個分析實時數據流的半自動化平台——實時預測，這正是我們用於監察經濟狀況的技術。為此，我們參考 Giannone *et al.*（2008）的計量經濟框架，並應用動態因子模型，檢驗月度和高頻數據在實時預測澳門實質 GDP 增長所發揮的作用。

在預測能力方面，經多個不同組別的變量所訓練的動態因子模型，能夠比如自回歸模型和簡易隨機遊走等單變量基準模型作出更精確的預測；準確度的提升在短期預測更為明顯。這顯示動態因子模型是由數據驅動，其能夠將大量變量和“大數據”的信息整合，此特性使該模型能適用於實時追蹤澳門的 GDP。此外，通過在 2018 年第一季至 2023 年第一季期間進行回測，證明隨着信息積累，模型的 GDP 實時預測將收斂至真實數值。

我們認為實時預測不僅是提供一個早期的 GDP 估算；其目的是創造一個能夠系統地讀取實時數據流的方法。我們在本文描述了一個專為此而設的計量經濟框架。準確來說，所有變量和潛在的經濟狀態均被納入至一個統一的系統之中，從而使我們能夠科學地識別經濟中的新信息。為闡明此概念，我們以實時預測 2022 年第二季澳門的 GDP 為例，展示動態因子模型如何提取信息，以及從實時數據流中追蹤變動的源頭。

最後，本文收集的數據集主要由“剛性”數據組成。因此，延伸至“軟性”數據，以及如谷歌趨勢等另類數據或能提供更豐富的資訊，並可能提高實時預測的質量。另一個未來研究方向則是探索如深度或機器學習等嶄新技術於實時預測的應用。

參考文獻

Ankargren, S. and U. Lindholm (2021), “Nowcasting Swedish GDP Growth,” National Institute of Economic Research Working Paper, No. 154.

Assunção, J. B. and P. A. Fernandes (2022), “Nowcasting GDP: An Application to Portugal,” Forecasting 2022, Vol. 4, 717-731.

Auroba, S. B., F. X. Diebold, M. A. Kose and M. E. Terrones (2011), “Globalization, the Business Cycle, and Macroeconomic Monitoring,” IMF Working Paper, WP/11/25.

Bańbura, M., D. Giannone, M. Modugno and L. Reichlin (2013), “Now-Casting and the Real-Time Data Flow,” Handbook of Economic Forecasting, Vol. 2, 195-237.

Bańbura, M., D. Giannone and L. Reichlin (2011), “Nowcasting,” The Oxford Handbook of Economic Forecasting, 193-224.

Bańbura, M. and M. Modugno (2010), “Maximum Likelihood Estimation of Factor Models on Data Sets with Arbitrary Pattern of Missing Data,” European Central Bank Working Paper Series, No. 1189.

Bańbura, M. and L. Saiz (2008), “Short-Term Forecasts of Economic Activity in the Euro Area,” European Central Bank Monthly Bulletin, Issue 2/2020, 69-74.

Barbaglia, L., L. Frattarolo, L. Onorante, F. M. Pericoli, M. Ratto and L. T. Pezzoli (2022), “Testing Big Data in a Big Crisis: Nowcasting under COVID-19,” European Commission JRC Working Papers in Economics and Finance, No. 2022/6.

Barhoumi, K., O. Darné and L. Ferrara (2009), “Are Disaggregate Data Useful for Factor Analysis in Forecasting French GDP?” Banque de France Working Paper Series, No. 232.

Bernanke, B. S. and J. Boivin (2003), “Monetary Policy in a Data-Rich Environment,” Journal of Monetary Economics, Vol. 50, Issue 3, 525-546.

Bok, B., D. Caratelli, D. Giannone, A. M. Sbordone and A. Tambalotti (2017), “Macroeconomic Nowcasting and Forecasting with Big Data,” Federal Reserve Bank of New York Staff Reports, No. 830.

Buono, D., G. Kapetanios, M. Marcellino, G. L. Mazzi and F. Papailias (2017), “The Challenge of Big Data,” Handbook on Rapid Estimates, 449-462.

Burns, A. F. and W. C. Mitchell (1946), “Measuring Business Cycles,” NBER Book Series Studies in Business Cycles.

Cepni, O., I. E. Guney and N. R. Swanson (2018), “Nowcasting and Forecasting GDP in Emerging Markets Using Global Financial and Macroeconomic Diffusion Indexes,” International Journal of Forecasting, Vol. 35, Issue 2, 555-572.

Chernis, T. and R. Sekkel (2017), “A Dynamic Factor Model for Nowcasting Canadian GDP Growth,” Bank of Canada Staff Working Paper, No. 2017-2.

Cristea, R. G. (2020), “Can Alternative Data Improve the Accuracy of Dynamic Factor Model Nowcasts?” Cambridge Working Papers in Economics, Faculty of Economics, University of Cambridge, No. 20108.

Dauphin, J. F., K. Dybczak, M. Maneely, M. T. Sanjani, N. Suphaphiphat, Y. Wang and H. Zhang (2022), “Nowcasting GDP – A Scalable Approach Using DFM, Machine Learning and Novel Data, Applied to European Economies,” IMF Working Paper, WP/22/52.

Dempster, A. P., N. M. Laird and D. B. Rubin (1977), “Maximum Likelihood Estimation from Incomplete Data via the EM Algorithm,” Journal of the Royal Statistical Society, Vol. 39, No. 1, 1-38.

Doz, C., D. Giannone and L. Reichlin (2011), “A Two-Step Estimator for Large Approximate Dynamic Factor Models Based on Kalman Filtering,” Journal of Econometrics, Vol. 164, Issue 1, 188-205.

Doz, C., D. Giannone and L. Reichlin (2012), “A Quasi Maximum Likelihood Approach for Large, Approximate Dynamic Factor Models,” The Review of Economics and Statistics, Vol. 94, No. 4, 1014-1024.

Engle, R. and M. Watson (1981), “A One-Factor Multivariate Time Series Model of Metropolitan Wage Rates,” Journal of the American Statistical Association, Vol. 76, Issue 376, 774-781.

Evans, M. D. D, (2005), “Where Are We Now? Real-Time Estimates of the Macroeconomy,” International Journal of Central Banking, Vol. 1, Issue 2, September.

Forni, M., M. Hallin, M. Lippi and L. Reichlin (2000), “The Generalized Dynamic Factor Model: Identification and Estimation,” The Review of Economics and Statistics, Vol. 82, Issue 4, 540-554.

FRBNY (2016), Nowcasting Report, September.

Giannone, D., L. Reichlin and L. Sala (2004), “Monetary Policy in Real Time,” NBER Macroeconomics Annual, Vol. 19, 161-200.

Giannone, D., L. Reichlin and D. Small (2008), “Nowcasting: The Real-Time Information Content of Macroeconomic Data,” Journal of Monetary Economics, Vol. 55, Issue 4, 665-676.

Harvey, A. C. (1990), “Forecasting, Structural Time Series Models and the Kalman Filter,” Cambridge Books, Cambridge University Press.

Hopp, D. (2022), “Economic Nowcasting with Long Short-Term Memory Artificial Neural Networks (LSTM),” Journal of Official Statistics, Vol. 38, Issue 3, 847-873.

Jansen, W. J., X. Jin and J. Winter (2016), “Forecasting and Nowcasting Real GDP: Comparing Statistical Models and Subjective Forecasts,” International Journal of Forecasting, Vol. 32, Issue 2, 411-436.

Jungbacker, B. and S. J. Koopman (2008), “Likelihood-based Analysis for Dynamic Factor Models,” Tinbergen Institute Discussion Paper, TI 2008-007/4.

Jungbacker, B., S. J. Koopman and M. Wel (2009), “Dynamic Factor Analysis in the Presence of Missing Data,” Tinbergen Institute Discussion Paper, TI 2009-010/4.

Kant, D., A. Pick and J. Winter (2022), “Nowcasting GDP using Machine Learning Methods,” DNB Working Paper, No. 754.

Kunovac, D. and B. Špalat (2014), “Nowcasting GDP Using Available Monthly Indicators,” Croatian National Bank Working Papers, W-39.

Luciani, M. and L. Ricci (2014), “Nowcasting Norway,” International Journal of Central Banking, Vol. 10, Issue 4, 215-248.

Marcellino, M. and V. Sivec (2021), “Nowcasting GDP Growth in a Small Open Economy,” National Institute Economic Review, Vol. 256, 127-161.

Mariano, R. S. and Y. Murasawa (2003), “A New Coincident Index of Business Cycles Based on Monthly and Quarterly Series,” Journal of Applied Econometrics, Vol. 18, No. 4, 427-443.

Matheson, T. D. (2011), “New Indicators for Tracking Growth in Real Time,” IMF Working Paper, WP/11/43.

Murray, L. M., A. Lee and P. E. Jacob (2013), “Rethinking Resampling in the Particle Filter on Graphics Processing Units,” January.

Rünstler, G., K. Barhoumi, S. Benk, R. Crisadoro, A. D. Reijer, A. Jakaitiene, P. Jelonek, A. Rua, K. Ruth and C. V. Nieuwenhuyze (2009), “Short-Term Forecasting of GDP Using Large Datasets: A Pseudo Real-Time Forecast Evaluation Exercise,” Journal of Forecasting, Vol. 28, Issue 7, 595-611.

Sargent, T. J. and C. A. Sims (1977), “Business Cycle Modeling Without Pretending to Have Too Much A-Priori Economic Theory,” Federal Reserve Bank of Minneapolis Working Papers, No. 55.

Schumacher, C. and J. Breitung (2008), “Real-Time Forecasting of German GDP Based on a Large Factor Model with Monthly and Quarterly Data,” International Journal of Forecasting, Vol. 24, Issue 3, 386-398.

Shumway, R. H., and D. S. Stoffer (1982), “An Approach to Time Series Smoothing and Forecasting Using the EM Algorithm,” Journal of Time Series Analysis, Vol. 3, Issue 4, 253-264.

Sonntag, M. and J. Little (2017), “Nowcasting the Economy to Develop a Better Investment Strategy,” HSBC Global Asset Management White Paper, January.

Stock, J. H. and M. W. Watson (1989), “New Indexes of Coincident and Leading Economic Indicators,” NBER Macroeconomics Annual, Vol. 4, 351-409.

Stock, J. H. and M. W. Watson (2002), “Forecasting Using Principal Components from a Large Number of Predictors,” Journal of the American Statistical Association, Vol. 97, No. 460, 1167-1179.

Stock, J. H. and M. W. Watson (2002), “Macroeconomic Forecasting Using Diffusion Indexes,” Journal of Business and Economic Statistics, Vol. 20, Issue 2, 147-162.

Watson, M. W. (2004), “[Monetary Policy in Real Time]: Comment” NBER Macroeconomics Annual, Vol. 19, 216-221.

Watson, M. W. and R. F. Engle (1983), “Alternative Algorithms for the Estimation of Dynamic Factor, Mimic and Varying Coefficient Regression Models,” Journal of Econometrics, Vol. 23, Issue 3, 385-400.

Woloszko, N. (2020), “Tracking Activity in Real Time with Google Trends,” OECD Economics Department Working Papers, No. 1634.

附錄 1：主成分分析法

在數學裡，主成分分析法可以定義為一個正交線性變換，其將原始數據序列變換到一個新的坐標系統之中，使得該組數據的某些標量投影的第 r 大方差落在第 r 個主成分之上。設 C 為原始數據集 Y 的協方差矩陣，並列出 $D = V^{-1}CV$ ，當中的 V 代表特徵向量，而 D 則是相應的對角特徵值矩陣。那麼數據集 Y 的主成分可通過以下的程序計算：

- (i) 按照特徵值的降序對 D 的列進行排序；
- (ii) 將上述排序套用至特徵向量 V 的列；
- (iii) 首 r 個主成分對應線性投影 $Y \cdot V$ 的首 r 個維度。

附錄 2：卡爾曼濾波和平滑器

在方程（1）—（2）的狀態空間表達式的框架下，假設給定過去的觀測值 $(Y_1, \dots, Y_{t-1})$ 和模型參數 θ ，那麼卡爾曼濾波可以為狀態向量 F_t 生成一個最小均方線性估計量，以及一個均方誤差矩陣 $P_{t|s}$ ：

$$F_{t|s} = \mathbb{E}[F_t | Y_1, \dots, Y_s; \theta], \quad P_{t|s} = \text{Var}(F_{t|s} - F_t | Y_1, \dots, Y_s; \theta)。$$

每當接收新信息的時候，卡爾曼濾波將通過下列的遞迴機制對先驗估計作出更新：

- (i) 新息殘差： $v_t = Y_t - \Lambda F_{t|t-1}$ ；
- (ii) 新息協方差： $G_t = \Lambda P_{t|t-1} \Lambda' + \Sigma_E$ ；
- (iii) 最優“卡爾曼收益”： $K_t = \Phi P_{t|t-1} \Lambda'$ ；
- (iv) 後驗狀態估計： $\hat{F}_{t+1|t} = \Phi F_{t|t-1} + K_t G_t^{-1} v_t$ ；
- (v) 後驗協方差估計： $\hat{P}_{t+1|t} = \Phi P_{t|t-1} \Phi' - K_t G_t^{-1} K_t' + \Sigma_U$ 。

另一方面，卡爾曼平滑器定義為卡爾曼濾波的反向傳導，估計量可使用類似的方法計算，惟相關計算需要進行反向遞迴：

- (i) “移除條件”： $L_t = G_{t|t} \Lambda' G_{t+1|t}^{-1}$ ；
- (ii) 平滑化的狀態估計： $\hat{F}_{t|T} = \hat{F}_{t|t} + L_t (\hat{F}_{t+1|T} - \hat{F}_{t+1|t})$ ；
- (iii) 平滑化的協方差估計： $P_{t|T} = P_{t|t} + L_t (P_{t+1|T} - P_{t+1|t}) L_t'$ 。

附錄 3：最大期望演算法

最大期望演算法的概念在於，估算部分不完整的數據並列出似然函數，以及迭代兩個步驟：在求期望值的步驟，我們根據數據和以上次迭代得出的參數作為條件，計算對數似然函數的期望值，而在求最大化的步驟，我們對期望值進行再優化。在某些閾值或調控條件下，最大期望演算法將收斂至一個（局部）最大值。

求期望值 – 遵從高斯分佈的對數似然函數 $l(Y, F; \theta)$ 的表達式如下：

$$l(Y, F; \theta) = -\frac{TN}{2} \log(2\pi) - \frac{1}{2} \sum_{t=1}^T \log |G_t| - \frac{1}{2} \sum_{t=1}^T v_t G_t^{-1} v_t'.$$

求最大化 – Watson 和 Engle (1983) 推導得出方程 (1) — (2) 中的模型參數的最大似然估計可以通過以下迭代程序計算：

$$\begin{aligned} \Lambda(j+1) &= \left(\sum_{t=1}^T \mathbb{E}_{\theta(j)} [Y_t F_t' | \Omega_T] \right) \left(\sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_t F_t' | \Omega_T] \right)^{-1}, \\ \Phi(j+1) &= \left(\sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_t F_t' | \Omega_T] \right) \left(\sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_{t-1} F_{t-1}' | \Omega_T] \right)^{-1}, \\ \Sigma_E(j+1) &= \text{diag} \left(\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\theta(j)} [Y_t Y_t' | \Omega_T] - \Lambda(j+1) \sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_t Y_t' | \Omega_T] \right), \\ \Sigma_U(j+1) &= \left(\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_t F_t' | \Omega_T] - \Phi(j+1) \sum_{t=1}^T \mathbb{E}_{\theta(j)} [F_{t-1} F_t' | \Omega_T] \right). \end{aligned}$$

當中涉及因子的條件矩的部分可使用卡爾曼濾波和平滑器計算 (Shumway 和 Stoffer 1982)。最後，關於最大期望演算法的一個補強版本可參考 (Bańbura *et al.* 2013)，該版本適用於任意數據缺失和特異部分存在自相關性的情況。