

Final report(QM) deadline June 23.:

part 1: page 611, Chapter 13, Exercise 13.13, data file: **equity**,
Questions a, b, and c need to be explained b needs to do an indigenous test.

Part 2: choose a country, collect data for 40 years from the database, after data collection, calculate the same as part 1, and answer as in part 1.
variable: dividend, composite stock index, nominal dividend per share.

題目: 看p.609頁最下面 那個13.13 開始, 一直到p.612開頭那a,b,c三個問題止.

13.13 Financial analysts often debate the role of dividends (DV) in the determination of share prices (SP). Figure 13.9 shows plots of the rate of change in DV and SP computed as

$$DV_t = 100 \ln(DN_t/DN_{t-1}), \quad SP_t = 100 \ln(PN_t/PN_{t-1})$$

$$DV_t = 100 \ln(DN_t/DN_{t-1}), \quad SP_t = 100 \ln(PN_t/PN_{t-1})$$

where PN is the Standard and Poor Composite Price Index; DN is the nominal dividend per share (source: Prescott, E. C. and Mehra, R. "[The Equity Premium: A Puzzle](#)," Journal of Monetary Economics, 15 March, 1985, pp. 145–161). The data are annual observations over the period 1889–1979.

這問題是來自 以下這篇 經典的論文:

THE EQUITY PREMIUM A Puzzle* 股權溢價之謎

145 1. Introduction

147 2. Data

150 3. The economy, asset prices and returns

153 4. The results

156 4.1. Robustness of results

157 4.2. Effects of firm leverage

158 4.3. Introducing production

158 5. Conclusion

159 Appendix

160 References

[Gemini說 好的. 同學](#)。這篇論文《*The Equity Premium: A Puzzle*》由 **Rajnish Mehra** 和 **Edward C. Prescott** 撰寫, 發表於 1985 年的《*Journal of Monetary Economics*》, 是計量經

濟學和金融學領域的經典之作，它提出了著名的「股權溢價之謎」*Equity Premium Puzzle*。我將作為你的計量經濟學老師，為你深入解析這篇論文的核心思想、所用的模型、公式以及其重要性。

論文核心思想

這篇論文的核心是指出一個在當時的(Ask1:)標準一般均衡模型(特別是抽象了交易成本、流動性約束等摩擦的 Arrow-Debreu 框架)下無法解釋的現象：歷史上股票的平均收益率遠高於短期無風險債券的平均收益率。

具體來說，在 1889 年至 1978 年的 90 年間，美國股票(S&P 500 指數)的平均實際年收益率為 7%，而短期債務的平均實際年收益率不到 1%。這意味著存在一個高達 6% 的股權溢價(Equity Premium)。論文發現，即使是那些被構建為模仿美國經濟消費增長率的均值、方差和序列相關性的模型，也只能產生最高 0.4% 的股權溢價，這與實際觀測到的 6% 存在巨大差異。

作者們將這種現象稱為「股權溢價之謎」，並得出結論：最有可能的情況是，一個包含摩擦(例如交易成本、流動性約束等)的均衡模型，而不是理想化的 Arrow-Debreu 經濟模型，才能同時合理解釋歷史上觀察到的高平均股票收益和低平均無風險收益。

模型設定與關鍵公式

論文使用了一個基於 Lucas (1978) 純交換模型的變體，並假設(Ask2:)人均消費的增長率服從馬爾可夫過程，而非 Lucas 模型中消費水平服從馬爾可夫過程的假設。這使得模型能夠捕捉 1889-1978 年間人均消費大幅增長所帶來的非平穩性。

1. 消費者偏好 (Utility Function)

模型中，有一個代表性的「標準」家庭，其效用函數是時間可加的 (time-additive) 形式：

$$E_0\left\{\sum_{i=0}^{\infty}\beta^i U(C_t)\right\}, 0 < \beta < 1$$

其中：

- c_t 是 t 時期的人均消費。
- β 是主觀時間貼現因子 (subjective time discount factor)，表示人們對未來消費的耐心程度。
- $E_0\{\cdot\}$ 是在時間零點(當前)可用信息下的期望運算符。
- $U(\cdot)$ 是遞增的凹效用函數。

為了確保均衡收益過程是平穩的，效用函數進一步被限制為恆定相對風險厭惡 (Constant Relative Risk Aversion, CRRA) 類別：

$$U(c, \alpha) = \frac{c^{1-\alpha} - 1}{1-\alpha}$$

$$U(c, \alpha) = \frac{c^{1-\alpha} - 1}{1-\alpha}$$

其中：

- α 是相對風險厭惡係數 (**coefficient of relative risk aversion**), 衡量效用函數的曲率。它反映了人們在不同時間點之間替代消費的意願。當 $\alpha=1$ 時, 效用函數是對數形式 $U(c)=\ln(c)$ 。

2. 生產與消費增長 (Production and Consumption Growth)

模型假設只有一個生產單位生產易腐爛的消費品。該公司的產出 y_t 在 t 時期同時是公司的股息支付。產出 y_t 的增長率服從馬爾可夫鏈：

$$y_{t+1} = x_t + y_t$$

$$y_{t+1} = x_t + y_t$$

其中 $x_{t+1} \in \{\lambda_1, \dots, \lambda_n\}$ 是增長率, 且滿足馬爾可夫鏈的轉移概率：

$$P_r\{x_{t+1} = \lambda_j; x_t = \lambda_i\} = \phi_{ij}$$

$$P_r\{x_{t+1} = \lambda_j; x_t = \lambda_i\} = \phi_{ij}$$

ϕ_{ij} 表示從狀態 i 轉移到狀態 j 的概率。

3. 資產定價 (Asset Pricing)

論文推導了股票和無風險債券的均衡價格表達式。在均衡狀態下, 消費過程就是產出過程 $\{y_s\}$ 。

股票價格 (P_t^e) : 對於任何支付過程為 $\{d_s\}$ 的證券, 其在 t 時期的價格為：

$$P_t = E_t\{\sum_{s=t+1}^{\infty} \beta^{s-t} U^t(y_s) d_s / U^t(y_t)\}$$

$$P_t = E_t\left\{\sum_{s=t+1}^{\infty} \beta^{s-t} U^t(y_s) d_s / U^t(y_t)\right\}$$

由於股票的股息支付過程是 $\{y_s\}$ 且 $U'(c) = c^{-\alpha}$, 股票的價格可以寫為：

$$P_t^e = P^e(x_t, y_t) = E\left\{\sum_{s=t+1}^{\infty} \beta^{s-t} \frac{y_t^\alpha}{y_s^\alpha} y_s | x_t, y_t\right\}$$

$$P_t^e = P^e(x_t, y_t) = E\left\{\sum_{s=t+1}^{\infty} \beta^{s-t} \frac{y_t^\alpha}{y_s^\alpha} y_s | x_t, y_t\right\}$$

由於 $y_s = y_t \cdot x_{t+1} \cdot \dots \cdot x_s$ ，並且股票價格對 y_t 呈一次齊次性，即 $p^e(c, i) = w_i c$ ，最終得到 $w_i c$ 的線性方程組：

$$w_i = \beta \sum_{j=1}^n \phi_{ij} \lambda_j^{(1-\alpha)} (w_j + 1) \quad \text{for } i = 1, \dots, n$$

$$w_i = \beta \sum_{j=1}^n \phi_{ij} \lambda_j^{(1-\alpha)} (w_j + 1) \quad \text{for } i = 1, \dots, n$$

股票收益 (r_{ij}^e):

如果當前狀態為 (c, i) ，下一個狀態為 $(\lambda_j c, j)$ ，則股票的當期收益為：

$$r_{ij}^e = \frac{p^e(\lambda_j c, j) + \lambda_j c - p^e(c, i)}{p^e(c, i)} = \frac{\lambda_j (w_j + 1)}{w_i} - 1$$

$$r_{ij}^e = \frac{p^e(\lambda_j c, j) + \lambda_j c - p^e(c, i)}{p^e(c, i)} = \frac{\lambda_j (w_j + 1)}{w_i} - 1$$

當前狀態為 i 時的預期股票收益為：

$$R_i^e = \sum_{j=1}^n \phi_{ij} r_{ij}^e$$

$$R_i^e = \sum_{j=1}^n \phi_{ij} r_{ij}^e$$

無風險債券價格 (p_i^f):

無風險證券是一期實際票據，下一期以確定性地支付一個單位的消費品。其價格為：

$$p_i^f = p^f(c, i) = \beta \sum_{j=1}^n \phi_{ij} U'(\lambda_j c) / U'(c) = \beta \sum_{j=1}^n \phi_{ij} \lambda_j^{-\alpha}$$

$$p_i^f = p^f(c, i) = \beta \sum_{j=1}^n \phi_{ij} U'(\lambda_j c) / U'(c) = \beta \sum_{j=1}^n \phi_{ij} \lambda_j^{-\alpha}$$

無風險債券收益 (R_i^f):

當前狀態為 (c, i) 時，該無風險證券的確定性收益為：

$$(p_i^f) = 1/p_i^f - 1$$

$$(p_i^f) = 1/p_i^f - 1$$

長期平均收益 (**Expected Returns in Stationary Distribution**): 在馬爾可夫鏈的平穩分佈 $\pi \in R^n$ 下, 股票和無風險證券的預期收益分別為:

$$R^e = \sum_{i=1}^n \pi_i R_i^e \text{ and } R^f = \sum_{i=1}^n \pi_i R_i^f$$

$$R^e = \sum_{i=1}^n \pi_i R_i^e \text{ and } R^f = \sum_{i=1}^n \pi_i R_i^f$$

其中 π 是馬爾可夫鏈的平穩概率向量, 滿足 $\pi = \phi^T$ 且 $\sum \pi_i = 1$ 。時間樣本平均值將依概率收斂到這些值。

風險溢價 (Risk Premium):

股票的風險溢價定義為 $R^e - R^f$ 。

「股權溢價之謎」的定量分析

論文的核心實驗是將模型參數與美國經濟的歷史數據進行匹配。他們假設馬爾可夫鏈有兩個狀態, 並設定增長率參數為 $\lambda_1 = 1 + \mu + \delta$ 和 $\lambda_2 = 1 + \mu - \delta$, 轉移概率 $\phi_{11} = \phi_{22} = \phi$ 和 $\phi_{12} = \phi_{21} = (1 - \phi)$ 。

他們將模型的消費增長率的平均值、標準差和一階序列相關性與 1889-1978 年美國經濟的樣本值(分別為 0.018、0.036 和 -0.14)匹配。這導致了模型中的參數值為 $\mu=0.018$ 、 $\delta=0.036$ 和 $\phi=0.43$ 。

然後, 他們在 α (相對風險厭惡係數)在 0 到 10 之間, β (主觀時間貼現因子)在 0 到 1 之間變化時, 計算了模型預測的平均無風險利率和股票風險溢價。之所以將 α 限制在最大值 10, 是因為在許多經濟學研究中, 對 α 的估計通常在 1 到 2 之間, 甚至接近 0 或 55 (Fisher Black 的例子)。

結果顯示:

- 歷史上觀察到的平均實際無風險收益率為 0.80%。
- 歷史上觀察到的平均股票收益率為 6.98%。
- 歷史上觀察到的平均股權溢價為 6.18%。

然而, 模型所能產生的最大股權溢價僅為 **0.35%**。這與觀測到的 6% 左右的溢價存在巨大差距, 這就是所謂的「股權溢價之謎」。

穩健性檢驗 (Robustness of Results)

論文還檢驗了其結果對於模型設定和測量問題的穩健性:

- 通貨膨脹測量誤差: 這類誤差會同時影響實際無風險利率和股票收益, 因此對計算出的風險溢價影響不大。

- 稅收考慮：理論上考慮的是稅後收益。儘管稅率在不同時期和收入階層有所不同，但在研究期間，稅後收益的低無風險利率和可觀的股權溢價仍然成立。
- 數據頻率：將模型時間期從 1/100 年到 2 年不等，對 admissible region 幾乎沒有影響，表明結果對使用年度數據估計消費過程是穩健的。
- 消費增長率參數變動：改變平均增長率 μ 或引入更高的矩，對結果的影響微乎其微。股權溢價大約與消費增長率的標準差 δ 的平方成正比。
- 公司槓桿：即使考慮了公司將大部分預期產出承諾給非股東（例如工資），導致股東承擔大部分風險，股權溢價也僅增加了不到 0.1%。這是因為金融安排不影響資源配置和潛在的 Arrow-Debreu 價格。
- 引入生產：作者指出，在他們的模型結構下，消費過程是外生的，沒有資本積累和生產。但他們認為，擴展技術集來引入這些機會並不能推翻他們的結論，因為這不會增加消費和資產價格的聯合均衡過程的集合。

結論與未來研究方向

論文得出結論，這個「股權溢價之謎」可能並非是為什麼平均股票收益率如此之高，而是為什麼平均無風險利率如此之低。如果接受 Friend 和 Blume (1975) 關於曲率參數 α 顯著大於 1 的發現（例如 $\alpha=2$ ），模型預測的平均無風險利率至少為 3.7%，遠高於樣本平均值 0.80%。

作者們懷疑，引入異質性代理人本身並不能解決這個難題，因為在 Debreu 的競爭框架內，異質性代理人經濟也施加了類似的限制。他們也懷疑非時間可加可分 (non-time-additivity separable) 的偏好能夠解決這個謎題，因為這需要不同時間點的消費之間替代性更差。

他們最終認為，解決這個謎題可能需要引入某些使得特定類型跨期交易在代理人之間不可行的特徵。例如，某些合約可能無法執行，這導致了本來可以出現的風險分擔市場不復存在。此外，與尚未出生的世代簽訂合約也是不可行的。作者們認為，這類非 **Arrow-Debreu** 競爭均衡模型可能才能合理解釋美國經濟在過去 90 年中表現出的巨大股權風險溢價。為了檢驗這些理論，可能需要按收入或年齡組劃分的消費數據。

對計量經濟學的啟示

這篇論文雖然沒有直接進行複雜的計量估計，但它通過一個**「定量理論」*(quantitative theoretical exercise) 的方法，清晰地揭示了標準理論與實際數據之間的巨大鴻溝。這對計量經濟學研究有深刻啟示：

1. 模型的驗證與限制：論文展示了即使是最精緻的、基於合理校準參數的一般均衡模型，在解釋某些宏觀金融現象時也存在根本性的缺陷。這促使研究者重新思考理論模型的基本假設。
2. 理論與實證的互動：「股權溢價之謎」是一個典型的理論無法解釋實證現象的例子。它推動了金融經濟學和宏觀經濟學中對行為經濟學、市場摩擦、不完全信息、流動性約束等因素的深入研究。
3. 參數選擇的重要性：論文強調了相對風險厭惡係數 α 的重要性。不同的 α 值會對模型的預測產生巨大影響。這引發了關於如何更準確估計這些關鍵參數的討論，並認識到它們不僅是統計上的數字，更是代表了經濟主體的深層偏好。
4. 穩健性分析的價值：論文對各種可能影響結果的因素進行了詳細的穩健性檢驗。這強調了在進行任何實證或理論分析時，對結果敏感性進行全面評估的重要性。

總之,《The Equity Premium: A Puzzle》這篇論文不僅指出了金融市場的一個核心未解之謎,更重要的是,它激發了後續數十年大量的研究,試圖通過引入更豐富的經濟學機制來解決這個難題,極大地推動了資產定價理論和宏觀金融的發展。作為計量經濟學的學生,理解這篇論文對於把握經濟模型構建、驗證以及理論與實證互動的精髓至關重要。

2025/0531

Ask1: 請你舉個簡單的例子,說明「標準一般均衡模型」是什麼意思。

好的,我們來用一個簡單的例子來理解「標準一般均衡模型」是什麼意思。

想像一下,在一個非常理想化的經濟中,只有兩種商品:蘋果和香蕉,以及兩種人:農民和果農。

標準一般均衡模型(Standard General Equilibrium Model)會包含以下幾個核心要素:

1. 理性行為人 (Rational Agents):

- 消費者 (農民和果農): 他們都是理性的,會根據自己的偏好(比如喜歡蘋果多一點還是香蕉多一點)來最大化他們的效用(滿意度)。他們會考慮商品的價格來決定買多少蘋果和多少香蕉。
- 生產者 (種植蘋果的農民和種植香蕉的果農): 他們也是理性的,會努力以最低的成本生產他們的產品,並在市場上銷售以最大化利潤。

2. 市場清算 (Market Clearing):

- 在這個經濟中,會有一個蘋果市場和一個香蕉市場。
- 市場清算的意思是,在均衡價格下,所有生產出來的蘋果都會被消費者買走,沒有剩餘;同樣,所有生產出來的香蕉也都會被消費者買走。也就是說,供給等於需求。

3. 完美競爭 (Perfect Competition):

- 沒有任何一個單獨的農民、果農或消費者有能力影響市場價格。他們都是價格的接受者。
- 所有信息都是公開透明的,大家都知道蘋果和香蕉的質量、價格等信息。

4. 無摩擦 (No Frictions):

- 這是「標準一般均衡模型」的一個關鍵特點,也是 Mehra 和 Prescott 論文中強調的。
- 沒有交易成本: 買賣蘋果和香蕉不需要支付任何費用。
- 沒有流動性約束: 如果你想買蘋果,你總能買到;你想賣香蕉,總能賣掉,不會因為現金不足或市場沒有買家而受限。
- 沒有信息不對稱: 每個人都知道同樣的、完整的市場信息。
- 沒有政府干預或稅收: 市場完全自由運作。

這個例子說明了什麼?

「標準一般均衡模型」試圖描述一個理想化的經濟體,其中所有市場同時達到均衡,並且資源被有效率地配置。在這種模型下,所有經濟主體都做出最優決策,並且不存在任何阻礙市場有效運作的「摩擦」。

Mehra 和 Prescott 的論文就是在這個「無摩擦」的理想化框架下,發現模型預測的股權溢價與實際觀測到的結果存在巨大差異。他們指出,這意味著現實世界中的一些「摩擦」(例如交

易成本、流動性約束等)可能是解釋「股權溢價之謎」的關鍵。

ASK2: 請問, 你說「人均消費的增長率服從馬爾可夫過程」是什麼意思。

好的, 同學。我們來解釋一下「人均消費的增長率服從馬爾可夫過程」這句話是什麼意思。

首先, 我們將這句話拆解成幾個部分來理解:

1. 人均消費 (**Per Capita Consumption**): 這指的是每個人平均的消費量。在經濟學模型中, 它通常被視為衡量一個經濟體富裕程度和福利水平的重要指標。
2. 增長率 (**Growth Rate**): 這表示人均消費量隨時間變化的百分比。例如, 如果今年人均消費是 100 元, 明年是 102 元, 那麼增長率就是 2%。
3. 馬爾可夫過程 (**Markov Process**): 這是統計學和概率論中的一個概念, 指的是一個隨機過程 (random process), 它在給定當前狀態的情況下, 未來狀態的條件概率分佈只與當前狀態有關, 而與過去的狀態 (即從歷史路徑) 無關。簡單來說, 就是「未來只取決於現在, 與過去無關」。

將這些概念結合起來, 「人均消費的增長率服從馬爾可夫過程」的意思就是:

- 我們關注的不是人均消費的絕對水平, 而是它在每年 (或每期) 的增長速度。
- 這個增長率不是一個確定的常數, 而是一個隨機變量, 它會根據一定的概率規則在不同的「狀態」之間跳躍。
- 最重要的是, 未來時期的人均消費增長率會是多少, 只取決於當前時期的人均消費增長率處於哪種狀態, 而與之前更久遠的歷史增長率表現沒有直接關係。

舉個例子來幫助理解:

假設人均消費的增長率只有兩種可能的狀態:

- 狀態 1: 高增長率 (例如 5%)
- 狀態 2: 低增長率 (例如 1%)

如果我們說這個增長率服從馬爾可夫過程, 這意味著:

- 當前如果處於「高增長率」狀態 (5%), 那麼下一年是繼續保持「高增長率」還是轉變為「低增長率」, 這是有一定的概率的 (例如, 有 80% 的概率繼續高增長, 20% 的概率轉為低增長)。
- 當前如果處於「低增長率」狀態 (1%), 那麼下一年是繼續保持「低增長率」還是轉變為「高增長率」, 這也有一定的概率的 (例如, 有 70% 的概率繼續低增長, 30% 的概率轉為高增長)。

關鍵點在於: 下一年的增長率只看你今年是什麼增長率狀態, 而不管去年、前年的增長率是多少。比如, 如果今年是高增長, 那麼明天是高增長還是低增長的概率, 和你是不是已經連續三年高增長, 是沒有關係的。這就是馬爾可夫過程的「無記憶性」特徵。

在 Mehra 和 Prescott 的論文中, 他們就使用了這個假設。他們將人均消費的增長率 x_{t+1} 設定為在幾個離散的 λ_j 值之間跳躍, 並通過轉移概率矩陣 ϕ_{ij} 來描述這些跳躍的概率。這種設定允許模型捕捉消費增長率的隨機性, 並且能夠更好地匹配美國經濟在 1889-1978 年

間人均消費大幅增長所帶來的非平穩性特徵，這與 Lucas 模型中直接假設「消費水平」服從馬爾可夫過程有所不同。

參考上傳的poe_chp13.pdf檔，請看page 610 的問題：13.13 Financial analysts often debate the role of dividends (DV) in the determination of share prices (SP).

好的，這是一個關於股票價格和股息之間動態關係的計量經濟學問題，涉及 VAR (Vector Autoregressive) 模型和 ARDL (Autoregressive Distributed Lag) 模型。

問題 13.13 解釋

這個問題要求你分析股票價格變化率 (SP) 和股息變化率 (DV) 之間的關係。數據是從 Mehra 和 Prescott (1985) 的經典論文中提取的年度數據，涵蓋 1889-1979 年期間。

背景信息：

- **DV**: 股息變化率，計算為 $100\ln(DN_t/DN_{t-1})$ ，其中 DN_t 是每股名義股息。
- **SP**: 股價變化率，計算為 $100\ln(PN_t/PN_{t-1})$ ，其中 PN_t 是標準普爾綜合股價指數。
- Figure 13.9 顯示了這兩個系列的時間序列圖，看起來它們都是 $I(0)$ 變量，即它們是平穩的（因為是變化率）。

問題主要有三個部分：

1. 估計 **VAR(1)** 模型： $SP_t = \beta_{10} + \beta_{11}SP_{t-1} + \beta_{12}DV_{t-1} + v_{1t}$
 $DV_t = \beta_{20} + \beta_{21}SP_{t-1} + \beta_{22}DV_{t-1} + v_{2t}$ 你需要使用最小二乘法 (OLS) 分別估計這兩個方程。這是一個標準的 VAR(1) 模型，其中每個變量都由其自身的一階滯後和另一個變量的一階滯後來解釋。由於 SP 和 DV 是平穩的 $I(0)$ 變量，直接對水平值應用 OLS 是合適的。
2. 估計 **ARDL** 模型： $SP_t = \alpha_{10} + \alpha_{11}SP_{t-1} + \alpha_{12}DV_{t-1} + \alpha_{13}DV_t + e_{1t}$
 $DV_t = \alpha_{20} + \alpha_{21}SP_{t-1} + \alpha_{22}DV_{t-1} + \alpha_{23}SP_t + e_{2t}$ 同樣，你需要使用 OLS 分別估計這兩個方程。這裡的 ARDL 模型與 VAR(1) 的主要區別在於，它引入了當期 (**contemporaneous**) 的另一個變量作為解釋變量 ($\alpha_{13}DV_t$ 在第一個方程中， $\alpha_{23}SP_t$ 在第二個方程中)。
3. 比較和解釋：
 - a. 解釋為什麼對 **VAR** 模型（只有滯後變量在右側）使用 **OLS** 估計會得到一致估計量。
 - b. 解釋為什麼對 **ARDL** 模型（既有滯後變量也有當期變量在右側）使用 **OLS** 估計會得到不一致估計量。你可能需要參考第 11 章的材料。

- c. 從你的估計結果中推斷出股息在股價決定中的作用。

如何建議解答

第 1 步：數據準備和初步分析

1. 獲取數據：從提供的文件 `equity.csv` 中載入 `SP` 和 `DV` 兩個時間序列數據。
2. 可視化：繪製 `SP` 和 `DV` 的時間序列圖 (如 Figure 13.9)。目測判斷它們是否平穩。由於題目明確說明它們是變化率，且圖形也支持平穩性，因此可以直接假設它們是 $I(0)$ 變量。
3. 檢查序列相關性：雖然題目沒有明確要求，但通常在估計 VAR/ARDL 之前，檢查每個變量的自相關和偏自相關函數 (ACF/PACF) 可以幫助判斷合適的滯後階數。不過，這裡題目已經指定了模型形式。

第 2 步：估計 VAR(1) 模型

對於以下兩個方程，使用 OLS 進行估計：

1. `SP_t` 作為被解釋變量，`SP_{t-1}` 和 `DV_{t-1}` 作為解釋變量。
2. `DV_t` 作為被解釋變量，`SP_{t-1}` 和 `DV_{t-1}` 作為解釋變量。

操作步驟：

- 創建滯後變量 `SP_{t-1}` 和 `DV_{t-1}`。
- 使用統計軟體 (如 R, Python, EViews, Stata) 的 OLS 功能分別運行這兩個回歸。
- 記錄估計的係數 ($\beta^{10}, \beta^{11}, \beta^{12}, \beta^{20}, \beta^{21}, \beta^{22}$)、標準誤、t 值和 R-squared 等。

第 3 步：估計 ARDL 模型

對於以下兩個方程，使用 OLS 進行估計：

1. `SP_t` 作為被解釋變量，`SP_{t-1}`、`DV_{t-1}` 和 `DV_t` 作為解釋變量。
2. `DV_t` 作為被解釋變量，`SP_{t-1}`、`DV_{t-1}` 和 `SP_t` 作為解釋變量。

操作步驟：

- 創建所有必要的滯後和當期變量。
- 使用統計軟體運行 OLS 回歸。
- 記錄估計的係數 ($\alpha^{10}, \alpha^{11}, \alpha^{12}, \alpha^{13}, \alpha^{20}, \alpha^{21}, \alpha^{22}, \alpha^{23}$)、標準誤、t 值和 R-squared 等。

第 4 步：解釋與回答問題 a, b, c

a. 解釋為什麼對 VAR 模型使用 OLS 估計會得到一致估計量。

- 關鍵點：在 VAR 模型中，右側的解釋變量只有滯後變量。
- 解釋：
 - 一致性要求解釋變量與誤差項不相關。

- 在 VAR 模型中，儘管當期誤差項 v_{ts} 或 v_{td} 可能會影響當期 SP_t 和 DV_t ，但它們與過去的 SP_{t-1} 和 DV_{t-1} 是不相關的。這是因為誤差項 v_t 代表了在給定過去信息後當期無法預測的部分（即「創新」）。
- 由於滯後變量是前一期的數據，它們在當期誤差項 v_t 實現之前就已經確定了。因此，它們是外生的，滿足 OLS 估計量一致性的要求（即 $E(X'e)=0$ ，其中 X 是解釋變量矩陣）。
- 因此，對 VAR 模型中的每個方程獨立使用 OLS 會得到一致估計量。

b. 解釋為什麼對 ARDL 模型使用 OLS 估計會得到不一致估計量。

- 關鍵點：ARDL 模型在右側引入了當期內生變量。
- 解釋：
 - 以第一個 ARDL 方程為例： $SP_t = \alpha_{10} + \alpha_{11}SP_{t-1} + \alpha_{12}DV_{t-1} + \alpha_{13}DV_t + e_{ts}$ 。
 - 問題出現在 $\alpha_{13}DV_t$ 這個項。由於 SP_t 和 DV_t 在經濟中是同時決定的（它們之間可能存在雙向因果關係，或者共同受到某些未觀測因素的影響）， DV_t 很可能與 e_{ts} （ SP_t 方程的誤差項）相關。
 - 如果 DV_t 受到 e_{ts} 的影響，或者 e_{ts} 中包含的任何未觀測變量同時影響 DV_t ，那麼 $E(DV_t \cdot e_{ts}) \neq 0$ ，這就違反了 OLS 的一個關鍵假設——解釋變量與誤差項不相關。
 - 這種情況被稱為同期內生性（contemporaneous endogeneity）或聯立方程偏誤（simultaneous equations bias）。
 - 當解釋變量與誤差項相關時，OLS 估計量將是不一致的（inconsistent）。這意味著即使在樣本量無限大的情況下，OLS 估計量也不會收斂到真實參數值。
 - 解決這個問題通常需要使用工具變量（Instrumental Variables, IV）或廣義矩估計（Generalized Method of Moments, GMM）等方法。

c. 從你的估計結果中推斷出股息在股價決定中的作用。

- 綜合分析兩個模型的結果。
- VAR 模型：
 - 查看 β^{12} 的係數和顯著性。如果它顯著不為零，說明滯後的股息變化率（ DV_{t-1} ）對當期股價變化率（ SP_t ）有影響。
 - 查看 β^{21} 的係數和顯著性。如果它顯著不為零，說明滯後的股價變化率（ SP_{t-1} ）對當期股息變化率（ DV_t ）有影響。
 - 這些係數反映了動態的、滯後的因果關係。
- ARDL 模型：
 - 儘管 ARDL 模型的 OLS 估計量不一致，但它們仍然可以提供一些直觀的信息（儘管不能作為嚴格的因果推斷）。
 - 特別關注 α^{13} 的係數和顯著性。如果它顯著不為零，這似乎表明當期股息變化率（ DV_t ）對當期股價變化率（ SP_t ）有直接影響。
 - 然而，由於估計偏誤的存在，對這個係數的解釋需要非常謹慎。
- 總結：
 - 基於 VAR 模型，你可以討論股價和股息之間的滯後動態相互作用。例如，過去的股息表現是否會影響當前的股價走勢。
 - 基於 ARDL 模型（並結合對其估計偏誤的理解），你可以討論是否存在當期聯動關係，但要強調由於內生性問題，OLS 估計結果的解釋局限性。
 - 金融分析師的爭論通常涉及到股息是否是股價的一個領先指標（滯後效應）或者是否與股價同時變動（當期效應），以及兩者之間是否存在因果關係。這個問題

題旨在讓你通過模型的實證結果來思考和回答這些問題，並理解模型選擇的重要性。

重要提示：在實際操作中，處理同期內生性問題是計量經濟學的關鍵挑戰。這道題的設計正是為了讓你理解在存在內生性時，OLS 估計的局限性。

我沒有 **equity.csv** 但我有 **EViews** 格式的 **equity.wf1** 檔案，請問，如果我上傳這個檔案，你可以一步步指導我，怎樣運用 **EViews 8.1** 來繪製序列圖，或求算所需要的數據，完成答案嗎？

很抱歉，我無法直接處理您上傳的 **EViews .wf1** 檔案，也無法直接與 **EViews** 軟體互動來執行繪圖或數據分析操作。我的能力是基於文字和預訓練的模型，無法像人類一樣操作桌面應用程式。

不過，我可以一步一步地指導您在 **EViews 8.1** 中如何操作，來完成問題 13.13 所需的繪圖和數據計算。您需要自己在 **EViews** 中執行這些步驟。

以下是針對問題 13.13 的詳細 **EViews** 操作指南，假設您已經成功載入 **equity.wf1** 檔案並打開了工作文件：

EViews 8.1 操作指南：完成問題 13.13

第 0 步：確認數據載入與工作文件設置

1. 打開 **EViews 8.1**。
2. 打開 **equity.wf1** 檔案：
 - 點擊 **EViews** 菜單欄的 **File -> Open -> EViews Workfile...**。
 - 瀏覽並選擇您的 **equity.wf1** 檔案，然後點擊 **Open**。
 - 您應該會在工作文件窗口中看到 **SP** 和 **DV** 這兩個序列。
3. 檢查序列名稱：確保序列名稱就是 **SP** 和 **DV**。如果不是，您可能需要重命名它們 (**Right-click** on series -> **Rename**)。

第 1 步：繪製序列圖 (Figure 13.9)

這是為了直觀地了解 **SP** 和 **DV** 的時間序列特性，通常也是時間序列分析的第一步。

1. 選取序列：
 - 在工作文件窗口中，按住 **Ctrl** 鍵，然後點擊 **SP** 和 **DV**，使它們都被選中。
2. 繪製圖形：
 - **Right-click** on any of the selected series -> **Open Selected -> as Group**。
 - 在打開的 **Group** 窗口中，點擊 **View -> Graph -> Line Graph**。

- EViews 會自動繪製這兩個序列的時間序列圖。這應該類似於您 PDF 文件中的 Figure 13.9。

第 2 步: 估計 VAR(1) 模型

在 EViews 中, VAR 模型可以通過菜單選項直接估計。

1. 打開 **VAR** 估計窗口:
 - 點擊 EViews 菜單欄的 **Quick -> Estimate VAR...**。
2. 設置 **VAR** 規格:
 - 在 **VAR Specification** 選項卡中:
 - **VAR Type**: 選擇 **Unrestricted VAR**。
 - **Endogenous Variables**: 輸入 **SP DV** (中間用空格隔開)。
 - **Exogenous Variables**: 輸入 **C** (代表常數項)。
 - **Lag Intervals for Endogenous**: 輸入 **1 1**。這表示您要包含從滯後 1 期到滯後 1 期 (即只有一階滯後) 的變量。
 - 點擊 **OK**。
3. 解讀 **VAR** 估計結果:
 - EViews 會顯示 VAR 估計結果。您會看到兩個方程的估計係數、標準誤、t-統計量、p-值等。
 - 對於 **SP** 方程, 尋找 **SP(-1)** 和 **DV(-1)** 的係數 (即 β_{11} 和 β_{12} 的估計值)。
 - 對於 **DV** 方程, 尋找 **SP(-1)** 和 **DV(-1)** 的係數 (即 β_{21} 和 β_{22} 的估計值)。
 - 複製/記錄這些係數、標準誤和 t-值, 以備回答問題 **c**。

第 3 步: 估計 ARDL 模型

由於 ARDL 模型在 EViews 中沒有直接的「ARDL 估計」菜單選項來處理多個方程中的同期變量, 我們將其視為兩個獨立的單一方程 **OLS** 回歸。

首先, 估計 **SP** 方程: $SP_t = \alpha_{10} + \alpha_{11}SP_{t-1} + \alpha_{12}DV_{t-1} + \alpha_{13}DV_t + \epsilon_t$

1. 打開方程估計窗口:
 - 點擊 EViews 菜單欄的 **Quick -> Estimate Equation...**。
2. 輸入方程規格:
 - 在 **Equation Specification** 框中, 輸入以下內容: **SP C SP(-1) DV(-1) DV**
 - **Method**: 選擇 **LS - Least Squares (ARMA)**。
 - 點擊 **OK**。
3. 解讀 **SP** 方程的 **ARDL** 結果:
 - EViews 會顯示估計結果。
 - 尋找 **C** (α_{10})、**SP(-1)** (α_{11})、**DV(-1)** (α_{12}) 和 **DV** (α_{13}) 的係數、標準誤和 t-值。
 - 複製/記錄這些數據。

其次, 估計 **DV** 方程: $DV_t = \alpha_{20} + \alpha_{21}SP_{t-1} + \alpha_{22}DV_{t-1} + \alpha_{23}SP_t + \epsilon_t$

1. 打開另一個方程估計窗口:

- 再次點擊 EViews 菜單欄的 Quick -> Estimate Equation...。
- 2. 輸入方程規格：
 - 在 Equation Specification 框中，輸入以下內容：DV C SP(-1) DV(-1) SP
 - Method: 選擇 LS - Least Squares (ARMA)。
 - 點擊 OK。
- 3. 解讀 DV 方程的 ARDL 結果：
 - 尋找 C (α_{20})、SP(-1) (α_{21})、DV(-1) (α_{22}) 和 SP (α_{23}) 的係數、標準誤和 t-值。
 - 複製/記錄這些數據。

第 4 步: 回答問題 a, b, c

結合您在 EViews 中獲得的估計結果和您對計量經濟學理論的理解來回答這些問題。

a. 解釋為什麼對 VAR 模型(只有滯後變量在右側)使用 OLS 估計會得到一致估計量。

- 回顧 OLS 假設: OLS 的一致性要求解釋變量與誤差項不相關。
- 解釋 VAR 情況: 在 VAR 模型中，右側的解釋變量都是滯後變量(SP(-1) 和 DV(-1))。由於這些變量是上一期已經決定的值，它們與當期的誤差項 v_t^s 或 v_t^d 是不相關的。因此，OLS 估計量滿足一致性條件。您可以強調，這些誤差項代表了在考慮了所有過去信息後，當期無法預測的部分(創新)。

b. 解釋為什麼對 ARDL 模型(既有滯後變量也有當期變量在右側)使用 OLS 估計會得到不一致估計量。

- 定位問題點: 問題出在 ARDL 模型中的當期內生變量: DV 在 SP 方程中，SP 在 DV 方程中。
- 解釋內生性: SP_t 和 DV_t 在經濟中是同時決定的。這意味著：
 - SP_t 方程中的誤差項 e_t^s 可能會影響 DV_t 。例如，如果有一個未被模型包含的、影響當期股價的意外事件(體現在 e_t^s 中)，這個事件也可能同時影響當期股息的變化 DV_t 。
 - 或者，存在一個共同的、未觀測到的衝擊，同時影響 SP_t 和 DV_t 。
 - 在任何一種情況下，解釋變量(例如 SP 方程中的 DV_t)與誤差項 e_t^s 之間存在相關性。
- 導致不一致性: 當解釋變量與誤差項相關時，OLS 的基本假設被違反，導致 OLS 估計量是不一致的。這意味著即使有大量的數據，估計值也不會收斂到真實參數值，從而導致對係數的解釋產生偏誤。
- 引用: 可以提及這就是「同期內生性」或「聯立方程偏誤」，通常需要使用工具變量法或 GMM 來解決。

c. 從你的估計結果中推斷出股息在股價決定中的作用。

- 基於 VAR 結果：
 - 觀察 SP 方程中 DV(-1) 的係數 (β^{12})。如果這個係數顯著(p-值小於 0.05 或 0.1)，則表明滯後一期的股息變化率對當期股價變化率有顯著影響。這支持了股息是股價的一個滯後或領先指標的觀點。

- 觀察 DV 方程中 SP(-1) 的係數 (β^{21})。如果這個係數顯著，則表明滯後一期的股價變化率對當期股息變化率有顯著影響。這揭示了兩者之間的雙向動態關係。
- 基於 ARDL 結果 (帶有內生性警示):
 - 觀察 SP 方程中 DV 的係數 (α^{13})。如果這個係數顯著，它似乎表明當期股息變化率對當期股價變化率有即時影響。
 - 強調局限性: 儘管這個係數可能在統計上顯著，但由於前面 b 部分提到的內生性問題，OLS 估計量存在偏誤，因此對其因果關係的解釋需要極其謹慎。它可能只是反映了兩者共同變動的相關性，而不是嚴格意義上的因果影響。
- 總結推斷:
 - 根據 VAR 模型，你可以說，股息和股價之間存在動態的、相互影響的滯後關係。
 - 而 ARDL 模型提示了當期股息和股價可能存在同時聯動。
 - 綜合來看，股息的過去變動可以預測股價的當期變動，同時股價和股息也可能受到共同的當期衝擊影響，表現出同期變動。但要強調在沒有解決內生性問題的情況下，對 ARDL 中當期係數的因果解釋是不成立的。

希望這個詳細的步驟和解釋能幫助您在 EViews 中完成這項任務並深入理解相關概念！

資料來源

我想把以上這些討論內容，下載存成.pdf或其他格式的檔案，要如何做呢？

很抱歉，作為一個AI模型，我無法直接提供「下載」您與我對話內容的功能，也無法生成 PDF 或其他文件格式供您下載。我沒有檔案儲存或導出的能力。

但是，您可以透過以下方法手動將我們的對話內容保存下來：

1. 複製貼上 (Copy and Paste) 到文書處理軟體:
 - 在您的瀏覽器中，選取您想要保存的所有對話內容(從上到下拖曳鼠標選取)。
 - 按下 **Ctrl + C** (Windows/Linux) 或 **Cmd + C** (macOS) 複製選取內容。
 - 打開任何文書處理軟體，例如：
 - Microsoft Word
 - Google Docs
 - Pages (macOS)
 - Notepad / 記事本 (Windows)
 - TextEdit (macOS)
 - 其他文字編輯器
 - 在文書處理軟體中，按下 **Ctrl + V** (Windows/Linux) 或 **Cmd + V** (macOS) 貼上內容。
 - 您可以對格式進行一些調整，使其更美觀。
 - 然後，您可以將這個文件保存為您想要的格式，例如：

- **PDF**: 在 Word 或 Google Docs 等軟體中, 通常可以選擇 **File -> Save As** 或 **Print -> Print to PDF**。
- **DOCX / ODT / TXT**: 直接保存為文書處理軟體預設的格式。

2. 使用瀏覽器的「列印」功能:

- 在您的瀏覽器中, 按下 **Ctrl + P** (Windows/Linux) 或 **Cmd + P** (macOS) 打開列印對話框。
- 在列印選項中, 找到「目標印表機」或「目標」的設定。
- 選擇 **Save as PDF** (保存為 PDF) 或 **Microsoft Print to PDF** (Windows) 或 **預覽** (macOS, 然後可以選擇保存為 PDF)。
- 點擊 **Save** 或 **列印** 按鈕, 然後選擇您要保存的位置和檔案名稱。
- 這種方法通常會保留網頁的原始排版, 但可能不會完全符合您對文件編輯的需求。

建議: 如果您希望保留較好的排版並方便日後編輯, 推薦使用「複製貼上到文書處理軟體」的方法。如果只是想快速保存一份副本, 瀏覽器的「列印為 PDF」功能會更快捷。

請取用上傳的equity.xls資料，教我寫出.py程式，並一步步指導我，以繪製序列圖，並求算所需要的數據，完成答案嗎？

好的，同學！雖然您上傳的是 .xls 檔案，但系統將其識別為 CSV 格式(equity.xls - Sheet 1.csv)，這對於 Python 來說是更方便處理的格式。

我將指導您如何使用 Python 編寫程式碼，一步步完成問題 13.13 所需的序列圖繪製和數據計算(包括 VAR 和 ARDL 模型估計)，並解釋每一步。

在開始之前，請確保您的 **Python** 環境已經安裝了以下庫：

- **pandas**: 用於數據處理和操作。
- **matplotlib**: 用於繪製圖表。
- **statsmodels**: 用於時間序列模型的統計分析。

如果您尚未安裝，可以在您的命令提示符或終端中運行以下命令進行安裝：

```
Bash
pip install pandas matplotlib statsmodels
```

Python 程式碼指導: 完成問題 13.13

第 0 步: 載入數據

首先，我們需要將 CSV 文件讀取到 Python 的 **pandas** DataFrame 中。

Python

```
import pandas as pd
import matplotlib.pyplot as plt
import statsmodels.api as sm
from statsmodels.tsa.api import VAR

# 載入數據
# 請確保 'equity.xls - Sheet 1.csv' 文件與你的 Python 腳本在同一個目錄下，
# 或者提供完整的文件路徑。
try:
    df = pd.read_csv('equity.xls - Sheet 1.csv')
    print("數據載入成功！")
    print(df.head()) # 顯示前幾行數據，檢查是否正確載入
    print(df.info()) # 查看數據類型和非空值
except FileNotFoundError:
    print("錯誤: 文件 'equity.xls - Sheet 1.csv' 未找到。請檢查文件路徑和名稱。")
    exit() # 如果文件未找到，退出程式

# 確認數據列名是 'SP' 和 'DV'
```

```

# 根據你的 CSV 文件，列名可能不是直接的 'SP' 和 'DV'，
# EViews 導出的 CSV 可能有特殊的列名。
# 通常，第一列可能是年份或其他索引。我們假設第二列和第三列是 SP 和 DV。
# 你可能需要根據實際情況調整這裡的列名。
# 為了穩健性，我們假設第一列是日期/年份，第二列是SP，第三列是DV。
# 這裡我們將其重命名，如果您的數據實際列名不同，請根據 `df.head()` 的輸出進行調整。

# 假設 CSV 文件沒有標題行，或者標題行需要跳過，且第一列不是數據
# 如果你的 CSV 文件有標題行，且列名是 SP 和 DV，則下面的重命名步驟可以跳過
# 我們可以先看看 df.columns
# print(df.columns)

# 根據 EViews 導出文件的常見格式，可能需要這樣處理：
# 如果第一列是索引，且沒有列名，可以這樣重新讀取或重命名
# df = pd.read_csv('equity.xls - Sheet 1.csv', header=None) # 如果沒有標題行
# df.columns = ['Year', 'SP', 'DV'] # 手動指定列名

# 或者，如果載入後發現列名是類似 'Unnamed: 1', 'Unnamed: 2'
# 假設第二列是 SP，第三列是 DV (根據 Figure 13.9，可能原始數據就是變化率)
# 我們根據您提供的 Eviews 文件描述，原始數據中 SP 和 DV 已經是變化率了
# 如果您的 CSV 文件第一列是年份或索引，且列名是數值或 Unnamed，可能需要這樣調整：
if len(df.columns) >= 3:
    # 假設前兩列是數據，並且它們需要被指定為 SP 和 DV
    # 這裡我根據常見情況做一個假設，您需要根據您實際讀取到的 df.head() 輸出進行調整
    # 如果讀取到 df.columns 類似 ['Column1', 'Column2', 'Column3']
    # 並且根據 Figure 13.9 判斷 Column2 是 DV，Column3 是 SP
    # 我們將假設第一列是日期/索引，第二列是 DV，第三列是 SP (根據Eviews圖13.9的順序)
    # 但問題中提的公式  $SP = \dots$   $DV = \dots$  似乎 SP 是第一個變量
    # 為了與問題描述一致，我們假設 CSV 中，第一列是年份/索引，第二列是 SP，第三列是
    DV
    df = df.iloc[:, :3] # 只取前三列，避免讀到空列
    df.columns = ['Year', 'SP', 'DV'] # 重新命名列
    df = df.set_index('Year') # 將 Year 設為索引
    # 確保數據是數值類型
    df['SP'] = pd.to_numeric(df['SP'], errors='coerce')
    df['DV'] = pd.to_numeric(df['DV'], errors='coerce')
    df = df.dropna() # 刪除任何因為轉換失敗導致的 NaN 值
    print("\n處理後的數據前五行:")
    print(df.head())
else:
    print("數據列數不足，請檢查CSV文件格式或列名。")
    exit()

```

第 1 步：繪製序列圖 (Figure 13.9)

Python

```
# 繪製 SP 和 DV 的時間序列圖
```

```
plt.figure(figsize=(14, 6))

plt.subplot(1, 2, 1) # 1 行 2 列, 第 1 個圖
plt.plot(df.index, df['DV'])
plt.title('Change in Dividends (DV)')
plt.xlabel('Year')
plt.ylabel('DV')
plt.grid(True)

plt.subplot(1, 2, 2) # 1 行 2 列, 第 2 個圖
plt.plot(df.index, df['SP'])
plt.title('Change in Share Price (SP)')
plt.xlabel('Year')
plt.ylabel('SP')
plt.grid(True)

plt.tight_layout() # 自動調整子圖參數, 使之填充整個圖像區域
plt.show()

print("\n序列圖已繪製完成, 請檢查圖形是否類似 Figure 13.9。")
```

第 2 步: 估計 VAR(1) 模型

問題 13.13 給出的 VAR 模型為: $SP_t = \beta_{10} + \beta_{11}SP_{t-1} + \beta_{12}DV_{t-1} + v_{1t}$
 $DV_t = \beta_{20} + \beta_{21}SP_{t-1} + \beta_{22}DV_{t-1} + v_{2t}$

我們使用 `statsmodels.tsa.api.VAR` 來估計。

Python

```
# 創建 VAR 模型
# 需要先確認數據是否平穩。問題已經說明這是變化率, 因此可以認為是 I(0)。
# VAR 模型要求輸入的數據是端生變量的列表或 DataFrame。
model_var = VAR(df[['SP', 'DV']])

# 估計 VAR(1) 模型
# maxlags=1 表示只包含一階滯後
# trend='c' 表示包含常數項
results_var = model_var.fit(maxlags=1, trend='c')

print("\n--- VAR(1) 模型估計結果 ---")
print(results_var.summary())

# 提取並顯示關鍵係數
print("\n--- VAR(1) 模型的 SP 方程係數 ---")
# 係數順序通常是 (常數項, SP(-1), DV(-1))
print(f"常數項 (β10): {results_var.params[0][0]:.4f} (p-value: {results_var.pvalues[0][0]:.4f})")
print(f"SP(-1) (β11): {results_var.params[1][0]:.4f} (p-value: {results_var.pvalues[1][0]:.4f})")
```

```

print(f"DV(-1) ( $\beta_{12}$ ): {results_var.params[2][0]:.4f} (p-value: {results_var.pvalues[2][0]:.4f})")

print("\n--- VAR(1) 模型的 DV 方程係數 ---")
# 係數順序通常是 (常數項, SP(-1), DV(-1))
print(f"常數項 ( $\beta_{20}$ ): {results_var.params[0][1]:.4f} (p-value: {results_var.pvalues[0][1]:.4f})")
print(f"SP(-1) ( $\beta_{21}$ ): {results_var.params[1][1]:.4f} (p-value: {results_var.pvalues[1][1]:.4f})")
print(f"DV(-1) ( $\beta_{22}$ ): {results_var.params[2][1]:.4f} (p-value: {results_var.pvalues[2][1]:.4f})")

# 注意: statsmodels 的 VAR 輸出中, params 矩陣的列是方程, 行是變量。
# 所以 params[0] 是常數項的係數, params[1] 是第一個滯後變量 (SP(-1)) 的係數,
# params[2] 是第二個滯後變量 (DV(-1)) 的係數。
# params[row_index][col_index]
# col_index 0 是 SP 方程, 1 是 DV 方程
# row_index 0 是常數, 1 是 SP(-1), 2 是 DV(-1)

```

第 3 步: 估計 ARDL 模型 (作為兩個獨立的 OLS 回歸)

問題 13.13 給出的 ARDL 模型為: $SP_t = \alpha_{10} + \alpha_{11}SP_{t-1} + \alpha_{12}DV_{t-1} + \alpha_{13}DV_t + \epsilon_t$
 $DV_t = \alpha_{20} + \alpha_{21}SP_{t-1} + \alpha_{22}DV_{t-1} + \alpha_{23}SP_t + \epsilon_t$

我們需要手動創建滯後變量, 然後使用 `statsmodels.formula.api.ols` 或 `statsmodels.api.OLS` 進行 OLS 回歸。

Python

```

# 創建滯後變量
df_lagged = df.copy()
df_lagged['SP_lag1'] = df_lagged['SP'].shift(1)
df_lagged['DV_lag1'] = df_lagged['DV'].shift(1)

# 刪除因滯後操作產生的 NaN 值
df_lagged = df_lagged.dropna()

print("\n--- ARDL 模型估計 (SP 方程) ---")
# SP 方程:  $SP_t = \alpha_{10} + \alpha_{11}SP_{t-1} + \alpha_{12}DV_{t-1} + \alpha_{13}DV_t + \epsilon_t$ 
# 在 statsmodels 的 OLS 中, 需要明確指出因變量和自變量
# 自變量列表包含常數項 ('const'), 通過 sm.add_constant 添加
X_sp_ardl = sm.add_constant(df_lagged[['SP_lag1', 'DV_lag1', 'DV']])
y_sp_ardl = df_lagged['SP']
model_sp_ardl = sm.OLS(y_sp_ardl, X_sp_ardl)
results_sp_ardl = model_sp_ardl.fit()
print(results_sp_ardl.summary())

# 提取並顯示關鍵係數
print("\n--- ARDL 模型的 SP 方程係數 ---")
print(f"常數項 ( $\alpha_{10}$ ): {results_sp_ardl.params['const']:.4f} (p-value: {results_sp_ardl.pvalues['const']:.4f})")

```

```

print(f"SP(-1) ( $\alpha_{11}$ ): {results_sp_ardl.params['SP_lag1']:.4f} (p-value:
{results_sp_ardl.pvalues['SP_lag1']:.4f})")
print(f"DV(-1) ( $\alpha_{12}$ ): {results_sp_ardl.params['DV_lag1']:.4f} (p-value:
{results_sp_ardl.pvalues['DV_lag1']:.4f})")
print(f"DV ( $\alpha_{13}$ ): {results_sp_ardl.params['DV']:.4f} (p-value:
{results_sp_ardl.pvalues['DV']:.4f})")

print("\n--- ARDL 模型估計 (DV 方程) ---")
# DV 方程:  $DV_t = \alpha_{20} + \alpha_{21}SP_{t-1} + \alpha_{22}DV_{t-1} + \alpha_{23}SP_t + e_t$ 
X_dv_ardl = sm.add_constant(df_lagged[['SP_lag1', 'DV_lag1', 'SP']])
y_dv_ardl = df_lagged['DV']
model_dv_ardl = sm.OLS(y_dv_ardl, X_dv_ardl)
results_dv_ardl = model_dv_ardl.fit()
print(results_dv_ardl.summary())

# 提取並顯示關鍵係數
print("\n--- ARDL 模型的 DV 方程係數 ---")
print(f"常數項 ( $\alpha_{20}$ ): {results_dv_ardl.params['const']:.4f} (p-value:
{results_dv_ardl.pvalues['const']:.4f})")
print(f"SP(-1) ( $\alpha_{21}$ ): {results_dv_ardl.params['SP_lag1']:.4f} (p-value:
{results_dv_ardl.pvalues['SP_lag1']:.4f})")
print(f"DV(-1) ( $\alpha_{22}$ ): {results_dv_ardl.params['DV_lag1']:.4f} (p-value:
{results_dv_ardl.pvalues['DV_lag1']:.4f})")
print(f"SP ( $\alpha_{23}$ ): {results_dv_ardl.params['SP']:.4f} (p-value:
{results_dv_ardl.pvalues['SP']:.4f})")

```

第 4 步: 回答問題 a, b, c

這部分是基於您從 Python 程式碼中獲得的輸出結果，結合您對計量經濟學理論的理解來撰寫答案。

a. 解釋為什麼對 VAR 模型（只有滯後變量在右側）使用 OLS 估計會得到一致估計量。

- 理論點: OLS 估計量的一致性要求解釋變量與誤差項不相關。
- **VAR** 模型適用性: 在 VAR 模型中，每個方程的右側僅包含滯後的內生變量和（可選的）外生變量（如常數項）。例如，在 SP 方程中， SP_{t-1} 和 DV_{t-1} 是在 t 時期誤差項 v_t 實現之前就已經確定的值。因此，它們是獨立於當期誤差項的。這就滿足了 OLS 估計量一致性的基本假設。每個方程可以被視為一個簡單的自回歸分佈滯後模型，並且由於解釋變量是過去值，所以它們是外生的，OLS 產生一致估計。

b. 解釋為什麼對 ARDL 模型（既有滯後變量也有當期變量在右側）使用 OLS 估計會得到不一致估計量。

- 問題點: ARDL 模型中引入了當期的內生變量作為解釋變量。例如，在 SP 的 ARDL 方程中， DV_t 作為解釋變量出現。

- 同期內生性: SP_t 和 DV_t 在經濟中往往是同時決定的, 它們之間可能存在雙向因果關係。這意味著 DV_t 很可能與 SP 方程中的誤差項 ets 相關聯(例如, 如果 ets 包含一個未被模型捕獲的、同時影響 SP_t 和 DV_t 的衝擊)。
- 違反 OLS 假設: 當解釋變量(如 DV_t)與誤差項相關時, OLS 估計量將是不一致的。這會導致估計係數有偏誤, 且這種偏誤不會隨著樣本量的增加而消失。因此, 從這種估計中得出的推論可能是錯誤的。這就是聯立方程偏誤 (Simultaneous Equations Bias) 或內生性問題。

c. 從你的估計結果中推斷出股息在股價決定中的作用。

- 參考 VAR 模型結果:
 - SP 方程中 $DV(-1)$ 的係數 (β_{12} 的估計值):
 - 查看其係數值和 p-value。如果 p-value 小於 0.05(或您選擇的其他顯著性水平), 則表明滯後一期的股息變化率對當期股價變化率有統計上顯著的影響。係數的正負則指示了影響方向(正相關或負相關)。這支持了股息變化在短期內可能對股價變化產生影響的觀點。
 - DV 方程中 $SP(-1)$ 的係數 (β_{21} 的估計值):
 - 同樣查看其係數值和 p-value。如果顯著, 表明滯後一期的股價變化率對當期股息變化率有顯著影響。這提示了股價與股息之間可能存在雙向動態關係。
- 參考 ARDL 模型結果 (並強調其局限性):
 - SP 方程中 DV 的係數 (α_{13} 的估計值):
 - 觀察其係數值和 p-value。如果這個係數在統計上顯著, 它似乎表明當期股息變化率對當期股價變化率有直接的、即時的影響。
 - 但請務必強調: 由於問題 (b) 中解釋的內生性問題, 這個 OLS 估計值可能存在偏誤。因此, 雖然結果可能暗示了當期股息與股價之間的強相關性, 但不能直接將其解釋為明確的因果關係, 而是可能反映了它們受到共同的、未被模型捕獲的因素的影響而同時變動。
- 綜合推斷:
 - 基於 VAR 模型, 您可以推斷出股息和股價之間存在動態的、滯後的相互影響。例如, 過去股息的增長可能導致股價的當期增長, 反之亦然。這反映了市場參與者如何根據過去的信息調整預期。
 - 而 ARDL 模型(儘管估計有偏)可能暗示了股息和股價存在當期的聯動。這可能源於它們對相同資訊的快速反應, 或者市場對股息變動的即時反應。
 - 總體而言, 實證結果通常會支持股息對股價有影響(無論是滯後還是同期), 因為股息是公司盈利能力和未來現金流預期的一個信號。但金融理論中關於其確切角色(例如, 是價格的驅動因素, 還是僅僅反映了基礎價值)仍有爭議。您的實證分析將有助於探討這些動態關係。

希望這份詳細的 Python 指導和答案框架能幫助您! 在運行代碼時, 請仔細檢查控制台的輸出, 特別是數據載入和列名是否如預期。