

數學建模競賽寫作課程

W2：摘要表——主張、證據與決策

Kenneth Sok Kin Cheng

修訂日期：2026 年 8 月 20 日

學習目標

完成本節兩小時課堂後，學生應能夠：

- 區分競賽摘要表、引言、學術摘要、方法清單與團隊工作日記；
- 在撰寫正文前建立證據清單；
- 將題目中的每項主要任務對應至模型、結果與決策；
- 以任務、策略、發現、信任／決策四個功能區塊組織摘要；
- 選取少量與決策相關的數字，並連同單位、基準及適用條件報告；
- 在精簡的主張中區分運算／程式核證、模型效度驗證、敏感度分析與不確定性分析；
- 在不削弱研究成果的前提下提出附有條件的建議；
- 在保留任務覆蓋與證據鏈的同時精簡草稿；
- 從答案可見度、證據、一致性與精確度進行快速同儕審閱。

1 競賽摘要表應完成甚麼

摘要並非把每個章節等比例縮短。它是一份決策簡報：讀者毋須翻閱完整報告，也應能辨認任務、建模路徑、主要發現，以及建議背後信心水平或限制。

官方指引也支持這項重點。現行 HiMCM 指示要求解答以摘要表開首，並要求團隊清楚呈現主要構想與結果；現行 IM²C USA 指示則指出評審十分重視摘要，而且最重要的結論應予突出。這些來源都沒有支持「評審恰好只看七分鐘」或「摘要佔 60% 分數」之類未經證實的說法。我們應依據官方要求，而非競賽傳聞。

啟動活動：20 秒答案測試

閱讀一頁摘要二十秒，然後把它蓋起來。在不回看的情況下回答：

1. 團隊須作出甚麼決策，或求出甚麼數值？
2. 團隊的主要建議或數值答案是甚麼？
3. 哪個模型或方法產生這些結果？
4. 哪些證據令建議可信？

若讀者無法憑記憶回答至少三項，這份摘要便尚未發揮摘要應有的功能。

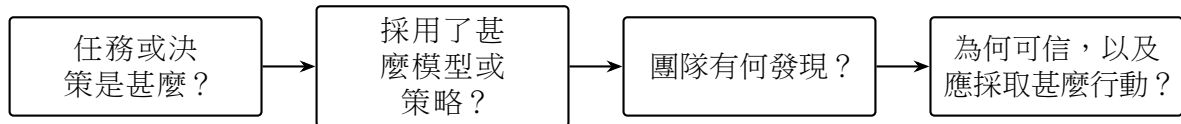


圖 1：摘要應回答讀者的四個問題。措辭與次序可以不同，但四項功能都應清晰可見。

模型檢視：摘要表、引言與學術摘要並不相同

- 引言鋪陳背景、解釋問題及交代報告結構，因此可以稍後才呈現數值答案。
- 研究論文的學術摘要通常遵循學術慣例，以精簡形式交代目的、方法、結果與結論。
- 競賽摘要表更接近一份供決策者閱讀的簡報：它應直接陳述答案、辨明建模架構，並說明答案為何站得住腳。

直接把引言複製為摘要，通常會造成背景資訊過多、結果過少。

2 不要從正文起筆：先建立證據清單

摘要薄弱，往往不只是句子寫得不好，而是證據取捨失當。團隊在判斷哪些任務、結果、比較與限制必須放進一頁之前，便急於寫成完整段落。

五欄證據清單

起草前，為每項重要任務填寫一列：

任務	核心答案	方法	證據 / 核證	條件 / 限制
必須回答甚麼？	數值、排名、路線、閾值或建議	產生結果的模型或演算法	基準、模型效度驗證、不確定性或敏感度分析	答案在甚麼情況下可能改變？

沒有核心答案的一列，表示分析尚未完成；沒有證據的一列，表示主張缺乏支持；沒有適用條件的一列，則可能反映過度自信。

例題 1：學校巴士案例的證據清單

一所學校須為 226 名學生重新設計路線，所用巴士每輛可載 60 人。目標是在最長預定車程不超過 45 分鐘的前提下，縮短乘車時間。

任務	答案	方法	證據 / 核證	條件
選擇車隊規模	四輛巴士	容量下界加路線規劃模型	$226/60 = 3.77$ ，故至少需要四輛	假設所有登記乘客當日均乘車
設計路線	四條路線，總長 87.4 公里	先分群、再作車輛路線規劃的啟發式方法與局部搜尋	現有路線總長 103.6 公里	不考慮道路封閉
縮短乘車時間	平均 28.2 分鐘；最長 44.1 分鐘	時變最短路徑	現有平均為 42.9 分鐘；所有規劃路線均短於 45 分鐘	使用早上行車時間估計
測試需求變動	在 20 個受測需求情境中，四輛巴士在 18 個情境仍足夠	各地區學生數作 ± 10 擾動後重新分配	在 17 個情境中，最長車程仍短於 45 分鐘	兩個高需求情境需要第五輛巴士

摘要毋須塞進表內所有數字；它只須保留能回答任務、證明改善幅度，以及揭示決策邊界的數字。

檢查點：檢查點：還欠甚麼？
某團隊寫道：「優化後的路線長度為 87.4 公里。」請根據證據清單指出至少三項欠缺的資訊。可追問：相對哪個基準？用了多少輛巴士？須符合甚麼容量與時間限制？結果有多穩定？

模型檢視：任務覆蓋先於文字優雅
即使摘要行文出色，只要遺漏一項必要任務，仍是不完整。把題目的任務動詞抄出來——估計、設計、比較、建議、測試——再檢查每個動詞是否直接出現，或是否已有明確結果句作答。

3 四個功能區塊——保留實用彈性

四區塊構想是一項有用的設計工具，卻不應被當作僵硬的段落規則。

建議的四區塊架構

區塊	主要功能	必須包含
1.任務 (Task)	引導讀者	決策情境、所需輸出、重要約束
2. 策略 (Strategy)	解釋模型架構	基準至改良的次序、具名方法、資料的用途
3. 發現 (Findings)	提供答案	先寫主要結果，再寫單位、比較與任務覆蓋
4. 信任與決策 (Trust and decision)	校準信心並連接行動	模型效度驗證／敏感度證據、閾值或注意事項、建議

典型的一頁摘要可以粗略分配 2-3 句給任務、3-5 句給策略、4-6 句給發現，以及 2-4 句給信任／決策。這些只是起草時可參考的範圍，並非硬性規則。

任務 (Task)：點明決策、輸出與約束——不是交代競賽歷史。

策略 (Strategy)：顯示模型如何銜接——不是羅列軟件和方程式。

發現 (Findings)：由答案寫起，再交代改善、閾值或不確定性。

信任 / 決策 (Trust/Decision)：交代核查、測試範圍與下一步行動。

圖 2：一個穩定的起草架構。若答案非常迫切，而且毋須額外背景即可理解，也可採用結果先行的開首。

模型檢視：何時適合以結果開首

若題目明顯以決策為導向，摘要可以答案開首：

我們建議一般上課日安排四輛巴士；預測乘客數達 236 人時開始提早檢視車隊；確認乘客數達 241 人（即超過二百四十人）時啟用第五輛應變巴士。

下一句仍須交代任務與適用條件。結果先行只是改變證據的呈現次序，並不等於可以刪去必要背景。

4 撰寫任務與策略區塊

4.1 任務區塊：精簡問題，而非刪光背景

有力的任務區塊應令交付成果一目了然。除非某項背景會改變建模選擇，否則毋須在此詳述歷史。

例題 2：薄弱與改良的任務區塊

薄弱

校車在很多地方都很重要，很多學生每天都乘搭校車。本文研究校車路線規劃；這是一個複雜問題。

改良

學校須以每輛容量 60 人的巴士接載 226 名學生，同時縮短車程，並確保每條預定路線的車程少於 45 分鐘。因此，我們尋找最小可行車隊、符合容量限制的路線組合，以及應對每日需求不確定性的啟動規則。

改良版本在兩句內點明決策變量、約束與所需輸出。

任務區塊句式

[決策者] 須為 [系統 / 群體] 作出 [決策]，並符合 [主要約束]。我們估計 / 設計 / 比較 [所需輸出]，以 [效能指標] 為主要準則。

4.2 策略區塊：先交代架構，後交代技術細節

策略區塊應解釋為何需要多種方法，以及它們如何互相銜接。軟件清單並不是建模策略。

例題 3：方法清單與模型架構

薄弱

我們使用 *Python*、分群、*Dijkstra* 演算法、最佳化、模擬與敏感度分析。

改良

我們先由容量限制導出車隊規模下界，再把鄰近車站分群，並在以時間加權的道路網絡上求最短路徑，以建立可行路線。其後以局部搜尋在路線之間交換車站，目標是縮短最長車程，而不只減少總距離。最後，我們在擾動後的地區需求下重新執行分配，找出四輛巴士不再足夠的條件。

改良版本呈現的是一連串有因果關係的建模決定。軟件只屬實作細節，可從摘要省略。

檢查點：檢查點：說明每次修正

逐項說明後一階段補救了前一階段的甚麼弱點：

1. 容量下界 → 可行路線規劃；
2. 最小總距離路線 → 針對最長車程的局部搜尋；
3. 名義需求 → 擾動需求下重新執行。

加入改良步驟，應是因為基準模型遺漏了重要因素，而不是因為第二種演算法聽起來較先進。

5 寫出有證據支撐的發現

發現區塊不是存放所有輸出的倉庫，而應呈現有層次的主張。

結果的三個層次

1. 決策結果：題目要求的路線、數量、排名、閾值或建議。
2. 效能證據：相對基準或要求的改善。
3. 穩健性邊界：決策在甚麼測試範圍內維持不變，以及何時改變。

摘要通常需要三個層次，但每層只選少量最有解釋力的數字。

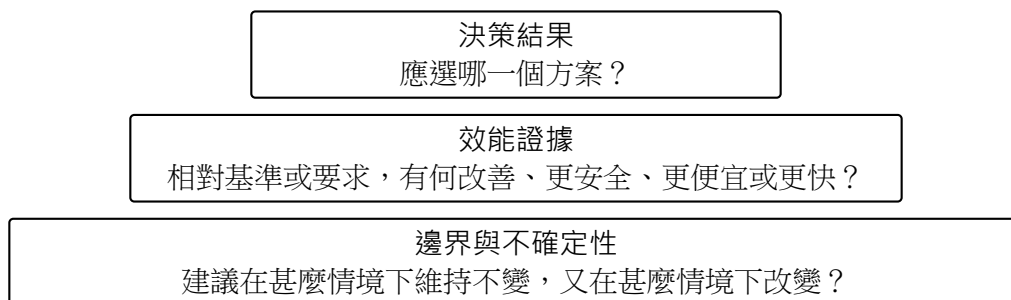


圖 3: 結果的選取應由決策推進至量化支持，再交代測試邊界。

量化主張必須有語境

有用的結果句通常包含下列數項：

- 數值與單位：28.2分鐘、87.4 公里、四輛巴士；
- 基準或目標：由 42.9 分鐘下降、低於 45 分鐘限制；
- 範圍：早上路線、226名登記學生、一般需求；
- 不確定性或測試範圍：20個情境中的 18 個、各地區 ± 10 名學生；
- 決策含義：維持四輛巴士、達到啟動條件時加入第五輛。

例題 4：逐步加強結果句

版本	句子
過於含糊	優化後的路線比現有路線好得多。
只有數值	平均車程為 28.2 分鐘。
加入比較	平均車程由 42.9 分鐘降至 28.2 分鐘，減少 34.3%。
可供決策	採用四輛巴士時，平均車程由 42.9 分鐘降至 28.2 分鐘，最長路線為 44.1 分鐘，符合 45 分鐘要求。
校準主張	採用四輛巴士時，平均車程由 42.9 分鐘降至 28.2 分鐘，最長路線為 44.1 分鐘；四輛巴士在 20 個擾動需求情境中的 18 個仍然可行，另有兩個高需求情境須使用第五輛。

模型檢視：精簡選數

數字越多，不代表證據越強。可選約三至六個各有用途的核心數字，分別回答答案、改善幅度、約束及穩健性邊界；路線細節、參數表和次要指標則留在正文。

檢查點：檢查點：補充分母

「方法的成功率為 90%」並不完整。請重寫句子，交代成功率來自哪些情境或重複試驗、輸入範圍為何，以及以甚麼閾值判斷成功。

6 交代可信度，同時避免誇大

最後一個區塊應告訴讀者發現為何可信，以及這種信任的界線。

四種不同的證據動詞

術語	回答的問題	摘要用語
運算 / 程式核證	我們有否正確實作或求解模型？	“手算的四站小型案例與程式輸出一致。”
模型效度驗證	模型有否重現相關現實或公認基準？	“預測車程與兩個留出早上的觀測值平均相差 2.4 分鐘。”
敏感度分析	輸入或假設改變時，結果如何變化？	“四輛巴士方案在 20 個擾動需求情境中的 18 個仍然可行。”
不確定性分析	估計值周圍有甚麼範圍或概率？	“在所有重複試驗中，模擬所得第 90 百分位等待時間介乎 7.8–9.1 分鐘。”

若唯一檢查只是程式順利執行，便不應寫成“validated”；若沒有說明測試範圍與結果指標，也不應聲稱結果“robust”。

例題 5：加入限制而不自我否定

薄弱表述

模型有很多不切實際的假設，所以結果可能不準確。

改良表述

四輛巴士的建議在已測試的各區 ± 10 名學生擾動下保持穩定，但測試未涵蓋道路

封閉和異常集中的缺席。因此，我們建議平日安排四輛巴士，預測乘客數達 236 人時提早檢視車隊，並在確認乘客數達 241 人（即超過二百四十人）或主要路線無法使用時啟用第五輛應變巴士。

改良版本清楚交代已測範圍、未測條件與實際應對方式。

精確用語庫

避免使用	有相應證據時可改用
模型證明了……	模型預測……／在所述假設下，分析顯示……
解決方案是最佳的。	解決方案是所述線性模型的最優解。／它是已測候選方案中找到的最佳結果。
結果是穩健的。	當 p 由基準值的 0.8 倍變至 1.2 倍時，建議維持不變。
模型是準確的。	留出觀測值上的平均絕對模型效度驗證誤差為 2.4 分鐘。
方法很有效率。	對 250 個節點的測試案例，執行時間由 48 秒降至 7 秒。

7 完整寫作轉化

本節把例題 1 的證據清單轉化為一份完整的競賽式摘要。案例只供說明；重點是句子設計，而非巴士路線本身的數學計算。

初稿

本文研究校車路線問題。我們使用分群、最短路徑、一種最佳化演算法、*Python* 和敏感度分析。我們發現模型表現良好，並能減少路線。結果顯示四輛巴士很好，學校應採用我們的模型。研究存在若干限制，但未來工作可以改善。

模型檢視：診斷

草稿只提及研究主題，沒有實際約束；它羅列工具，卻沒有交代建模架構；它也沒有按任務給出單位、基準、閾值或結果。「表現良好」、「很好」和「限制」均是缺乏證據的標籤，而建議亦沒有啟動條件。

完整示範摘要

任務。學校須以每輛容量 60 人的巴士運送 226 名登記學生，同時縮短車程，並確保每條早上路線少於 45 分鐘。我們求出最小可行車隊、建立路線方案，並判斷何種需求不確定性會令系統需要額外容量。

策略。我們先從容量限制求得下界，證明至少需要四輛巴士。其後把相鄰車站分群，並以道路網絡上的時間加權最短路徑連接各站。由於只減少總距離可能造成某條路線過長，局部搜尋階段會在維持容量可行的前提下，於不同巴士路線之間交換車站，以縮短最長車程。最後，我們在二十個需求擾動情境下重新執行分配；每個情境把各地區乘客數增加或減少最多十人。發現。我們建議一般上課日採用四條路線，總長 87.4 公里。相對現有 103.6 公里的方案，路線總長減少 15.6%。預測平均乘車時間由 42.9 分鐘降至 28.2 分鐘，而最長預定車程為 44.1 分鐘，符合少於 45 分鐘的要求。四輛巴士在 20 個需求擾動情境中的 18 個仍然可行；在兩個需求最集中的高需求案例中，至少一條路線會超出容量，因而需要第五輛巴士。

信任與決策。手工追蹤的小型案例重現了路線成本計算，而模型車程亦已與最佳化前的早上行車紀錄比較。因此，這項方案適用於一般需求，但不應視為無條件成立。我們建議平日安排四輛巴士，預測乘客數達 236 人時開始提早檢視車隊，並在確認乘客數達 241 人（即超過二

百四十人)、地區需求異常集中,或主要道路封閉令行車時間矩陣失效時,啟用第五輛應變巴士。

檢查點：句子功能檢查

請在完整摘要中分別畫線標示：

1. 陳述核心答案的句子；
2. 交代相對基準改善幅度的句子；
3. 報告約束測試的句子；
4. 報告敏感度分析的句子；
5. 把無條件建議改成附條件建議的句子。

若兩項不同的必要任務只能依靠同一句含糊句子作答，便須增加具體資訊。

8 精簡與修訂

一頁限制帶來的是取捨問題。目標不是平均縮短每句，而是移除價值較低的材料，同時保護完整證據鏈。

三遍修訂

1. 完整性檢查：每項主要任務都有答案嗎？建議及重要條件是否明確？
2. 證據檢查：每項主要主張都有數值、比較、核查或測試範圍支持嗎？
3. 精簡檢查：刪去歷史敘述、工作過程、重複的方法細節、空泛形容詞、軟件名稱、方程式及填充式的未來工作。

不要從精簡開始；簡短但不完整的摘要，仍然不完整。

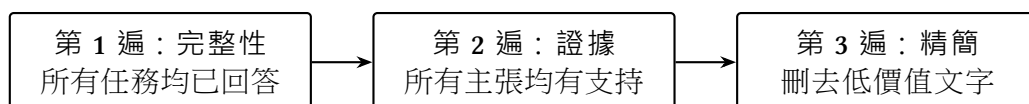


圖 4: 修訂次序很重要。在核對任務覆蓋前刪改，可能會誤刪唯一回答某項要求的句子。

例題 6：精簡而不丟失證據

精簡前——44 字

為判斷我們提出的路線模型在各地區所分配的學生人數改變時能否繼續有效運作，我們進行了一項涵蓋多個替代需求情況的敏感度分析。

精簡後——19 字

我們在二十個情境下重新執行分配，各地區需求擾動為 ± 10 名學生。

較短的句子仍保留行動、樣本數、所改變的輸入及擾動範圍。

模型檢視：必須保護的高價值細節

保留任務答案、決策閾值、單位、基準、模型效度驗證指標、測試範圍及建議的適用條件。可以刪去競賽年份、團隊工作流程、程式語言、方程式、冗長參數清單，以及「未來工作可以改善模型」一類空泛句子。

9 寫作工作坊與同儕審閱

核心路線：由證據寫成摘要

運用校車證據清單，自行撰寫一份 12–16 句的摘要，不要照抄上述完整版本。要求如下：

- 回答全部四項任務；
- 包含至少四個有單位的數值或量化比較；
- 包含一句運算／程式核證或模型效度驗證陳述；
- 包含一項敏感度邊界；
- 以附有條件的建議作結。

挑戰路線：兩種開首

分別撰寫摘要首四句的兩個版本：

1. 任務先行；
2. 結果先行。

與同伴交換草稿。哪個版本既令答案更易辨認，又不會令背景含糊？

延伸路線：摘要中的指標衝突

假設最短路線方案把總距離降至 82 公里，卻令一名學生的車程達 51 分鐘；平衡方案的總距離為 87.4 公里，但所有車程均少於 45 分鐘。請用兩句解釋為何建議採用平衡方案，並說清目標、約束與取捨，不要在沒有定義的意義下稱任何方案為「最佳」。

兩分鐘同儕審閱——課堂工具，並非官方評分準則

每項給予 0、1 或 2 分。

準則	問題
任務覆蓋	每項主要要求都能與答案對應嗎？
答案可見度	能在十秒內找到核心建議嗎？
模型架構	讀者能理解各方法如何銜接，以及為何需要改良嗎？
量化證據	重要主張有單位、基準、閾值或不確定性支持嗎？
信任校準	有否準確描述運算／程式核證、模型效度驗證或敏感度分析，並附上邊界或注意事項？
精簡與一致性	每句均有作用，而且所有數字與正文一致嗎？

評分後只寫兩項意見：最有價值的一句，以及最重要的一項修訂。

檢查點：課末檢核

按每項提示各寫一句：

1. 你所解決的一個建模問題之核心答案；
2. 令答案具有意義的一個基準或閾值；
3. 支持讀者信任答案的一項證據；
4. 令建議改變的一項條件。

這四句正是日後撰寫「發現」與「信任／決策」區塊的核心。

10 進入 W3 前的準備

W2 重點整理

1. 撰寫正文前，先建立證據清單。
2. 把四個區塊視為四項功能：任務（Task）、策略（Strategy）、發現（Findings）、信任／決策（Trust/Decision）。
3. 先寫核心答案，再寫次要輸出。
4. 報告數字時交代單位、基準、閾值與範圍。
5. 準確區分運算／程式核證（verification）、模型效度驗證（validation）、敏感度分析（sensitivity analysis）與不確定性分析（uncertainty analysis）。
6. 限制若能轉化為條件、啟動規則或警告，便具有決策價值。
7. 按完整性 → 證據 → 精簡的次序修訂。
8. 摘要必須與正文、表格及結論完全一致。

在 **W3**：假設、效度驗證與敏感度分析，我們會建立本節「可信度句子」背後所需的證據。學生將設計可供檢驗的假設，區分模型的運算／程式核證與現實世界的模型效度驗證，並把參數變動轉化為與決策相關的敏感度結論。

附錄 A：證據清單工作紙

任務（Task）	核心答案 （Headline answer）	方法（Method）	證據／核證 （Evidence/check）	條件／限制（Condition/caveat）

附錄 B：句式

- 任務 (**Task**)：[決策者] 須在 [約束] 下 [作出決策]；因此，我們估計 / 設計 / 比較 [輸出]。
- 架構 (**Architecture**)：我們先以 [基準] 開始，再加入 [改良] 以補救 [具名弱點]，最後進行 [測試 / 決策階段]。
- 核心結果 (**Headline result**)：我們建議 [行動 / 數量]，其在 [範圍] 下產生 [附有單位的指標]。
- 比較 (**Comparison**)：相對 [基準]，[指標] 由 [舊值] 變為 [新值]，改善 [百分比或絕對值]。
- 模型效度驗證 (**Validation**)：相對 [留出資料 / 基準]，模型達到 [誤差或一致程度量]。
- 敏感度分析 (**Sensitivity**)：當 [輸入] 在 [範圍] 內改變，建議維持不變；但達到 [閾值] 時，建議便會改變。
- 附條件決策 (**Conditional decision**)：在 [條件] 下採用 [行動 A]；達到 [啟動條件] 時改用 [行動 B]。

附錄 C：最終一致性核對表

- 題目要求的每項任務均在摘要及正文出現。
- 核心數字與最終結果表完全一致。
- 單位、分母、時間範圍及情境數均清楚可見。
- 「最優 (optimal)」、「已驗證 (validated)」、「穩健 (robust)」與「準確 (accurate)」均有恰當證據支持。
- 除非影響可重現性或解釋，否則不列軟件名稱。
- 除非方程式本身就是核心結果，否則不在摘要放入方程式。
- 建議須包含決策者、行動及相關條件。
- 摘要不包含正文沒有出現的引文或主張，也不附參考文獻表。
- 頁面符合當前官方競賽樣式與識別規定。

官方來源註記。本節涉及競賽規則的陳述，已於 2026 年七月按 COMAP 2026 HiMCM/MidMCM Instructions 及 2026 IM²C USA Instructions 核對。規則、樣式及地區要求或會改變；每次參賽前應重新查閱最新官方頁面。

模型審核卡

在信任或傳達模型之前先作審核			
目的	模型要回答甚麼問題或支援甚麼決策？	資料	哪些輸入屬觀測、推導或模擬資料？
假設	哪些簡化可能改變答案？	單位	量綱與尺度是否一致？
運算 / 程式核證	方程、程式、限制與極限情況是否已核對？	模型效度驗證	是否用獨立證據檢驗預定用途下的表現？
不確定性	測量、參數、求解器及模型結構誤差如何影響結論？	倫理	誰受益、誰承擔風險，公平準則是甚麼？
可重現性	是否記錄資料來源、版本、種子與轉換？	溝通	主張是否有條件、可追溯且與決策相關？

延伸閱讀

- HiMCM / MidMCM 競賽指引（COMAP）
- 組織研究文章（IEEE Author Center）

來源、資料脈絡與授權

本講義依循 SIAM / COMAP GAIMME 報告對反覆建模與評核的觀點。競賽相關說明於 2026 年 8 月 20 日依據官方 COMAP HiMCM 規則核對；提交前仍須查閱當時最新規則。除非另有明確來源，課堂資料與數值情境均屬合成或模擬教學例子，不構成任何真實機構的實證主張。第三方參考資料沿用各自條款。Kenneth Sok Kin Cheng 的原創教材採用 CC BY-NC-SA 4.0 授權。

學生作答指引與檢核準則
封閉式計算題請列出定義、代入、單位及至少一項獨立檢查；本課的完整例題可作方法提示，但須自行完成數值。開放建模題按五項檢核：目的與輸出是否明確、資料與來源是否交代、假設與適用範圍是否合理、運算／程式核證與模型效度驗證是否可追蹤，以及不確定性、限制與建議是否互相一致。每項須附具體證據，而非只寫「已完成」。