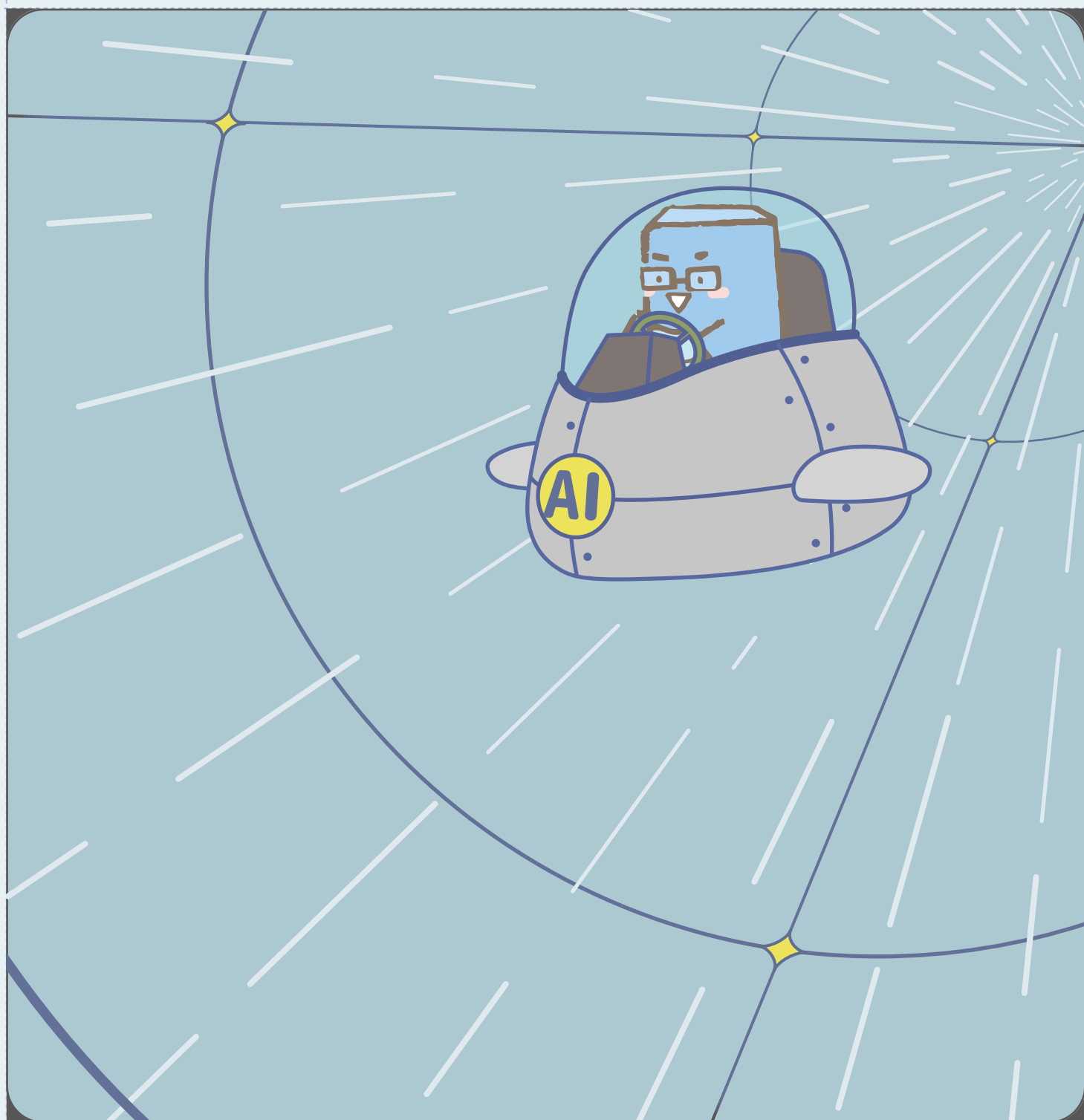




高中生 智慧駕馭AI

2026 年 9 月



部長序

親愛的同學們：

2022年底 ChatGPT 問世後，生成式人工智慧（Generative AI，GenAI）迅速改變了人們獲取知識、創作內容與解決問題的方式。短短幾年間，AI 已從新興技術走進日常生活。如今，AI 正從回應問題、生成內容，進一步發展為能依照目標規劃步驟、運用不同工具，並協助完成多階段任務的「AI 代理人」。當 AI 能做的事情愈來愈多，我們需要學習的便不只是如何下指令，更包括如何設定目標、分配任務、檢查過程、修正結果，以及判斷哪些決定不能交給 AI。

2024年9月，聯合國教科文組織（UNESCO）發布《學生人工智慧素養框架》，提出學生在 AI 時代應具備的核心能力，強調批判思考、倫理判斷、合作創新，以及負責任地使用 AI。科技的進步，不只是能力的競賽，更是價值與判斷力的考驗。

教育部特別規劃了三本專為中小學學生設計的學習手冊——國小生聰明用 AI、國中生靈活用 AI、高中生智慧駕馭 AI；希望透過三個學習階段，引領我們從多角度認識 AI、學習人機協作，並逐步建立安全、適切且負責任的使用能力。

這本《中小學 AI 之學習應用手冊——高中生智慧駕馭 AI》，以培養具備「AI 共學力」的青年公民為目標，從 AI 的運作原理、提問策略與資訊判讀出發，進一步探討偏誤、公平性、著作權、隱私保護及倫理等重要議題。

AI 能協助我們更有效率地找到答案，但無法取代我們思考；能生成大量內容，但無法替我們決定價值。AI 代理人可以協助規劃與執行，卻不應取代人的目標、判斷與責任。真正決定科技如何影響未來的，始終是使用它的人。面對快速演進的科技時代，希望這本手冊能陪伴我們建立扎實的知識基礎、培養批判思考與責任意識；在使用 AI 的同時，期許我們能保持對世界的好奇、對事實的追求，以及對人的關懷，成為能將 AI 當成優秀助手，與 AI 協作、能監督 AI，也能在必要時拒絕 AI 建議的自主學習者與新世代公民。

教育部部長



謹識

2026年9月

● 主編序

人工智慧 (Artificial Intelligence，簡稱 AI) 正快速改變世界，也正在改變學習。

從生成式 AI 協助完成文字創作、圖像生成與知識整理，到智慧系統參與教育、醫療、交通、產業與公共治理，人類正進入一個由數據、演算法與智慧系統共同形塑的新時代。對今日的學生而言，AI 不再只是未來科技，而是已經存在於日常生活、學習歷程與社會互動中的重要夥伴。

面對這波科技變革，我國近年積極推動數位轉型與 AI 教育發展。教育部透過「推動中小學數位學習精進方案」持續建構數位學習環境，並結合因材網及 AI 學習夥伴等智慧學習工具，協助學生發展自主學習能力。同時發布《臺灣中小學教師與學生 AI 素養框架》，引導學生逐步建立理解、應用、評估與創新運用 AI 的核心能力，培養立足未來世界所需的重要素養。

另一方面，立法院三讀通過《人工智慧基本法》，揭示我國 AI 發展應以人為核心，並兼顧公平、透明、安全、隱私保護、責任治理與永續發展等核心價值。這項重要立法不僅代表臺灣邁向 AI 治理的新里程碑，也提醒教育工作者：AI 教育的目的，不只是培養能操作科技的人，更是培養能理解科技、善用科技、監督科技，並願意承擔社會責任的公民。

本手冊從學生的真實生活經驗出發，循序漸進引導讀者理解人工智慧的本質與限制。第一章介紹 AI 與大型語言模型的運作原理，理解演算法、數據與預測的基本概念；第二章進一步學習如何透過有效提問與人機協作，善用 AI 作為學習與創作的夥伴；第三章聚焦於資訊辨識與數位思辨，探討深偽、幻覺、假訊息與過濾氣泡等議題，培養學生查證與批判思考能力；第四章則從學術誠信、個資保護與 AI 倫理出發，引導學生思考科技發展背後的人文價值與責任。第五章進一步探討 AI 如何成為個人化學習的重要助力，協助學生建立自主學習與終身學習能力；第六章則將視野擴展至 AI 時代的數位公民素養，帶領學生理解演算法對個人認知、公共參與及民主社會的影響，培養在智慧社會中與 AI 共存共榮的能力。

這六個章節不只是 AI 知識的介紹，更是一條從「認識 AI」到「善用 AI」、從「理解 AI」到「反思 AI」、從「使用 AI」到「與 AI 共創」的學習旅程。

我們相信，未來真正重要的能力，不只是操作工具的能力，而是提出問題、辨識真偽、理解價值、解決問題，以及與他人合作創新的能力。AI 可以協助我們搜尋資訊、生成內容與分析資料，但無法取代人類的好奇心、判斷力、同理心與責任感

因此，本手冊最終希望培養的，不是依賴 AI 的人，而是能夠駕馭 AI 的人；不是被演算法牽引的人，而是能夠保持自主思考的人；更不是將決定權交給科技的人，而是能夠運用科技創造公共價值的人。

期待每一位讀者都能透過本手冊的學習，在快速變遷的 AI 時代中持續保有探索世界的熱情、面對未知的勇氣，以及獨立思考的能力，成為兼具科技素養、人文關懷與社會責任的未來公民。

中小學 AI 之學習應用手冊編輯群 謹識

2026年9月

AI時代，最重要的能力不是 會用AI，而是會思考

親愛的同學們：

你是否曾經和AI聊天、請AI幫忙寫文章、畫圖片，或解決功課上的問題？

也許你已經發現，現在的AI好像越來越厲害了。它會回答問題、寫故事、翻譯語言、生成圖片，甚至還能陪你聊天。有時候，它回答得又快又完整，讓人忍不住覺得：「哇！AI是不是什麼都知道？」不過，我們想先告訴你一件很重要的事情：

如果AI可以幫你完成很多事情，那麼未來人才最重要的能力，究竟是什麼？

是更快地找到答案嗎？是更熟練地操作工具嗎？還是擁有一顆願意思考、懂得判斷、能夠分辨真假的大腦？我們相信，答案是第三個。

人工智慧正在改變世界，但人類的價值從未因此消失。AI可以幫助我們搜尋資訊，卻無法代替我們判斷資訊是否正確；AI可以產生文章，卻無法代替我們思考文章背後的意義；AI可以模仿語言與創作風格，卻無法取代人類的創意、情感、責任與價值選擇。因此，教育部推動AI教育時，特別強調「以人為中心的思維」。我國《人工智慧基本法》也指出，人工智慧的發展應以增進人類福祉為目的。換句話說，科技存在的價值，不是讓人類停止思考，而是幫助人類變得更有能力。

這本手冊正是基於這樣的理念而編寫。

在這裡，你不會只學到AI工具怎麼使用，更重要的是學習如何理解AI、善用AI、評估AI，以及與AI合作學習。你將認識人工智慧的運作原理，了解AI為什麼會出錯、為什麼會產生幻覺；你將學習如何設計更有效的提問，讓AI成為學習夥

伴；你也會接觸深偽技術、假訊息、過濾氣泡與演算法偏見等重要議題，學習在資訊爆炸的時代保持清醒的判斷力。此外，你還將思考 AI 倫理、數位公民責任，以及如何運用 AI 協助自己持續學習與成長。

閱讀這本手冊時，我們有幾個小建議送給你。

- 一、不要急著追求快解。**看到問題時，先試著自己思考，再與 AI 對話。因為學習最重要的收穫，往往不在答案本身，而是在思考的過程。
- 二、把 AI 當成學習夥伴，而不是答案製造機。**當 AI 提供內容時，試著追問：「為什麼？」、「有沒有不同觀點？」、「這個資訊可靠嗎？」。好的學習者不是照單全收，而是懂得提問與驗證。
- 三、保持懷疑，也保持好奇。**AI 有時候會犯錯，甚至一本正經地說出不存在的事情；但同時，它也能成為探索知識與創意的重要工具。保持好奇，能讓你持續學習；保持懷疑，則能讓你避免被誤導。
- 四、動手實作。**本手冊設計了許多案例、活動與挑戰任務。請不要只閱讀文字，更要親自操作、驗證與反思。因為真正的 AI 素養，不是知道概念，而是能在真實生活中做出正確判斷。

未來的世界，AI 將無所不在。但我們相信，最有價值的從來不是 AI 本身，而是懂得運用 AI 的人。期待你在閱讀這本手冊的過程中，保持好奇、勇於提問、樂於查證，不只是學會使用 AI，更能學會與 AI 合作、與 AI 共學，並始終保有屬於自己的獨立思考與價值判斷。願你在快速變動的 AI 時代裡，成為一位能夠理解科技、善用科技、引領科技，而不是被科技引領的人。

現在，就讓我們一起出發吧！

認識吱可思小夥伴



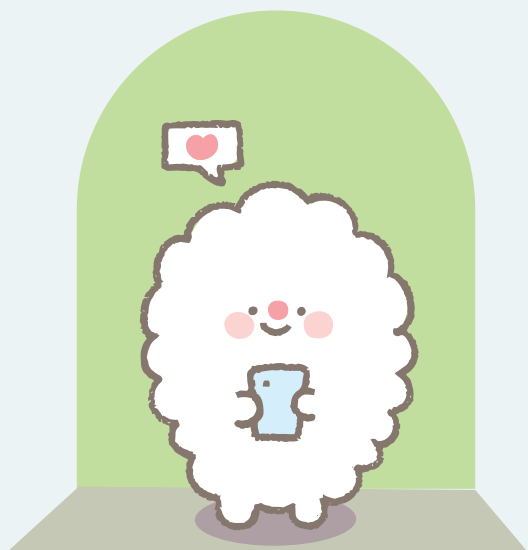
吱吱

喜歡各種嘗試，愛交朋友，有點自戀。有經營直播頻道，富正義感但有時候口快了一些。



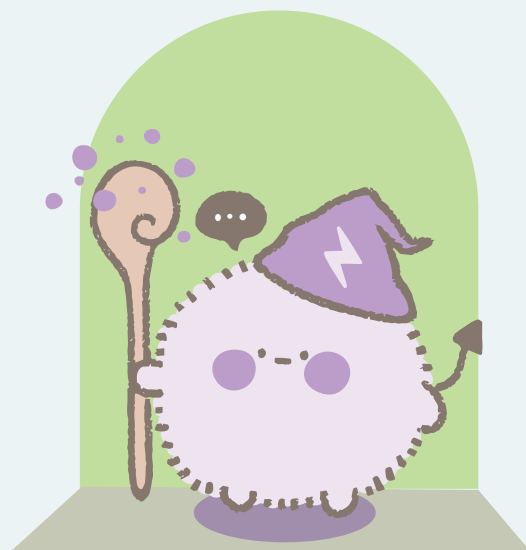
可可

沉穩謹慎，有點神經質。屬於網海中潛水型的人，總是ㄅㄅ，重視隱私。



思思

愛作白日夢，非常容易相信別人，有禮貌，被打了還會跟你說謝謝的那種人。很容易出包，還好身邊有兩位好朋友。



歐匿

可愛又神祕。特殊技法是決定控制資訊的匿名性與否，來獲得更多的觸及率。

目錄

	壹	01	AI 真的懂你嗎？
18	貳		AI 協作與學習效率
	參	40	事實查核與破解幻覺
57	肆		AI 倫理與學術誠信
	伍	72	AI 助你高效學習
97	陸		AI 世代下的數位公民

最後面附錄有「互動式教材」的QRcode可以掃描喔！
一起來動動手、動動腦、想一想，幫助你更有效率學習手冊內容！

AI 真的 懂你嗎？

從數據到決策、認識演算法的底層邏輯

「這不是在教你變成 AI 工程師，
而是教你在 AI 時代保有判斷力。」

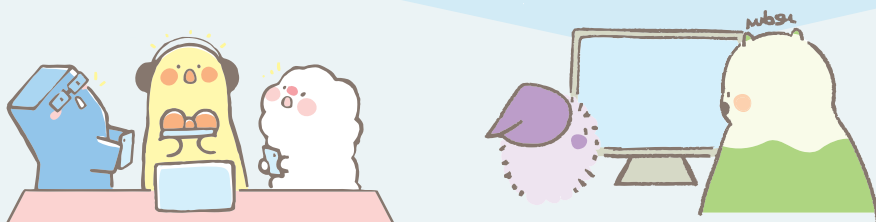
AI 很像懂你的想法，但其實是在預測你的期待。

這一章，教你看穿 AI 的底層邏輯，讓你成為駕馭 AI 的人，而不是被 AI 帶著走的人。



📌 本章學習目標：

- 了解機器學習如何從數據中找到規律
- 看懂 LLM 「預測下一個詞」的機率機制
- 判斷數據品質如何影響 AI 的輸出結果
- 學會使用「AI 回答風險檢查表」批判性評估 AI 生成內容
- 認識認知卸載風險，避免思考被 AI 取代
- 了解通用 AI 與教育 AI 的差異
- 掌握批判型 Prompt，學會挑戰 AI



一、你的日常，AI 的數據

你上週在 IG 上的哪些貼文點了讚？YouTube 推了什麼影片給你？Netflix 怎麼知道你喜歡哪類影集？

這些行為，對你來說只是日常，但在 AI 的世界裡，它們共同構成了一份關於你的**數位足跡 (Digital Footprint)**。AI 不認識你這個人，它認識的是你留下的數據紀錄。

🎯 你有沒有過這種經驗？

📺 IG Reels 情境

你只是點了一次健身影片，之後 Reels 都充滿健身內容？

🔗 這就是演算法正在幫你「分類」。

📺 YouTube 同溫層情境

當演算法一直推送你喜歡的內容，你可能越來越少看到不同觀點，形成「同溫層」。

🔗 你的世界觀，可能在不知不覺中被框限。



🔗 圖片生成提示詞參考：

Simple infographic showing a brain connected to a computer neural network, with data flowing in and predictions coming out, flat colorful illustration, blue and teal color scheme, educational style for middle school students, no text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

🔄 人類 vs AI：誰認識你？

👤 人類朋友認識的你

- 你的個性、心情、成長背景
- 你說的話背後的「意義」
- 你今天狀態好不好
- 你真正在意什麼

🤖 AI「認識」的你

- 你按讚的數據紀錄
- 你搜尋過的關鍵字
- 你觀看影片的秒數
- 你的消費與點擊行為

💡 核心一句話

你在 AI 眼中 = 一組數字，而 AI 不是在「了解你」，它是在用你留下的數據紀錄預測你接下來可能做什麼。

二、AI 怎麼「學習」？機器學習的本質

(一) 傳統程式設計 vs 機器學習

⚙️ 傳統程式設計

人類寫規則 → 電腦執行
例如：如果郵件包含「免費」和「立即點擊」→ 標記為垃圾郵件

✖ 問題：規則一改，系統就失效

🤖 機器學習

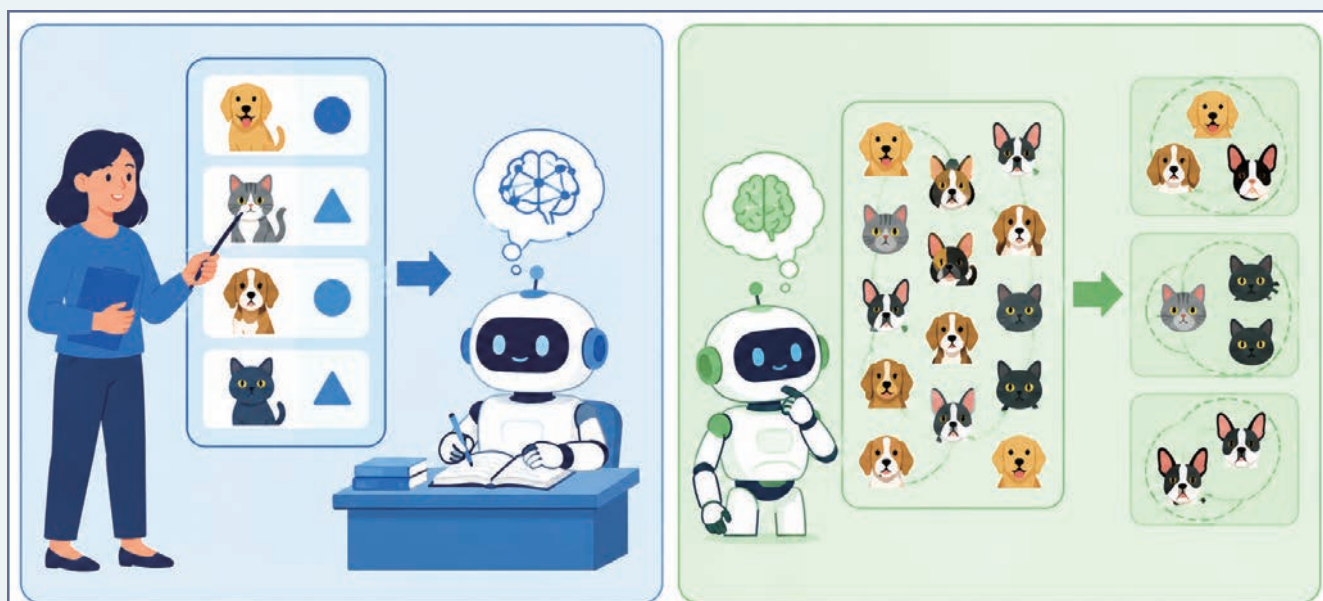
數據 → AI 自動找規律 → 對新數據做預測
例如：給 AI 幾萬封標記好的郵件，它自己學會辨別垃圾郵件

☑ 優勢：會自動適應新情況

💡 核心一句話

AI 沒有真正「理解」任何東西——它只是在數據中找到讓預測誤差最小的數學規律。

(二) 監督式學習 vs 非監督式學習



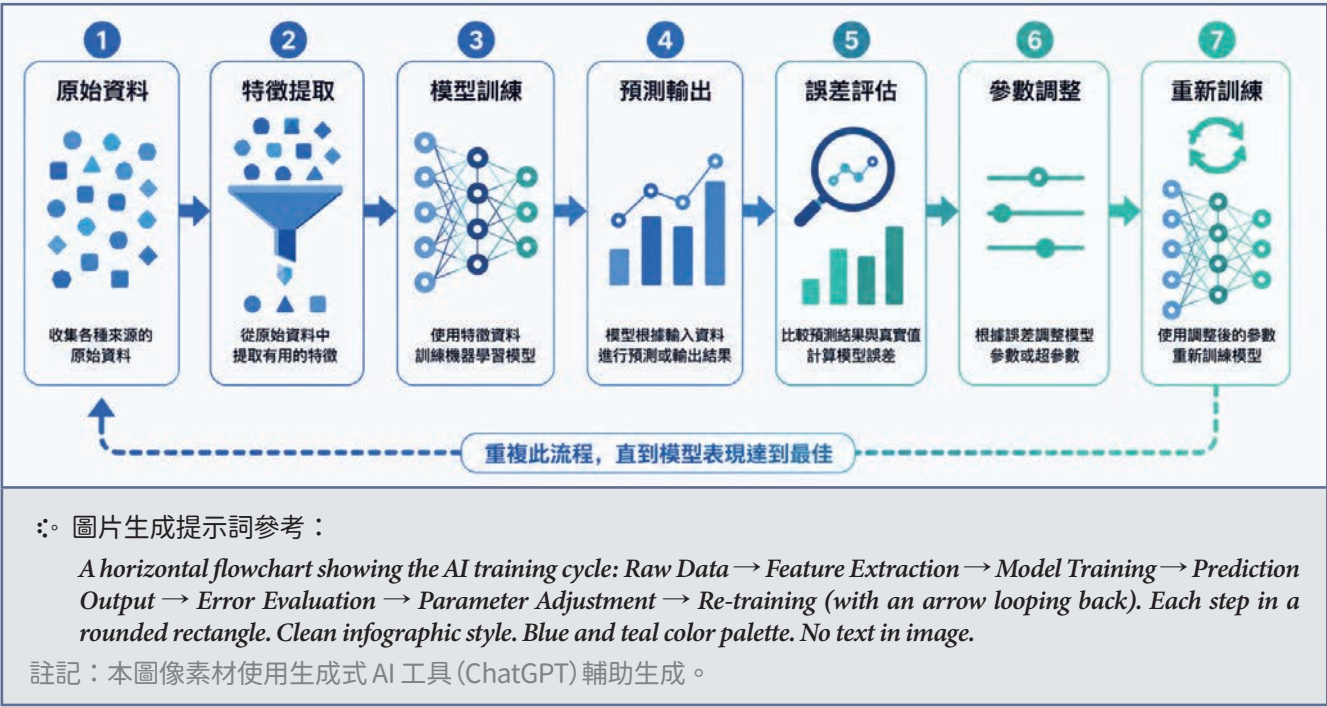
∴ 圖片生成提示詞參考：

Two side-by-side diagrams. Left: Supervised learning shown as a teacher giving labeled examples (dog/cat images with labels) to a student robot. Right: Unsupervised learning shown as a robot discovering patterns in unlabeled data clusters on its own. Flat illustration style, blue and green colors, no text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

項目	監督式學習	非監督式學習
名詞解釋	有老師改作業——給「問題」和「標準答案」，AI 學習兩者之間的規律	沒有老師、沒有答案，AI 自己在數據中找隱藏的分群結構
生活例子	人臉辨識、健保診斷輔助、信用評分	串流平臺自動分類音樂、電商消費者分群、社群推薦演算法
高中生注意！	要確認「標準答案」的來源，及標籤本身是否包含偏見	AI 不知不覺把你分進某個「族群」，你的資訊世界可能因此被框限

(三) AI 訓練流程圖



🗣️ 用一句話說明 AI 訓練

AI 就像反覆練習的學生——每次做錯就調整，直到預測越來越準，但它「記住」的是統計規律，不是真正的「理解」。

(四) 各群科的機器學習應用範例

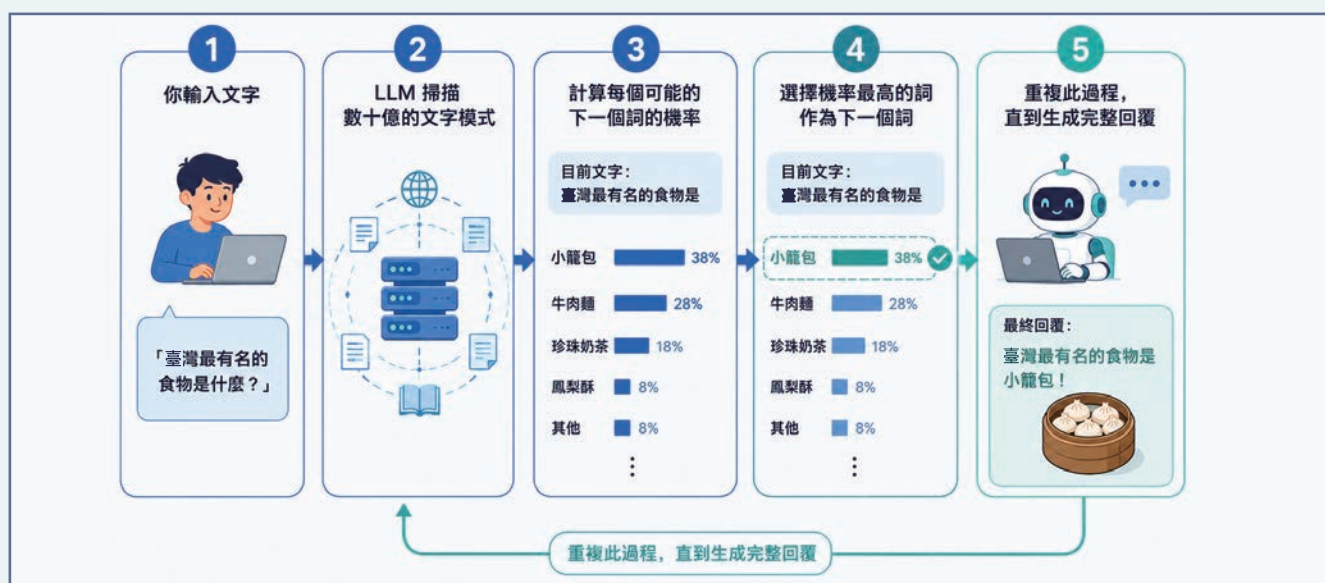
群科	AI 機器學習應用	思考方向
🏢 商管群	AI 分析消費者購買行為，自動分群推薦商品	這份分群的數據準確嗎？有偏見嗎？
💻 資電群	AI 輔助程式除錯，自動學習常見錯誤模式	AI 建議的解法，你理解背後的邏輯嗎？
🎨 設計群	AI 根據風格偏好自動生成視覺草圖	AI 的「美感」是從哪些數據訓練出來的？
🍴 餐旅群	AI 分析顧客評論，學習推薦菜單與行銷方向	顧客評論數據是否代表所有消費者？

三、LLM：預測下一個詞的機器

(一) LLM 的核心秘密——一個超簡單的概念

ChatGPT、Claude、Gemini 這些大型語言模型 (Large Language Model, LLM)，最根本的運作原理出奇地簡單：「**給定目前已有的文字，預測下一個最可能出現的詞。**」就這樣，不斷重複這個步驟，就會產生看起來流暢自然的回答。

(二) LLM 運作流程圖



∴ 圖片生成提示詞參考：

A step-by-step flowchart showing how LLM works: Step 1: You input text ('臺灣最有名的食物是什麼?') → Step 2: LLM scans billions of text patterns → Step 3: Calculates probability for each possible next word → Step 4: Selects the highest probability word → Step 5: Repeats until response is complete. Show probability percentages as bar charts next to candidate words. Clean flat design. Blue and teal palette. No Chinese text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

(三) 人類 vs AI：思考方式大不同

人類回答問題時	LLM 回答問題時
● 理解問題的「意義」 ● 根據知識推理 ● 知道自己說的是否正確 ● 有情感和意圖 ● 能判斷道德對錯	● 識別文字的統計模式 ● 根據機率關聯生成文字 ● 不知道自己說的是否正確 ● 沒有情感，只有數學運算 ● 無法真正判斷道德對錯

✂ AI 不是真的懂你，而是在預測下一句最可能出現的話。

(四) 幻覺——AI 為什麼會「說謊」？



∴ 圖片生成提示詞參考：

A step-by-step flowchart showing how LLM works: Step 1: You input text ('臺灣最有名的食物是什麼?') → Step 2: LLM scans billions of text patterns → Step 3: Calculates probability for each possible next word → Step 4: Selects the highest probability word → Step 5: Repeats until response is complete. Show probability percentages as bar charts next to candidate words. Clean flat design. Blue and teal palette. No Chinese text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

📖 真實案例：小論文的幻覺陷阱

小安使用 AI 幫忙整理小論文的參考資料，AI 提供了一篇標題看起來很專業、格式也正確的「研究論文」，連作者姓名和期刊名稱都有。但小安實際查詢後，這篇論文根本不存在——AI 把看似合理的詞語拼湊在一起，產生了一個「假引用」。

🔗 **AI 可能會「捏造引用」，這不是故意說謊，而是 LLM 的本質：**它只負責生成高機率的詞語序列，不管那是不是真實存在的資訊。

✅ **解決方法：**所有引用資料都必須自行查證並閱讀比對。先確認作者、出版來源、發表日期與原文內容，再以 Google 學術搜尋 (Google Scholar)、臺灣博碩士論文知識加值系統、圖書館資料庫等可靠來源交叉比對；不要只看搜尋摘要或 AI 整理的結果。

(五) 其他重要 LLM 的特性

特性	說明	對你的影響
溫度	「溫度」是部分 AI 工具提供的設定值，用來控制回答的隨機程度。數值較低時，AI 較傾向選擇常見、穩定的說法；數值較高時，回答會有較多變化與創意，但也較可能偏離題意或出現不精確內容。	同一個問題問兩次，可能得到不同答案，這是正常設計。
知識截止日期	LLM 的訓練數據有最後時間點，之後的事它不知道。	問 AI 最新新聞、最新法規要非常小心，結果可能已過時。
繁體中文資料稀缺	全球 LLM 訓練數據多為英文，繁體中文和臺灣在地資料較少。	詢問臺灣本土文化、法律、俚語時，AI 可能出錯。

TW 臺灣專屬資源：TAIDE



國家實驗研究院開發了 TAIDE (Trustworthy AI Dialogue Engine)，專注於建立以繁體中文和臺灣在地數據為核心的語言模型，降低文化偏誤。這是目前最適合臺灣學生使用的在地化 AI 工具之一。

四、垃圾進，垃圾出——數據品質的重要性

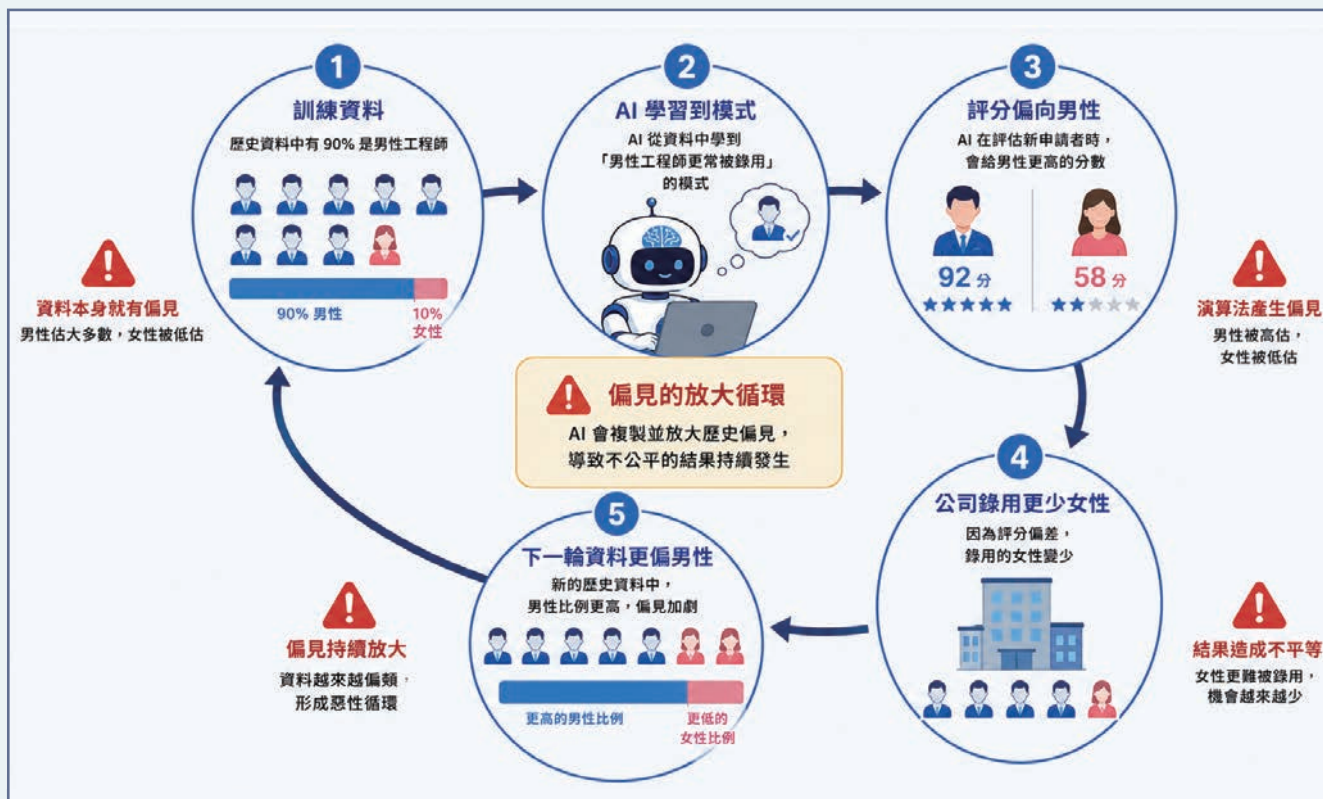
數據問題類型	說明	真實案例
數量不足	訓練樣本太少，AI 無法學到普遍規律	只用100張照片訓練人臉辨識，準確率極低
代表性不足	數據只涵蓋部分群體	主要用西方人臉訓練的系統，對亞裔準確率較低
標籤錯誤	「正確答案」本身有誤	醫療數據記錄錯誤，AI 學到錯誤診斷模式
歷史偏見	數據反映了過去的不公平	招募 AI 因歷史數據偏男性，而學會偏好男性履歷
繁中資料稀缺	訓練數據多為英文	LLM 處理臺灣本土文化、俚語時表現較差

GIGO 法則：Garbage In, Garbage Out

電腦科學界的鐵律：

不論 AI 模型的架構多先進，若訓練數據本身有問題，輸出結果一定也有問題

偏見如何被放大——演算產出的歧視



📌 圖片生成提示詞參考：

A circular diagram showing the amplification of bias: Training data (90% male engineers) → AI learns pattern → Scores male applicants higher → Company hires fewer women → Next round of data becomes even more male-biased → loops back to training data. Show each step in a circle with arrows connecting them. Warning icons. Flat illustration style. No text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

🔗 AI 不只會「重現」數據中的偏見，還可能「放大」它

這就是「演算法歧視 (Algorithmic Discrimination)」的形成機制。當你使用 AI 時，這份偏見可能已經內嵌在它的每一個回答裡。

五、AI 回答風險檢查表——判斷前先確認一遍

每次看到 AI 的回答，先暫停一下，用這份檢查表評估看看：

每次使用 AI 前，問自己這些問題	確認
① 這個回答是否涉及最新資訊（近六個月內）？	☐
② AI 是否提供了可查核的來源？	☐
③ 這個主題是否涉及法律或醫療建議？	☐
④ 這份回答是否可能含有偏見（性別、種族、文化）？	☐
⑤ 這個說法是否有數據支持，還是只是聽起來合理？	☐
⑥ 這個資訊是否需要交叉查證（至少兩個可信來源）？	☐
⑦ 如果這個資訊是錯的，後果嚴重嗎？	☐

最終判斷欄	
☒ 哪些部分我接受？ （列出你認為可靠、已查證的內容） 	☒ 哪些部分我保留？ （列出需要再查證或存疑的內容）
 為什麼我這樣判斷？（寫下你的理由） 	

六、通用 AI vs 教育 AI —— 你在用哪一種？

不同 AI 的設計目的不同，使用前要先知道你用的是哪一種：

類型	代表工具	設計重點	可能的風險
 通用 AI	ChatGPT、Gemini、Claude、Perplexity、Bing AI	回答廣泛問題，適用各種情境	回答廣博但不一定精準，可能有幻覺或偏見
 教育 AI	教育部因材網e度、Khanmigo	針對教學情境設計，配合課程	功能較窄，但內容較符合學習需求

✎ 高中生使用 AI 的建議優先順序

- 學習概念理解 → 教育 AI 或通用 AI 搭配教科書使用
- 小論文資料查找 → 以學術資料庫為主（不可只用 AI）
- 練習與即時反饋 → 通用 AI（但需自行查證重要資訊）
- 創意與圖像設計 → 圖像生成 AI（注意著作權與倫理）

七、認知卸載——思考不能外包

什麼是認知卸載？



📌 圖片生成提示詞參考：

Two side-by-side illustrations. Left: A student actively thinking, brain glowing, writing notes while looking at an AI screen — labeled 'Learning WITH AI'. Right: A student passively copying from AI screen directly to paper, brain dark and empty — labeled 'Outsourcing to AI'. Contrasting colors: green and warm for left, gray and cold for right. Flat illustration. No text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

☑ AI 幫你學習（正確用法）

- 看懂概念後，能用自己的話解釋
- 能舉出新的例子說明
- 沒有 AI 幫忙還是能解決問題
- 能反駁 AI 說的話
- 會質疑 AI 的答案

✗ AI 幫你完成（錯誤用法）

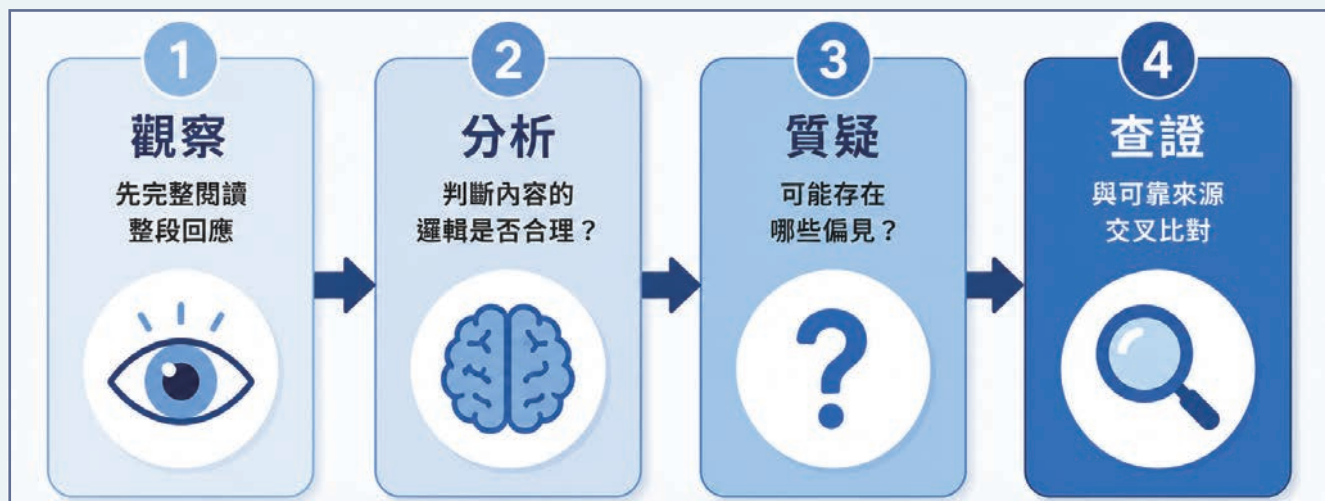
- 直接複製 AI 的回答
- 沒有真正理解
- 離開 AI 就不會
- 直接接受 AI 觀點
- 從不懷疑 AI 說的話

⚠ 虛假精熟 (False Mastery)

你完成了任務，但學習並未真正內化——你以為自己學會了，其實只是讓 AI 做了；考試的時候，你才發現什麼都不會。

✍ AI 寫了報告 ≠ 你學會了報告的主題。

八、AI 判讀四步驟——不要被「看起來合理」騙了



∴ 圖片生成提示詞參考：

A horizontal step-by-step flowchart with 4 steps: Step 1 (eye icon): 'Observe - Read the full response first'; Step 2 (brain icon): 'Analyze - Is the logic sound?'; Step 3 (question mark icon): 'Question - Where might biases exist?'; Step 4 (magnifying glass icon): 'Verify - Cross-check with reliable sources'. Each step in a rounded rectangle connected by arrows. Blue gradient from light to dark. Icons only, no text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

步驟	你要做的	問自己這個問題
① 觀察	讀完整個 AI 回答後，不要立刻反應	AI 說了什麼？ 我真的讀完了嗎？
② 分析	檢視邏輯是否成立、數據來源是否存在	這個結論合理嗎？ 它有引用具體來源嗎？
③ 質疑	找出可能的偏誤、被省略的觀點、不確定的部分	哪些地方可能有偏見？ 有沒有不同立場的聲音？
④ 查證	查證關鍵事實，回到原始資料或可信來源	我能從其他可靠來源確認這件事嗎？

✍ 觀察、分析、質疑、查證，再做負責任的決策。
這個流程不只適用於 AI，也適用於新聞、社群貼文與網路謠言。

九、批判型提示詞 (Prompt) —— 不只會用 AI，更要挑戰 AI

什麼是批判型提示詞 (Prompt)？

高中階段的重點不是「會問 AI」，而是「會挑戰 AI」，批判型提示詞讓你主動引導 AI 揭示自己的侷限與不確定性。

🕒 一般提示詞 (被動使用)	🔗 批判型提示詞 (主動思考)
「幫我解釋氣候變遷。」 → AI 給你一段答案，你接受。 但是，你不知道 AI 說的哪些部分是確定的，哪些是不確定的。	「解釋氣候變遷後，請列出： ● 你最不確定的地方 ● 可能存在偏見的部分 ● 需要查證的資訊 ● 不同立場可能怎麼看」 → 多了一段 AI 協助反思的步驟。

🔗 萬用批判型提示詞模板

在你回答完後，請你分析這個回答：

- ① 你最不確定的地方是哪裡？
- ② 你的回答中可能存在哪些偏見或立場傾向？
- ③ 哪些資訊需要讀者自行查證？
- ④ 如果換一個完全不同立場的人，他們會怎麼看這個問題？

高中生常用情境的批判型提示詞

情境	批判型提示詞這樣說
小論文寫作	「幫我分析這個論點後，請指出哪些主張可能存在爭議，以及我應該補充哪些反駁觀點。」
學習歷程撰寫	「幫我描述這份學習心得後，請告訴我哪些部分聽起來太空泛，需要加入更具體的例子。」
查詢資料	「提供這個主題的資訊後，請說明你的回答有哪些部分是截止日期前的舊資訊，哪些可能已改變。」
考試準備	「解釋這個概念後，請設計三個能測試我是否真正理解（而非背起來）的問題。」

十、媒體識讀與法律責任——AI 生成 ≠ 不用負責

(一) 深偽與假訊息

⚠️ 看起來很真實的假影片

深偽技術 (Deepfake) 製造的影像



別被騙了！
深偽影片可能用來散布假訊息、操縱輿論或詐騙。


看起來很真實


語音也能偽造


容易被散播


用途不良

🔍 如何辨識深偽影片？

仔細檢查這些可疑的破綻



- 邊緣不自然**
臉部或頭髮邊緣模糊、有鋸齒
- 眨眼不正常**
眨眼頻率異常或完全不眨眼
- 嘴型對不上**
嘴型與語音不同步或出現扭曲
- 畫面有破綻**
局部模糊、變形或像素異常

查證才安心！


查來源
確認消息來源是否可靠


多方比對
與其他媒體或原始影片比對


注意時間
確認影片發布時間是否合理


使用查證工具
善用事實查核網站或反深偽工具

保持懷疑
不輕信、不轉傳
守護資訊安全！

📌 圖片生成提示詞參考參考：

Split illustration: Left panel shows a realistic-looking deepfake video of a public figure with a red 'FAKE' watermark overlay. Right panel shows a magnifying glass revealing digital artifacts and manipulation marks in the image. Warning icons and fact-check symbols. Flat illustration style, red and orange warning colors. No real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

🎞️ 深偽影片	📰 AI 假新聞
AI 生成的名人影片，看起來像真的，但實際從未發生。近年已有多起利用深偽影片詐騙、散布假訊息的案例。	AI 可以快速生成大量「看起來像新聞」的內容——有標題、有段落、有引用，但全部是捏造的。
🔗 轉發前要查證，「看起來真實」≠真實發生。	🔗 查證方法：看原始出處 + 比對多個可信媒體。

(二) 高中生使用 AI 的法律責任——AI 生成 ≠ 不用負責

🔗 人類使用 AI 生成內容，仍須負責

AI 不是法律或道德責任的主體。使用、修改、發布或轉傳 AI 生成內容的人，必須先查證內容、尊重他人權利，並為造成的影響承擔責任。

高中生常見情境	可能涉及的法律問題	後果
用 AI 生成同學照片惡搞	肖像權、名譽權、數位霸凌	可能面臨學校記過、民事賠償，嚴重者涉及刑事責任
用 AI 生成不實資訊並轉傳	誹謗罪、散布不實訊息	即使你只是轉傳，仍可能承擔法律責任
用 AI 生成的圖片當作業繳交	學術誠信問題、著作權爭議	可能被判定作弊，並被要求重做
未標示 AI 生成而發布影音	欺騙公眾、著作權不透明	在臺灣，部分平臺已要求強制標示 AI 生成內容

使用 AI 的基本原則

- 不得生成真實人物的虛假影像或言論
- 分享 AI 生成內容時，應清楚標示「AI 生成」
- 不得使用 AI 散布謠言或騷擾他人
- 學術作品使用 AI 必須依學校規定揭露
- 上傳他人照片進行 AI 加工前，需取得本人同意

學習檢核表

學習目標	我能做到 	還需練習 
1. 我知道監督式學習和非監督式學習的核心差異是什麼。	☐	☐
2. 我知道 LLM「預測下一個詞」的機率機制是怎麼運作的。	☐	☐
3. 我知道什麼是「幻覺」以及它為什麼會發生。	☐	☐
4. 我知道 GIGO 法則如何影響 AI 的輸出品質。	☐	☐
5. 我知道通用 AI 和教育 AI 有什麼不同，也知道何時該用哪個。	☐	☐
6. 我知道什麼是「認知卸載」，也知道如何避免陷入「認知卸載」。	☐	☐
7. 我知道使用 AI 生成內容可能涉及的法律責任。	☐	☐
8. 看到深偽影片或 AI 假新聞，我知道如何辨識和查證。	☐	☐

挑戰任務

任務一 | 解剖一個 AI 推薦系統

選擇你日常使用的任何一個推薦平臺（IG Reels、YouTube、Spotify、Netflix）。

- 它可能使用了什麼類型的機器學習（監督式或非監督式）？
- 它大概在使用什麼數據紀錄來訓練？
- 這份數據可能存在什麼偏見或盲點？
- 這個系統有沒有讓你的資訊世界變窄？請舉一個具體例子。

任務二 | 測試 LLM 的幻覺

設計三組問題，測試 AI 的極限：

- 問一個你知道正確答案的具體事實，看 AI 是否答對
- 請 AI 提供一篇學術論文的引用，再去查那篇論文是否真實存在
- 問 AI 臺灣特有的文化或法律問題，評估它的正確性

記錄結果，並寫下：這個測試改變了你對 LLM 的信任程度嗎？

任務三 | 批判型提示詞實戰

選擇你目前最困惑的一個學科概念，用批判型提示詞和 AI 對話至少三輪：

- 第一輪：請 AI 解釋這個概念
- 第二輪：請 AI 說明它最不確定的地方與可能的偏見
- 第三輪：請 AI 設計三個能測試你是否真正理解的問題

最後離開 AI，用自己的話寫下你的理解（不能複製 AI 的回答）。

🌐 延伸閱讀與資源

資源名稱	說明	如何取得
TAIDE 計畫官網	臺灣在地化 LLM，了解繁中 AI 的發展	
《What Is ChatGPT Doing?》 ／Stephen Wolfram	LLM 技術最深入淺出的英文解析（免費線上閱讀）	
《臺灣中小學教師與學生 AI 素養框架》	教育部官方 AI 素養架構文件	
Deepfake 查核工具 （FakeCatcher 等）	幫你辨識深偽影片的實用工具	

📖 本章核心觀念回顧

1. AI 不是真的懂你，而是在預測——這是 LLM 的本質。
2. 數據決定 AI 的「世界觀」—— GIGO、偏見、繁中資料稀缺都會影響輸出。
3. 幻覺、截止日期、繁中偏誤——三大使用風險，每次查詢都要留意。
4. 風險檢查表和四步驟判讀——你的 AI 使用安全網，養成習慣。
5. 批判型提示詞——不只會用 AI，更要挑戰 AI。
6. AI 生成 ≠ 不用負責——法律責任不會因為「是 AI 做的」而消失。

💡 最重要的一句話

AI 是工具，判斷力是你的。在 AI 時代，最重要的不是讓 AI 替你想，而是你能清楚知道 AI 在做什麼、說什麼、以及哪些地方你需要保持懷疑。



AI 協作 與學習效率

讓 AI 成爲你的學習夥伴

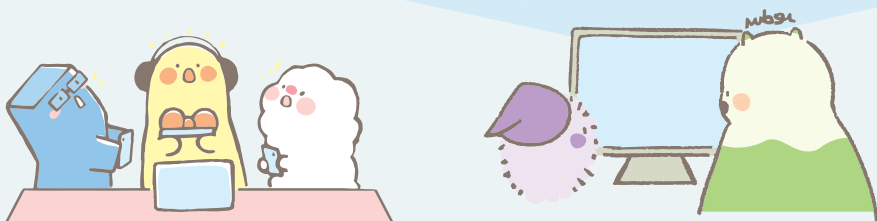
「AI 的價值，不是替你思考，
而是幫助你思考得更深入。」

AI 很擅長整理、追問、模擬對話，但你才是決定方向的人。
這一章，教你用 AI 提升學習效率，讓你學得更快、想得更深。



📌 本章學習目標：

- 理解 AI 的基本運作方式
- 學會用「逐步推理」讓 AI 給出更好的答案
- 把 AI 用在你真實的學習情境：學習歷程、小論文、英文、程式、創作
- 認識 AI 的風險：版權、個資、學術誠信、深偽影片
- 培養「慢用 AI」思維，讓 AI 輔助思考，而不是取代思考





📌 圖片生成提示詞參考：

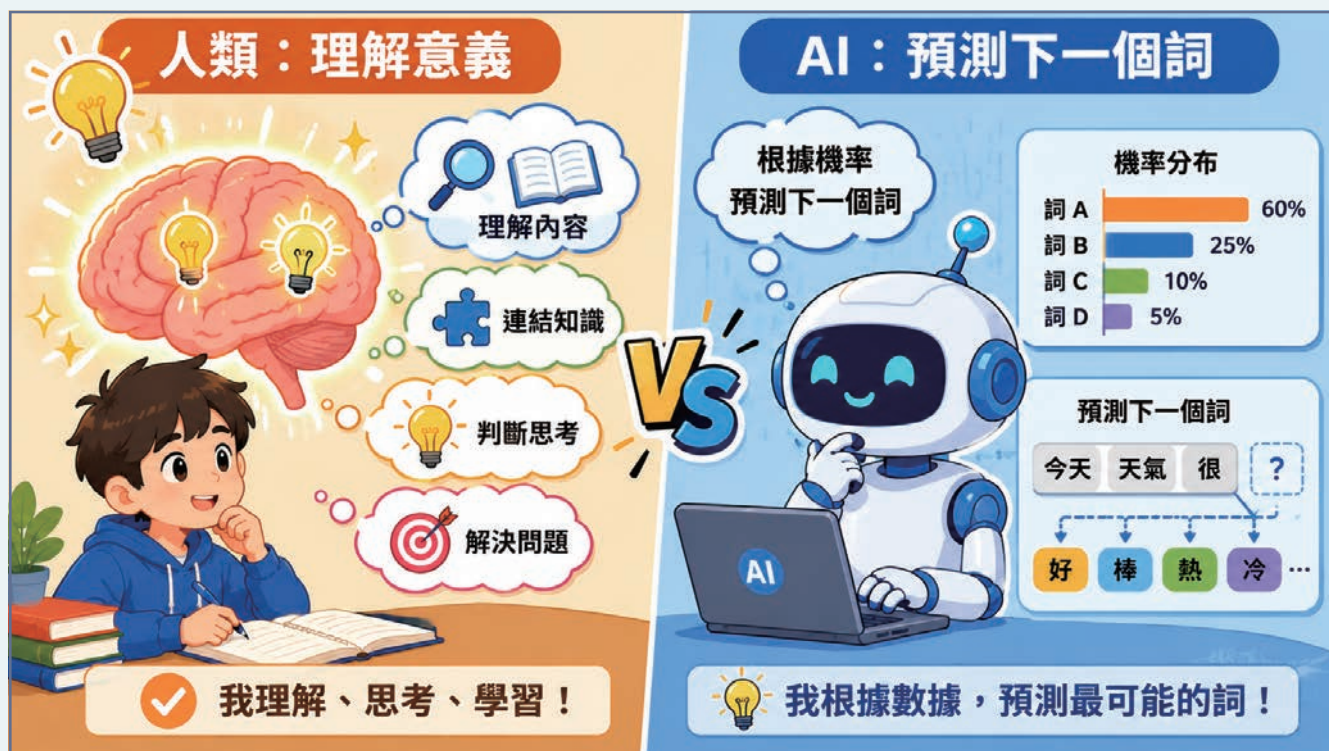
A high school student sitting at a desk, looking thoughtfully at a laptop screen showing an AI chat interface. Books, notes, and a coffee cup nearby. Warm, cozy study room atmosphere. Flat illustration style. Soft blue and orange tones. No text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

一、AI 到底是怎麼「想」的？

在開始學 AI 協作技巧之前，先搞懂一件事：AI 和人類的「思考」方式根本不一樣。

🧑 人類如何思考？	🤖 AI 如何「思考」？
● 真的「理解」意思 ● 能從生活經驗連結 ● 知道自己不懂的地方 ● 有情感與直覺	● 在海量文字中學習「什麼詞接著出現」 ● 預測最合適的下一個字 ● 不「理解」，只是「預測」 ● 沒有意識，也不知道自己的答案對不對



✎ 圖片生成提示詞參考：

Side-by-side comparison infographic: Left panel shows a human brain with light bulbs and thought bubbles labeled "understand meaning". Right panel shows an AI robot with probability charts and text prediction arrows labeled "predict next word". Clean flat design. Blue and orange colors. No text except simple labels.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

💡 AI 的白話運作方式

想像你讀了幾千萬本書和網頁，學會了「在這種情況下，下一句通常是什麼」的規律。

AI 就是這樣——它接受了超大量的文字訓練，學會預測「接下來最可能出現的內容」。

所以：

- 你問的問題越清楚，AI 就越容易預測出你需要的答案。
- AI 並不真的知道答案「對不對」——它只是預測最合理的文字。
- AI 可能很自信地說出錯的事情（稱為「幻覺」）。

🔑 重點

AI 的答案聽起來很有信心，不代表一定正確！你的查證能力比 AI 的「信心」更重要。

二、讓 AI 幫你想更深：逐步推理的力量

你有沒有發現，有時候 AI 給的答案太跳躍、理由說不清楚？其實，你可以要求 AI 「一步一步想」。這不是什麼神奇技巧，就是清楚跟 AI 說明你要它怎麼做。

(一) 為什麼「一步步說明」很重要？

AI 每回答一句話，前面已說的內容就成為它繼續回答的依據。當它被要求「先分析步驟」，每一步的推理會讓下一步更嚴謹——就像你解數學題，先列出公式再計算，比直接猜答案準確多了。

✕ 直接問答案	☑ 要求逐步思考
「我應該選理組還是文組？」 → AI 可能給一個模糊的答案	「請分別從『課程內容』、『未來出路』、『我的興趣』三個角度，各給我一個分析，再給出整合建議。」 → 分析更清楚、理由更完整

(二) 逐步推理的三種層次

▣ 層次一：加一句「請一步步說明」

最簡單的做法就是在問題後面加上這句話，AI 的回答品質通常會明顯提升。

📄 提示詞範例

一個水池有 3 個入水管，分別能在 2 小時、3 小時、4 小時注滿水池，如果同時開啟三個入水管，多久可以注滿？

請一步步說明你的推理過程，最後再給出答案。

🔗 為什麼有效？

這樣做讓 AI 把「推理過程」明確寫出來，你才能看清楚它是怎麼算的，也能找出哪裡算錯了。

AI 不是計算機，它在數學計算上可能出錯，逐步說明讓你「看得出」它哪裡錯了。

▶ 層次二：給 AI 一個示範，讓它照做

如果你想要 AI 用特定格式分析，可以先給它一個「示範」，它就能照著格式幫你解決新問題。

提示詞範例（小論文分析）

以下是一個分析示範：

問題：社群媒體是否影響青少年睡眠？

步驟 1 — 確證論點：「影響睡眠」是研究假設。

步驟 2 — 找反例：有些研究顯示影響有限。

步驟 3 — 整合：整體而言，影響因使用習慣而異。

請用相同格式，分析這個題目：

「AI 技術是否會讓高中生的寫作能力退步？」

▶ 層次三：多角度思考，自己來判斷

對於複雜的問題，可以請 AI 從不同角度各分析一次，然後你自己做最後的判斷，而不是讓 AI 直接給你結論。

提示詞範例（選系志願）

關於「我要選資工系還是電機系」這個問題，請分別從三個角度給我分析：

- 角度一：課程內容的差異
- 角度二：未來就業市場
- 角度三：適合喜歡寫程式但也對硬體有興趣的人

分析完三個角度後，不要直接給結論，讓我自己比較判斷。

🔧 動手做：逐步推理實驗

選一個你最近遇到的難題（數學、歷史、選系……都可以）：

- 先直接問 AI，記下答案
- 再加上「請一步步說明推理過程」，記下第二次答案
- 比較兩次：哪個更有說服力？你找到哪裡不同了？

把你的觀察寫下來，這就是批判性評估 AI 的第一步！

🔍 我的判斷在哪裡？

每次使用 AI 完成任務後，請在下表標示你各步驟的判斷及你自己決定的地方。

📍 第一步：提出問題

我為什麼選擇問 AI 這個問題？我想解決的真正問題是什麼？

我的紀錄： _____

📍 第二步：選擇資料

我從 AI 的回答中，選擇了哪些資訊？為什麼選這些，而不是其他的？

我的紀錄： _____

📍 第三步：修正觀點

AI 的回答讓我改變了什麼想法？哪些地方我不同意或覺得需要再查證？

我的紀錄： _____

📍 第四步：最後決策

我最終的判斷是什麼？這個決定是我做的，不是 AI 替我決定的。

我的紀錄： _____

三、高中生 AI 協作實戰：你真實會用到的情境

以下是高中生最常使用 AI 的場景，每個情境都提供實用的提示詞範例，以及你必須注意的事項。

情境一 | 學習歷程 & 小論文



∴ 圖片生成提示詞參考：

Illustration of a high school student writing a learning portfolio (學習歷程) on paper, with an AI chatbot interface on a nearby screen suggesting ideas. The student is thoughtfully writing, not just copying. Clean flat style. Green and blue tones.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

▶ 學習歷程：AI 幫你整理，但反思要自己來

學習歷程最核心的是：你「真正」參與了什麼、思考了什麼、學到了什麼，AI 可以幫你整理文字，但那份經歷是你的，不是 AI 的。

📝 有效的學習歷程提示詞範例

我參加了學校的科展，負責的部分是○○○(如：實驗設計)，以下是我的活動記錄：

(貼上你自己的筆記或描述)

請幫我整理成學習歷程的格式，包括：

● 活動描述(簡短) ● 我在其中做了什麼 ● 我遇到的挑戰

請用「我」的第一人稱，根據我提供的內容整理，不要添加我沒有提到的事情。

⚠ 學習歷程的 AI 陷阱

- ✗ 把 AI 寫的心得直接貼上，但自己根本沒參與活動
- ✗ AI 幫你「創造」你沒有經歷過的反思
- ✓ 請 AI 幫忙整理你的筆記，你再修改成自己的語氣，這才是正確的用法。

▶ 小論文：AI 找方向，你來查核

小論文最重要的是：資料來源必須真實存在，若要引用必須由你自己查證。AI 可能「編造」不存在的論文或資料！

📝 小論文研究提示詞範例

我是高中生，正在寫一篇小論文，主題是「臺灣年輕人的政治參與」。

請幫我：

- 列出這個主題的 3-5 個可能研究角度
- 說明每個角度的主要研究問題是什麼
- 建議我可以在哪些地方找到可靠的資料（例如哪個類型的網站或機構）

注意：不要幫我虛構論文標題或數據，我需要自己去查核真實來源。

🔍 查核 AI 資料的步驟

- AI 提到某個「研究發現」→ 去 Google 學術搜尋(Google Scholar) 或政府資料庫確認。
 - AI 給了某個數字 → 找原始出處，不要直接引用 AI 的話。
 - AI 推薦某篇論文 → 確認這篇論文真實存在(AI 可能編造!)
- 🔗 AI 查資料，你來查核——這個分工才對。

情境二 | 英文作文 & 面試模擬



∴ 圖片生成提示詞參考：

Illustration of a high school student practicing an English mock interview with an AI on a tablet screen. The student is speaking confidently, speech bubbles showing English phrases. Friendly and encouraging atmosphere. Flat illustration style. Blue and green tones.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

▣ 英文作文：AI 幫你修改，不是幫你寫

📝 英文作文協助提示詞範例

我是高中生，英文程度中等，以下是我自己寫的英文段落：

(貼上你自己的文字)

請幫我：

- 指出文法錯誤，並說明為什麼錯
- 建議更自然的用法，但保留我的原意
- 不要完全改寫，讓我能夠從錯誤中學習

🔑 記住：AI 幫你修改，不是幫你代寫

先自己寫，哪怕很粗糙，再請 AI 協助修改。每次修改後，問 AI 「為什麼這樣改比較好？」——這才是學習的正確態度。

► 英文面試模擬：AI 當你的練習夥伴

升大學面試、英語口說考試，都可以用 AI 來練習。

📄 英文面試練習提示詞範例

你現在是大學英文面試的面試官，請用英文問我問題，每次問一個，等我回答後再繼續。我的面試主題是：「我為什麼對資訊工程有興趣」，如果我說錯了，請用中文提醒我哪裡可以說得更好。

重要：一次只問一個問題，讓我練習即時回答。

情境三 | 程式學習（不同群科都適用！）

AI 讓學程式的門檻變低了，但這不代表你可以「只複製程式碼，不理解邏輯」，不管哪個群科，AI 都可以是你學程式的好夥伴——重點是你真的學懂。



∴ 圖片生成提示詞參考：

Illustration of diverse high school students from different academic tracks (普高, 商管, 設計, 電機, 餐旅, 外語) each using AI on their laptop for different tasks: data analysis, visual design, debugging, menu planning, language practice. Connected by a network pattern. Flat illustration. Colorful, vibrant.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

程式設計實戰：用 AI 做一款貪食蛇小遊戲

AI 讓寫程式變得比較容易，但真正重要的不是「複製一份會動的程式碼」，而是你能不能說出：它怎麼動、哪裡錯、怎麼改得更好。

💡 什麼是 Vibe Coding ？

先用自然語言說出你想創造的「感覺」，再和 AI 一起修改、測試、理解；不是叫 AI 全部寫好，而是你和 AI 共同把作品做出來，同時把邏輯學起來。

用 AI 一步步打造貪食蛇遊戲！ ✨

從基礎 → 進階 → 高階，越做越厲害！

1. 基礎版
做出會玩的貪食蛇

- 鍵盤方向鍵控制
- 顯示分數
- 有開始、重來按鈕
- 玩法說明

我們先做一個簡易的版本，掌握蛇的移動和吃食物！

2. 進階版
加入挑戰性與更多機制

- 分數越高，速度越快
- 顯示等級
- 撞牆或撞到自己會扣生命
- 加入障礙物

太棒了！我們來加上速度、生命和關卡，讓遊戲更有挑戰！

3. 高階版
打造專屬風格與特色玩法

- 自訂主題與背景
- 手機觸控操作
- 加入音效與特效
- 特殊道具與關卡

發揮創意，做出屬於你自己的貪食蛇遊戲，超有成就感！

學習重點：

- 變數與資料
- 條件判斷
- 迴圈控制
- 碰撞偵測
- 函式設計
- 難度設計
- 狀態管理
- 模組化思維
- 美術與音效
- 使用者體驗
- 觸控事件
- 創意發想

📌 圖片生成提示詞參考：

Illustration of a high school student building a Snake game step by step with AI on a laptop. Three versions shown side by side: basic Snake game, advanced version with speed levels and lives, and a themed custom version. Retro pixel art style mixed with modern flat design. Blue, green, purple tones. No text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

▶ 自學任務：三步驟完成貪食蛇 (Vibe Coding)

層級	任務目標	學生要做到	QRcode
● 基礎	做出可以玩的貪食蛇	有開始、重來、得分、玩法說明 https://btosichen.github.io/AI-Vibe-Coding/snake_01.html	 (自學網頁教材)
● 進階	改出更好玩的機制	加入等級、速度變化、生命值或障礙物 https://btosichen.github.io/AI-Vibe-Coding/snake_02.html	
● 高階	做出自己的風格	加入主題風格、音效、手機控制或特殊關卡 https://btosichen.github.io/AI-Vibe-Coding/snake_03.html	

▶ 基礎版提示詞

📝 基礎版提示詞範例

我是高中生，剛開始學 HTML、CSS、JavaScript，請幫我設計一款簡單的「貪食蛇」網頁小遊戲。

請符合以下條件：

- 所有程式碼整合成一個 HTML 檔案
- 可以用鍵盤方向鍵控制
- 要有開始按鈕、重新開始按鈕
- 要顯示分數
- 要有簡單玩法說明
- 程式碼中請加入中文註解，讓我看得懂每一段在做什麼
- 不要使用太複雜的語法

▶ 進階版提示詞

📝 進階版提示詞範例

請在剛才的貪食蛇遊戲中，加入以下功能：

- 分數越高，蛇移動速度越快
- 加入生命值，撞牆或撞到自己會扣生命
- 加入關卡等級顯示
- 請說明你修改了哪些程式碼，以及每個功能的邏輯

▶ 高階版提示詞（選一個主題風格！）

📝 一、高階版提示詞範例

請把貪食蛇遊戲改成我自己的主題風格：

主題（選一個）：未來科技風／可愛像素風／餐旅美食風／外語闖關風

請幫我加入：

- 遊戲封面
- 音效或提示音
- 手機觸控按鈕
- 特殊道具或隱藏加分物
- 請保留中文註解，並告訴我哪些地方可以自己再修改

💡 二、各群科都能玩的主題建議範例

群科	推薦主題方向
設計群	可愛像素風：設計專屬角色造型與配色
餐旅群	餐旅美食風：蛇變成筷子，食物變成各式菜餚
外語群	外語闖關風：吃到正確單字得分，吃錯的扣分
電機與資電群	未來科技風：加入電路板視覺效果與技術細節
商管與管理群	加上排行榜功能，練習數據統計與顯示邏輯

▶ Debug 提示詞

📄 Debug 提示詞範例（貪食蛇版）

我的貪食蛇遊戲出現問題，請不要直接重寫全部程式。

請幫我：

- 找出錯誤原因
- 用高中生聽得懂的方式解釋
- 告訴我是哪一段程式造成問題
- 提供最小修改版本
- 告訴我下次怎麼避免

AI 可以幫你把想法變成程式，但真正的學習是你能看得懂、改得動、說得出來。

不是「我叫 AI 寫好了」，而是「我和 AI 一起把作品做出來，也把邏輯學起來」。

這樣才是真正的升級，不是把腦袋外包。

情境四 | 科展 & 專題研究



📌 圖片生成提示詞參考：

Illustration of a high school student working on a science project, using a tablet to consult an AI while also reading physical research papers. Beakers, charts, and a poster board visible. Flat illustration style. Purple and blue tones. No text in image.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

AI 可以在研究的多個階段幫助你，但科學研究的核心——批判性思考與嚴謹驗證——永遠需要你自己完成。

研究階段	AI 可以幫你	你必須做的
確立研究問題	腦力激盪，列出相關問題	判斷哪個問題有價值、可行
文獻回顧	解釋摘要、翻譯英文論文	親自閱讀原文，批判評估
數據分析	協助寫分析程式	驗證數據來源，確認方法正確
撰寫與修改	協助潤飾文句，檢查邏輯	確保內容原創，正確引用

情境五 | AI 與創作：改變的是方式，不是你的創意

AI 最大的影響之一，是改變了創作的方式。設計海報、寫歌詞、剪影片等，AI 讓過去需要大量時間的工作變快了，但「快」不等於「好」，你的品味、觀點、風格，才是讓作品有靈魂的東西。

✧ AI 是創意夥伴，你才是創作的主角！✧

AI 提供靈感與建議，我們加入想法與判斷，讓作品更有價值！

1 設計社製作海報

AI 協助產生草圖與配色建議



AI 生成草圖

環保行動 從我做起

我的創作決定

- ✓ 選擇更有力量與標語
- ✓ 調整顏色與字體
- ✓ 加入手繪繪圖
- ✓ 強調行動呼籲！

AI 給靈感，我來做決定！

2 音樂創作與編曲

AI 協助生成旋律與和弦建議



AI 旋律建議

AI 和弦建議：C, Am, F, G

我的創作決定

- ✓ 調整旋律節奏
- ✓ 更換和弦進行
- ✓ 加入自己的風格
- ✓ 設計副歌與橋段！

AI 給素材，我來創造感動！

3 影片剪輯與字幕

AI 協助辨識語音與生成字幕



AI 生成字幕

00:08 勇敢追夢，勇敢追夢！
00:12 勇敢追夢，勇敢追夢！
00:15 勇敢追夢，勇敢追夢！
00:17 勇敢追夢，勇敢追夢！

我的創作決定

- ✓ 修正用詞與標點
- ✓ 調整字幕時間點
- ✓ 設計字體與樣式
- ✓ 加入轉場與特效！

AI 幫處理，我來說好故事！

重點提醒：AI 可以幫助你更快開始，但真正有價值的創意，來自你的想法與選擇！ ★

∴ 圖片生成提示詞參考：

Illustration showing three creative scenarios: 1) A student in a design club working on a poster with AI-generated draft on screen. 2) A student using AI to generate melody ideas on a music app. 3) A student using AI for subtitle generation while video editing. All three show human creative decisions layered on top of AI assistance. Flat colorful illustration.

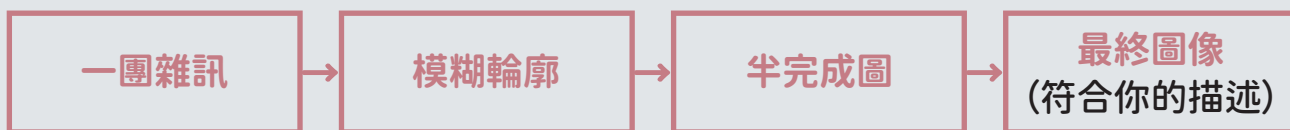
註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

創作情境	AI 的角色 / 你的角色
🎨 設計海報	AI 生成草圖概念 → 你選出最有說服力的設計 → 你決定配色、字體、訴求重點。 💡 AI 加速了「發想」，但海報說什麼、傳達什麼價值是你的責任。
🎵 音樂創作	AI 生成旋律草稿 → 你選出符合你情緒的片段 → 你決定歌曲走向與意涵。 💡 旋律可以由 AI 給靈感，但歌曲的故事是你的。
🎬 影片剪輯	AI 協助字幕生成、腳本整理 → 你決定剪接節奏、哪段要留哪段要剪。 💡 AI 處理重複性工作，創作節奏和意境由你掌控。
📝 社群貼文	AI 協助草稿文案 → 你確認是否真實代表你的觀點 → 你決定是否發布。 💡 你的價值觀與觀點，不應完全交給 AI。

▶ AI 如何生成圖片？

🔗 AI 生成圖片的過程

想像你拿著一張滿是雜訊的模糊圖，然後 AI 根據你的文字描述，一步一步把圖像「修清楚」。



這就是為什麼給 AI 越詳細的描述，AI 越能生成符合你預期圖片的原因。

💡 AI 並不真正「理解」圖像的意義，只是透過學習到的文字與圖像對應關係來生成圖片。

四、AI 風險：你必須知道的事

AI 很好用，但使用前你需要了解哪些事情可能讓你陷入麻煩，包括可能涉及的法律問題。

(一) 版權與個資風險

使用 AI 要小心！這些行為可能違法或侵權！

聰明使用 AI，保護自己也尊重他人！

1 上傳個人照片要注意！
未經同意上傳個資，可能洩漏隱私！

你的照片
可能被儲存、分享
或用於訓練 AI 模型

⚠️ 提醒：上傳前請確認平臺的隱私政策！

2 未經授權使用他人作品！
讓 AI 生成別人作品的風格，可能侵權！

他人作品

⚠️ 提醒：未經授權使用他人作品，可能違反著作權法！

3 深偽影片可能觸法！
用 AI 變造影像或聲音，可能涉及法律責任！

深偽影片
(Deep Fake)

- ✗ 誤導他人
- ✗ 侵犯名譽
- ✗ 可能構成犯罪

⚠️ 提醒：製作或散布深偽影片，可能觸犯法律！

要怎麼做？

保護自己
不隨意上傳個資或照片

尊重他人
使用素材要取得授權

遵守法律
不做違法或有為的行為

負責任地使用 AI
善用 AI，創造正向價值

AI 是工具，
使用方式決定
你是**創造者**，
還是**違規者**！

📌 圖片生成提示詞參考：

Warning illustration showing three scenarios: 1) A person uploading a photo to AI platform with a privacy warning icon. 2) An AI generating an image from someone else's artwork with a copyright symbol. 3) A deep-fake video icon with a legal gavel. Red warning colors mixed with educational blue. Flat vector style. No real faces.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

⚠️ 風險情境	可能涉及的問題
上傳他人照片到 AI 平臺	肖像權、個人資料保護、對方的同意權
用 AI 模仿知名藝術家的風格生成圖	可能涉及著作權爭議 (各國法律不同)
AI 生成圖未標示「AI 製作」就發布	可能誤導他人，涉及不實資訊
把別人的私人資訊輸入 AI	個資保護法；AI 平臺可能儲存或使用該資料

🔑 使用 AI 創作的的基本守則

- 能生成，不代表可以任意使用
- 發布 AI 圖像前，標示「AI 生成」是誠信的表現
- 上傳他人照片前，先取得對方同意
- 商業用途的 AI 圖像，需確認平臺的授權條款

(二) 深偽 (Deepfake) 與法律責任

⚠️ Deepfake 不是玩笑，可能是犯罪

📖 情境

小志用 AI 工具把同學的臉換到一段搞笑影片上，在班級群組分享，覺得只是好玩。

但他可能涉及：

- 名譽權侵害：影片讓同學在他人眼中形象受損
- 數位霸凌：未經同意使用他人形象
- 不實影像散播：影片讓人誤以為是真實情況

即使「只是好玩」、「沒有惡意」，法律責任不會因為你的動機而消失。

🔗 AI 生成 ≠ 不用負責

- 你用 AI 生成的內容，法律責任還是在你身上
- 未經同意使用他人臉部或聲音生成影片，可能違法
- 在臺灣，散播深偽影片已有相關法律規範，情節嚴重可能負刑事責任

(三) 學術誠信：常見違規案例

以下是高中生最常見的 AI 相關學術誠信問題——很多同學可能不知道這些算違規。

✖ 違規行為	問題在哪裡？
AI 幫寫小論文，直接貼上交作業	這是你的作業，不是 AI 的——你沒有實際學習，也沒有獨立完成
引用 AI 「提供」的論文，未核實	AI 可能捏造不存在的論文！你引用了假資料，報告的可信度歸零
AI 幫生成研究數據或結果	捏造研究結果，是最嚴重的學術不誠信行為之一
AI 幫寫學習歷程心得，自己沒有參與	學習歷程是你真實的成長紀錄，AI 無法代替你的經歷

🔗 AI 可以協助整理，但不能取代你的學習與誠實。

🔗 使用 AI 的底線：你必須真正理解、參與、思考過你交出去的東西。

(四) 在作業中揭露使用 AI 範例

誠實揭露 AI 的使用，是學術誠信的一部分。以下是一個你可以參考的範本：

≡ AI 使用揭露範例(可直接修改使用)

AI 工具使用聲明

本作業在以下環節使用了 AI 工具 (Claude / ChatGPT) 協助：

- 整理研究方向與初步資料搜尋
- 修正部分中文語句的表達
- 協助英文段落的語法檢查

以下內容由本人獨立完成：

- 資料的核實與引用查驗
- 分析與論點的形成
- 最終撰寫與修改

本作業的核心內容、分析與結論，均代表本人的思考與理解。

五、慢用 AI：讓你保持思考的主導權

「慢用 AI」不是要你少用 AI，而是要你「用得有意識」。AI 可以讓你更快完成事情，但如果你從來不自己思考，你的能力其實在慢慢消失。

★ 用對 AI，學得更多！★

快不代表好，慢一點，思考更多，學習才會真正發生！

不動腦的用法

快的 AI



複製貼上完成

- 不理解內容
- 沒有自己的想法
- 學習效果差

❌ 結果：壓力大、收穫少，能力沒提升！

被動依賴 AI
變得越來越懶

不理解內容
容易出錯

缺乏思考
無法解決問題

聰明學習的用法

慢的 AI

1 先思考
釐清問題與重點



我會先想：

- ✓ 我要表達什麼？
- ✓ 有哪些重點？
- ✓ 我知道什麼？
- ✓ 還需要哪些資料？

2 善用 AI
詢問、比較、學習



我會這樣做：

- ✓ 問題精準的問題
- ✓ 比較不同答案
- ✓ 吸收有用的資訊
- ✓ 標記來源與想法

3 整理與修改
加入自己的想法



我會這樣做：

- ✓ 整理重點
- ✓ 用自己的話寫出來
- ✓ 加入實例或想法
- ✓ 檢查與修正內容

✅ 結果：理解深、記得住，能力大提升！

主動思考
培養判斷力

理解內容
學得更扎實

表達自己的想法
更有創意

真正提升能力
面對未來挑戰

★ AI 是好幫手，但學習的主角永遠是你自己！★

∴ 圖片生成提示詞參考：

Illustration comparing two scenarios: Left "Fast AI" shows a student mindlessly copy-pasting from an AI screen, looking stressed and passive. Right "Slow AI" shows a student thinking with a pencil first, then consulting AI, then revising thoughtfully. Clean flat style. Left side in gray tones, right side in warm green and blue.

註記：本圖像素材使用生成式 AI 工具 (ChatGPT) 輔助生成。

✗ 快 AI (要避免)	☑ 慢用 AI (要培養)
直接複製 AI 答案，連讀都沒讀	先讀懂 AI 的回答，再用自己的話重新說
遇到問題第一步就問 AI	先自己想 5 分鐘，寫下初步想法，再問 AI
立刻交作業，沒有查證	用 AI 生成草稿後，查核重要資訊，再修改繳交
AI 說什麼就是什麼	對 AI 的答案保持懷疑，主動查核重要內容
所有事情都讓 AI 決定	用 AI 收集資訊和角度，自己做最後的判斷和決定

設計你的 AI 使用個人公約

參考以下範本，根據自己的狀況設計你的 AI 使用原則：

我的 AI 使用個人公約

- ☒ 考試和作業的核心內容，我會自己獨立完成
- ☒ 使用 AI 協助後，我會讀懂輸出，不盲目複製
- ☒ 提交作業時，我會誠實說明 AI 協助的部分
- ☒ AI 提供的重要資訊，我會交叉比對其他可信來源
- ☒ 不把他人的個人資訊或私密內容輸入 AI
- ☒ 定期問自己：AI 正在幫助我學習，還是取代我學習？

六、AI 回答風險檢查表

每次使用 AI 的答案前，用這個檢查表確認一下。特別是重要的作業、報告或需要引用資料的情況。

檢查問題	☑ 確認	⚠ 待查
AI 有沒有提供引用來源？來源是否真實存在？	☐	☐
這個資訊是否可能過時？(AI 的訓練有截止日期)	☐	☐
這個內容是否涉及版權問題？(文字、圖像、音樂)	☐	☐
AI 的回答是否可能有偏見或立場？	☐	☐
我有沒有把含有個資的內容輸入 AI？	☐	☐
我真的理解 AI 說的內容了嗎？還是只是看懂文字？	☐	☐
如果我要引用這個資訊，我有自己查核過嗎？	☐	☐

七、本章總結與自我評量

本章重點摘要 – 記住這五件事

- AI 是「預測」，不是「理解」：它很有信心，但不一定正確。
- 「請一步步說明」讓 AI 的回答更清楚，也讓你更容易找出錯誤。
- AI 幫你整理和修改學習歷程、小論文、英文作文，核心思考還是你的。
- AI 能生成，但不代表可以任意使用，注意版權、個資和深偽影片的風險問題。
- 慢用 AI 思維：先自己想、再問 AI、最後查核，這才是真正的學習。

情境式自我評量（不要問 AI！請自己思考）

情境一

AI 幫你生成了一篇英文作文，你修改了其中大約 30% 後準備交出去。

思考問題：

- 這算是你自己的作品嗎？為什麼？
- 你需要揭露 AI 的協助嗎？如何揭露？
- 如果作文的內容有事實錯誤，責任是誰的？

你的想法：_____

情境二

你在做科展，請 AI 幫你找「臺灣青少年使用手機與睡眠時數」的研究資料。AI 給了你三篇論文的標題和摘要。

思考問題：

- 你會直接引用這三篇論文嗎？為什麼？
- 如果你引用了，後來發現其中一篇不存在，會發生什麼事？
- 你應該如何正確使用這份 AI 給的資料？

你的想法：_____

情境三

你用 AI 工具把朋友的照片和另一個電影角色的臉合成，覺得很有趣，想在 Instagram 分享。

思考問題：

- 你在發布之前應該考慮哪些問題？
- 朋友如果不知道，有什麼問題？
- 電影角色的圖像有什麼版權問題？

你的想法：_____

學習檢核表

學習目標	我能做到 ☑	還需練習 🔄
1. 我能用自己的話說明 AI 和人類思考的差異。	☐	☐
2. 我知道怎麼用「一步步說明」讓 AI 的推理更清楚。	☐	☐
3. 我能說明學習歷程和小論文中，哪些地方不能讓 AI 代替。	☐	☐
4. 我了解版權、個資和深偽影片的基本風險。	☐	☐
5. 我知道什麼是「幻覺」，知道為什麼要查核 AI 的資料。	☐	☐
6. 我能設計一份 AI 使用揭露聲明。	☐	☐
7. 我理解「慢用 AI」的概念，並能說出至少三個具體做法。	☐	☐

延伸資源

推薦延伸探索

- Google Colab (colab.google) —— 免費的雲端 Python 環境，適合高中生練習
- Gemini Canvas —— AI 協作寫 HTML 遊戲、互動網頁
- 教育部因材網 e 度 AI 學習夥伴 —— 符合臺灣課綱的 AI 學習工具
- Canva AI (canva.com) —— 適合初學者的 AI 輔助設計工具
- ChatGPT / Claude —— 進行多輪對話練習，體驗「慢用 AI」思維
- 臺灣事實查核中心 (tfc-taiwan.org.tw) —— 練習查核 AI 提供資訊的好工具

本章核心觀念回顧

1. AI 不是真的懂你，而是在預測：AI 會根據大量資料預測最可能的答案。
2. AI 說得有道理，但不一定是對的：AI 可能出錯，因此要學會查證。
3. 問得越清楚，答案越有幫助：好的提問，才能得到好的回答。
4. 一步一步問，效果更好：逐步推理能讓 AI 回答更完整、更有邏輯。
5. AI 是學習夥伴，不是代寫工具：AI 可以協助學習，但不能代替你的思考。
6. 最後的判斷權在自己手上：AI 提供建議，決定仍要由你負責。
7. AI 能加快工作，但不能取代創意：真正有價值的想法與觀點來自你自己。
8. 善用 AI，也要負責任地使用 AI：遵守誠信原則，尊重版權、隱私與法律。

最重要的一句話

AI 的價值，不是替你思考，而是幫助你思考得更深入。



事實查核 與破解幻覺

識破假象，奪回思考自主權

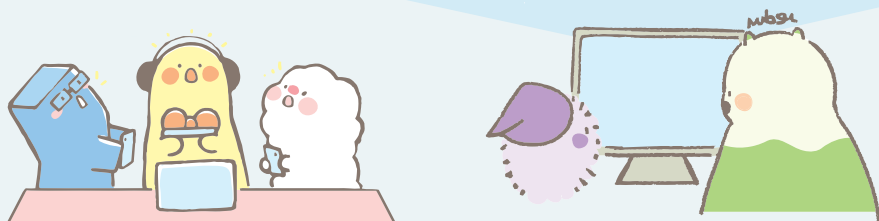
「識破假象的能力，
是這個時代重要的素養之一。」

在 AI 爆發的時代，許多之前專屬人類的生物特色與能力，
部分已可讓 AI 為我們代勞，但千萬不可讓我們的大腦被外包了。
如何活出人的獨特性，而不是活成「AI 運作就是我」，也考驗著這個時代的我們。



📌 本章學習目標：

- 防範深偽 (Deepfake)：分析假訊息的傳播機制與對民主社會的影響
- 對抗幻覺：建立嚴謹的檢驗機制，學習如何將 AI 作為「草稿員」而非「最終仲裁者」
- 演算法偏見與回聲室效應：了解 AI 推薦系統如何重塑我們的世界觀



一、防範深偽 (Deepfake)

你有沒有想過，若有人用你的臉和聲音製作出一段你從未拍過的影片，並在網路上瘋傳，你該怎麼辦？深偽徹底顛覆了「眼見為憑」這個我們自古以來判斷真假的方式。



∴ 圖片生成提示詞參考：

Split illustration: Left panel shows a realistic-looking deepfake video of a public figure with a red FAKE watermark overlay. Right panel shows a magnifying glass revealing digital artifacts and manipulation marks in the image. Warning icons and fact-check symbols. Flat illustration style, red and orange warning colors. No real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

🔍 什麼是深偽 (Deepfake) ？

深偽這項技術能快速生成假影片，捏造出任何人說出他們「從未說過的話」、去「從未去過的地方」、做「從沒做過的事情」的影片。以前的我們常用眼見為憑或用影片來判斷某件事是否為真，但深偽徹底顛覆了這個認知。

因此，當眼見再也不能都真為憑時，提升媒體識讀和事實查核能力，就變得相當重要，也成為數位公民的必修課題。

關於 AI 生成的影片、圖片或音樂等成果，現在都會出現 AI 公司的標誌、商標（LOGO）或嵌入「數位隱形浮水印」，讓 AI 生成的影片、圖片或音樂可簡易被辨識是否為 AI 生成；但這不代表它一定就是假訊息，我們該看的是內容有無進行不實的扭曲或誤導，以及內容**有無進行事實查核**。

AI 是優化表達方式的利器，它讓複雜的資訊變得更易閱讀、更美觀。只要核心資訊是正確的，AI 的輔助反而提升了資訊傳遞的價值。倘若關於影片、圖片或音樂，不易辨識判斷是否為 AI 生成，亦可將影片、圖片或音樂上傳 AI 平臺，讓其判斷 AI 生成的機率有多少。

🔒 數位隱形浮水印

📍 小提醒：Google Gemini 生成影片、圖片、音樂或文字等，會嵌入「數位隱形浮水印」，除了提升 AI 生成內容的透明度與可追溯性，目前這技術也已成爲跨產業應對深偽問題的重要工具。

除了隱形浮水印，全球主流 AI 廠牌（如 OpenAI、Microsoft、Adobe、Midjourney 等）已普遍在輸出檔案中寫入**內容憑證（Content Credentials, CR）**。這其實是一種**後設資料（Metadata）**；簡單來說，就是「關於檔案的數位履歷資訊」，如同你用數位相機或手機拍照，照片檔案內部會記載拍照時間與機型等相關資訊。AI 生成檔案的後設資料，就像是一張「數位身分證」，詳細記錄著這張圖片、影片或音樂是由哪個 AI 模型生成或進行哪些修改歷程。

倘若使用者進行「截圖」，這張身分證（後設資料）就可能喪失部分資訊。但只要與**數位隱形浮水印**搭配，就能發揮雙重防護的效果。另外，認識假訊息的常見類型特徵、傳播模式與阻斷傳播機制，也是現代公民該具備的數位素養。

📋 假訊息常見類型一覽表

類型	定義	目的 或 特徵 或 釋例
指桑罵槐	以諷刺或惡搞呈現，可能無意造成傷害，但具有誤導可能性。	通常出現在幽默搞笑網站，被不明就裡的人當真，進行轉傳分享。
張冠李戴	提供錯誤連結，或者標題、圖片與實際內容不符。	典型的「標題黨」，利用聳動標題騙點擊流量。常見於內容農場的網站。
斷章取義	撰寫誤導性內容，刻意曲解資訊，用來抹黑他人或某件事情。	擷取部分事實，但隱藏關鍵背景，用以誤導觀眾產生偏見。
移花接木	內容真實但與錯誤的時間或地點、背景結合，提供以假亂真的情境。	例如引用十年前的真實演習畫面，假裝說成昨天爆發戰爭。常見於社群媒體，如：LINE 群組或臉書等。

類型	定義	目的或特徵或釋例
冒名頂替	假冒真實的新聞來源、政府機構或名人等的冒充內容。	偽裝成臉書或推特藍勾勾認證帳號發布假訊息，或假冒為知名新聞臺發布假新聞。
偷梁換柱	操弄撰寫內容，對真實資訊進行竄改或修改圖文內容。	透過修圖、剪接，改編照片或影片的原意呈現。
無中生有	百分百憑空捏造，旨在欺騙或傷害的完全虛構內容。	編造完全沒有事實依據的假訊息，並進行轉傳、分享。常見於刻意發動資訊戰。

⚠ 假訊息為何傳播得比真相快？

為什麼假訊息傳播得比真相快？



媒體識讀，求證是關鍵！

註記：本圖像素材使用生成式AI工具(Gemini)輔助生成。

假訊息傳播快，是因為它具有以下特質：

- 情緒激化：內容刻意引發憤怒、恐懼、仇恨或同情，讓人不假思索就轉傳。
- 新奇聳動：誇張離奇的標題最容易吸引點擊與分享。
- 八卦題材：可成為群體間的社交談論話題。
- 訴諸直覺：內容貼近既有偏見，看起來「很合理」所以未查證。
- 從眾效應：看到朋友都在轉傳，就覺得「大家都說一定是真的」。

最有效阻止假訊息傳播的方法，就是對收到的訊息都要做到「勿輕信、要審視、勤查證」，求證無誤後，才能進行轉傳、分享。

原則	說明
勿輕信	保持冷靜、理性、中立、包容與客觀的立場，不要輕易被訊息左右，而失去了自己的判斷力。
要審視	確認資訊的來源是否可靠、內容或影像是否真實正確、評論是否客觀且多元，同時也要小心「慣性思維」的直覺反應與避免「同溫層效應」的影響。
勤查證	遇到可疑或有爭議的訊息，主動進行查核與驗證，並能立即阻斷假訊息繼續傳播。

其他查證相關資源和查核平臺介紹請參閱附錄四：查證相關資源和查核平臺介紹。

血 對民主社會的衝擊

看透假訊息！對民主社會的衝擊

假訊息透過刻意捏造、散播不實資訊，誤導大眾，達成特定目標與利益。

1. 人民與社會層面



引起大眾心理焦慮，引發社會不安。

例如：量販店散布不實衛生紙即將漲價的假訊息，引發恐慌性搶購衛生紙囤貨行為。

2. 經濟與商業利益



干擾股市、市場行情或打擊競爭對手，獲取不當商業利益。

例如：散布雞瘟假訊息，企圖藉機哄抬蛋價。

3. 政治目的與操弄



抹黑、造謠對手或斷章取義，影響選民投票意願風向。

例如：斷章取義的訊息散布來影響選民，企圖達成某些目的。

4. 軍事與戰略用途



進行認知作戰，擾亂社會運作或干擾政府決策，影響民心士氣。

例如：使用假訊息影響民心士氣。

提升識讀力，守護我們的民主！

註記：本圖像素材使用生成式AI工具(Gemini)輔助生成。

若讓假訊息恣意於社會亂竄，可能撕裂社會安寧或影響國家運作，它的目標是讓我們「對彼此失去信心」。民主防衛韌性不在於消滅所有異音，而是強化大眾對民主制度的信任。當人民能保持冷靜、多方求證，並在面對滔天巨浪的威脅時，能及時警醒同心面對問題，共同尋求最有利和最佳解決辦法，齊心爲了家國的安定，凝聚團結力與共識，方能克服風險。這正是由公民意識築成的防線，也是假訊息和認知作戰無法逾越的民主社會長城。

🔍 事實查核方法：SIFT 查核法

SIFT 事實查核法：破解假訊息四步法 (Mike Caulfield 提出)

S 冷靜審視，暫停分享。
面對令人震驚或憤怒的訊息時，先冷靜審視，並暫停分享的衝動。思考其來源是否可靠。

I 調查發布者背景。
了解訊息發布者或媒體的背景與立場。善用 Google AI 搜尋工具與搜尋「來源」，確認發布者是否為權威可信。

F: Find better coverage (尋找「佐證」)
F 比對權威媒體報導，尋找更多可靠來源佐證。
利用權威媒體或事實查核機構進行交叉驗證。

T 確認影像與言論脈絡，追溯原始出處脈絡。
確認影像或言論是否被斷章取義，或由 AI 生成。

活用 SIFT，成為資訊達人！

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

SIFT 查核法由資訊素養專家柯斐德 (Mike Caulfield) 提出，當面對真假難辨的訊息時，可使用這四個簡單步驟協助判斷：

步驟	說明
S — Stop (暫停)	當你看到令人震驚或憤怒的訊息時，先「冷靜審視」，並暫停分享的衝動，思考其來源是否可靠。
I — Investigate the source (調查來源)	了解這則訊息來自哪裡，是可信的媒體、政府機構，還是來路不明的帳號？
F — Find better coverage (尋找佐證)	用關鍵字搜尋，看看有沒有其他可靠的媒體或事實查核機構也有相關的報導？
T — Trace claims, quotes and media to the original context (追溯原始脈絡)	確認影像或言論是否被斷章取義，或由 AI 生成？

二、對抗幻覺 (AI Hallucination)

💡 小叮嚀

你是否曾經請 AI 幫你查資料，結果發現它「引用」的論文根本不存在？這不是 AI 在說謊，而是它的運作本質所造成的——它統計預測，而非邏輯思考。

對抗幻覺，可以建立嚴謹的檢驗機制，學習如何將 AI 作為「草稿員」而非「最終仲裁者」。

建立工作流：

1. 草稿員原則



將AI視為協助助手，而非自己不加思索，直接大腦外包給AI。

2. 多方交叉驗證法



針對AI產出的事實資訊，必須嚴格執行查核。至少比對兩個以上的權威管道。

3. 提示詞加入要求條件



輸入提示詞...(要求AI「如果不知道，請回答不知道」，並要求註明資料來源，以利後續的人工溯源查證)

例如要求AI「如果不知道，請回答不知道」，並要求註明資料來源，以利後續的人工溯源查證。

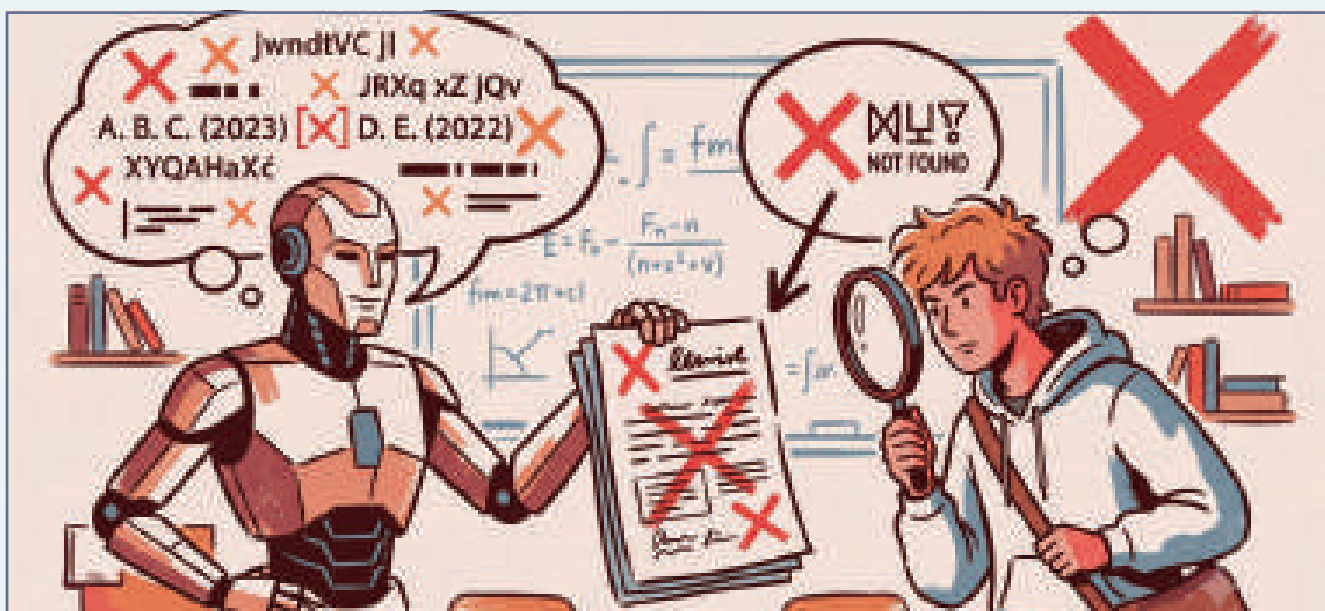
聰明使用AI，不被誤導！



🔍 什麼是幻覺 (AI Hallucination) ?

幻覺是指 AI 模型產出「邏輯通順但事實錯誤或不存在」的內容。為什麼會發生「幻覺」？因為 AI 的運作本質是機率預測，而非邏輯思考；它是在預測「下一個最有可能出現的字詞」。

當它為了滿足文法邏輯，硬湊出機率最高但未經考證的資料時，幻覺就產生了。



🔗 圖片生成提示詞參考：

Illustration of an AI robot confidently presenting a document labeled Fake Research Paper. A student looks at it with a magnifying glass and finds the paper does not exist. Show a thought bubble from the AI saying high probability words while generating nonsense citations. Warning style with red X marks. Flat illustration. No text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

📌 AI 可能會「捏造引用」

這不是故意說謊，而是 LLM 的本質：它是在預測「下一個最有可能出現的詞」，而不是真的在查資料庫確認真假。

📌 如何將 AI 當作「草稿員」？

- 把 AI 的回答當成初稿，不是最終答案。
- 主動針對 AI 給的每一項資料，去查核可信來源。
- 對引用的論文或書目，必須在圖書館或學術資料庫查到原文才能使用。
- 培養「懷疑的能力」與「查核的自律」，才是 AI 時代最重要的競爭力。

三、過濾氣泡與回聲室效應

你有沒有發現，你的社群動態幾乎都是你「認同」的觀點？你以為在自由瀏覽網路，但其實演算法早已為你量身打造了一個「資訊牢籠」，產生資訊接收的偏聽現象。



📌 圖片生成提示詞參考：

Illustration of a teenager trapped inside a transparent bubble floating in a digital space. Inside the bubble: social media feeds showing only similar content, like symbols, and agreeing faces. Outside the bubble: diverse opinions, varied news, different perspectives that cannot reach inside. Flat illustration, blue and purple color palette. No real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

🔍 過濾氣泡 (Filter Bubble)

社群媒體推薦演算法的目標是「留住你的注意力」，它會不斷推播你喜歡、認同的內容，自動過濾掉挑戰你觀點的資訊，最終讓你被包圍在一個舒適的「氣泡」中。

🔊 回聲室效應 (Echo Chamber)

在氣泡內，你會發現身邊的人意見都跟你一樣，這會放大你的偏見，讓你誤以為「全世界都這麼想」。

這可能導致社會極端對立，不同立場的人失去溝通的可能性。

💡 小叮嚀

AI 時代最強大的競爭力，不是你多會使用工具，而是你是否具備「懷疑的能力」與「查核的自律」。別讓演算法定義你的世界觀，真相需要自己去追查。

💡 突圍策略：三招打破資訊牢籠

- ①意識到演算法的不中立：演算法並不中立，它是為了商業利益（提高觸動點擊率）而設計的。
- ②勿「偏食」單一媒體並啟用無痕模式：刻意追蹤與自己立場不同的深度媒體，並使用「無痕模式」瀏覽網路資訊，可降低演算法的偏見產生。
- ③批判性提問自己：當看到一則讓你非常認同的貼文時，問問自己：「這則資訊有沒有漏掉哪些不同的觀點？」時刻警醒尊重多元，才能讓思考更趨向周延、完整。

另外，事實查核平臺MyGoPen提供防詐工具「麥擱騙 | 詐騙網址幫手」，也是很好的查核輔助工具。

網址：<https://check.mygopen.com/>

四、數位時代的法律紅線：深偽與性影像犯罪

⚠ 嚴正警告：AI 不當使用的法律後果

隨著 AI 技術普及，利用深偽製作不實性影像的案件激增。別讓一時的不慎變成一輩子的汙點。

法律後果	說明
● 刑法重罪	製作或散佈深偽性影像，最高可處5年有期徒刑；若涉及營利或恐嚇，刑期更重。
💰 民事賠償	受害者可以要求高額的精神撫慰金，這可能會讓你與你的家人面臨嚴重的財務負擔。
💻 數位足跡永久存在	網路是有記憶的，一旦進入司法程序，你的升學、未來就業（尤其是公職或大型企業）都可能受影響。
「尊重，是數位公民的基本素養；法律，是惡作劇的終點線。」	

SOS 應對口訣：【截、求、告】

如果你正面臨深偽影像威脅或恐懼，請記住這三個字：

-  **截（截圖）**：把對方的威脅訊息、照片網頁等資訊全部截圖，保留證據以利求助舉證用。
-  **求（求助）**：找你信任的大人（如：雙親、老師、親友……）或聯繫 iWIN 或性影像處理中心等相關單位。
-  **告（告發）**：到警察局報案。爲了保護你自己擺脫傷害，並讓錯誤的行為被制止。

步驟	具體行動	目的
1. 冷靜蒐證	將訊息截圖，保留網址、社群帳號、威脅訊息等相關證據。	建立法律上舉證用的證據鏈。
2. 尋求下架	聯繫「iWIN」或「性影像處理中心」。	阻止影像等數位訊息進一步擴散。
3. 法律訴訟	前往警局報案，啟動刑事調查。	讓加害者付出代價，並尋求保護。
4. 心理支持	尋求家人、校園輔導室或諮商師協助。	修復事件帶來的心理創傷。

五、媒體識讀實戰活動

🔗 活動說明

YouTube 上會有許多影片，但內容都正確嗎？都能相信嗎？

邊看影片邊使用「附錄二：媒體識讀實戰查核指引學習單」和「附錄三：數字編碼假訊息查核表」進行事實查核。透過假訊息查核表檢核訊息後，你得到的警報顏色結果是什麼？請將判斷結果紀錄於「警報顏色判定欄」，並實踐查核行動的建議。

📺 查核範例影片

標題：臺灣最多農藥水果Top10！第1名你天天吃！

🕒 時代大動脈 2026年6月27日

影片連結：<https://www.youtube.com/shorts/nJ1GfKz7ebg>

#臺灣水果 #農藥殘留 #食安危機 #Shorts #臺灣

提醒：以上影片也可到MyGoPen網站找尋平臺留存影片與查核文章。

網址：<https://www.mygopen.com/2026/07/fruit.html>



🚨 警報顏色判定表

- ☐ 🚨 **紅色警報（高度可疑）**：這則訊息極有可能是假的或充滿誤導性。請不要分享，也可提醒轉傳者要多小心。
- ☐ ⚠️ **黃色警報（部分真實，需再查核求證）**：訊息可能包含部分事實，但整體結論被誇大或變造誤導他人，盡可能就別分享。若分享時，能提醒收訊者務必特別小心，最好能附上你檢視驗證收集到的相關事實查核資訊。
- ☐ ✅ **綠色警報（可信度高）**：這則訊息有可靠的來源和多方佐證，可做為參考；但永遠要記得，保持批判思考和檢視查核習慣。

媒體識讀實戰活動分析參考說明：

🔍 媒體識讀查核要訣：如何一眼看穿謠言？

疑點分析：

影片具有**內容農場**的常見招數：

- 這類影片通常具備「**內容農場**」的典型特徵，透過心理戰術吸引點擊並誤導大眾。

媒體識讀查核要訣：如何一眼看穿謠言？

- **聳動且恐懼的標題**：使用「最多農藥」、「慢性毒死」、「致命組合」、「99%的人都中招」等極端詞彙，旨在激發讀者的恐懼，誘使人們急於點擊或轉發。
- **匿名、雜亂或缺乏公信力的頻道**：發布者是具公信力的媒體或官方機構，還是風格雜亂的匿名頻道？同時檢視平臺其他內容是否標示AI公司標誌或商標。
- **時間線模糊**：影片是否頻繁出現「最近」、「剛剛發生的」、「最新公告」，卻始終沒有標明事件或數據發生的具體時間？
- **舊聞新炒**：刻意拿數年前已經澄清過舊事件或舊數據，重新剪輯後當作「當前危機」來發布，製造人爲恐慌。
- **資訊真假夾雜**：在正確訊息中，刻意混入錯誤的導向結論，降低閱聽者的防備心。
- **違背或無專業常識**：是否頻繁宣稱違背自然科學的現象或醫學等常識？或無法拿出公信力機構的實驗數據或採用錯誤科學的論述或胡說八道。
- **宣稱或推廣未經證實的「神奇偏方」**：是否強烈推薦某種偏方或方法，來宣稱驚人療效或功效，卻始終沒有醫學臨床實驗佐證或缺乏衛生主管機關的認證或無科學實驗證明？
- **訴諸權威但缺乏具體身分**：宣稱「專科醫師公開」或「名醫警告」，卻沒有具體的姓名、所屬醫院或相關研究報告來源，是典型的「訴諸權威」謬誤。
- **鎖定特定族群(長輩)**：強調「長輩中招」，利用長輩對健康的高度關注與資訊落差，增加傳播率。
- **視覺與聽覺特徵**：常使用AI配音(語調較平淡且機械化)、庫存素材拼湊(畫面與口說內容不完全對齊)以及高飽和度的縮圖。或者緊湊、緊張或驚悚的配樂節奏，來挑動你的情緒。
- **用語與字體破綻**：影片中出現簡體字或非本土慣用語或缺乏彈性且過於絕對的斷言話語，可能是境外詐騙集團或有心人士……等的設局。
- **斷章取義或扭曲假造**：是否大量剪輯專家發言、新聞畫面或官方報告，卻只截取其中幾句話或特定片段，刻意隱去完整的前後文脈絡與原始背景或加以扭曲假造來增添話語說服力？

→ 面對這類「謠言」，可運用以下四個查核步驟：

- **S：Stop (暫停)**。看到標題很嚇人、語氣很絕對的影片，先別急著相信或轉發。
- **I：Investigate the source (調查來源)**。調查來源是否爲具有公信力單位或專業醫療機構的頻道，還是匿名、雜亂的影音頻道？審視內容的邏輯：如果真的會「毒死人」，爲何主流新聞媒體或衛生福利部還沒有發布緊急公告？
- **F：Find better coverage (尋找佐證)**。利用關鍵字搜尋，例如搜尋「最多農藥」，通常第一時間就能看到事實查核平臺的結果。
- **T：Trace claims, quotes and media to the original context (追溯原始脈絡)**。確認影像或言論是否被斷章取義或由AI生成。

事實查核重點整理（以範例影片為例）

1. 非近期新事件：引用2021年的舊數據編造發布

- 網傳的名單（百香果、西瓜、木瓜、荔枝、草莓、龍眼、印度棗、洋香瓜、香瓜、枇杷）源自農業部《110年度水果農產品農藥殘留監測研究成果報告》，此為數年前的舊資料，被內容農場重新編造、剪輯成影片發布。

2. 「全果採樣」不等於民衆「全果全吃下肚」

- 主管機關進行農藥檢測時是採用「全果（連皮帶肉）」採樣進行檢驗。
- 民衆實際食用時大多會清洗、剝皮或削皮（如百香果、荔枝、木瓜、西瓜等），因此「檢驗農藥超標」不等於民衆會把水果上的農藥全都吃下肚。

3. 專家說明：網傳病蟲害與用藥說明錯誤百出

中興大學植物病理學系名譽教授曾德賜教授指出，傳言內容對這些水果的生長、用藥，以及常見的病害、蟲害根本毫無概念和實務經驗。

- 百香果(第1名)：傳言稱「果面皺摺易藏污垢多農藥」。專家表示，百香果在樹上成熟時表面是光滑的，變皺是採收後乾掉所致，其皮厚，農藥難滲透到果肉。
- 西瓜(第2名)：傳言稱「易受白粉病威脅」。專家表示，西瓜主要是炭疽病與蔓割病，白粉病非主要問題。
- 木瓜(第3名)：傳言稱「果實易受蚜蟲侵襲」。專家表示，蚜蟲叮的是嫩葉而非果實(果皮上的乳汁蚜蟲受不了)，且目前多採用網室栽培防蟲。
- 荔枝(第4名)：傳言稱「為對抗地老虎噴藥」。專家表示荔枝根本沒有「地老虎」害蟲，主要是蒂頭部位會有「細蛾」，而有細蛾的荔枝也代表農藥用量較少。
- 草莓(第5名)：傳言稱草莓貼近地面且皮薄多汁，為防灰黴病用藥達8至10次。專家表示，灰黴病是在盛產期後，因人力無法管理田間所致；果農若需用藥會先採光成熟果實，並遵照3至7天以上的安全採收期讓農藥自然降解。
- 龍眼(第6名)：傳言稱曾抽驗發現農藥殘留超標，引發食品安全疑慮。專家表示，龍眼雖然與荔枝有部分相似的病蟲害問題，但龍眼比荔枝好種許多，比起其他水果算是非常容易種植的。
- 印度棗(第7名)：傳言稱面臨20幾種病蟲害，整個生長期用藥高達30至40次。專家表示，使用農藥也需成本，農民不大可能用藥那麼多次，只要做好田間管理，專家自己都會連皮食用。
- 洋香瓜(第8名)：傳言稱皮薄糖分高易病害，生長期大量噴灑殺菌劑以維持網紋外觀。專家表示，洋香瓜皮其實很厚，且主要的白粉病與黑點根腐病使用微生物防治法效果就很好，根本不需要大量用藥。
- 香瓜(第9名)：傳言稱使用「二甲戊靈」。專家表示，二甲戊靈是種子萌發前的用藥，或整地時用來除雜草；若直接噴在果樹或果實上會致其死亡，果農不可能這樣使用。
- 枇杷(第10名)：傳言稱因高溫高濕易感炭疽病，果農在花期至膨大期高頻率用藥。專家表示，炭疽病對枇杷影響不大，主要病害是真菌引起的「灰斑病」，且現今果農多已採用「套袋」來防護。

事實查核重點整理（以範例影片為例）

4. 從挑選到保存！降低水果農藥殘留的2個實用小撇步

- 認明標章，源頭把關：優先選購具有「產銷履歷標章」或「有機認證」的水果，可確保農產品來源安全與安心享用。
- 適度靜置，自然分解：耐放的水果買回後，可放置1～2天，利用農藥在自然環境下的降解或植物本身的代謝作用來分解部分農藥。但切勿擺放過久，以免流失新鮮度與營養價值。

5. 正確清洗水果的方法與勿信偏方

- 偏方無效：使用小蘇打、鹽水、醋、麵粉等中和農藥缺乏科學實證，過度清洗浸泡，還可能減損營養與風味。
- 「流動清水」最有效：
 - (1) 浸泡：先在清水中浸泡約3分鐘，溶解去除接觸型農藥。
 - (2) 沖洗：使用「流動的清水」徹底沖洗，表面凹凸處(如芭樂)可用軟毛刷輕刷。
 - (3) 先洗後切：清洗乾淨後再進行去皮、切塊或去蒂頭，可避免刀具將表皮農藥帶入果肉造成交叉污染。

總結小叮嚀：

- 小心內容農場的錯假訊息：內容農場的文章或影片或訊息大量出現，主要是由**經濟利益與搜尋引擎機制**共同催生的產物，資訊內容多為抄襲、斷章取義或憑空造假，內容正確性的查核對其而言，反而不重要了；其主要能快速吸引廣大閱聽者點擊或轉傳傳播賺取廣告流量。隨著資訊傳播或有心人士操弄等的演變，現今則可能牽扯認知戰的傳播，因此具備媒體識讀錯假訊息的能力，就變得很重要了。
- 不明影片別亂傳：轉發未經證實的食安議題或醫療謠言，可能會引發不必要的恐慌，甚至造成果農損失、物價衝擊或誤導正確觀念與使用方法。
- 查證工具要善用：遇到疑慮時，請多利用「MyGoPen」、「Cofacts 真的假的」或「臺灣事實查核中心」進行查證。
- 做到「勿輕信、要審視、勤查證」來對抗錯假訊息。

關於「事實查核重點整理」的內容，節錄與編修自MyGoPen文章。

2026/07/22【錯誤】網傳臺灣最多農藥的10種水果？食安危機影片？缺乏實務觀念的不實描述！專家詳解。

網址：<https://www.mygopen.com/2026/07/fruit.html>

 小叮嚀

關於更多查核資訊分析等實戰學習活動，可到臺灣事實查核中心平臺網站找尋查核報告或MyGoPen網站：謠言澄清區查閱及練習，以精進自己的事實查核力。

🔗 臺灣事實查核中心：<https://tfc-taiwan.org.tw/fact-check-reports-all/>

🔗 MyGoPen：<https://www.mygopen.com/search/label/謠言>

別讓錯假訊息佔領你的認知

非農業專業背景的民衆，面對海量的網路訊息，確實容易被錯假訊息誤導。除了到圖書館查閱專業書籍，善用具公信力的官方與專業機構資源，也是進行事實查核的最佳途徑。

以下整理影片中提及的知識延伸與專業機構查核資源：

◎ 延伸學習與知識補給

● 印度棗：

即為日常生活中常見的蜜棗、棗子或棗仔，外皮呈鮮綠色，口感爽脆多汁。

● 洋香瓜：

1.作物特性：一年生蔓性草本作物。莖部呈藤蔓狀，可匍匐於地面或攀附網架生長。開黃花。雌花授粉後，子房會膨大發育成果實。常見種類有網紋綠肉、網紋橙肉、光皮及哈密瓜等。

2.網紋形成原理：當果實表皮生長並逐漸硬化後，內部果肉仍持續發育膨大，導致表皮產生微細裂痕。此時果實會啟動自我修復機制，形成「結痂樣態」的木栓化組織，即為所見網紋，這是植物體重要的保護機制。可能因氣候或水分因素，導致裂痕過深、自我修復不及，病原體就會趁機入侵，致使果實病變、腐壞。

◎ 權威資源與事實查核資料

● 農藥殘留與清洗食安指南：

1.衛生福利部食品藥物管理署(Taiwan Food and Drug Administration)

《藥物食品安全週報》第 995 期(2024/10/11)：避免農藥殘留，蔬果清洗有撇步！

網址：<https://www.fda.gov.tw/Tc/PublishOtherEpaperContent.aspxid=1530&tid=4915&r=1460310396>

2.有機農業全球資訊網(Taiwan Organic Information Portal)

文章(2018/12/18)：如何清洗蔬果吃得安心

網址：<https://info.organic.org.tw/3609/>

3.農業部農業藥物試驗所 技術諮詢交流服務網
問答集

網址：<https://mbox.acri.gov.tw/TA02.asp>

● 農業知識與食農教育資源相關平臺：

1. **農業知識入口網(農業部)** | 網站：<https://kmweb.moa.gov.tw/index.php>
整合研究團隊資源，針對消費者、生產者與學者進行分眾知識，核心包含：
 - 專業知識庫與分類主題館：除分享臺灣農業精進成果外，專家系統化整理農、林、漁、牧各類特定作物的知識、圖片與生活資訊。
 - 優質農家經驗：透過主題報導與產銷專欄，介紹推動臺灣農業蛻變與優秀農經驗。
 - 彙整多樣化資源：利用「影音專區」與「資源推薦」，集中管理多媒體檔案與實用系統，方便民衆學習。
 - 農業知識家論壇：打造關心農業發展者的互動討論與知識共享平臺。
2. **食農教育資訊整合平臺** | 學習資源專區網址：<https://fae.moa.gov.tw/ws.php?id=4>
3. **農業部動植物防疫檢疫署**(前身爲行政院農業委員會動植物防疫檢疫局) | 電子書網址：https://www.aphia.gov.tw/publish/sweet/sweet_index.html
4. **農業部動植物防疫檢疫署** | 時事澄清專區網址：https://www.aphia.gov.tw/theme_list.php?theme=NewInfoListWS&sub_theme=explain

挑戰任務

初階任務(觀察與記錄型)

連續三天，記錄你在社群媒體上看到的一則讓你「很想轉傳」的內容，並用 SIFT 查核法逐一分析，記錄查核結果。

進階任務(分析與比較型)

選一個你常用的社群平臺(如 IG、YouTube)，分析「演算法推薦給你的內容」與「你主動搜尋的內容」有何不同？觀察並記錄「過濾氣泡」的存在跡象。

高階任務(創作與應用型)

設計一份屬於你班級的「假訊息識讀小手冊」(至少 4 頁)，說明 3 種常見假訊息類型、SIFT 查核法步驟，以及如何面對幻覺；可配合 AI 工具輔助排版。

學習檢核表

學習目標	我能做到 ✓	還需練習 🔄
1. 我可以說明深偽這項技術如何顛覆「眼見為憑」，及其對民主社會的潛在衝擊。	☐	☐
2. 我可以分辨至少三種假訊息類型（如：張冠李戴、斷章取義、移花接木或冒名頂替）。	☐	☐
3. 我可以覺察自己是否正被「情緒挑動」、「好奇心」或「從眾效應」影響而相信假訊息。	☐	☐
4. 我可以運用 SIFT 查核法（暫停、調查來源、尋找佐證、追溯脈絡）對可疑訊息進行查核。	☐	☐
5. 我可以解釋幻覺的成因，並說明如何將 AI 當作「草稿員」而非「最終仲裁者」。	☐	☐
6. 我可以說明過濾氣泡與回聲室效應如何影響個人世界觀，並提出突圍策略。	☐	☐
7. 我可以辨識深偽涉及的法律責任，知道遇到威脅時的應對口訣（截、求、告）。	☐	☐
8. 我可以使用「AI 回應風險檢核表」、「媒體識讀實戰查核指引學習單」與「數字編碼假訊息查核表」對實際案例進行分析。	☐	☐

本章核心觀念回顧

1. 深偽徹底顛覆了「眼見為憑」，任何影音都可能是 AI 合成。
2. 假訊息類型多樣（移花接木、張冠李戴、無中生有等），情緒是它最大的武器。
3. SIFT 查核法是面對可疑訊息的基本工具：暫停、調查來源、尋找佐證、追溯脈絡。
4. 幻覺是 AI 的本質，要將 AI 視為「草稿員」，而非「最終仲裁者」。
5. 演算法的過濾氣泡讓你看不到全貌，要主動打破資訊牢籠。
6. 利用深偽製作性影像是違法的，面對威脅要記得「截、求、告」。

最重要的一句話

真相需要自己去追查，別讓演算法定義你的世界觀。



AI 倫理 與學術誠信

從便利與努力之間找到平衡

「科技雖然做得到，
但人們仍需要思考：這樣做真的合適嗎？」

當 AI 越來越聰明，我們的生活也開始變得越來越方便。

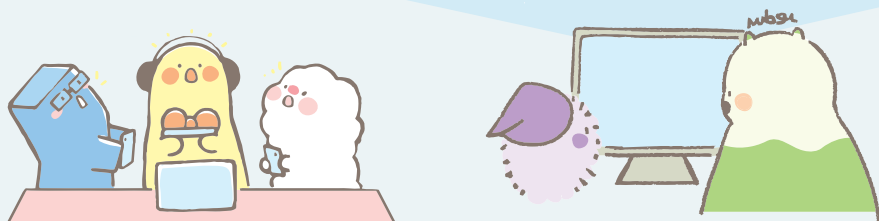
然而，學習 AI 不只是學會操作工具，更重要的是培養判斷與反思能力。

當我們開始願意停下來思考科技可能帶來的影響，才能讓 AI 真正對人與社會有幫助。



📌 本章學習目標：

- 理解學術誠信的意涵，學會負責任地將 AI 作為學習助手而非代寫工具
- 掌握數位隱私的保護原則，辨識哪些資訊不應輸入 AI 平臺
- 能構思 AI 系統設計中的倫理問題，思考公平、透明、責任與安全的平衡
- 建立與 AI 協作的倫理共識，培養團隊中的即時倫理判斷能力
- 將 AI 倫理內化為日常習慣，養成「倫理反射」的思考模式





👉 圖片生成提示詞：

Infographic of a robot and a student sitting on a large balance scale, comparing technology ethics, flat colorful illustration style, educational with clear visuals. The left scale pan holds convenient technology items, the right holds responsibility items like a scale and lock, central is a thinking lightbulb. Back blackboard features Chinese and English text for "Technology Ethics" and related concepts. Detailed vector art style.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

一、學術誠信與負責任的 AI 使用

當 AI 可以幫你寫作業時，你真的有學到東西嗎？你有沒有想過，如果 AI 可以替你完成作業，那學習的意義到底是什麼？效率很重要，但誠實面對自己的學習過程，更是高中生在 AI 時代最重要的能力之一。

😊 情境 | 小瑜的報告

高中的小瑜最近常使用 AI 工具幫忙寫小論文。剛開始，她覺得非常方便，輸入題目後，幾秒鐘就能得到一篇完整內容。有一次，她甚至直接把 AI 生成的文章貼進作業裡交出去。雖然老師沒有立刻發現，但當同學問她報告內容時，她卻答不太出來。

那一刻，小瑜才開始思考：如果只是把 AI 當成「代寫工具」，自己真的有完成學習嗎？

現在的 AI 工具越來越普遍，能協助查資料、翻譯、整理重點、產生圖片，甚至幫忙規劃學習流程。善用 AI 的確能提高效率，不過 AI 提供的答案不一定完全正確，如果沒有經過查證就直接使用，可能會誤導自己或他人。

(一) 學術誠信是什麼？

學術誠信不只是「不要作弊」，更要誠實面對自己的學習過程。如果把 AI 生成內容直接當成自己的成果，卻沒有說明來源或加入自己的理解，就可能變成抄襲；真正好的 AI 使用方式，是把它當成「學習助手」，而不是「代替自己思考的人」。

(二) 怎麼負責任地使用 AI 輔助學習？

☑ 負責任的 AI 學習方式	✗ 違反學術誠信的做法
請 AI 解釋難懂概念，再用自己的話重新整理	直接複製 AI 輸出當作作業繳交
請 AI 提供不同觀點，再透過自己的思考取捨	把 AI 生成的文章說成完全是自己寫的
用 AI 協助整理大綱，再自己展開內容	在禁止使用 AI 的考試或作業中偷偷使用
引用 AI 協助的部分時，清楚標示來源	引用 AI 捏造的資料，沒有查證就繳交
把 AI 當提問練習的對象，加深自己的理解	讓 AI 做完整份專題，自己完全沒有思考

(三) 我與 AI 的學習分工表

在使用 AI 協助學習前，用這張表確認你的分工是否合理：

倫理與誠信自我提問	我自己處理(☑)	AI 可協作的項目
☐ 是否直接複製貼上？	☐ 理解內容	☐ 提供範例
☐ 是否誠實註明 AI 協助？	☐ 判斷資訊正確性	☐ 協助整理架構
☐ 是否侵犯他人作品？	☐ 加入自己的觀點	☐ 解釋概念
☐ 是否輸入敏感個資？	☐ 個人反思	☐ 提供不同觀點
☐ 是否過度依賴 AI？	☐ 最終決定	☐ 幫忙修正文句

(四) 改寫 AI 生成內容，詮釋你的理解

 練習：請用 AI 查詢你目前正在學習的一個主題，再依照以下格式重新整理

AI 原文(摘錄)：_____

我的理解(用自己的話重新說明)：_____

我補充的觀點或疑問：_____

AI 說的哪一點讓我最有共鳴或最有疑問？_____

💡 小叮嚀

- 💡 AI 提供的內容不等於「正確答案」，使用前一定要查證。
- 💡 學術誠信的核心是：你有沒有真的理解，而不是你有沒有用 AI。
- 💡 清楚標示「AI 協助的部分」，是對老師和讀者最基本的尊重。

📖 學術誠信 (Academic Integrity)

學術誠信是高等教育的核心價值，在 AI 工具普及後面臨新挑戰。美國多所大學已更新學術誠信政策，要求學生「披露 AI 使用情形」，學生需要思考：AI 在我的學習中是「鷹架」還是「替代品」？建議搭配「後設認知反思」活動，讓學生書寫自己的學習過程，而不只是繳交作品。

☑ 本節重點整理

- 學術誠信 = 誠實面對學習過程，不只是「不作弊」。
- AI 是學習助手，不是代替你思考的人。
- 使用 AI 輔助時，要標示來源、加入自己的理解。
- 分工合理：理解、判斷、反思、最終決定由人來負責。
- 改寫練習是將 AI 知識轉化為自己理解的最佳方法。

二、數位隱私與個人資料保護

當你把問題丟給 AI 時，也把自己的資訊送出去了嗎？你每次輸入給 AI 的文字，都可能在某個伺服器裡留下記錄。「數位足跡」無形中積累，一旦包含個人資料，就可能帶來意想不到的風險。

👤 情境 | 阿哲的班級名單

阿哲最近很常使用 AI 工具。有一天，他爲了方便，直接把班上同學的完整名單和成績表貼進聊天視窗，希望 AI 幫忙做統計。結果朋友看到後提醒他：「這些不是別人的個人資料嗎？」阿哲這才發現，原來使用 AI 不只是「會操作」而已，還需要注意隱私保護。

當我們上傳資料給 AI，也代表這些內容可能會被系統處理或儲存。每一次輸入，都可能留下「數位足跡」，如果沒有留意，就可能不小心暴露自己的個資，甚至影響別人的隱私。

(一) 哪些資訊不應該輸入 AI？

✕ 絕對不應輸入 AI 的資訊	為什麼危險？
姓名 + 學校 + 學號的組合	可拼湊出完整的個人身份，被有心人士利用
他人照片（未經同意）	涉及肖像權，也可能被用來製作深偽
帳號密碼或登入憑證	資料可能被儲存或外洩，導致帳號被盜
班級名單、成績表、聯絡資料	屬他人個資，未經同意提供可能違反個資法
照片中包含可辨識地點的背景資訊	可能暴露常出沒的位置，帶來安全風險
私人對話或他人的聊天紀錄	侵犯他人隱私，可能構成法律責任

(二) 去識別化：保護資料的基本方法

💡 什麼是去識別化（De-identification）？

去識別化是指移除或替換資料中能夠辨識個人身份的資訊，讓資料在使用時不會直接連結到特定個人。

去識別化方法	說明	範例
匿名化	完全移除識別資訊，用代號取代	學號 → 「學生A」
資料遮蔽	將部分資訊用符號遮蓋	王○○ / 0912-XXX-XXX
資料聚合	只提供統計數字，不提供個別資料	平均成績 78 分
虛構化	用虛構的姓名和細節代替真實資訊	小明的案例

(三) 我的 AI 隱私守則

📝 立下你的 AI 隱私承諾

我不會上傳包含 _____ 的資料。
 當必要上傳時，我會用 _____ 方法，來去識別化。
 我上傳前會想一想：
 「這些資訊適合被分享嗎？如果被其他人看到會怎麼樣？」
 「被知道這些資訊可能有什麼風險？」

立書人：_____ 時間：____年____月____日

情境小劇場 | 阿哲遇到了個資問題

阿哲：「我把全班名單丟給 AI，讓它幫我分析學習狀況，應該沒關係吧？」

AI 助理：「等一下！這份名單包含同學的真實姓名和成績，這是他們的個人資料喔，你有得到他們的同意嗎？」

阿哲：「沒有……我只是想省時間……」

AI 助理：「沒關係，你可以先去識別化——把名字換成代號，成績用區間替代，這樣分析一樣準確，而且保護了同學的隱私！」

阿哲：「原來如此！我馬上修改，謝謝你的提醒。」

💡 小叮嚀

- 💡 每次把資料輸入 AI 前先想一想：「這份資料屬於誰？我有權限分享嗎？」
- 💡 臺灣《個人資料保護法》規定，未經同意蒐集他人個資是違法行為。
- 💡 數位足跡一旦留下就很難完全清除；輸入前思考，比事後後悔更重要。

📖 個人資料保護法

個人資料保護法是臺灣保護個人隱私的主要法律，適用範圍包括學生使用 AI 工具時不當分享他人資料的行為。同學可以閱讀 AI 服務的「隱私權政策」，了解資料被儲存、使用的範圍。歐盟的一般資料保護規則 (General Data Protection Regulation, GDPR) 提供更嚴格的框架，可作為進階延伸討論，建議同學可實際嘗試將真實資料「去識別化」後再輸入，建立習慣。

☑ 本節重點整理

- 每次輸入 AI 都可能留下數位足跡，個人資料要謹慎處理。
- 絕對不應輸入：他人姓名＋學號組合、密碼、班級名單、未經同意的照片等。
- 去識別化的四種方法：匿名化、資料遮蔽、資料聚合、虛構化。
- 使用前問自己：「這份資料屬於誰？我有權限分享嗎？」。
- 臺灣個資法保護每個人的隱私，未經同意分享他人個資可能違法。

三、AI 系統設計的六大倫理原則

如果 AI 能幫人做決定，我們該先教它什麼？AI 不是「自己懂得對錯」——它的判斷來自人類的設計與資料，如果設計本身有偏見，AI 的輸出也會不公平。這一節，我們來想想：設計 AI 時，應該把什麼放進去？

情境 | 校園 AI 助手設計

資訊課上，老師請大家設計一個「校園 AI 助手」，阿廷的組別想做一個能自動偵測學生情緒的 AI，希望幫助老師更快關心需要協助的同學。大家討論得很熱烈，但老師突然問了一句：「如果 AI 判斷錯誤，會發生什麼事？」教室瞬間安靜了下來。

現在的 AI 已能完成許多複雜任務，但 AI 並不是「自己懂得對錯」，它做出的判斷往往來自人類提供的資料與設計方式，如果資料本身有偏見，就可能導致不公平的結果；如果系統過度蒐集個人資料，也可能侵犯隱私。

(一) 設計 AI 的六大倫理原則

倫理原則	設計 AI 時需要問的問題
隱私	這個功能真的需要這些資料嗎？資料會被如何儲存與使用？
公平	這個 AI 對每個人都一樣公平嗎？會不會對某些群體產生偏見？
責任	如果 AI 出了問題，誰應該負責？有沒有申訴或修正的機制？
透明	人們知道 AI 是怎麼做決定的嗎？決策過程可以被理解嗎？
安全	如果 AI 出錯，會造成什麼影響？最壞的情況是什麼？
人類自主	使用這個 AI 後，人還有自己選擇的空間嗎？

(二) AI 倫理的兩難困境

很多 AI 倫理問題，不是「對」與「錯」那麼簡單。不同的價值之間常常會互相拉扯，沒有完美的解法，只有更負責任的思考。

倫理兩難	價值衝突	思辨問題
情緒偵測 AI	效率(快速發現需要幫助的學生) vs. 隱私(持續監控情緒)	爲了幫助學生，可以在不告知的情況下偵測情緒嗎？
AI 批改作文	效率(即時回饋) vs. 學習自主(AI 能取代老師的判斷嗎？)	如果 AI 評分和老師不同，應該以誰的標準爲準？
推薦系統	個人化(更好的推薦) vs. 資料蒐集(要知道越多才推得越準)	爲了更好的服務，你願意讓系統知道你的多少事？
AI 輔助決策	效率(AI 分析更快) vs. 公平(AI 的訓練資料有沒有偏見？)	當 AI 推薦的結果對某個群體不公平時，應該如何處理？

(三) 情境判斷討論：自動偵測學生情緒的 AI

思考「覺察學生情緒 AI」這個案例，針對六大倫理原則逐一討論：

情境	六大倫理原則分析	倫理兩難討論
自動偵測學生情緒的 AI 目的：幫助老師更快發現需要關心的學生 做法：透過攝影機、表情分析或語音偵測情緒	◎ 隱私：這個功能真的需要這些資料嗎？ ◎ 公平：AI 對不同族群的準確度一樣嗎？ ◎ 責任：如果 AI 判斷錯誤，誰來處理？ ◎ 透明：學生知道自己被偵測了嗎？ ◎ 安全：判斷錯誤會對學生造成什麼影響？ ◎ 人類自主：老師還是最後決定者嗎？	● 爲了提升學生安全，是否可以犧牲部分隱私？ ● 如果情緒偵測結果可能對學生造成標籤化，應該怎麼設計保護機制？ ● 學生和家長需要被告知並同意嗎？

💡 小叮嚀

- 💡 AI 不會自己「懂得」什麼是對的，它的價值觀來自設計者和訓練資料。
- 💡 倫理原則之間常常互相衝突；沒有標準答案，但不能逃避思考。
- 💡 設計 AI 系統時，受影響的每一種人的聲音都應該被考慮進去。

📖 《AI 倫理建議書》

UNESCO 於 2021 年發布《AI 倫理建議書》，提出以人爲本、包容、透明、問責、安全等核心原則，是目前最具影響力的國際 AI 倫理框架之一。中華民國國家科學及技術委員會也發布了《人工智慧科研發展指引》，可作為本章討論的本土參考文本。同學可以「倫理稽核員」角色演練，針對一個實際 AI 產品逐一評估六大原則，培養批判性思維。

☑ 本節重點整理

- AI 設計的六大倫理原則：隱私、公平、責任、透明、安全、人類自主。
- 倫理兩難沒有標準答案，但思考不同價值的平衡是設計 AI 時不能忽略的一步。
- 真實的 AI 倫理問題常涉及效率 vs. 隱私、個人化 vs. 資料蒐集等衝突。
- 情緒偵測等 AI 系統需要從多角色視角思考其影響。
- 人類自主是最不能被 AI 取代的核心價值。

四、與 AI 協作的即時倫理判斷

當 AI 成為你的合作夥伴時，誰來提醒「這樣真的可以嗎？」在與 AI 協作時，倫理問題往往不是最後才出現，而是可能隨時發生；真正讓合作變得可靠的，仍然是人們彼此之間的溝通、判斷與共同承擔。

😊 情境 | 小琳的設計比賽

小琳參加了學校的創新設計比賽，和同學想做一個能幫助學生安排讀書計畫的 AI 工具。剛開始，大家只專注在功能設計，不過在測試過程中，有同學突然問：「如果 AI 一直叫成績比較差的人多讀書，會不會讓他們壓力更大？」這個問題，讓小琳開始重新思考，原來 AI 不只是工具，它同時也會影響人的感受與選擇。

(一) 與 AI 協作時的即時倫理檢查

在協作流程的每個節點，都可以問這些問題：

協作階段	倫理自我提問
使用資料前	這些資料涉及他人隱私嗎？是否已去識別化？
引用 AI 內容前	這個資訊查證過了嗎？AI 的來源可靠嗎？
設計功能時	這個功能對不同背景的使用者都公平嗎？
發布成果前	是否清楚標示哪些部分是 AI 協助完成的？
遇到分歧時	是否願意停下來討論，而不是假裝問題不存在？

(二) 建立團隊的 AI 倫理共識

團隊應該共同討論的問題	建議的共識原則
哪些資料不能隨便上傳？	個資、他人作品、機密資料一律不輸入
AI 生成的內容要不要標示？	有 AI 參與就說明，標示方式事先統一
如果 AI 出錯該怎麼處理？	建立查證流程，指定負責人確認關鍵內容
是否有人完全依賴 AI 而不動腦？	每位成員要能說明自己在理解什麼、負責什麼
倫理問題由誰負責提出？	任何人都有責任提出倫理疑慮，不能等別人先說

(三) 任務討論：不同角色的 AI 倫理視角

以「自動偵測學生情緒 AI」為主題，你們可以分組扮演不同角色，討論各角色最重視的倫理問題：

團隊應該共同討論的問題	工程師 (開發者)	支援者 (輔導老師)	管理者 (教師)	使用者 (學生)
AI 設計目的與功能 (這個情緒偵測 AI 的目的是什麼？哪些功能是必要的？)				
AI 設計的限制與挑戰 (系統可能有哪些錯誤？誰會受到影響？如何減少傷害？)				
覺察學生情緒 AI 的倫理問題 (這個系統對不同角色各有什麼倫理顧慮？你的角色最重視什麼？)				

討論後，你們可以互相分享：

1. 不同角色之間的倫理顧慮有哪些共同點？有哪些根本的衝突？
2. 如果要設計這個系統，你的組別會優先保護哪一個價值？原因是什麼？

💡 小叮嚀

- 💡 倫理問題不會等到「完成後」才出現，在協作過程的每個節點都需要思考。
- 💡 「任何人」都有責任提出倫理疑慮，不要等別人先說，也不要假裝沒看見。
- 💡 好的 AI 產品背後，一定有願意停下來討論的人。

📖 設計倫理 (Ethics by Design)

設計倫理或負責任創新 (Responsible Innovation) 是目前科技教育的重要趨勢。角色扮演 (Role-playing) 是培養多元視角的有效方法，從不同利害關係人的立場思考，能有效打破「技術中立」的迷思。沒有一個角色的考量是可以完全忽略的，設計 AI 需要多元聲音的參與。

☑ 本節重點整理

- 在與 AI 協作的每個節點都要進行即時倫理檢查。
- 團隊倫理共識五大原則：資料限制、標示規範、查證流程、分工清楚、共同責任。
- 不同角色 (開發者、教師、學生) 對同一個 AI 系統有不同的倫理顧慮。
- 「任何人都有責任提出倫理疑慮」是健康協作文化的核心。
- 倫理問題需要在設計過程中持續討論，不是完成後才處理。

五、當使用 AI 成為日常後，你還會停下來思考嗎？

從規則到習慣——培養「倫理反射」

科技越方便，我們越容易「想都不想就用了」。但真正成熟的 AI 使用者，會把倫理判斷變成一種自然的習慣——在按下「送出」之前，先停下來想一想。

🎧 情境 | 阿傑的錄音

阿傑最近越來越習慣使用 AI，寫報告時，他請 AI 整理重點；做簡報時，他用 AI 生成圖片；就連安排讀書計畫，也會先問 AI 的建議。有一次，阿傑直接把老師上課的錄音貼給 AI 分析，快速生成重點整理，雖然只是想省時間，但也不禁開始反思：「這樣的操作，我學到的是什麼？」「老師同意我錄音了嗎？」「這段錄音輸入 AI，合適嗎？」

很多時候，人們使用 AI 時最先想到的是「快不快、方不方便」。但真正成熟的科技使用習慣，是在行動前先思考：「這樣會不會影響別人？資料使用是否合理？內容是否需要標示？」

(一) 什麼是「倫理反思」？

💡 倫理反思 (Ethical Reflex)

指的是在使用 AI 之前，不需要刻意提醒，就會自然地停下來問自己「這樣做合適嗎？」；就像開車前自動繫安全帶一樣，是一種內化的習慣。

科技越方便，人們越容易忽略背後的影響。AI 倫理不是等到出事才討論，而是應該成爲一種自然反應，包含在接收 AI 生成內容時也要主動查證。

(二) 從規則到習慣：AI 倫理的三個層次

層次	樣貌	如何到達這個層次
📋 知道規則 (Rule)	知道「不能直接複製 AI 的文章繳交」，但有時還是會做	謹記規定、了解後果
💡 內化價值 (Value)	理解學術誠信的意義，願意誠實標示 AI 協助，即使沒人監督	討論倫理兩難、反思自己的行爲
⚡ 成爲習慣 (Habit)	使用 AI 前自動想到倫理問題，倫理判斷成爲工作流程的一部分	反覆練習、在團隊中建立共同文化

(三) 什麼是「倫理摩擦 (Ethical Friction)」？

💡 倫理摩擦 (Ethical Friction)

指的是個人、組織或社會在面對多元價值觀時，因爲道德標準、規範或利益不同而產生的衝突與碰撞。這通常發生在傳統價值與現代創新（如科技發展）交會之際，或者在不同文化的互動中。

◆ 倫理摩擦不是問題，而是成長的機會。

合作過程中難免遇到不同意見：有些人想追求效率，有些人更重視公平或隱私。這些「倫理摩擦」其實很正常，真正重要的，是我們願不願意停下來討論，而不是假裝問題不存在。

🔍 AI 倫理日常自我檢查清單

每次使用 AI 前後，用這張清單問問自己：

使用前	使用後
☐ 我輸入的資料涉及他人隱私嗎？	☐ AI 說的這件事，我查證過了嗎？
☐ 這個場合允許我使用 AI 嗎？	☐ 如果有人問我這份作品內容，我能回答嗎？
☐ 我使用 AI 是爲了學習，還是爲了逃避學習？	☐ 我有沒有遵守這個場合（學校、比賽、工作）對 AI 使用的規定？
☐ 我輸入的 Prompt 會讓 AI 產生有害內容嗎？	☐ 這份作品的呈現，對閱讀它的人公平嗎？

😊 情境小劇場 | 阿傑的倫理反射練習

阿傑（心想）：「等一下……老師同意我錄音了嗎？這段錄音包含老師的聲音，算是老師的個人資料嗎？」

AI 助理：「你願意先停下來想，這就是倫理反射的開始！老師的授課錄音確實涉及智慧財產和同意問題。建議你先問過老師，或改用自己的筆記來整理重點。」

阿傑：「你說得對。我改用我自己整理的手寫筆記來問你問題，這樣應該沒問題了！」

AI 助理：「很棒！這才是把 AI 當學習助手的正確方式，你的倫理反射正在慢慢建立起來了！」

💡 小叮嚀

- 💡 「倫理反射」不是一天就能養成的，但每一次「停下來想一想」都是練習。
- 💡 遇到倫理摩擦時，停下來討論比假裝沒問題更有價值。
- 💡 從「知道規則」到「成爲習慣」，需要反覆練習和團隊文化的支持。

📄 倫理思考自動化

心理學研究顯示，習慣的形成需要重複的行為與一致的環境觸發。在 AI 倫理教育中，教師可以協助學生建立「使用前停下來問三個問題」的固定流程，讓倫理思考逐漸自動化。

諾貝爾經濟學獎得主卡尼曼（Daniel Kahneman）的「快思慢想」理論提醒我們：快速便利的工具更容易讓人跳過慢思考，AI 倫理教育正是在訓練這種「慢思考」習慣。

☑ 本節重點整理

- 倫理反射 = 使用 AI 前自然地停下來問「這樣做合適嗎？」
- AI 倫理的三個層次：知道規則 → 內化價值 → 成為習慣。
- 倫理摩擦是正常的，不同價值觀的碰撞是成長的機會。
- 「AI 倫理日常自我檢查清單」可幫助建立使用前後的倫理習慣。
- 從規則到習慣需要反覆練習與團隊文化的共同支持。

🚀 挑戰任務

難度	任務描述
🌱 初階 觀察記錄	詢問並記錄一位你信任的大人：「你覺得高中生使用 AI 最需要注意什麼倫理問題？」聽聽不同世代的觀點。
🌿 進階 分析比較	選擇一個你常用的 AI 工具（如 ChatGPT、Google Gemini 等），閱讀其隱私權政策，找出哪些資料會被蒐集，並評估是否符合六大倫理原則。
🌳 高階 創作應用	設計一份「班級 AI 使用倫理公約」：包含禁止事項、標示規範、查證流程，以及倫理疑慮的處理機制。讓全班共同討論後通過並實際遵守。

📊 學習檢核表

學習目標	我能做到 ☑	還需練習 🔄
1. 我能說明「學術誠信」和「善用 AI」之間的正確關係。	☐	☐
2. 我知道哪些資訊不應該輸入 AI，以及去識別化的基本方法。	☐	☐
3. 我能用六大倫理原則分析一個 AI 系統的設計問題。	☐	☐
4. 我了解「倫理兩難」的概念，能舉出至少一個真實例子討論。	☐	☐
5. 我能為一個 AI 協作團隊提出具體的倫理共識原則。	☐	☐
6. 我了解「倫理反射」是什麼，並能描述自己目前在哪個層次。	☐	☐
7. 我能使用「AI 倫理日常自我檢查清單」評估自己的使用行為。	☐	☐
8. 遇到倫理摩擦時，我願意停下來討論，而不是假裝沒看見。	☐	☐

📖 本章核心觀念回顧

1. 學術誠信不只是「不作弊」，而是誠實面對學習過程；AI 是助手，不是代替你思考的人。
2. 個人資料一旦輸入 AI 系統，就可能被處理或儲存，要養成「輸入前先想清楚」的習慣。
3. 去識別化是保護資料的基本方法：匿名化、資料遮蔽、聚合、虛構化。
4. AI 設計的六大倫理原則：隱私、公平、責任、透明、安全、人類自主。
5. 倫理兩難沒有標準答案，但思考不同價值的平衡，是設計 AI 時不能忽略的一步。
6. 與 AI 協作時，需要建立團隊的倫理共識，任何人都有責任提出倫理疑慮。
7. 「倫理反射」是目標：讓倫理判斷成為使用 AI 前的自然反應，而非事後的亡羊補牢。

💡 最重要的一句話

科技雖然做得到，但人們仍需要思考：這樣做真的合適嗎？



AI 助你 高效學習

讓AI變成你最強的學習神隊友

「讓AI神隊友開啟你的高效學習能力。」

想讓AI變成你最強的學習神隊友嗎？

它可以幫你整理資料、幫助你更深入思考。

只要你學會怎麼下達清楚的提示詞，AI就能聽懂你的需求，給出最貼心、最精準的回覆。

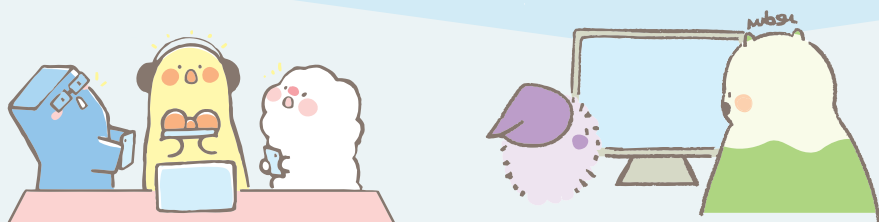


📌 本章學習目標：

- 明白AI是陪伴學習與思考的夥伴，而非答案機
- 學會用角色、任務、背景、格式寫出有效提示詞
- 能用AI輔助各科任務，如作文與幾何圖解
- 堅持先動腦並主動查證，不直接採用AI答案
- 運用AI進行辯論練習、生涯探索與小組討論

🎯 萬用提示詞公式：角色R + 任務T + 背景B + 格式F

- 角色 (Role)：你希望AI扮演的身份、職業或人物
- 任務 (Task)：你希望AI具體執行的指令或要完成的動作，是提示詞的核心，必須清晰、明確
- 背景 (Background)：你提供給AI的情境、相關資訊或先前對話
- 格式 (Format)：你希望AI輸出答案的特定結構或樣式



【準備篇】如何聰明選用 AI 工具？

1. 檢視使用條款與隱私權保護

在使用任何 AI 工具前，務必先閱讀其隱私權政策與年齡限制，了解哪些資料會被收集以及是否符合倫理原則。請特別留意，你的對話紀錄是否會被拿去訓練未來的 AI 模型？切記，絕對不要將個人隱私資料（如身分證字號、住址、密碼等）輸入給 AI，保護好自己的數位安全。

2. 選對工具做對事

不同的 AI 工具擁有不同的超能力。如果你需要的是靈感發想、文章摘要或多語言翻譯，擅長文字處理的語言模型是最佳選擇；但如果你需要生成精美的簡報配圖，或是處理複雜的數學運算與圖表，就需要挑選專門針對圖像生成或邏輯運算優化的 AI 工具。先釐清自己的任務目的，才能找到對的神隊友。

3. 留意模型版本與更新狀態

AI 的大腦進化得非常快！同一個工具往往會有不同版本的語言模型，新版本的 AI 通常理解力更強、出錯率較低。此外，也要注意該 AI 的知識庫更新到什麼時候——有些 AI 無法連上網路，提供的可能是舊資訊。隨時留意工具的更新動態，能幫助你更準確地評估它給出的答案是否符合現況。

掌握好指令公式與挑選工具的技巧後，接下來，你會看到不同科目的真實例子和練習，一步步學會怎麼運用 AI，讓它變成陪你一起學習、一起創作的超級夥伴。當你愈會用它，你的思考力、創造力也會一起升級！

一、怎麼讓 AI 變成語文課程的神隊友？



∴ 圖片生成提示詞參考：

A teacher-student scene with a friendly AI robot sitting beside a high school student at a desk with open Chinese textbooks. The robot is pointing at a text passage while the student takes notes. Flat illustration, educational, blue and orange tones, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

(一) 寫好提示詞 (Prompt)，讓 AI 秒懂你！

只要輸入一篇網路新聞或社論，AI 就能幫你自動抓出重點、整理摘要，還會丟出一些讓你思考的問題，像是：「作者的主要想法是什麼？」、「這篇報導有沒有偏見呢？」

如果你要寫議題作文或練習論證，AI 也能幫你準備正反兩方的觀點、相關資料，甚至連名言佳句都能幫你找好，讓你更容易整理想法、建立有邏輯的論點，寫出更有說服力的文章。

(二) 文言文白話翻譯與賞析

AI 可以幫你把文言文翻成白話文，還會順便解釋詞語、句型和修辭，甚至幫你補上那個時代的背景和文化小知識，讓你更容易看懂古文，不再覺得難。

更酷的是，AI 還能讓古文角色「穿越」到現代！像是把《水滸傳》或《紅樓夢》裡的人物放進現在的生活情境，幫你生成對話、劇本或小故事，讓你一邊玩創意、一邊更懂文學在說什麼。

(三) 口語表達與溝通

只要輸入你的簡報主題，AI 就能幫你自動生出大綱，還會一起準備圖片或視覺素材，讓你的簡報更有吸引力。如果你要練習演講，AI 也能幫你當教練，它會給你發音、語調、語速，甚至內容結構的建議，幫你愈講愈順、愈講愈有自信。

(四) 寫作靈感與結構

你只要輸入一個寫作主題，AI 就能幫你想出好多不同的方向、架構跟關鍵詞，讓你不再卡關、靈感滿滿！它能陪你從模仿開始，慢慢練到能自己創作出有風格的文章。

如果你想挑戰寫詩或散文，AI 也可以幫忙提供意象、比喻和情境描寫的例子；你還能和 AI 一起創作詩句或段落，在互動的過程中學到新的寫作技巧，讓創意像泉水一樣冒出來。

(五) 跨領域議題探究

你也可以把 AI 和不同科目結合起來，像是歷史或地理。例如，你可以請 AI 幫你整理某個歷史事件相關的文學作品，或是分析不同時代人們寫信的方式有什麼不一樣，這樣就能從不同角度來比較、發現有趣的跨領域連結。

(六) 練習時間

短文練習題目：一場暴雨——從日常到變遷的凝視（至少600字）

夏日的午後，原本炙熱的陽光被迅速集結的烏雲吞噬；放學的途中，豆大的雨點毫無預警地猛烈砸下。這場景我們或許熟悉，但你是否曾留意，近年來新聞報導中的「暴雨」似乎有些不同？

「短延時強降雨」、「破紀錄的時雨量」……這些詞彙頻繁出現，伴隨而來的是都市部分區域瞬間變為汪洋、山區土石滑落的驚心畫面。我們腳下的這片土地，正經歷著氣候變遷帶來的嚴峻考驗。天氣，好像越來越極端了。

這一次，請你寫的「一場暴雨」，不只是一場個人記憶中的雨，更是一面映照出時代問題的鏡子。請你以「一場暴雨」為核心，連結你的個人經驗與對臺灣社會時事、氣候變遷的觀察與思考。

想一想

如何使用AI幫助思考寫作觀點？

◆ AI 使用指南：以教育部因材網e度的寫作精靈為例

通用型學習夥伴
教學影片
操作說明
注音音：



嗨～我是您的AI教學夥伴e度。

你要先跟我說題目，我會引導你把內容一段一段講出來、寫出來，最後一起修改變更順。先告訴我：你想寫的主題是什麼？例如：我的暑假、一次難忘的經驗、我最喜歡的動物。

慢速
正常
快速

請同學進入因材網後，點按畫面右下角e度圖案，並選擇[寫作精靈]。

寫作精靈 切换模式會重新對話喔!

練習一：輸入較模糊的提示詞，如：請幫我寫一篇600字，關於一場暴雨的作文，此時e度會邀請你提供更多的任務細節，表達想法和觀點。
請幫我寫一篇600字，關於一場暴雨的作文



這是一個很有趣的主題。我們可以一起來構思它的內容。你覺得這場暴雨是你親身經歷的，還是你想要創作一個虛構的故事呢？或者你更想要描述暴雨帶來的影響，比如對自然、城市或人們生活的改變？

慢速
正常
快速
00:46:24

練習二：若是能換個較詳細的提示詞(角色+任務+背景+格式)，e度會給你更豐富的引導及回答，留給你試試。

🔑 詳細提示詞範例(角色+任務+背景+格式)

角色：你是一位關心臺灣環境議題的記者。

任務：提供我幾個撰寫「一場暴雨」作文的切入點。

背景：請結合近年來「短延時強降雨」、「都市熱島效應」等新聞關鍵字。

格式：請從「一個高中生的視角」出發，分別以「個人經歷」、「社會觀察」、「環境反思」三個角度提供具體想法。

通用型學習夥伴

教學影片

操作說明



嗨～我是您的AI教學夥伴e度。

你要先跟我說題目，我會引導你把內容一段一段講出來、寫出來，最後一起修改變更順。
先告訴我：你想寫的主題是什麼？例如：我的暑假、一次難忘的經驗、我最喜歡的動物。

⏮ 慢速

▶ 正常

▶▶ 快速

**輸入較詳細的提示詞(角色+任務+背景+格式)，
e度寫作精靈會給你更豐富的引導及回答。**

寫作精靈



切換模式會重新對話喔!

角色：你是一位關心臺灣環境議題的記者。

任務：提供我幾個撰寫「一場暴雨」作文的切入點。

背景：請結合近年來「短延時強降雨」、「都市熱島效應」等新聞關鍵字。

格式：請從「一個高中生的視角」出發，分別以「個人經歷」、「社會觀察」、「環境反思」三個角度提供具體想法。

雖然e度很厲害，但有時候腦袋也會打結～重要的事情，還是要靠你聰明的腦袋再檢查一遍喔！

⚡ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的寫作主體性。

同時，請找出 AI 回答中一個需要確認的地方，再用課本、老師提供的資料或可信網站查證。

學習主題：背單字——高中英文參考字彙表 Level 3

背單字是學英文的基本功，但你是不是常常覺得——背了又忘、忘了又背？用死記硬背的方式真的超無聊，背到後來根本沒動力。不過現在有 AI 幫忙，你可以用更聰明、更有趣的方式記單字！它能幫你做練習、出小測驗、甚至用例句陪你複習，讓背單字不再那麼痛苦，效率直接升級。

想一想

如何使用 AI 幫助更有效率地背單字？

◆ 請自己選擇適合的 AI 來練習

► 較模糊的提示詞 → 請幫我練習高中英文參考字彙表 Level 3 單字。

🔑 詳細提示詞範例（角色＋任務＋背景＋格式）

角色：你是一位高中英文老師，同時也是精通「主動回想」與「間隔重覆」技巧的記憶專家。

- **任務A：**請針對這些單字，提供我一些結合「情境例句」、「字源拆解」或「諧音聯想」的記憶方法。
- **任務B：**請針對這些單字，設計一個互動測驗。

背景：我已經學過高中英文參考字彙表 Level 3 的前五個單元單字，但常常混淆相似字和詞性變化。

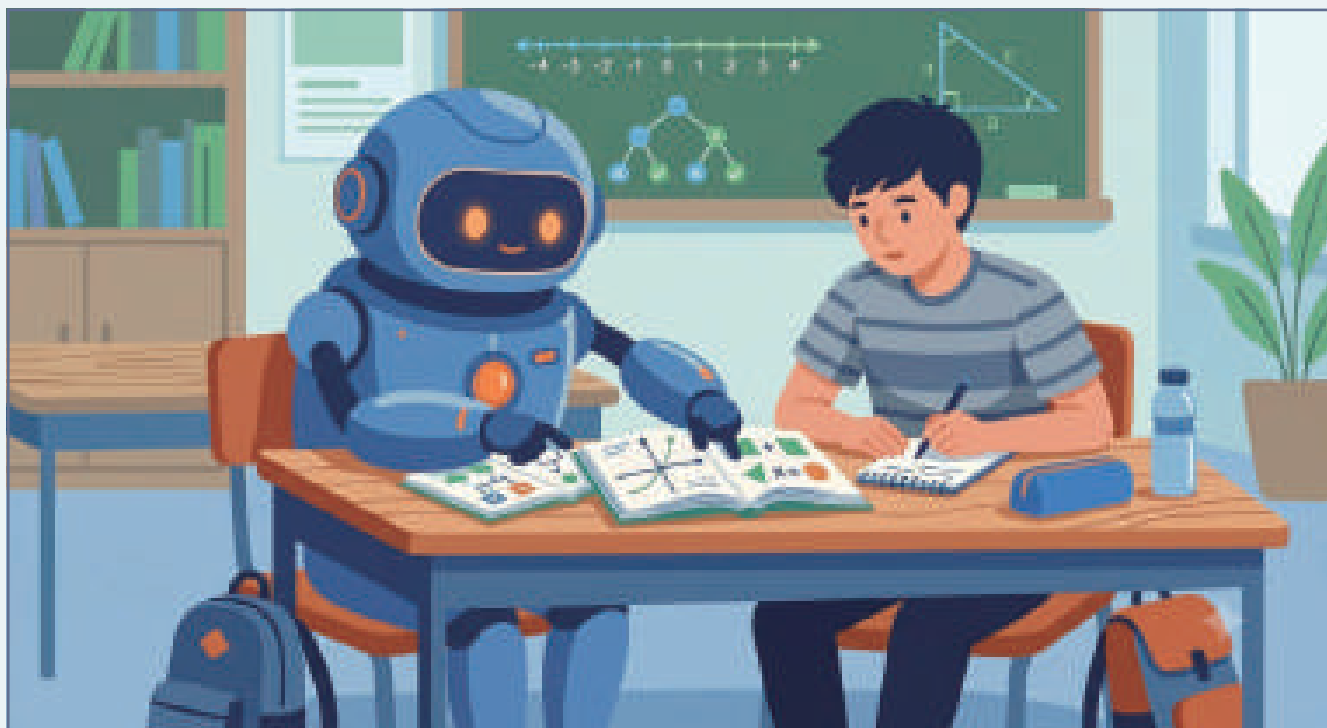
- **格式A：**記憶方法請以表格形式呈現，包含單字、中文解釋以及至少兩種記憶技巧。
- **格式B：**互動測驗請以測驗形式呈現。

⚠ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的寫作主體性。

同時，請找出 AI 回答中一個需要確認的地方，再用課本、老師提供的資料或可信網站查證（例如確認單字與翻譯是否正確）。

二、怎麼讓 AI 變成數學課程的神隊友？



✿ 圖片生成提示詞參考：

A friendly AI robot tutoring a student on math problems. On a whiteboard behind them: number line, right triangle with labeled sides, and a tree diagram for combinations. Flat illustration, clean educational style, blue and green tones, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

(一) 解二次不等式

當你在學二次不等式時，如果想更清楚地了解圖解（拋物線）到底代表什麼，你可以把不等式輸入到 AI 裡，請它一步步幫你解題，還會幫你畫出函數圖形與 x 軸的關係圖！

AI 不只會標出解的範圍（例如圖形在 x 軸上方或下方的區間），還會用簡單的文字告訴你「為什麼判別式和正負號會決定圖形的走向」，幫你更容易看懂不等式的幾何意思。

你也可以多試幾個不同的二次不等式（像是兩交點、相切或無解的情況），看看 AI 畫出來的拋物線有什麼變化，從中自己歸納出規律，就能更快掌握這個單元的重點。

(二) 學習平面向量與幾何意義

學習平面向量時，想驗證兩個向量之間的幾何關係嗎？你只需將兩個向量的座標元件輸入 AI，請 AI 判斷它們是否互相垂直或平行。AI 不僅能精準給出判斷結果，還能進一步在座標平面上還原這兩個向量的實際樣貌，甚至幫你計算出它們的內積與夾角餘弦值。

你可以隨意改變向量座標，讓 AI 快速驗證多組數據，從中強化對向量運算與幾何意義的直觀理解！

(三) 學習排列組合與機率

計算複雜的排列組合問題，你可以將問題情境（例如從 10 個人中選 3 人擔任不同職務）輸入 AI，請 AI 列出所有可能的排列或組合方式。AI 能提供公式與計算過程，甚至能用樹狀圖視覺化呈現，幫助理解每個步驟的意義。

此外，也可將不同情境輸入 AI，讓 AI 計算機率，協助你掌握排列組合在機率中的應用。

(四) 練習時間

學習主題：三角函數的基礎概念與應用

三角函數是數學中一個非常重要的主題，它描述了直角三角形中角與邊之間的關係。這些關係透過正弦（ $\sin$ ）、餘弦（ $\cos$ ）、正切（ $\tan$ ）等函數來表示，並廣泛應用於物理、工程、天文學等領域。

想一想

如何使用 AI 幫助理解與練習三角函數？

◆ 請自己選擇適合的 AI 來練習

► 較模糊的提示詞 → 請幫我解三角函數題目、三角函數是什麼？

🔑 詳細提示詞範例（角色＋任務＋背景＋格式）

角色：你是一位耐心的數學家教。

- **任務A**：請用生活化的例子解釋什麼是「弧度」以及它和「度數」的關係。
- **任務B**：請提供一個表格或圖示，清晰地說明 $\sin$ 、 $\cos$ 、 $\tan$ 在四個象限中的正負號變化，並簡要說明其原因。
- **任務C**：針對每個象限，請提供一個具體的例子，例如某個特定角度的三角函數值計算，幫助我理解。

背景：我正在學習高中一年級的三角函數，對於「弧度」這個概念感到有些困惑，也分不清楚 $\sin$ 、 $\cos$ 、 $\tan$ 在不同象限的正負號變化。

- 格式A：請以條列式說明，並在解釋弧度時，可以運用比喻。
- 格式B：提供互動練習題。

⚠ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的解題主體性。

同時，請找出 AI 回答中一個需要確認的地方，親自拿紙筆將 AI 的計算過程與答案「驗算」查證一遍。

三、怎麼讓 AI 變成自然課程的神隊友？



✎ 圖片生成提示詞參考：

A high school science lab setting with a friendly AI robot helping students analyze experimental data on a tablet. Beakers and lab equipment visible. Flat illustration, blue and green scientific colors, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

(一) 複雜概念白話文與生活比喻

高中自然科充滿了抽象的理論和複雜的機制（例如：化學平衡、物理的電磁感應、生物的中心法則等）。如果你對課本上生硬的定義感到困惑，可以請 AI 把這些概念翻譯成白話文，或者請它舉個生活中的例子；透過生動的比喻，原本看不見、摸不到的科學原理就會變得像看故事一樣容易理解。

(二) 實驗數據分析與報告撰寫建議

做完自然科實驗，面對一堆數據不知道怎麼寫報告嗎？你可以把數據趨勢告訴 AI，請它協助你找出變因之間的關聯，或是幫你整理探究與實作的報告架構（例如：研究動機、實驗方法、結果分析、結論等）。它甚至能引導你思考實驗中可能出現的誤差來源，讓你的報告更具科學嚴謹性。

(三) 客製化解題教練與素養題庫

遇到解不出來的物理或化學計算題，不要只叫 AI 給答案！你可以請 AI 扮演「蘇格拉底」，一步步引導你思考解題觀念。

此外，現在大考越來越強調科學素養，你也可以請 AI 針對特定單元（例如：氣候變遷、基因編輯、綠能科技）出幾題情境素養題，幫助你檢驗學習成果。

(四) 練習時間

學習主題：化學反應的拔河比賽——化學平衡與勒沙特列原理

你有沒有想過，為什麼打開汽水瓶蓋時，氣泡會不斷冒出來，但鎖緊瓶蓋後，氣泡就不見了呢？這其實是一場看不見的化學拔河比賽！

在化學的世界裡，很多反應不是單向進行的，而是像在玩拔河一樣，正向和逆向反應同時發生。當兩邊的力量一樣大時，看起來就像靜止不動，這就是「化學平衡」。

而「勒沙特列原理」就像是一個討厭改變的固執裁判，只要你對系統施加壓力、改變溫度或濃度（比如打開瓶蓋讓壓力變小），這個裁判就會想辦法讓反應往反方向移動，努力抵消你的改變！學會看懂這個原理，就能輕鬆破解高中化學的大魔王。

想一想

如何使用 AI 幫助我理解複雜的科學機制並進行測驗？

◆ AI 使用指南：以教育部因材網e度的「自然探究精靈」為例

練習一：輸入較模糊的提示詞試一試。

►較模糊的提示詞 → 什麼是化學平衡？解釋勒沙特列原理。幫我出三題化學題。

通用型學習夥伴 教學影片 操作說明 注音：(?)

請告訴我一個你想研究的主題，我會帶你走 5E 探究：先觀察與探索 → 解釋原理 → 延伸挑戰 → 最後用小測驗確認你有沒有學會！

慢速 正常 快速

解釋勒沙特列原理

慢速 正常 快速

勒沙特列原理(Le Chatelier's Principle)指出，當一個化學平衡系統受到外界條件(如濃度、壓力或溫度)的改變時，系統會自動調整以減少這種改變的影響，從而重新達到新的平衡。你對這個原理的應用有什麼具體的問題嗎？或者你想要舉個例子來更好地理解嗎？

慢速 正常 快速

自然探究精靈 勿換模式會重新對話喔!

練習二：輸入較詳細的提示詞(角色+任務+背景+格式)試試，比較一下與練習一的回答，是不是更詳盡了？

🔑 詳細提示詞範例(角色+任務+背景+格式)

角色：你是一位幽默且擅長用生活化比喻教學的高中化學老師。

任務：請用「一個非常忙碌的吃到飽餐廳供餐區」作為比喻，詳細解釋「勒沙特列原理」中「系統會盡力抵消外加改變」的核心概念。

具體說明：在一個「放熱反應」中，如果我(1)增加反應物的濃度、(2)升高系統的溫度，化學平衡會如何移動？

解釋完後，請根據這個單元出 3 題符合學測趨勢的「情境素養選擇題」(例如結合打開碳酸飲料、哈伯法製氨或游泳池的酸鹼平衡等生活應用)，讓我測試自己的理解程度。

背景：我是一名高中生，正在學習化學科的「化學平衡與勒沙特列原理」，但我對於濃度、溫度、壓力改變時，平衡到底會往左還是往右移動總是搞不清楚，只能死背。

● **格式A**：生活比喻解說，請善用列點，並用粗體字標示專有名詞與比喻的對應關係。

● **格式B**：3 題選擇題，請將解答與詳解獨立放在最後面，讓我能先自己思考。

通用型學習夥伴

教學影片

操作說明

注音  



第一部分：生活比喻解說

讓我們用一個「吃到飽餐廳的供餐區」來比喻 勒沙特列原理，來幫你理解化學平衡如何因外部改變而調整。

1. 化學平衡系統 = 吃到飽餐廳的供餐區

想像一個吃到飽餐廳的供餐區，廚房負責補充食物(反應物)，顧客則不斷取用食物(生成物)。這個供餐區的目標是保持餐盤的數量穩定，讓所有人都能吃到食物。

2. 增加反應物濃度 = 廚房一次補充了很多食物

如果廚房突然送來了超多的食物(增加反應物)，那麼供餐區就會顯得「太擁擠」，顧客會加速取用食物，直到餐盤的數量重新達到平衡。

對應化學平衡的移動：系統會向生成物方向移動，因為要「消耗掉」多餘的反應物。

3. 升高系統溫度(放熱反應) = 餐廳變得太熱，顧客開始離開

自然探究精靈 

切換模式會重新對話喔！

⚡ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的解題主體性。

同時，請找出 AI 回答中一個需要確認的地方，再用課本、老師提供的資料或可信網站查證「實驗證據與科學定律」。

四、怎麼讓 AI 變成社會課程的神隊友？



∴ 圖片生成提示詞參考：

An AI robot standing next to a world map and historical timeline, pointing at different periods. A high school student stands beside looking curious. Flat illustration, warm earth tones and blue, educational style, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

(一) 世界文明的發展與交流

如果你想更了解歷史人物，可以讓 AI 帶你玩一場「時空訪談」！像是輸入秦始皇的資料，請 AI 扮演他，回答你關於統一文字、度量衡或修築長城的問題。你還可以把這段「訪談」錄音或轉成文字，放進你的探究報告裡。這樣不只更有臨場感，還能幫你更深入理解他的想法和對歷史的影響。

(二) 自然環境與人類活動

你可以用 AI 的圖像分析功能，來幫你比較同一個地區10年前和現在的衛星空照圖，AI 會自動幫你標出有哪些變化，比如哪裡變成城市、哪裡的森林被砍掉、或者哪塊農地變了用途。這樣一來，你就能清楚看到人類活動怎麼改變地景，還能在報告裡更具體地分析對環境造成的影響。

(三) 練習時間

學習主題：社會辯論角色扮演——是否應全面禁止校園內販售含糖飲料？

近年來，學童肥胖問題日益受到關注，含糖飲料被認為是其中一個重要的因子，為了學生的健康，許多國家和地區開始推動校園內禁止販售含糖飲料的政策。然而，也有人認為這限制了學生的選擇自由，且可能效果不彰，在校園中，我們是否應該全面禁止含糖飲料的販售呢？

這次的辯論角色扮演，將邀請同學從不同立場出發，透過資料蒐集、觀點組織與口語表達，訓練思辨能力與團隊合作精神。

想一想

如何使用 AI 幫助思考辯論觀點？

◆請自己選擇適合的 AI 來練習

►較模糊的提示詞 → 請幫我寫一篇關於校園禁售含糖飲料的贊成或反對論點。

🔑 詳細提示詞範例（角色＋任務＋背景＋格式）

角色：你是一位中學生，正在準備一場關於校園含糖飲料販售的辯論。

任務：提供我正反方在「是否應全面禁止校園內販售含糖飲料」這個議題上，可能使用的論點、支持證據與反駁策略。

背景：中學生肥胖率上升，含糖飲料是重要因素。學校希望促進學生健康，但也有學生認為自己有選擇食物的權利。

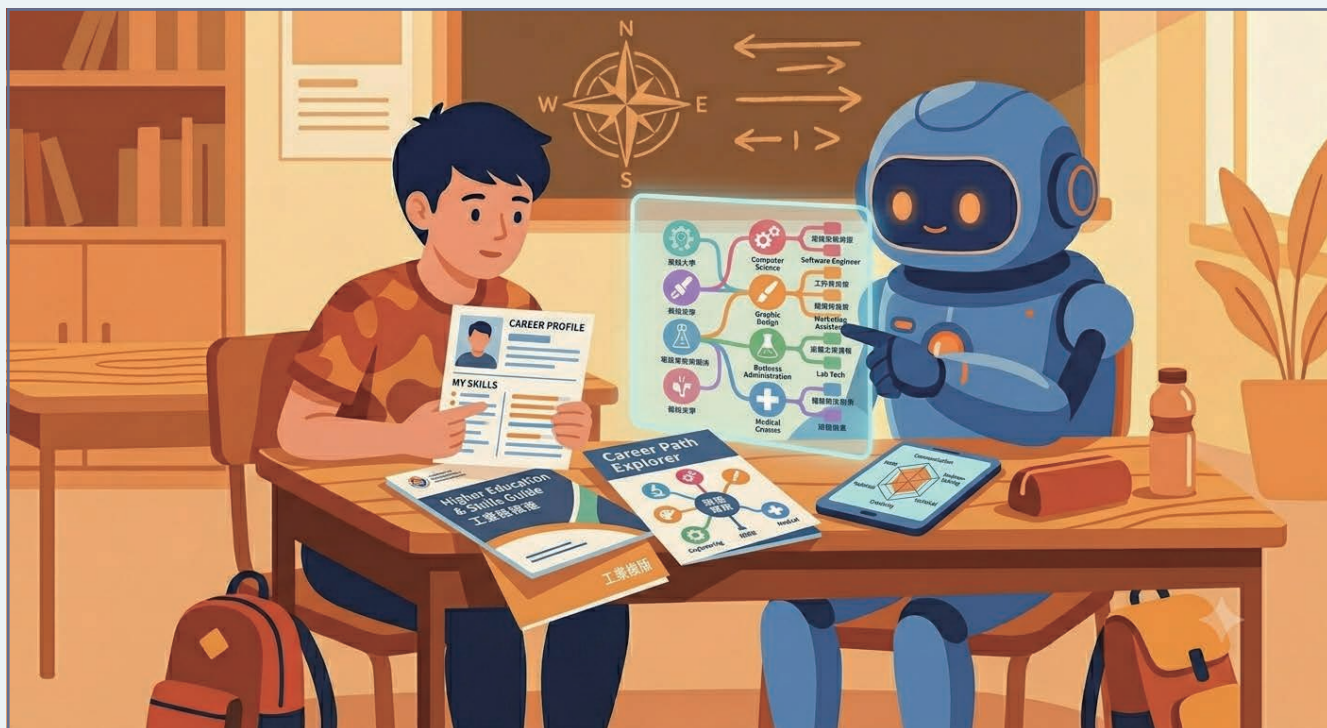
格式：請以條列式呈現，並分別針對「學生健康與飲食習慣」、「學生自主權與教育方式」、「政策可行性與替代方案」三個面向提供具體觀點。

⚡ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的探究主體性。

同時，請找出 AI 回答中一個需要確認的地方，再對照「課本或官方資料」查核相關資料是否正確。

五、怎麼讓 AI 變成生涯規劃的神隊友？



∴ 圖片生成提示詞參考：

A high school student with a compass and roadmap, with a friendly AI robot pointing at different career path branches on a digital screen. Flat illustration, warm optimistic colors, educational setting, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

製作學習歷程檔案時，AI 神隊友可在繁雜的資料中找出重點，並將其轉化為有說服力的文字，可協助的工作有以下四項：

(一) 統整與歸納資料

在準備學習歷程時，通常會有很多分散的檔案，例如專題報告、活動照片、社團紀錄等。你可以將所有相關的文字資料（如報告、心得）上傳檔案給 AI，請它分析內容，並歸納出核心主題或關鍵能力。

例如，AI 能從上傳的多份報告中，統整出在「解決問題」或「數據分析」方面展現的能力，幫助你快速掌握自己的優勢。

(二) 撰寫與優化自述

學習歷程的自述（如學習心得、課程學習成果反思）是展現個人特質的關鍵，你可以提供 AI 幾份草稿，請它協助潤飾語句。如果不知道如何起頭，也可以提供核心概念，請 AI 撰寫不同風格的開頭或結尾，例如「感性分享」或「數據化呈現」，作為參考與啟發。

(三) 設計與美化版面

一個好的學習歷程檔案，除了內容豐富，版面設計也很重要。你可以請 AI 提供版面設計的建議，例如，給 AI 幾個關鍵字（如「簡約」、「科技感」、「活潑」），請它推薦配色、字體與排版樣式；也可以利用 AI 繪圖工具，生成與主題相關的圖像或圖標，讓學習歷程檔案更具個人風格。

(四) 模擬面試與問答

在準備面試時，能預先掌握教授可能提出的問題，你可以將自己的學習歷程檔案提供給 AI，請它模擬大學教授的角色，提出可能的問題。AI 可以根據檔案內容詢問「為什麼選擇這個主題？」或「這個專題最大的啟發是什麼？」等問題，讓你可以練習回答，提前做好準備。

(五) 練習時間

學習主題：我的未來藍圖——結合個人興趣與未來趨勢的生涯探索計畫

請 AI 當你的「生涯探索領航員」，規劃未來更有方向。這個世界變化得超級快，每天都有我們想不到的新工作冒出來，生涯規劃不能只看當下，應從探索自我興趣與專長出發，結合世界趨勢，才能制訂一份既適合我們、又具未來發展性的生涯藍圖。

想一想

如何使用 AI 幫助我進行生涯探索？

◆ AI 使用指南－以教育部因材網e度的「一般模式」為例

- 較模糊的提示詞 → 我做一份生涯規劃報告、告訴我什麼工作最賺錢、我喜歡打電動，可以做什麼工作？

🔑 詳細提示詞範例（角色＋任務＋背景＋格式）

角色：你是一位專業的生涯諮詢師，特別擅長分析臺灣的產業趨勢與新興職業。

任務：請根據我的興趣、專長與個性，為我分析並推薦 3 個未來 10 年內在臺灣具有發展潛力的生涯方向（不一定是單一職業）。

背景A—興趣：喜歡看科幻電影、玩策略遊戲、關心環保與動物保育議題。

背景B—在校學科：國文和生物是我的強項，數學比較弱。

背景C—個性：我很有耐心，喜歡團隊合作，也樂於找出問題並解決它。

格式——針對每一個生涯方向，請用以下格式說明：

【方向名稱】：這個生涯領域的核心是什麼？

【連結點】：這個方向如何與我的興趣與專長結合？

【關鍵技能】：若要往這個方向發展，我現在可以開始培養哪些硬實力和軟實力？

【探索建議】：建議我可以透過哪些大學科系、線上課程、或課外活動來進一步探索這個領域？

◆AI 使用指南—以教育部因材網e度的「一般模式」回答為例，節錄如下圖

The screenshot displays the 'e-degree' AI interface. At the top, there are tabs for '通用型學習夥伴' (General Learning Partner), '教學影片' (Teaching Video), and '操作說明' (Operation Guide). A '注音' (Pinyin) toggle is on the right. The main content area, highlighted in green, discusses the core of 'Digital Content and Immersive Experience Design', listing elements like visual design, interactive design, narrative design, sound design, and UX design. Below this, a text input box is shown with a dropdown menu set to '一般模式' (General Mode). A red callout box points to the input area, stating: '能與 AI 多輪對話，進行生涯探索，對相關領域有更深入的了解。' (Can have multiple rounds of dialogue with AI for career exploration, gaining deeper understanding of related fields). To the right of the input box are icons for voice input, a search icon, and a submit button. At the bottom, a disclaimer reads: '雖然e度很厲害，但有時候腦袋也會打結～重要的事情，還是要靠你聰明的腦袋再檢查一遍喔！' (Although e-degree is very powerful, sometimes the brain can get stuck ~ important things still need to be checked by your own smart brain!).

⚡ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的探究主體性。

同時，請找出 AI 回答中一個需要確認的地方，再對照「課本或官方資料」查核相關資料是否正確。

六、怎麼讓 AI 變成小組討論的神隊友？



✧ 圖片生成提示詞參考：

Four students sitting around a table with a friendly AI robot as the fourth member. Speech bubbles showing different ideas. Digital whiteboard with pros/cons table visible. Flat illustration, collaborative team colors, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

(一) 擔任特定專長的「外掛組員」

小組討論時如果缺乏某方面專長（例如：沒人擅長數據分析，或沒人懂程式），你可以邀請 AI 擔任該領域的「外掛組員」。只要賦予它特定身分，它就能以專家的角度參與討論，補足團隊能力的缺口，提供專業建議。

(二) 突破盲點的「反方組員」

大家意見太一致時，反而容易忽略潛在問題。你可以請 AI 擔任「專門找碴的組員」或「魔鬼代言人」，針對小組的共識提出質疑和挑戰，幫你們提前踩雷，這能幫助你們的企劃或報告變得更嚴謹。

(三) 統整意見的「文書組員」

大家七嘴八舌討論完，卻沒人知道結論是什麼？請 AI 擔任你們的「文書處理組員」，把群組裡雜亂的對話紀錄丟給它，讓它快速歸納出重點、分類優缺點，並清晰列出下一步每個人該做什麼。

(四) 練習時間

學習主題：專題討論大挑戰——邀請 AI 成為你們的第四位組員

你有沒有過這種經驗？小組討論時，大家天馬行空聊了一節課，最後卻不知道下一步該做什麼？或者大家想法都一樣，交出去的報告卻被老師說缺乏深度與可行性？

在這種時候，你可以直接把 AI 當成虛擬第四位組員加入團隊。你不僅能指定它的個性（例如：務實冷靜、創意十足），還能賦予它專屬任務（例如：評估預算、提出反對意見）。

想一想

如何設定 AI 的個性和專長，讓它成為我們小組不可或缺的組員？

◆ AI 使用指南－以教育部因材網e度的「樂學精靈」為例

► 較模糊的提示詞 → 幫我們想一個專題題目、評價一下我們的企劃、幫忙整理我們的討論紀錄。

詳細提示詞範例（角色＋任務＋背景＋格式）

角色：你是我們小組的第四位組員「小智」，你的個性務實理性，非常擅長資料收集與評估「執行可行性」。說話語氣請像一般高中生一樣自然，把我們當成你的同學。

- **任務A：**以「組員」的身分，客觀分析我們三個人提案的優點與高中生在校園執行時會遇到的最大困難。
- **任務B：**提出一個融合大家想法，而且在學校裡最容易落地執行的綜合企劃案。
- **任務C：**用同學在 LINE 群組討論的口吻，寫一段話跟我們分享你的結論，並問我們一個關鍵問題來引導下一步的討論。

背景：我們是一個高中生專題小組，主題是「校園減塑行動方案」。目前組員A想推廣租借環保杯，組員B想辦減塑集點抽獎活動，組員C想拍宣傳微電影。我們下週就要交企劃書了，但不知道該選哪一個，或者該怎麼結合起來。

- **格式A：**3個提案的優缺點與執行難度分析（請用表格呈現，方便大家看）。
- **格式B：**小智的綜合企劃案（請條列說明）。
- **格式C：**群組對話訊息（請加上適當的表情符號Emoji，展現組員感）。

◆ AI 使用指南－以教育部因材網e度的「樂學精靈」為例

下圖為樂學精靈成為組員，與其多輪式對話，協助整理提案優缺點及執行難度表格。



The screenshot shows the '樂學精靈' (Happy Learning Spirit) interface. At the top left is a small cartoon character. Below it is a table with four columns: '提案' (Proposal), '優點' (Pros), '缺點' (Cons), and '執行難度 (高中生在校園執行)' (Execution Difficulty (High School Students on Campus)). The table contains three rows of proposals. Below the table are playback controls (慢速, 正常, 快速) and a timer (00:29:18). At the bottom, there is a dropdown menu labeled '樂學精靈' and a button labeled '切換模式會重新對話喔!'.

提案	優點	缺點	執行難度 (高中生在校園執行)
A. 租借環保杯	實際減少一次性垃圾，效果最直接。	需要資金購買環保杯，管理和清潔是個大工程。	高
B. 減塑集點抽獎	提高同學參與意願，增加活動趣味性。	獎品費用、點數系統設計和兌換需要人力。	中
C. 宣傳微電影	傳播廣泛，能有效提升減塑意識。	拍攝、剪輯需要專業技能和設備，耗時費力。	中

⚠ 防雷與查證任務：不能直接採用 AI 答案！

記得提示詞得先自己寫，再讓 AI 協助優化，保持你的探究主體性。
同時，請找出 AI 回答中一個需要確認的地方，再對照「課本或官方資料」查核內容是否正確。

七、延伸學習：認識 AI 的進階型態——人工智慧代理人 (AI Agent)

在前面的內容中，我們學會了如何使用「角色 + 任務 + 背景 + 格式」的指令來讓 AI 成為神隊友；但隨著科技進步，AI 已經從「一問一答的對話工具」進化成了更聰明、更具主動性的**人工智慧代理人 (AI Agent)**。

(一) 什麼是 Agentic AI 與 AI Agent? 它們和一般的 AI 有什麼不同?

近年科技界熱門的 Agentic AI，指的是 AI 不再只是被動回答，而是具備了主動解決問題的自主行動力。當一個 AI 系統運用這種能力，能夠自行設定目標、拆解步驟並完成複雜任務時，我們就會稱它為 AI Agent (人工智慧代理人)——就像是一個能為你獨立執行專案的專屬數位夥伴。

如果你把一般的 AI 當成圖書館員 (你問一個問題，它幫你找出一本書或一段答案)，那麼 AI Agent 就是你的「專屬學習秘書」。

- **一般的 AI (被動)**：你輸入一個指令，它給出一個回答，然後對話就結束了。如果你需要更多資訊，你必須再想下一個指令。
- **AI Agent (主動，即展現 Agentic 能力)**：它具備「目標導向」與「自主規劃」的能力。只要你給它一個最終大目標，它能自己把任務拆解成好幾個步驟、主動判斷需要使用什麼工具，甚至能在過程中發現錯誤並自我修正，直到幫你完成目標為止。

(二) AI Agent 的三大核心能力

要被稱為一個稱職的 AI Agent，通常具備以下三個特徵：

- **大腦 (記憶與推理)**：它能記住你們之前的對話脈絡，並根據你的學習習慣進行邏輯推理。例如，它記得你數學的「空間向量」特別弱，在後續規劃進度時，就會自動幫你加強這部分的練習題。
- **感知 (接收環境資訊)**：它不只依賴你輸入文字，還能讀取你上傳的講義 PDF、分析你拍下的錯題照片，甚至能連網搜尋最新的新聞時事來補充教材。
- **行動 (使用工具與執行)**：它可以串接不同的數位工具。例如，它能自己打開瀏覽器查證資料、自動幫你整理好的重點畫成心智圖，或是直接幫你將英文單字卡匯入到複習軟體中。

(三) 目前常用的 AI Agent 服務

在我們的日常學習與生活中，其實已經有許多觸手可及的 AI Agent 工具，它們能幫我們做更複雜的自動化任務：

- **Custom GPTs (ChatGPT)**：這是讓使用者能自行客製化的專屬 AI 特助。你可以自由上傳特定課本或資料庫，並規定它的說話語氣與解題邏輯。除了自己動手訂製外，還能在 GPT Store 下載全世界創作者分享的各領域專業代理人。
- **Gems (Google Gemini)**：這是 Google 推出的個人化 AI 代理人功能。它最大的優勢是能與你個人的 Google 雲端硬碟、Gmail 及文件進行深度串接。例如你可以建立一個「專題 Gem」，讓它主動翻閱你收集的雲端資料並自動生成報告大綱。

- **Claude Code (Anthropic)**：這是由 Anthropic 推出的程式開發專屬 AI 代理工具。它最大的特色是能直接深入你的電腦開發環境(如終端機)中運作，而不侷限於網頁聊天視窗。只要給予指令，它就能自主閱讀你的專案資料夾、撰寫程式碼並執行測試。例如在製作資訊課專題時，你可以指派它自動幫你抓出程式碼中的錯誤(Debug)，或是依照需求直接建構出網頁的基本架構。
- **Microsoft Copilot Studio**：這是微軟針對工作流自動化所設計的 Agent 開發平台。它能深入整合 Office 系列軟體，自主幫你處理繁瑣的行政庶務，例如你可以指派它自動分類電子郵件、擷取行事曆空檔，並主動排定小組開會時間等。
- **OpenClaw (小龍蝦)**：在GitHub上爆紅、以小龍蝦為吉祥物的開源自主型 AI 代理人系統。它能常駐在你的電腦後台，並直接串接到Telegram等通訊軟體。只要你在聊天視窗下達大方向指令，它就能在你的實體電腦上自主開瀏覽器查資料、讀寫檔案，甚至自己寫程式來解決問題。

(四) 專屬你的學習 Agent 應用情境

在未來的學習中，你可以開始嘗試將單純的對話框，轉換成不同職責的 AI Agent：

- **【語言交換 Agent】**：不只是幫你翻譯，你可以設定它每天早上主動傳一句英文生活對話給你，並根據你回覆的語法錯誤，自動幫你整理成一本專屬錯題本。
- **【專題研究 Agent】**：當你需要做一份社會科探究報告時，你只要告訴它我要研究臺灣校園減塑政策，它就能自主幫你完成：(1) 搜尋近三年新聞、(2) 整理正反方論點表格、(3) 擬定一份供你參考的訪談問卷草稿。
- **【大考複習 Agent】**：把你的模擬考成績單給它，它能自動為你排定考前30天的專屬讀書計畫，並在每天提醒你該完成的進度。

⚠ 防雷與查證任務：不能直接採用 AI 答案！

請記得，提示詞與任務核心必須由你親自建構，再交由 AI Agent 進行自動化拆解與優化，這樣才能確保你在人機協作中主導學習的主體性。

同時，當 AI Agent 進行多輪對話或輸出複雜成果時，請務必挑出至少一個關鍵論點或數據，親自對照課本、老師資料或權威官方網站進行二次查證，嚴防幻覺與錯誤引導。

(五) AI Agent 的風險

雖然 AI Agent 能夠像虛擬小助手一樣，自動幫我們完成多步驟的複雜任務(例如自動搜集資料、寫草稿並發送郵件)，但這種高度自動化也帶來了不可忽視的風險。在享受便利的同時，我們必須保持警覺：

- **過度依賴與能力退化**：當 AI Agent 幫我們把規劃、執行甚至決策都做完時，我們可能會失去自己思考和解決問題的能力。如果連如何下判斷都交給 AI，大腦的批判性思維就會逐漸退化，讓我們變成只會按確認鍵的操作員。
- **權限濫用與安全漏洞**：AI Agent 通常需要存取我們的個人資料、社群帳號或信箱才能完成任務。如果它遭到駭客攻擊，或者因為系統錯誤而把私密資訊發送給錯誤的人，就會造成嚴重的隱私外洩。把太多鑰匙交給 AI，就等於暴露在更大的資安風險中。
- **責任歸屬模糊**：如果 AI Agent 在自動執行任務時犯了錯——例如在網路上發表了不當言論、侵犯了他人的版權，或者是買錯了昂貴的商品——誰該負責？是開發 AI 的公司？還是啟動 AI 的使用者？當 AI 能夠自己做決定時，犯錯的責任歸屬將變得非常難以釐清。
- **放大偏見與錯誤**：AI 偶爾會產生幻覺。當 AI Agent 把這些錯誤資訊自動串聯到下一步行動時（例如把錯誤的數據直接寫進要繳交的報告並發送出去），錯誤的影響力就會被瞬間放大，甚至造成連鎖反應。

挑戰任務

 **任務說明**：請選擇你最感興趣的一個學科（例如國文、數學、自然、社會或英文），設計一個「AI 如何幫我學得更好」的任務範例，並包含以下四個部分：

項目	說明	你的答案
①角色	你希望 AI 扮演誰？（例如：數學家教、作文老師、辯論教練）	
②任務	希望 AI 幫你解決什麼問題？（例如：講解觀念、改作文、模擬實驗）	
③背景	你目前在學什麼？（例如：二次不等式、氣候變遷、化學平衡）	
④格式	你希望 AI 以什麼方式呈現？（例如：表格、對話、流程圖、摘要）	

挑戰加分題

完成後，把你的 AI 對話成果分享給同學，看看大家的神隊友怎麼幫忙！

學習檢核表

學習目標	我能做到 ☑	還需練習 🔄
1. 我知道 AI 不只是幫我找答案，而是能陪我一起學習、思考、創作的好夥伴。	☐	☐
2. 我會用「角色、任務、背景、格式」四個重點，寫出清楚又有效的 AI 指令。	☐	☐
3. 我能完成一個在課堂上運用 AI 的小任務（像是用 AI 幫我改作文、畫圖、算數學或設計探究報告）。	☐	☐
4. 我知道 AI 也能幫我做更進階的任務，像是練習辯論、設計實驗或規劃學習歷程。	☐	☐
5. 我知道 AI 是工具，真正要動腦、下決定的還是我自己。	☐	☐
6. 我了解 AI Agent 具備自主拆解任務的能力，但我知道定義最終目標、指派工作方向的人還是我自己。	☐	☐

本章核心觀念回顧

1. 在各學科學習中善用 AI 神隊友，了解不同科目中運用 AI 的具體範例。
2. 好的提示詞應包含四個要素：角色、任務、背景和格式。這能讓 AI 更精準地理解我們的需求，並提供更符合期待的回答。
3. 學習進階技巧：你也可以試試請 AI 擔任「提示詞教練」，只要給它一個目標，它就能協助你打造出結構更完整、更有效的提示詞。
4. 防雷與查證任務：不能直接採用 AI 答案！記得提示詞得先自己寫，再讓 AI 協助優化，保持你的學習主體性。同時，請找出 AI 回答中一個需要確認的地方，再用課本、老師提供的資料或可信網站查證。

最重要的一句話

「當你愈會用它，你的思考力、創造力也會一起升級！」



AI 世代下的數位公民

從「數位時代」邁向「AI 世代」

「當 AI 越來越懂你， 你還清楚自己為什麼會這樣想嗎？」

以前談到數位公民，人們主要關心的是如何安全使用網路、遵守網路禮儀、保護個人資料。

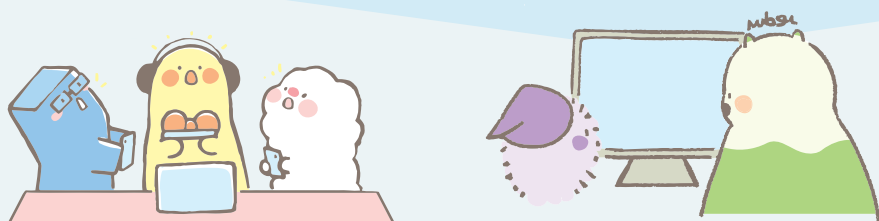
但在 AI 世代，情況根本性地不一樣了——

數位環境不只是「被使用的工具」，而是會主動回應人的系統。

真正重要的，不只是會不會使用 AI，而是能不能理解 AI 如何影響我們的生活、想法與決定。

📌 本章學習目標：

- 理解數位時代與 AI 世代的本質差異，辨識智能系統對生活的影響
- 掌握 AI 世代數位公民的三個新核心能力：理解機制、協商關係、承擔共責
- 深化四大數位公民素養在 AI 世代的新意涵：隱私、媒體判讀、公民參與、數位身分
- 能分析演算法、推薦系統與生成式 AI 對個人思考與社會的影響
- 培養適應能力與批判能力的持續共存，成為能與智能系統共處的公民





✿ 圖片生成提示詞參考：

A teacher-student scene with a friendly AI robot sitting beside a high school student at a desk with open Chinese textbooks. The robot is pointing at a text passage while the student takes notes. Flat illustration, educational, blue and orange tones, no real faces, no text in image.

註記：本圖像素材使用生成式 AI 工具 (Gemini) 輔助生成。

數位時代 vs. AI 世代——什麼改變了？

數位時代的公民，面對的是一個由人建造、由人操作的數位環境；網路是工具，社群媒體是平臺，內容由人產製、決策由人做出。

AI 世代下的數位公民，面對的環境根本不同——環境本身開始有了能動性。演算法決定你看到什麼、生成式 AI 產製你閱讀的內容、推薦系統塑造你的偏好、自動化系統影響你的機會，公民不只是在環境中行動，而是與一個會回應、會學習、會影響的系統共同生活。

數位時代公民面對的環境	AI 世代下數位公民面對的環境
網路是工具，由人建立與操作	環境本身有了能動性，會主動回應與學習
內容由人產製、由人選擇閱讀	生成式 AI 快速產製文章、圖片與影片
資訊由人主動搜尋	推薦系統主動決定你先看到哪些內容
個資保護主要防範惡意攻擊	資料被平臺用來分析、預測與訓練模型
假訊息靠查證來源辨別	深偽讓「眼見為憑」越來越不可靠
數位身分靠個人形象管理	AI 能模仿聲音、生成照片、偽造影片
人是環境中的行動者	人與一個會學習、會影響的系統共同生活

一、當 AI 開始替你決定時，你還知道自己為什麼會這樣想嗎？

你有沒有發現，社群媒體越用越像一個「只推你喜歡的東西」的世界？這不是巧合，而是演算法的設計結果。當推薦越來越精準，你看到的世界也越來越窄。

情境 | 小涵的演算法世界

放學回家的路上，小涵一邊滑手機，一邊看著社群平臺推薦的新影片。她最近明明只是搜尋過一次環保議題，結果接下來幾天，平臺不斷推薦跟環保相關的內容。過沒多久，她開始覺得：「大家對這個議題的看法，是不是都跟我一樣？有人想法不同嗎？」但仔細一看，她才發現，自己看到的內容，其實是演算法幫她挑選過的世界。

(一) 演算法推薦系統：你看到的，是真實的世界嗎？

演算法推薦系統的運作邏輯

平臺的 AI 系統會持續追蹤你的行為：你點了什麼、看了多久、跳過了什麼、按讚了哪些內容。這些資料被用來建立一套「關於你的模型」，預測你下一步最可能想看什麼，再把那些內容推到你面前。

後果就是過濾氣泡 (Filter Bubble) —— 當系統只推薦你「可能喜歡的」，你看到的世界就越來越像「你想看到的世界」，而不是「真實存在的世界」。

(二) 過濾氣泡的三個層次影響

影響層次	具體表現	對公民能力的衝擊
個人認知層次 (你怎麼看世界)	以為「大家都這樣想」，不知道自己的資訊圈其實很窄	難以理解不同立場，批判思考能力下降
社會討論層次 (公共議題的形成)	同溫層效應加劇，社會越來越難達成共識與對話	公共討論品質降低，極化現象加劇
民主參與層次 (誰的聲音被聽見)	演算法決定哪些議題被看見，哪些被埋沒	公民參與受到平臺規則影響，不再完全自主

(三) 隨堂練習：推薦內容觀察日記

📅 進行一天的觀察，記錄你在社群平臺上看到的推薦內容，分析演算法可能的運作邏輯：

這些推薦內容，反映了你最近關注的哪些主題？

有沒有哪些觀點或議題，你幾乎從未在平臺上看到？

如果平臺只推薦你喜歡的內容，你認為這對你的判斷力有什麼影響？

你打算如何主動打破自己的資訊氣泡？

情境小劇場 | 小涵與AI 助理的氣泡對話

小涵：「為什麼我的頁面裡全都是環保的影片？我只是搜尋過一次而已。」

AI 助理：「那是因為推薦演算法學到你對這個主題有興趣，就開始替你過濾內容了。問題是，你現在看到的現在只是它覺得你喜歡的東西，而不是世界上全部的聲音。」

小涵：「所以說，我的世界已經開始被它塑造了？」

AI 助理：「對，想打破這個氣泡，你可以主動搜尋不同觀點、follow不同立場的帳號，或者定期清除瀏覽記錄。最重要的是：意識到氣泡的存在，才能主動選擇。」

小涵：「原來如此，我以為我看的是真實的世界……」

⚠️ 小叮嚀

- 💡 「大家都這樣想」可能只是演算法讓你這樣以為，現實中有更多不同聲音。
- 💡 主動搜尋不同觀點、閱讀多元媒體，是打破資訊氣泡最直接的方法。
- 💡 演算法的目標是讓你「留在平臺上更久」，不一定是讓你「看到更全面的資訊」。

 教師補充

「過濾氣泡 (Filter Bubble)」概念由美國作家 Eli Pariser 在 2011 年的著作《The Filter Bubble》中提出，描述個人化演算法如何讓每個人活在專屬的資訊繭房，研究顯示，長期暴露於單一觀點的資訊環境，會加劇認知偏誤和政治極化。同學可實際測試：在同一臺電腦用不同帳號搜尋同一關鍵字，觀察結果差異。

 本節重點整理

- 演算法推薦系統根據你的行為資料，建立「關於你的模型」並持續篩選內容。
- 資訊氣泡影響三個層次：個人認知、社會討論、民主參與。
- 主動打破氣泡：搜尋不同觀點、接觸多元媒體、定期反思推薦內容。
- 演算法的目標是「延長停留時間」，不是「讓你看到全面資訊」。
- 意識到氣泡的存在，才是批判思考的第一步。

二、當 AI 越來越懂你時，你還知道該怎麼判斷嗎？

AI 世代的數位公民，不再只是「會操作科技的人」，而是能理解系統、與系統保持適當距離，並願意對自己的選擇負責的人。這需要三個超越「會操作工具」的核心能力。

 情境 | 阿哲的報告危機

阿哲發現自己越來越離不開 AI。寫作文時，他請 AI 修改句子；準備報告時，他用 AI 整理資料；就連不知道午餐吃什麼，也會問聊天機器人推薦。有一次在做職涯探索時，他直接問 AI：「請根據我的個性與嗜好，推薦我未來應該從事什麼工作？」看著 AI 給出的答案，他想到自己家裡的經濟狀況與外語能力，漸漸感覺到無力，也不禁思考：這個答案符合真實的我嗎？

那天之後，阿哲開始對 AI 生成的回答，產生了與以前不太一樣的看法。

(一) 第一個新能力：理解機制 (Understanding the System)

💡 什麼是理解機制？

AI 並不像人類一樣真正理解世界，它是透過大量資料進行預測與生成。有時候，AI 可能會說出看似合理、其實錯誤的內容（幻覺，AI Hallucination）；有些判斷也可能受到訓練資料偏差的影響。如果不了解 AI 的運作方式，人們就很容易把 AI 的答案當成「絕對正確」。

不理解機制的風險	理解機制後的應對方式
把 AI 的答案當作絕對正確的事實	把 AI 的答案當作起點，再用可靠來源查證
不知道 AI 可能產生「幻覺」	對具體數字、人名、事件特別謹慎查證
看到流暢的文字就以爲內容正確	文字流暢不代表內容可信，要問「來源是什麼？」
不了解訓練資料偏差對結果的影響	問自己「這個 AI 可能從哪些資料學習？可能有什麼偏見？」

(二) 第二個新能力：協商關係 (Negotiating the Relationship)

與 AI 的互動，更像是一種持續的合作與協商，而非單純的工具操作。什麼時候適合相信 AI？什麼時候需要保留懷疑？當 AI 的建議和自己的判斷不同時，該怎麼選擇？這些問題沒有固定答案，需要不斷思考與調整。

情境	完全依賴 AI 的問題	協商關係的做法
用 AI 寫作業	失去思考過程，被問到答不出來	請 AI 提供架構，自己理解後重新組織表達
用 AI 做決定	決策過程不透明，不知道爲什麼做這個選擇	把 AI 的建議當參考，最終決定由自己負責
用 AI 查資料	錯誤資料被當作事實使用	AI 提供線索，再用可靠來源確認
用 AI 創作	作品缺乏個人觀點，成爲「平均值的輸出」	用 AI 補強技術，保留自己的創意與立場

(三) 第三個新能力：承擔共責 (Sharing Responsibility)

💡 什麼是承擔共責？

只要人選擇使用 AI，就代表參與了這個結果。有些人會覺得：「這不是我做的，是 AI 產生的。」但其實，無論是作業內容、生成圖片，還是推薦決策，都不能因為有 AI 參與，就完全把責任推給它。AI 參與越多，人們反而越需要清楚意識到自己的選擇與影響。

情境	共責的意涵
用 AI 生成假圖片散播消息	不能說「都是 AI 的問題」——是你選擇了散播
把 AI 內容直接當成自己的報告	不能因此免除學術責任
使用有偏見的 AI 系統做決策	使用者也有責任意識到並指出問題
AI 參與越多	人們越需要清楚意識到自己的選擇與影響

📝 三個新能力的自我評估

【理解機制】我上次使用 AI 時，有沒有查證它給的答案？(☐ 有 ☐ 沒有 ☐ 部分查證)
我認為 AI 最容易在哪些地方出錯？

【協商關係】我通常在什麼情況下選擇相信 AI，什麼時候保留懷疑？

【承擔共責】如果我用 AI 完成了一份作業，我認為我對結果負有什麼責任？

⚠️ 小叮嚀

- ✂️ 「AI 說的」不等於「正確的」，流暢的語言和準確的內容是兩回事。
- ✂️ 協商關係的關鍵：AI 提供選項，最終決策由你來決定。
- ✂️ 使用 AI 是一種選擇，選擇了就要為結果承擔責任。

📝 本節重點整理

- 理解機制：了解 AI 的限制，不把它當作萬能神諭。
- 協商關係：與 AI 合作，保持判斷力，最終決定由人來負責。
- 承擔共責：使用 AI 就是參與了結果，不能以「AI 做的」為由推卸責任。
- 幻覺是真實存在的風險，流暢不等於正確。
- 三個新能力合在一起，才是真正的 AI 世代數位公民核心素養。

三、AI 會不會開始改變「你看到的世界」？

過去談數位公民的四大素養，在 AI 世代都有了更深層的意涵，不只是「怎麼用工具」，而是「如何與一個會主動影響你的系統共存」。

情境 | 小芸的三重發現

晚自習結束後，小芸滑著手機放鬆。社群媒體很快推薦了她喜歡的影片、新聞與短文，甚至連最近煩惱的升學問題，也出現了「最適合你的科系分析」，她覺得演算法好像越來越懂自己。

但有一天，她看到一段看似真實的名人影片，後來才知道那其實是 AI 生成的深偽影像；接著，她又發現朋友貼出的文章內容，居然是 AI 自動生成的；更讓她驚訝的是，當她搜尋某些議題後，社群媒體推薦的內容竟越來越單一，好像所有人都只剩下一種聲音。

(一) 網路安全與隱私：從防範攻擊到理解資料如何被使用

以前的隱私保護重點是防止駭客或陌生人偷走資料。現在，即使沒有惡意攻擊，我們的資料也可能被平臺拿來分析、預測與訓練 AI 模型——搜尋紀錄、停留時間、按讚習慣，都可能慢慢建立出一套「關於你」的系統理解。威脅不一定來自壞人，而可能來自那些看似方便、卻不夠透明的系統。

AI 世代需要問的隱私問題	具體行動
我的資料被如何使用？	閱讀平臺的隱私政策，了解資料蒐集範圍
系統正在分析關於我的什麼？	定期檢查帳號的廣告偏好設定，了解平臺的用戶標籤
我有哪些資料控制權？	行使資料刪除權 (Right to be Forgotten)，要求平臺說明
AI 訓練使用了我的資料嗎？	確認使用的 AI 服務是否使用用戶對話訓練模型

(二) 資訊素養與媒體判讀：從查證來源到辨識生成內容

生成式 AI 讓「眼見為憑」越來越不可靠。以前辨別假新聞，多半靠查證來源與比對資料；但現在，生成式 AI 能快速製作真假難辨的內容，甚至連圖片、聲音與影片都可能被偽造。

內容類型	AI 時代的風險	辨識與應對方法
文字內容	流暢的文字不代表內容正確，AI 可能自信地說錯（幻覺）	查找原始來源，確認作者資格，比對多個可靠媒體
圖片	AI 生成的圖片越來越難用肉眼辨識	使用圖片反搜工具，注意邊緣、手指、光影細節
深偽影片	名人或政治人物的假影片可能被用來操控輿論	注意嘴型與聲音的不一致，尋找原始影片來源
聲音複製	有人的聲音被 AI 複製，用來詐騙或散播假訊息	接到不尋常的電話，用約定好的暗號確認身份

練習：深偽辨識實驗（請老師提供一段可疑的影片或圖片）

步驟 1：我的直覺判斷（真 / 假）：_____

理由：_____

步驟 2：技術分析（光影、邊緣、臉部細節）：_____

步驟 3：來源查證（原始發布者是誰？是否有其他平臺轉載？）：_____

最終判斷（真 / 假）：_____

這個練習讓我學到：_____

（三）數位參與與公民行動：從敢不敢發聲到理解平臺如何影響聲音的流動

AI 世代的公民參與，不只是「敢不敢發聲」，現在很多公共討論都發生在社群平臺上，但平臺的 AI 系統會決定哪些內容被看見、哪些被忽略，有些言論可能被放大，有些則可能因演算法而被埋沒。

過去的數位參與	AI 世代的數位參與
學習如何在網路上表達意見	理解演算法如何影響哪些聲音被放大
辨別網路上的假訊息	思考平臺規則是否有利於某些立場
不散布仇恨言論	意識到自己的分享行為如何強化演算法偏好
保護個人帳號安全	思考平臺的商業模式如何影響公共討論品質
尊重不同意見	主動打破資訊氣泡，接觸多元觀點

(四) 數位身分與倫理責任：從個人形象管理到 AI 時代的自我認識

當 AI 可以「像你一樣說話」，「我是誰」這件事開始變得複雜。以前管理數位身分，多半是經營個人形象、尊重他人，但現在 AI 已經能模仿聲音、生成照片、偽造影片，甚至模擬人的說話方式。AI 世代的數位公民，需要建立更穩定的自我意識：哪些是自己的真實想法？哪些只是演算法推送後形成的偏好？

數位身分的 AI 世代新挑戰	應對方式
你的聲音可能被 AI 再製利用，去製作假影片	定期搜尋自己的名字，設置身份確認機制
你的思想可能被演算法「塑造」，非完全自主形成	定期問自己「這個想法是我本來就有的，還是被推薦出來的？」
你分享的內容可能在不知情的狀況下，被用來訓練 AI	了解你使用的平臺是否將用戶內容用於 AI 訓練
AI 生成的內容可能被誤認為是你的原創	在分享 AI 協助生成的內容時，清楚標示 AI 的參與

情境小劇場 | 小芸的深偽挑戰

小芸：「AI 助理，這段影片看起來很真，但感覺有點怪……我要怎麼判斷？」

AI 助理：「先用三步驟來看看：第一，注意嘴型和聲音有沒有不同步？第二，看看臉部邊緣和背景有沒有模糊或不自然？第三，去找找這段影片有沒有在學校的官方帳號發布過。」

小芸：「咦，嘴型好像有點對不上……而且學校官網完全沒有這個消息！」

AI 助理：「很好，這很可能是 AI 合成的影片。記得：眼見不一定為憑，任何重要資訊都要先查核來源。」

小芸：「我差點就轉傳出去了……謝謝你！以後我一定要先停下來想一想。」

⚠️ 小叮嚀

- 💡 隱私威脅不一定來自駭客，平臺的資料分析也可能悄悄影響你的生活。
- 💡 眼見不一定為憑：深偽技術讓圖片、影片、聲音都可能被偽造。
- 💡 你分享的每一篇文章，都在強化演算法對你的判斷，分享前先想一想。

📄 深偽 (Deepfake)

深偽讓人臉合成和模仿聲線的門檻大幅降低，已被用於政治操控、詐騙與名譽傷害。臺灣近年已有多起深偽詐騙案例，同學可使用 InVID、Google 圖片反搜等工具練習辨識；關於平臺演算法的公民影響，可延伸討論歐盟的《數位服務法》(Digital Services Act, DSA) 要求平臺公開演算法機制的規範，思考：如果你能設計平臺的推薦演算法，你會如何平衡商業利益與公共責任？

📝 本節重點整理

- 隱私升級：從防範攻擊到理解「我的資料在系統中是什麼樣子」
- 媒體素養升級：從查證來源到辨識深偽，眼見不一定為憑
- 公民參與升級：理解演算法如何影響哪些聲音被聽見
- 數位身分升級：在 AI 可複製人聲的時代，確認並守護自己的真實想法
- 四大素養都需要從「會使用工具」進化到「理解系統對我的影響」

四、AI 世代數位公民能力的統整

數位時代的公民教育，培養的是一個有能力在數位環境中生存的人，提問的核心是「怎麼用」；而 AI 世代的數位公民教育，需要培養的是一個有能力與智能系統共存，並持續追問「這對人類意味著什麼」的人。

🔧 適應能力 (Adaptive Capacity)

學會使用新工具、理解系統運作、在變化中找到有效的行動方式
缺乏適應能力 → 無法使用新工具，被時代淘汰

🔍 批判能力 (Critical Capacity)

保持懷疑、問「為什麼」、思考「對誰有利、對誰有害」、不放棄自己的判斷
缺乏批判能力 → 被系統操控，失去自主判斷

◆ 唯有兩種能力持續共存，才能面對不斷變化的世界。

能力面向	數位時代 (會……)	AI 世代 (還需要……)	關鍵問題
工具使用	安全有效地使用數位工具	理解工具背後的系統邏輯與限制	為什麼這個工具這樣設計？對我有什麼影響？
資訊判讀	查證來源、辨別真假新聞	辨識 AI 生成內容、懷疑「眼見為憑」	這個內容是誰（或什麼）產製的？目的是什麼？
隱私保護	保護帳號、不隨意公開個資	理解資料如何被系統分析與使用	我的資料在系統中是什麼樣子？誰在使用它？
公民參與	在網路上負責任地表達意見	理解平臺如何影響聲音的流動與公共討論	誰有權決定哪些聲音被聽見？
自我認識	管理個人數位形象	確認自己的真實想法，避免被演算法塑造	哪些想法是我的？哪些是演算法幫我形成的？
責任承擔	對自己的言行負責	理解使用 AI 也是一種選擇，需要承擔共責	我使用了 AI，我要為這個結果承擔什麼？

五、本章終極挑戰：AI 世代公民的思辨任務

📝 任務一

追蹤你的資訊氣泡至少連續三天，記錄你在同一個社群平臺上看到的內容類型。三天後，分析：
哪些主題被重複推薦？

哪些觀點或族群的聲音，你幾乎沒有看到？

如果要主動打破這個氣泡，你會怎麼做？

📝 任務二：AI 共責分析

選擇以下一個情境，分析「承擔共責」的意義：

- A. 有人用 AI 生成一位老師的假影片，說老師做了某件事
- B. 一篇由 AI 生成的假新聞在班上群組廣泛轉傳
- C. 有人使用 AI 幫忙做了整份期末專題，沒有說明

我選擇的情境：_____

在這個情境中，哪些人有責任？ 各自的責任是什麼？

如果你是旁觀者，你應該做什麼？

🚀 挑戰任務

難度	任務描述
🌱 初階 觀察記錄	連續三天記錄一個社群平臺的推薦內容，列出你看到的主題類型，分析哪些聲音被放大、哪些幾乎消失。
🔍 進階 分析比較	針對同一個新聞事件，分別在三個不同來源（主流媒體、獨立媒體、社群平臺）閱讀相關報導，比較敘述角度的差異，分析哪個版本的視角更完整。
🌳 高階 創作應用	設計一份「AI 世代數位公民手冊」（至少一頁），說明你認為高中生最需要掌握的三個 AI 時代公民素養，並舉出生活中的真實例子。

📊 學習檢核表

學習目標	我能做到 ☑	還需練習 🔄
1. 我能說明數位時代與 AI 世代的本質差異。	☐	☐
2. 我了解演算法推薦系統的運作邏輯以及資訊氣泡的三層影響。	☐	☐
3. 我能說明「理解機制」、「協商關係」、「承擔共責」三個新能力。	☐	☐
4. 我了解幻覺是什麼，以及面對 AI 內容應如何查證。	☐	☐
5. 我知道 AI 世代的隱私威脅不只來自惡意攻擊，也來自系統的資料分析。	☐	☐
6. 我能了解辨識深偽內容的基本方法，並知道「眼見為憑」已不再可靠。	☐	☐
7. 我了解平臺演算法如何影響公民參與及公共討論。	☐	☐
8. 我能分析使用 AI 時的共責問題，不會以「那是 AI 做的」為由推卸責任。	☐	☐
9. 我能描述「適應能力與批判能力共存」對 AI 時代公民的意義。	☐	☐

📖 本章核心觀念回顧

- AI 世代的數位環境本質改變：從「工具」變成有能動性、會學習的系統。
- 演算法推薦系統形成資訊氣泡，影響個人認知、社會討論與民主參與。
- 三個新核心能力：理解機制（懂 AI 怎麼運作）、協商關係（與 AI 保持適當距離）、承擔共責（使用 AI 就要負責）。
- 隱私保護升級：從防範攻擊到理解資料如何被系統分析與使用。
- 媒體素養升級：從查證來源到辨識 AI 生成內容，深偽讓「眼見為憑」不再可靠。
- 公民參與升級：從敢發聲到理解平臺演算法如何影響聲音的流動。
- 數位身分升級：在 AI 可複製人聲的時代，確認並守護自己的真實思想。
- 目標：適應能力與批判能力的持續共存，才能面對不斷變化的世界。

💡 最重要的一句話

真正重要的，不只是會不會使用 AI，
而是能不能理解 AI 如何影響我們的生活、想法與決定。



■ 附錄一：AI 回應風險檢核表

 AI 檢核偵探：文字／圖片／影片幻覺大體檢	
任務說明： AI 有時候會「一本正經地胡說八道」或製造出不符合現實的影音。請帶著懷疑的眼光，把 AI 生成的內容(文字、圖片或影片)當作「嫌疑犯的供詞」，逐一對照以下項目進行大體檢！	
SECTION 1：任務基本資料 Q1. 你的名字／組別：_____ Q2. 本次要檢核的 AI 內容類型(可多選) ☐ 文字(文章、報告、對話回答) ☐ 圖片(AI 繪圖、合成照片) ☐ 影片(AI 生成影片、深偽 Deepfake 影片)	
SECTION 2：【文字類】AI 內容檢核表	
維度	檢核項目(符合者打勾)
維度一：事實與證據(適用：歷史、地理、公民、日常新聞)	☐ 【幽靈文獻】 課本、圖書館、Google 權威網站完全查不到 AI 說的這個人、這本書或這條法律。 ☐ 【數字魔術】 出現聽起來很專業的精準數據，但問 AI 數據哪來的，它交不出真實報告。 ☐ 【時空錯亂】 歷史事件的時間點對不上，例如：清朝人物搭捷運，或把 20 世紀的事移到 19 世紀。
維度二：邏輯與常識(適用：數學、理化、生物、語文寫作)	☐ 【前後打臉】 文章前半段和後半段的觀點或設定互相矛盾。 ☐ 【物理大崩壞】 不符合科學常理，文字上出現不合邏輯的因果關係；圖影中出現東西往天上飄、水往高處流。 ☐ 【瞎編裝懂】 故意拿不存在的假概念問它，它沒有拒絕，反而一本正經地胡說八道。
維度三：多媒體、影音與社群(適用：視覺及表演藝術創作、生活科技、FB/IG/YT)	☐ 【複製人與異形手】 圖片或影片中，人的手指數量不對(多或少)、關節扭曲，或背景路人的臉全部長得一樣。 ☐ 【神祕精靈文】 畫面中招牌、課本封面、衣服上的文字，放大看全部變成扭曲怪異的符號。 ☐ 【物件融合】 畫面中的東西不自然地黏在一起，例如：眼鏡直接嵌進皮膚裡，或吸管穿透過杯子。
維度四：媒體與心理直覺(適用：日常媒體識讀、防詐騙)	☐ 【完美迎合】 這則訊息「太完美了」，剛好完全符合你的喜好，或內容極度誇張、聳動，讓人一眼就想轉傳。 ☐ 【無人認領】 找不到這份內容的「真人作者」或「出處來源」，看起來純粹是由 AI 自動生成的內容。
SECTION 3：【圖片類】AI 內容檢核表	
檢核項目(符合者打勾)	
☐ 【手指與關節】 手指數量不對(6 隻手指)、手指頭融合在一起，或手肘關節彎曲角度詭異。 ☐ 【五官與配件】 兩邊耳環不對稱、眼鏡框邊緣模糊、瞳孔形狀不是圓形，或雙眼視線方向不一致。 ☐ 【毛髮與皮膚】 頭髮像是一整塊片狀，或有些頭髮莫名其妙融入背景，皮膚質感滑膩、沒有毛孔。 ☐ 【光影與重力】 影子的方向和光源對不起來，或是物體飄在空中沒有合理的影子。 ☐ 【背景與文字】 背景的人臉都模糊變形、背景建築物的線條扭曲、圖片中的招牌文字變成看不懂的「外星文字」。 ☐ 【物件融合】 衣服的領口線條不見、背包帶子直接穿過身體，或杯子和手融為一體。	
SECTION 4：【影片類】AI 內容檢核表(深偽 Deepfake 檢測)	
檢核項目(符合者打勾)	
☐ 【對嘴不自然】 說話時，嘴型跟聲音對不上、有延遲感，或嘴唇閉合時聲音還在持續。 ☐ 【眼神與眨眼】 影片中的人幾乎不眨眼，或眨眼次數多到很不自然，眼神看起來呆滯、沒有焦點。 ☐ 【邊緣閃爍】 當人物轉頭或手在臉前揮過時，臉部邊緣、下巴或脖子會出現奇怪的邊框或殘影。 ☐ 【時空扭曲】 影片背景的物品會突然改變形狀、不正常地閃爍，或走過去的人突然消失。 ☐ 【動作卡頓】 人物的動作不連貫，前一秒手在下面，下一秒突然瞬間移動到上面。	
SECTION 5：🕵️ 偵探結案報告	
統計總分： 請數一數，你在上面所有檢核表中，總共勾選了幾個「 ☐ 」？ ☐ 0 個 →  綠燈：安全通過！這個 AI 內容高度可信，加上自我觀點，可安心使用。 ☐ 1 - 2 個 →  黃燈：半信半疑！AI 開始出現幻覺，資訊必須經過修改與查證才能用。 ☐ 3 個(含)以上 →  紅燈：嚴重幻覺！這個 AI 內容充滿漏洞，不可採用或需查核後，大範圍修正。	
你的結案心得(請用一句話總結，你發現了這個 AI 的什麼秘密或謊言?) 	

第一部分：媒體識讀評估分析

評估向度	評估項目	檢視或實踐指標	評估勾選			評估備註或 反思紀錄
			存 疑	不確定	信 任	
一、來源 審視：檢 核訊息從 哪來？	訊息來源可信度？	1.傳訊者或網站的可信度？ 2.傳統媒體或網路社群媒體成立時間和組織成員變革的歷史背景紀錄？(是否為短期成立？) 3.傳統媒體或網路社群媒體的公信力？				
	訊息作者是誰？	1.有無標明作者或單位？ 2.是否匿名或使用假名、捏造人物而無法查核？ 3.是否為知名專業人士被冒名盜用？				
	作者或其任職單位聯絡方式？	1.能否與作者真實聯繫？ 2.網站有沒登載關於我們或聯絡方式？				
二、內容 審視：檢 核訊息內 容有沒問 題？	標題與內文？	1.標題與內文有無矛盾？ 2.有沒聳動或煽動的文字意涵？ 3.內文用語是否為我國慣用語、有沒過多錯別字或不通順語句？				
	日期和時效性？	1.有沒具體日期？ 2.時效性正確嗎？ 3.是否為舊聞假冒近期事件？				
	引述證據及出處的可靠性？	無法追蹤引述的真實性，如： 1.引述網址連結可否查證？ 2.是否出現以下字眼，如：「據專家指出」、「內部人士透露」、「某人說」？				
	圖片或影片的真實性？	1.影像與訊息內容是否有一致性？ 2.影像是否為 AI 生成？(細節有無符合自然現象？)				
三、情 緒、思維 與同溫層 效應審視 ：檢核訊 息引起我 什麼感覺 和思緒？	情緒操弄言語？	情緒字眼： 1.企圖引起恐慌、激起憤怒、製造對立或仇恨？ 2.快轉發(傳)？				
	公正性、中立性及客觀性？	1.客觀與平衡報導？ 2.公正和多元觀點？ 3.中立與包容共榮？				
	慣性思維和同溫層效應檢視？	1.能否覺察有沒慣性思維反應？ 2.是否易讓自己輕信而誤判？ 3.是否偏向自己想聽的話或投己所好？				
四、審視判斷結果：		統計存疑、不確定與信任各項目勾選數量。 1.計算 (存疑數量 + 不確定數量)。 2.(存疑數量 + 不確定數量)與信任數量比大小。				
		當(存疑數量 + 不確定數量) ≥ 信任數量 => 事實查核進階行動	☐ 信任訊息。 ☐ 立即進行查核行動。			

第二部分：針對疑點進行事實查核進階行動(運用 SIFT 法)

步 驟	疑點與查核紀錄
S — Stop (停)	
I — Investigate the source (調查來源)	
F — Find better coverage (尋找佐證)	
T — Trace to original context (追溯原始脈絡)	

事實查核後結果判斷：該訊息是 ☐真實 ☐錯誤 ☐部分誤導

第三部分：反思—若為假訊息，他有何目的？可如何阻斷假訊息散佈？

■ 附錄三：數字編碼假訊息查核表

依照下表依序進行檢視，可協助你判斷訊息的真實性，最終得到紅、黃、綠三色警報判定。

查核步驟		查核指標	下一步查核行動、關鍵叮嚀或查核結果說明
步驟 1	1a	檢視訊息能讓我情緒波動 (如：生氣、害怕、激動等)，或標題聳動誇張、內容有明顯錯別字或非生活慣用語、排版很亂。	註記：高風險。 繼續前往步驟 2。
	1b	檢視訊息的語氣和內容看起來還算正常。	繼續前往步驟 2。
步驟 2	2a	檢視發布訊息的單位和作者，確認是我認識且可信賴的單位和作者。(例如：公視、中央社、政府或學校官網)。	來源可信。 請直接跳到步驟 4。
	2b	檢視發布訊息的單位和作者。確認是匿名的、我不認識的，或名字很奇怪的網站或粉絲專頁。	來源不明，需查核。 請繼續前往步驟 3。
步驟 3	3a	(啟動查核) 另外開啟瀏覽器搜尋該單位和作者姓名，發現單位或作者評價良好，具有專業背景。	來源可信度提高。 請繼續前往步驟 4。
	3b	(啟動查核) 另外開啟瀏覽器搜尋後發現，它是著名的內容農場，常發表偏激言論，或有造假紀錄，或根本查不到任何資訊，或粉絲專頁常更改名稱。	最終判斷：紅色警報 (高度可疑) 原因：來源不明或聲譽不佳，不應相信或分享。
步驟 4	4a	訊息內容能提出具體證據 (如：具公信力媒體或經權威期刊或報章雜誌等轉載，能提供報告連結、數據資料、原始照片或影片等)。	證據需驗證檢核。 請繼續前往步驟 5。
	4b	訊息內容只用「專家說」、「內部消息」、「聽說」等模糊字眼，無法提供任何證據；或專家是張冠李戴或虛構人物。	證據薄弱。 請跳到步驟 6 做最後確認。
步驟 5	5a	(檢核證據) 我點開連結或反向搜圖，發現證據真實且支持訊息的說法。	證據力強。 請繼續前往步驟 6。
	5b	(檢核證據) 我點開連結或反向搜圖，發現證據是被扭曲、誇大的斷章取義，甚至是把舊的圖文套用在新事件上，進行移花接木等捏造訊息。	最終判斷：黃色警報 (部分真實，需再查核求證) 原因：訊息可能混雜了真假內容，有誤導嫌疑。
步驟 6	6a	(多方比對) 我用訊息的關鍵字進行搜尋，發現多家獨立且可靠的新聞或官方機構或查核平臺，同時報導或刊登相同內容，該內容正確可信。	最終判斷：綠色警報 (可信度高) 原因：經過多方查證，訊息內容基本屬實。
	6b	(多方比對) 我用訊息的關鍵字搜尋後，找不到任何其他可靠的單位報導，或只找到澄清或闢謠的文章，或查核平臺已刊登內容不可信。	最終判斷：紅色警報 (高度可疑) 原因：缺乏佐證，很可能是虛構的，不應相信且勿分享。

■ 附錄四：查證相關資源、查核平臺與防詐網站介紹

(一) 查證相關資源介紹：

- 善用網路搜索進行反向追蹤，如 Google 以圖搜圖的功能，來辨識圖片和搜索可能的相關資訊。除了 Google，還可使用 TinEye。
網址：<https://tineye.com/>
- 關於以圖搜圖的技巧學習，可看 YouTube 說明影片：TFC 臺灣事實查核中心「查核工具箱 - 搜圖達人」。
網址：<https://www.youtube.com/watch?v=JPuqK3b7XaU>
- Google 事實查核工具：提供最近已被事實查證的英文與中文新聞或資料，特別注意中文包含繁體中文和簡體中文內容。另外還可輸入關鍵字與執行以圖搜圖的方式進行查詢。
網址：<https://toolbox.google.com/factcheck/explorer/search/list:recent;hl=en?authuser=0>

💡 小叮嚀：沒查到，不代表完全沒問題，只是還沒被查證完成。請保持懷疑，持續關注後續查核結果。

(二) 查核平臺介紹：

善用下列平臺進行事實資訊查核，也能得到真假訊息的相關資料，不妨多加運用。

- 臺灣事實查核中心 網址：<https://tfc-taiwan.org.tw/>
- MyGoPen 網址：<https://www.mygopen.com/>
- LINE 訊息查證 網址：<https://fact-checker.line.me/>
- Cofacts 真的假的 網址：<https://cofacts.tw/>
- 疾管家：疾管署官方 LINE，可針對所有傳染病及防疫相關問題進行提問，亦可查詢假訊息。
網址：<https://page.line.me/vqv2007o?openQrModal=true#~>
- 食藥署：提供實用食藥資訊和最新正確訊息。如：接收食藥署週報內容，包含圖文並茂的食藥醫粧資訊。同時還有資訊查核功能，如：「食藥闢謠機器人」查詢網路訊息真偽，有關食品、藥品、醫療器材及化粧品安全等傳言或迷思，都可查核，並逐一破解。
網址：<https://page.line.me/ero7519h?openQrModal=true>
- 數發部「網路詐騙通報查詢網」：看到可疑投資或名人廣告？先查再說！貼上網址，AI 立即查詢是否已被通報為詐騙。若發現新詐騙，立即回報守護大眾。受冒名者可通報下架；AI 也會自動偵測偽冒廣告並通知平臺。
網址：<https://fraudbuster.digiat.org.tw/accessibility/index>

(三) 數位防詐雙核心：165 全民防騙網 X 165 打詐儀錶板

- 165 全民防騙網：提供民眾網路報案、檢舉詐騙廣告，並定期公告「涉詐高危險業者」名單，方便大眾即時識別風險。
網址：<https://165.npa.gov.tw/#/>
- 165 打詐儀錶板：公開最新的詐騙數據、各種最新受騙案例與案例分析。下載專區提供最新的打詐圖卡，可轉給你的親友，預防受騙上當。
網址：<https://165dashboard.tw/>

■ 附錄五：高中生智慧駕馭 AI —— 互動式教材



Lorem ipsum

■ 附錄六：AI 數位學習相關資源

名稱	QR code	說明
《人工智慧基本法》		115年1月14日發布。立法目的是為建設智慧國家，促進以人爲本之人工智慧研發與人工智慧產業發展，建構人工智慧安全應用環境，落實數位平權，保障人民基本權利，增進社會福祉，提升國人生活品質，促進社會國家之永續發展，維護國家文化價值及提升國際競爭力，並確保技術應用符合社會倫理。
臺灣教師與學生 AI 素養框架		本框架提供中小學 AI 教育與學習的參考方向，培養師生面對人工智慧所需的知識、技能與態度。內容從認識 AI、適切運用 AI 到思考 AI 帶來的影響，引導學生理解 AI 的基本概念、可能與限制，並重視批判思考、問題解決、倫理責任及與 AI 協作的能力，幫助學生在學習與生活中負責任地運用 AI。
115年中小學校長數位與 AI 學習領導指引		115年校長數位與 AI 學習領導指引旨在協助校長應對數位與 AI 科技蓬勃發展，提供校長進行校園數位化轉型、推動數位學習與引導 AI 應用於教與學之策略，內容包含校長自身之校長 AI 素養、推動校園數位行政轉型、引導校園內應用數位與 AI 互動方式學習、了解與導入 AI Agent 用於教學及輔導與社區溝通、訊息辨識應對等領導層面，期望協助全國校長將數位與 AI 科技納入校園促進個人化學習，並建構安全且永續的環境。
115年中小學教師數位與 AI 教學指引		《115年中小學教師數位與 AI 教學指引》是教師迎向 AI 時代的實用教學指南，提供數位與 AI 素養架構、課程設計原則、教學應用案例、生成式 AI 與 AI 代理人使用方法，以及風險防護與自我檢核工具，協助教師從備課、教學到評量安心運用 AI，提升教學效能，培養學生自主學習、批判思考與負責任使用 AI 的能力。
115年中小學家長數位與 AI 學習知能指引		本書旨在協助家長因應數位科技與人工智慧快速發展趨勢，提升親子陪伴策略與數位學習知能。內容從數位學習政策、AI 素養出發，引導家長正確認識與善用生成式 AI，介紹適合孩子的數位學習資源，並涵蓋網路安全、訊息辨識、網路霸凌、詐騙防範、短影音使用及心理健康等重要議題。期望促進家長數位與 AI 陪伴能力，培養孩子自主學習及健康使用科技習慣，共同打造安全、健康、優質的家庭數位與 AI 學習環境。

名稱	QR code	說明
中小學 AI 之學習 應用手冊——國小 生聰明用 AI		本書專為國小三至六年級學生設計，從生活中的人工智慧出發，透過故事、圖像、互動任務與實作活動，引導學生認識 AI，學習如何與 AI 溝通，並逐步接觸可應用於問答、學習、探索與創作的教育型 AI 工具。內容同時強調資訊查證、安全使用、隱私保護與數位公民責任，培養學生的好奇心、判斷力、創造力及獨立思考能力，讓孩子建立正確、安心且負責任的 AI 使用觀念與習慣。
中小學 AI 之學習 應用手冊——國中 生靈活用 AI		本書旨在協助國中生建立實用且負責任的 AI 素養能力，內容從認識 AI 開始，引導學生掌握有效提問、提示設計、人機協作及自主學習方法，善用 AI 作為學習夥伴。書中並涵蓋資訊查證、幻覺、媒體識讀、深偽內容、隱私安全與 AI 倫理等議題，透過任務與生活情境練習，培養學生辨識錯假訊息、反思科技影響及靈活運用 AI 解決問題的能力，為未來學習、生活與生涯探索做好準備。
中小學使用生成式 人工智慧注意事項 2.1(教師、行政人 員及家長版)		教育部「AI 人才方舟計畫」入口網提供的官方下載資源，適合中小學教師、行政人員與家長參考。
中小學使用生成式 人工智慧注意事項 2.1(學生版)		教育部「AI 人才方舟計畫」入口網提供的官方下載資源，適合中小學學生參考。

教育部

中小學AI之學習應用手冊

總召集人

葉家宏

教育部資訊及科技教育司司長

副總召集人

郭伯臣

AI人才方舟計畫推動營運中心主持人
國立臺中教育大學校長

計畫主持人

王立仁

國立臺中教育大學助理教授

協同主持人

莊宗嚴

國立臺南大學教授

總編輯

王立仁

國立臺中教育大學助理教授

美術編輯

徐彥哲

桃園市立經國國民中學國文專任教師

研發團隊 (依姓氏筆畫數排序)

李秉芳	教師	高雄市立旗山國民中學
李美惠	教師	臺北市立仁愛國民中學
吳穎沕	教授	國立中央大學網路學習科技研究所
林建毅	教師	臺中市烏日區東園國民小學
林穎俊	教師	宜蘭縣宜蘭市中山國民小學
林靜怡	科長	教育部資訊及科技教育司
施春輝	教師	新北市新莊區丹鳳國民小學
洪詠善	研究員	國家教育研究院課程及教學研究中心
陳佑華	教師	臺北市立建國高級中學
陳政川	教師	臺北市立陽明高級中學
楊宗榮	校長	臺中市豐原區豐原國民小學
楊琬琳	副教授	國立成功大學師資培育中心
詹明峰	副教授	國立中央大學學習與教學研究所
魏秀玲	教師	臺中市南區樹義國民小學
羅玗貞	教師	臺北市立內湖高級中學

助理團隊

何意茹 鄭栢堯