

提示詞素養的概念建構、層次化框架與 制度化培育意涵：回顧式概念分析取向

張明文¹、陳國生²、陳建志³

¹ 新北市政府教育局局長

² 新北市政府教育局局本部輔導員

³ 國立臺北教育大學教育經營與管理學系副教授

摘要

隨著生成式人工智慧與大型語言模型快速進入教育場域，自然語言互動已成為影響學習歷程與知識建構的重要介面。作為人機互動的核心媒介，提示詞不僅影響生成內容的方向與品質，也關係到生成式 AI 在教育中能否被負責任地使用。然而，現有研究對提示詞素養的概念界定、能力內涵與發展層次仍呈現高度分散，致使相關成果不易累積，也限制其在課程、評量與教師專業發展上的應用。本研究採回顧式概念分析取向，整合數位素養、AI 素養與人機互動研究，以釐清提示詞素養的理論定位、核心能力構面與層次化理解。研究結果顯示，提示詞素養宜定位為建立在數位素養與 AI 素養之上的人機協作導向高階素養，其核心能力可操作化為提示設計、互動調節與批判評估三個相互連動的構面。進一步而言，本研究提出三層次提示詞素養框架，用以說明基礎素養、互動操作與高階反思之間的動態關係，並據此分析其在課程設計、教師專業發展、評量機制與治理支持上的制度化培育意涵。本研究提供一個具理論整合性的概念框架，可作為後續研究與教育實務之分析基礎。

關鍵詞：人工智慧素養；人機協作；大型語言模型；生成式人工智慧；提示詞素養

壹、緒論

一、研究背景與目的

近年來，生成式人工智慧（generative artificial intelligence，以下簡稱生成式 AI）快速進入教育場域，特別是以大型語言模型（large language models [LLMs]）為基礎的對話式系統（Brown et al., 2020），使學習者得以透過自然語言進行寫作輔助、概念釐清、資料整理與文本修訂。相較於過往以搜尋、推薦或自動評分為主的教育科技應用，生成式 AI 強調即時互動與內容生成，顯著改變了學習者與科技之間的關係，也對既有的教學設計、學習歷程與評量方式提出新的挑戰（Kasneci et al., 2023）。在此脈絡下，教育現場的關切已不再限於「是否允許使用」或「如何防範濫用」，而逐步轉向如何將生成式 AI 納入課程、教學與評量的能力架構，並同步回應倫理、透明與制度治理等要求（UNESCO, 2023, 2024）。

隨著生成式 AI 的普及，教育研究逐漸注意到「提示詞（prompt）」在影響生成結果品質與方向上的關鍵角色。不同於傳統指令式輸入，提示詞往往包含對任務目標、情境脈絡與輸出規格的語言化描述，並可在多輪互動中持續被修正與調整。此一特性使提示行為不再只是技術操作，而逐漸被視為一種影響學習歷程與認知參與的重要中介行動。人機互動研究亦指出，非 AI 專業使用者在與語言模型互動時，成效差異往往源自提示設計與互動策略的不同，而非單純的語言能力或工具熟悉度（Zamfirescu-Pereira et al., 2023）。因此，若欲理解生成式 AI 介入學習後的差異來源與可培育能力，提示行為及其教學與評量意涵，構成不可忽略的研究焦點。

在此背景下，近年文獻開始以「提示詞素養（prompt literacy）」或相關概念，嘗試描述學習者在生成式 AI 情境中提問、引導、修正與評估生成內容的能力。部分研究將其視為提示設計或提示工程（prompt engineering）的教育應用，強調如何撰寫結構清楚、資訊完整的提示（Chu & Chiu, 2025）；另有研究則將提示詞素養納入 AI 素養或數位素養的延伸，關注學

習者是否理解系統能力、限制與風險，並能負責任地使用生成式 AI (Ng et al., 2021; Qian, 2025)。

此外，亦有研究從學習歷程與互動觀點出發，指出提示詞素養涉及多輪對話中的策略性調節與批判評估，而非一次性輸入行為 (Tour & Zadorozhnyy, 2025)。然而，現有提示詞素養相關研究在概念界定、理論定位與操作化方式上仍呈現高度分散：不同研究分別聚焦於提示結構、互動行為、學習成效或使用態度，使研究成果難以對話與累積；甚至在實證研究中，提示詞素養有時被簡化為提示長度或具體程度，有時則被視為學習者對 AI 的整體態度，顯示其作為教育研究概念仍缺乏穩定的分析框架 (Plaatjies & van Wyk, 2025)。

值得注意的是，這種「概念與衡量的分散」正與國際教育治理趨勢形成張力。國際上，生成式 AI 的教育治理與能力框架已由工具使用議題，逐步轉向以人為本的能力建構、倫理規範與制度治理。例如，聯合國教育、科學及文化組織 (United Nations Educational, Scientific and Cultural Organization [UNESCO]) 針對生成式 AI 之教育與研究指引強調，國家層級需同步推進短期規範、長期政策與人力增能，並將透明、責任與人本價值納入教育系統治理 (UNESCO, 2023, 2024)。UNESCO 亦提出教師 AI 能力框架，將教師所需能力整合為人本思維、倫理、AI 基礎、AI 教學法與專業學習等構面，顯示教師能力架構已成為各國推進生成式 AI 教育應用的重要議題 (UNESCO, 2024)。

在學習者端，經濟合作暨發展組織 (Organisation for Economic Co-operation and Development [OECD], 2025) 與歐盟亦推進面向中小學的 AI 素養架構，嘗試以可操作的能力描述與課堂情境示例，支撐跨學科導入與學校層級政策 (OECD, 2025)。在法規治理層面，歐盟 AI Act 以風險分級思維規範 AI 系統，並對教育相關高風險用途（如可能影響入學、評量或教育機會之系統）要求更嚴格的治理與文件化，凸顯教育場域的 AI 使用已被視為涉及公平與基本權利的制度議題，而非單一教學創新 (European

Union [EU], 2024)。

基於上述脈絡，本研究聚焦於生成式 AI 教育情境下之提示詞素養，透過整合既有研究，釐清其概念界定、能力構面與制度意涵，並將公平與治理納入分析。當生成式 AI 的教育應用由工具使用議題轉向能力建構與制度治理議題時，提示行為及其素養化不宜僅被理解為技術技巧或個人偏好，而應被視為可納入課程、評量與治理設計的學習能力。藉由文獻整合與概念定位，本文嘗試為提示詞素養研究提供較穩定的分析架構，並作為教育實務與政策討論中課程、評量與治理設計之理論基礎。

二、研究取向與分析方式

本研究採回顧式概念分析取向，定位為理論性論文之文獻統整與概念建構工作，非樣本蒐集與統計推論為主的實證研究。分析範圍以近年生成式 AI、提示詞素養、數位素養、AI 素養與人機互動相關文獻為主，並輔以國際教育治理文件作為制度脈絡參照，以釐清提示詞素養在教育場域中的概念邊界與理論位置。

在文獻納入與統整原則上，本研究優先選取能直接回應提示詞素養概念界定、能力構面、層次化理解與教育應用意涵之研究，並對不同研究之核心主張、分析單位與操作方式進行交互比對。分析過程著重於辨識各研究之間可相互連結的概念節點，據以建構三構面與三層次之整合框架。

為提升分析信實性，本研究採取概念一致性、脈絡適切性與論述連貫性三項原則：其一，檢視不同研究對提示、互動與評估之用語是否具有可比較性；其二，將教育場域中的課程、評量與治理需求納入概念判準；其三，於文末明確標示本研究屬概念整合與研究議程建構性質，其框架仍有待後續實證研究持續檢驗與修正。

貳、相關議題論述

為回應提示詞素養概念分散、方法交代不足與應用意涵連結薄弱等問題，本部分依序從理論脈絡與概念定位、核心能力構面、實證研究脈絡、層次化框架，以及制度化培育意涵等面向進行整合，以建立後續結論所依循的完整論述鏈結。

一、提示詞素養的理論脈絡與概念定位

在生成式人工智慧快速進入教育場域之前，教育研究已長期透過「素養」概念理解學習者與科技、資訊及知識之間的關係。素養並非僅指工具操作能力，而是鑲嵌於社會文化脈絡中的讀寫與意義建構實踐；此一取向強調「如何在特定情境中運用符號資源、做出判斷並承擔後果」，而不只是「能不能使用工具」（Street, 1984）。進一步而言，多元溝通管道與語言文化多樣性，使素養視野擴展至多模態與設計導向的意義生成，凸顯學習者在文本、生產與詮釋之間的主動性（New London Group, 1996）。因此，若忽略素養理論的演進脈絡，提示詞素養容易被簡化為短期「技巧訓練」，難以被定位為具教育意義、可培育與可評估的核心能力。

（一）數位素養的擴展：從操作能力到批判取向

早期數位素養多以「使用數位工具」與「資訊操作技能」作為核心，評估焦點偏向學習者是否能有效搜尋、處理與呈現資訊。然而，網路平台化與內容生產門檻降低後，研究逐漸指出：僅具備操作能力不足以支撐有意義的學習與公民參與，素養更關乎對資訊品質、來源可信度、立場與影響的辨識與判斷。以新素養觀點來看，數位環境中的讀寫實踐同時涉及社群互動、規範理解與身分建構，學習者需要在參與、詮釋與產出之間展現能動性（agency）（Lankshear & Knobel, 2011）。這一階段的理論轉向，為後續 AI 素養與提示詞素養的發展奠定了關鍵立場：科技能力必須與批判思考、價值判斷與責任倫理相結合，而非停留在工具使用層次。

(二) AI 素養的興起：理解系統、限制與責任

隨著人工智慧在教育與社會中的應用擴張，研究者開始提出 AI 素養以回應新科技所帶來的理解落差與風險治理需求。相較於數位素養主要處理資訊與媒介環境，AI 素養更強調使用者對 AI 系統基本概念、運作邏輯、限制與倫理風險的理解，以及在應用情境中做出適切判斷的能力 (Ng et al., 2021; Long & Magerko, 2020)。在教育研究中的綜述觀點裡，AI 素養常被概念化為「理解、運用、評估、倫理」等面向的整合能力；在人機互動 (Human-Computer Interaction [HCI]) 取向的討論中，AI 素養也被界定為能「辨識並善用 AI」，同時以批判眼光評估其適切性與影響的能力。

對本研究而言，AI 素養之所以構成提示詞素養的重要前提，在於提示互動並非憑空發生：使用者若缺乏對模型能力邊界、常見錯誤型態與不確定性的基本理解，便難以提出有效提示，更難對生成結果進行後續查核與修正 (Qian, 2025)。因此，提示詞素養雖具有互動歷程上的獨特性，仍須建立在 AI 素養所提供的系統理解與責任判斷基礎之上。

(三) 生成式 AI 的出現與既有素養理論的張力

生成式 AI (尤其是大型語言模型) 以自然語言互動、即時生成內容的特性，使既有素養框架面臨新的張力。一方面，生成式 AI 降低了表達與產出的門檻，讓學習者得以以「對話式方式」進行資料彙整、構思與文本修訂；另一方面，生成內容的可塑性與高度擬真，模糊了「搜尋／使用」與「創作／生成」的界線，也提高了錯誤擴散、來源不透明與過度信任等風險。相關評論指出，教育應用不只涉及工具導入，更牽涉學習者與教師如何理解模型的侷限、偏誤與責任歸屬，以避免把輸出誤認為權威知識 (Kasneci et al., 2023)。在此情境下，若僅以 AI 素養回答「是否理解系統」，仍不足以解釋使用者如何在互動歷程中主動引導、反覆修正與評估生成結果；這也促使研究焦點進一步轉向「prompting 作為可訓練的互動實踐」，並為提示詞素養的提出提供理論動因。

（四）提示詞素養的理論定位

基於上述脈絡，提示詞素養可被理解為建立在數位素養與 AI 素養之上的「延伸能力」，而非取代既有素養。其核心不只是寫出更長或更精細的提示，而是能把學習目標與任務需求轉譯為可與模型協作的互動策略：包括界定問題、設定限制、提供脈絡、提出評估標準、進行多輪追問與修正，並對輸出進行品質判讀與責任歸因。

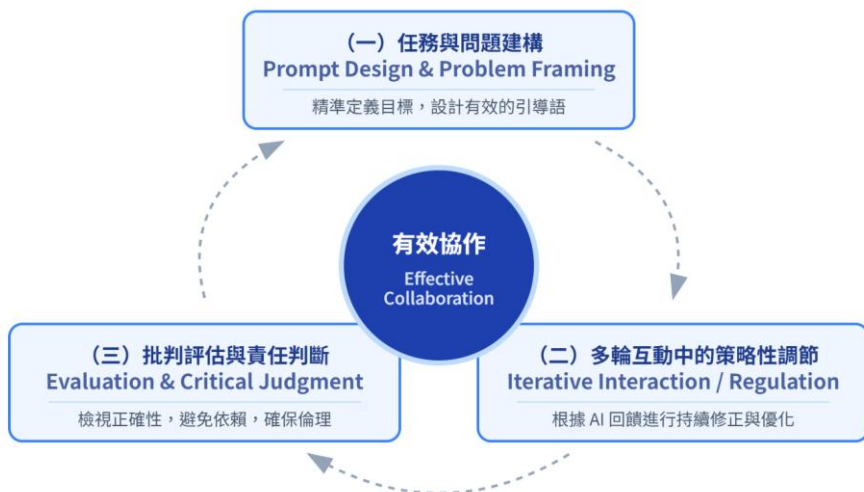
就概念建構而言，提示詞素養的重要貢獻在於把「人機互動」前移到學習歷程中心，使「如何提問、如何引導、如何驗證」得以成為教學分析與能力培育的核心單位（Tour & Zadorozhnyy, 2025）。相對地，若僅將提示詞素養視為 AI 素養的子項，容易忽略其在互動歷程與問題建構層面的獨特性；相關回顧亦指出，提示詞素養應同時涵蓋對系統侷限的理解、對生成內容非權威性的批判意識，以及可遷移的策略性互動能力（Plaatjies & van Wyk, 2025）。此外，若從批判與權力視角理解資料與技術實作，亦可提醒我們：任何看似中立的輸出都可能隱含價值選擇與不平等再製，因而提示詞素養的培育不宜脫離倫理與公平的討論（D'Ignazio & Klein, 2020）。

綜合而言，數位素養從操作取向走向批判取向，奠定了「科技使用必須連結判斷與責任」的基本立場（Lankshear & Knobel, 2011; Street, 1984）；AI 素養則進一步把焦點推向對 AI 系統概念、限制與倫理風險的理解與評估（Long & Magerko, 2020; Ng et al., 2021）。而生成式 AI 的互動式生成特性，使「如何以語言引導系統、如何在多輪互動中協作與驗證」成為新的能力核心，促成提示詞素養作為延伸能力的理論定位（Kasneci et al., 2023; Tour & Zadorozhnyy, 2025）。此一定位不僅使提示詞素養得以被視為具教育延續性與分析價值的能力結構，也對其能力構念與層次化框架的建構提供必要的理論基礎（Plaatjies & van Wyk, 2025; Qian, 2025）。

二、提示詞素養的核心能力構面

前述研究回顧已指出提示詞素養並非可被簡化為「把提示寫得更長、更詳細」的操作技巧，而是一種建立在素養理論、數位／AI 素養之上、並因生成式 AI 互動性而被凸顯的延伸素養（Long & Magerko, 2020; Ng et al., 2021; Tour & Zadorozhnyy, 2025）。本部分進一步聚焦於提示詞素養的核心能力構面，主張其可被操作化為三個相互連動的面向（如圖 1）。

圖 1
提示詞素養的核心能力構念



註：出自作者依據 Long & Magerko (2020)、Ng et al. (2021)、Kasneci et al. (2023)、Zamfirescu-Pereira et al. (2023) 與 Tour & Zadorozhnyy (2025) 理念繪製。

其內容分別為（一）任務與問題建構導向的提示設計（prompt design / problem framing）；（二）多輪互動中的策略性調節（iterative interaction / regulation）；以及（三）生成結果的批判評估與審慎判斷（evaluation / critical judgment）。這三個面向共同構成學習者在生成式 AI 情境中能否有效協作、避免過度依賴並產出可負責任的成果的關鍵機制（Kasneci et al., 2023;

Zamfirescu-Pereira et al., 2023)。

須先說明的是，圖 1 所呈現者為提示詞素養的橫向核心能力構面，重點在於說明學習者於人機互動歷程中需要動員哪些能力；下文圖 2 所呈現者則為提示詞素養的縱向發展層次，重點在於說明這些能力如何建立在不同素養基礎之上並逐步深化。兩圖並非重複，而是共同構成「構面—層次」的分析矩陣。

（一）提示設計：將學習目標轉譯為可協作的任務規格

提示設計首先涉及學習者能否將原本模糊或複雜的學習目標，轉譯為生成式 AI 可處理的「任務規格」。此一過程不只是語句措辭，而是包含問題界定、情境提供、限制條件設定、輸出格式規範與品質標準描述等結構化行動。從 AI 素養的觀點，使用者需理解系統能力與限制，才能做出有效的任務拆解與提示安排 (Long & Magerko, 2020; Ng et al., 2021)。然而，生成式 AI 的互動特性使「設計」不再是一次性輸入，而是把提示視為一種可迭代的設計物：使用者在互動中逐步校準問題、補足脈絡並明確化評量準則，以引導模型生成更符合意圖的回應 (Tour & Zadorozhnyy, 2025)。相關實證亦顯示，學習者在學術寫作任務中呈現的提示型態，常可區分出是否提供足夠背景與具體需求；較高 AI 素養者傾向採用較具情境與規格的提示策略，並能更有效運用生成式 AI 作為協作夥伴 (Kim et al., 2025)。

就操作化而言，提示設計至少可再區分為四個次面向：其一，任務拆解能力，係指將複雜問題轉化為模型可回應之步驟或子任務；其二，脈絡提供能力，係指能補入必要背景、對象特性與使用情境；其三，限制條件設定能力，係指能界定範圍、角色、格式、篇幅或判準，以避免輸出過度發散；其四，輸出規格指定能力，係指能明確要求回應的結構、語氣、證據形式或呈現格式。如此一來，提示設計便不再只是抽象概念，而可成為較明確的分析與教學基礎。

（二）互動調節：多輪對話中的策略迭代與後設控制

第二個能力面向是「互動調節」，亦即學習者在多輪對話中能否以策略性方式調整提示、追問、澄清與修正，以逐步提升輸出品質。這裡的核心不僅是「會提問」，更在於能否在互動過程中形成可辨識的調節循環：辨識不足、提出釐清問題、調整限制或示例、比對生成結果，再進一步修正。HCI 研究指出，非 AI 專業使用者在設計提示時經常失敗，其困難主要在於尚未建立將互動視為「逐步試探與校準」的策略框架；換言之，提示互動的有效性高度依賴使用者能否形成迭代心智模型與互動策略（Zamfirescu-Pereira et al., 2023）。在語言學習與寫作修訂情境中，研究亦觀察到學習者在 ChatGPT 輔助修訂時呈現不同的提示行為類型，並且行為與修訂品質之間存在關聯；更具策略性的學習者會明確表達修訂目標、要求具體建議、要求示例與理由，並在回應後再追問以對齊需求（Hwang et al., 2024）。同樣地，課堂互動研究也顯示，學生的參與程度與對工具的感知會影響其提示互動深度與持續性，進而影響學習成效與使用信任（Yang, 2025）。

（三）批判評估：把生成結果視為「可疑但可用」的中介資源

第三個能力面向是「批判評估與審慎判斷」，亦即學習者如何檢核、比較、修正與歸因生成內容。生成式 AI 的輸出具有高可讀性與擬真性，容易引發「權威錯覺」與過度信任，因此提示詞素養必須包含對輸出可靠度的判斷與驗證行動，如：要求來源與依據、交叉查核、辨識不確定性、檢視偏誤與不當推論等（Kasneci et al., 2023）。生成式 AI 的多模態擴展亦使評估議題更為重要；研究指出，當生成資源跨越文字、圖像等模態，學習者更需要具備辨識適切性、檢核一致性與修訂輸出的能力（Kang & Yi, 2023）。在高等教育情境的研究亦發現，較成熟的人機協作經驗通常伴隨「對輸出進行修改與精煉」而非直接提交；學習者會把 AI 產出視為草稿或素材，透過人類判斷進行再創作與責任承擔，顯示評估與修正是人機協作不可或缺的環節（Hwang & Lee, 2025）。此外，AI 素養綜述亦強調倫理與

評估面向的重要性，指出 AI 素養不應止於使用，而應包含對影響與倫理後果的辨識（Ng et al., 2021）。因此，在本研究的論述架構下，提示詞素養的評估面向可被界定為：能把 AI 產出當作「中介資源」而非「最終答案」，並以批判性標準完成查核、修正、引用與責任歸屬。

綜合而言，本部分將提示詞素養操作化為「設計、調節、評估」三構念：提示設計負責把任務規格化，互動調節負責在多輪對話中校準輸出，而批判評估負責確保生成結果的可靠度與責任可歸屬。這一能力鏈條同時回應 AI 素養強調的理解、使用、評估與倫理面向（Ng et al., 2021），並補足 AI 素養在生成式 AI 情境中對「互動策略」與「迭代調節」的不足（Zamfirescu-Pereira et al., 2023）。就教育實作而言，上述構念也提供後續進一步發展評量指標與教學介入設計的基礎：例如可透過學習者提示文本的結構、互動回合的調節行為，以及對生成內容的查核與修正紀錄，形成較具體的能力證據（Hwang et al., 2024; Hwang & Lee, 2025; Kim et al., 2025）。

三、實證研究對提示詞素養框架建構的啟示

既有實證研究顯示，提示詞素養已逐步由概念性主張轉為可觀察的教育研究對象。相關研究多不再僅以單一提示品質作為判準，而是從提示文本、互動歷程、修正行為與評估行動等面向加以操作化，據以分析學習者如何在生成式 AI 情境中界定任務、調整策略並檢核輸出品質（Hwang et al., 2024; Moorhouse et al., 2025）。此一發展顯示，提示詞素養並非單一技巧，而是可透過歷程性證據加以辨識的複合能力。

進一步而言，既有研究亦指出，不同學習者在提示設計深度、互動調節能力與生成結果評估上呈現明顯差異，而此類差異常與 AI 素養、任務理解、人機協作經驗及後設監控能力有關（Kim et al., 2025; Hwang & Lee, 2025）。同時，教學介入研究發現，若缺乏明確示範、比較、反思與回饋機制，學生多停留於表層使用；反之，若將提示結構、追問策略與生成內

容驗證納入教學設計，則較能促進較高層次的提示互動能力（Qian, 2025; Plaatjes & van Wyk, 2025）。

綜合上述，既有實證研究一方面說明提示詞素養具有可觀察的能力差異與教學可塑性，另一方面也顯示其表現並非單由單次提示技巧所決定，而是涉及基礎素養、互動策略與批判判準的整合運作。基於此，有必要進一步以層次化框架整合既有研究中的分散發現，以說明不同能力成分之間的關係及其在教育情境中的發展邏輯。

四、提示詞素養的層次化框架與層次關係

承接前述理論脈絡、核心構面與實證脈絡，本部分進一步透過概念分析與跨研究比較，說明為何提示詞素養需要以層次化方式加以理解，並據此建構三層次框架與層次間的動態關係，以避免其被誤解為僅是單一技巧清單。

透過前述文獻分析可知，提示詞素養並非單一操作技巧，而是涵蓋提示設計、多輪互動修正，以及生成內容評估與反思等相互關聯的能力集合。然而，僅以能力清單方式呈現，仍不足以回應不同研究情境中提示詞素養的發展差異，亦難以說明其在教育體系中如何被培育與評量。因此，有必要進一步將這些能力置於層次化結構中加以理解，使提示詞素養得以被視為一種具發展性與可制度化潛力的能力架構。

（一）層次化理解的必要性：從能力聚合到發展框架

既有研究雖已辨識多項與提示互動相關的能力，但多半未明確處理能力之間的關係。有些研究將提示詞素養等同於進階技巧，假定使用者只要掌握較複雜的提示語法即可提升表現（Chu & Chiu, 2025）；另一些研究則將其視為 AI 素養的自然延伸，未區分不同層次能力所需的前提條件（Tour & Zadorozhnyy, 2025），此類取向容易導致兩種問題：其一，忽略基礎數位與 AI 素養不足時，提示教學的成效有限；其二，將不同認知層次的能力混為一談，降低理論解釋力。因此，本研究主張以層次化視角理解提示詞

素養，將其視為建立在既有素養基礎之上的高階能力結構。此一取向不僅有助於澄清不同研究中能力界定的差異，也能為課程設計與教師專業發展提供較清楚的發展路徑。

（二）建構三層次提示詞素養框架

綜合前述文獻分析，本研究提出一個由低至高、彼此遞進且相互交織的三層次提示詞素養框架，如圖 2。此架構並非意圖取代既有素養概念，而是作為一種整合性視角，說明不同能力如何在生成式 AI 互動中共同運作。第一層次可視為基礎數位與 AI 素養層。此層次涵蓋使用者對數位工具的基本操作能力，以及對生成式 AI 運作原理的初步理解，包括對模型限制、資料來源不確定性與生成結果非權威性的認知（Qian, 2025）。缺乏此一層次能力時，使用者往往難以判斷提示失效的原因，亦無法進行有意義的修正。

第二層次為生成式 AI 互動與提示操作層，對應於前述核心構面所分析的提示設計與多輪互動能力。在此層次中，使用者能根據任務需求設計結構化提示，並在互動過程中即時調整策略，以逐步引導模型產出符合情境的內容（Chu & Chiu, 2025; Kim et al., 2025）。此層次能力強調操作性與策略性，亦是多數實證研究目前聚焦的核心。

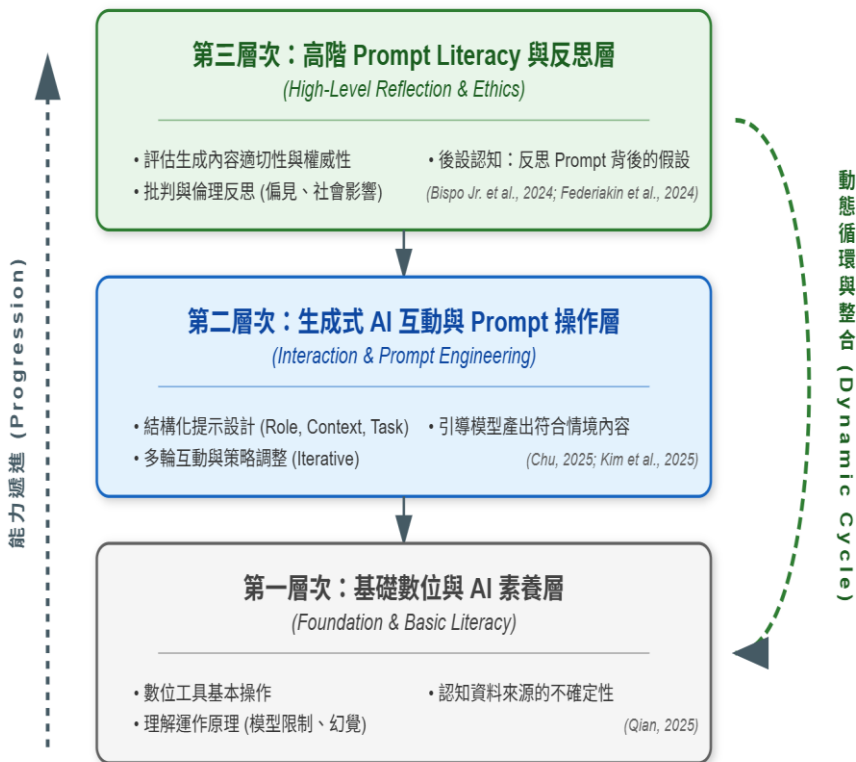
第三層次則為高階提示詞素養與反思層。此層次超越操作技巧，聚焦於使用者對生成內容的評估、比較與倫理反思能力。使用者不僅能判斷生成結果的適切性，亦能反思提示設計背後的假設，並考量生成內容在學術、社會與倫理層面的影響（Bispo Jr. et al., 2024; Federiakin et al., 2024）。此一層次亦與後設認知與批判思考能力高度相關，是提示詞素養能否支撐長期學習與負責任使用的關鍵。

（三）層次間的關係：必要條件、遞進發展與動態回饋

值得強調的是，三層次能力並非線性階梯，也非「學會高階能力即可取代低階能力」的關係。相反地，各層次能力在實際互動中往往同時被動

員，形成動態循環（Chen et al., 2024）。例如，高階評估與反思能力會回饋至提示設計策略的調整，而基礎 AI 素養的提升亦可能改變使用者對提示操作的理解方式。此一觀點回應了部分研究對「提示工程技術化」的疑慮，指出提示詞素養的核心價值不在於技巧累積，而在於不同層次能力之間的整合與轉化。相較於僅列舉能力向度的研究取向，層次化框架有助於解釋為何在相似教學情境中，學習者的提示使用表現仍呈現顯著差異。這些差異往往源自於不同層次的能力發展不均，而非單一技巧的掌握程度。

圖 2
三層次提示詞素養動態框架圖



註：出自作者依據 Qian (2025)、Kim et al. (2025)、Tour & Zadorozhnyy (2025)、Bispo Jr. et al. (2024) 與 Federiakin et al. (2024) 理念繪製。

換言之，第一層的基礎數位與 AI 素養並非與後兩層並列，而是第二層互動操作與第三層高階批判得以成立的必要條件；第二層則提供實際的人機協作經驗，使學習者能在反覆操作中累積對模型限制與回應品質的判斷；第三層所形成的批判反思，又會反過來深化對第一層與第二層的理解，避免三層次被誤解為靜態或平行分類。

（四）與既有研究的對照及補充意涵

將提示詞素養置於層次化結構中，有助於回應既有研究中的概念模糊問題。部分研究雖已提出提示詞素養的操作化指標，但未明確說明其所對應的能力層次（Tour & Zadorozhnyy, 2025）。本研究所提出的三層次架構，則可作為一種分析工具，用以重新詮釋既有研究成果，並比較不同研究在能力層次上的側重。

此外，此一框架亦為後續實證研究與制度化培育提供理論基礎。研究者可依據不同層次設計評量工具與教學介入，而教育實務者亦能據此區分短期技能訓練與長期能力培育的不同目標。此一層次化理解，使提示詞素養得以從零散的研究描述，轉化為可被累積與比較的學術概念。

綜合而言，本部分提出的三層次提示詞素養框架，建立在既有文獻對核心構面的分析之上，並進一步回應其概念分散與理論斷裂的問題。透過層次化視角，提示詞素養得以被理解為一種具發展性、整合性與制度化培育潛力的素養結構，而非單一技巧或短期教學策略。以此層次化框架作為整合與發展的支點，亦為探討提示詞素養在課程、教師專業與制度層面的培育與研究方向，奠定概念基礎。

五、提示詞素養的理論脈絡與制度化培育意涵

承接前述概念定位、三構面與三層次框架，本部分進一步討論提示詞素養如何由理論分析推導至教育實踐。本部分將重點置於說明課程設計、教師專業發展、評量機制與治理支持等面向，如何由前述概念框架推演而來，以凸顯制度化培育意涵與理論分析之間的連續性。

前述分析顯示，提示詞素養可被理解為一種具層次性與發展性的複合素養；因此，其制度化培育意涵必須回到本文所提出之三構面與三層次框架加以推導。具體而言，提示設計、互動調節與批判評估三構面，分別對應課程任務設計、學習歷程引導與結果判準建構；而基礎素養、互動操作與高階反思三層次，則提供教育體系規劃培育順序與支持機制的參照。

（一）課程設計：由提示技巧轉向能力導向學習

在課程層面，既有研究多以單元式或短期介入方式導入提示教學，著重於提示設計技巧的示範與練習（Qian, 2025）。此類課程設計雖能在短期內提升學習者的操作熟練度，但亦被批評為過度聚焦於工具使用，未能系統性培養高階的反思與評估能力。相對而言，部分研究開始主張將提示詞素養納入能力導向或素養導向課程設計中，強調與學科任務、學習目標及評量方式的連結（Qian, 2025）。在此取向下，提示互動不再被視為獨立技能，而是作為支持學習歷程的中介工具，使學習者在完成寫作、問題解決或創作任務時，同時發展其提示設計與後設反思能力。這類研究顯示，若缺乏層次化能力框架作為指引，課程設計往往難以在不同學習階段之間建立連續性。

（二）教師專業發展：由技術支援走向學習引導

制度化培育提示詞素養亦涉及教師專業角色的轉變。多數實證研究指出，教師在生成式 AI 教學中的角色，並非單純提供工具操作說明，而是協助學生辨識何種提示策略適用於特定學習情境，並引導其反思生成結果的品質、限制與責任歸屬。

然而，若缺乏系統性的教師專業發展支持，教學現場仍可能把提示教學簡化為範例套用或效率導向技巧，進而限制學生發展高階提示詞素養的可能性。因此，教師專業發展不宜被理解為附帶措施，而應與 AI 素養、課程設計與評量素養共同納入制度化培育架構之中。

（三）評量機制：由結果品質擴展至歷程與倫理判準

除課程與教師專業外，提示詞素養的制度化亦涉及評量與倫理層面的挑戰。若僅以生成結果品質作為評量指標，容易忽略提示互動歷程中的問題建構、互動調節與批判查核能力，也可能導致學習者過度追求表面成效而忽視責任歸屬。因而，制度化評量需同時關注提示文本、互動歷程、查核修正與倫理判準，才能較完整反映提示詞素養的能力內涵。

（四）治理支持：由資源配置延伸至制度規範

治理支持之所以重要，在於提示詞素養若僅依賴個別教師或個別課程的自發導入，往往難以形成穩定且公平的培育條件。教育體系需要同步處理資源可近用性、校內使用規範、學術誠信指引、資料與隱私保護，以及師生對生成式 AI 風險與責任的共同理解。唯有將這些支持條件制度化，提示詞素養才不致淪為少數高資源學習者的競爭優勢，而能成為教育系統中的共同能力目標。

由此可見，本研究所稱的制度化培育意涵，係指將提示詞素養納入課程設計、教師專業發展、評量機制與治理支持之整體規劃之中。如此，制度化培育意涵便不再只是文獻摘述，而能作為前述概念框架的教育實踐延伸。

（五）理論脈絡圖與後續研究方向

若從理論整合角度回看全文，提示詞素養之所以值得獨立概念化，正在於其同時連結素養理論、數位素養、AI 素養與生成式 AI 互動實踐。圖 3 是用以說明提示詞素養如何立基於數位素養與 AI 素養之上，並進一步連結問題建構、互動調節、批判查核、人機協作與公平治理等觀點，以凸顯其在理論整合上的學術貢獻，並非著重「大小素養」的靜態層級。

1. 評量工具與能力指標的持續建構

儘管近年已有研究嘗試透過提示文本分析、互動歷程編碼或修訂品質比較來操作化提示詞素養，但整體而言，系統性、可重複使用的評量工具仍相當有限。多數研究仍倚賴情境化分析或質性分類，較少發展具信效度

檢驗的量化指標或混合方法設計（Chu & Chiu, 2025; Koong Lin et al., 2025）。此一限制使得不同研究之間的比較困難，也影響提示詞素養作為教育成效指標的應用潛力。未來研究有必要在不過度簡化能力內涵的前提下，發展能同時捕捉「提示結構、互動策略與評估行為」的多維度評量架構。

2. 研究情境與對象的擴展

目前多數提示詞素養的實證研究集中於高等教育場域，特別是學術寫作、語言學習與部分專業培育情境。相對而言，中小學教育、技職教育與非正式學習場域的研究仍相當有限（Koong Lin et al., 2025）。此外，現有研究對學習者背景變項（如語言能力、學科知識、文化脈絡）與提示詞素養發展之間的關係探討不足，可能低估了不同學習族群在互動策略與評估能力上的差異。未來研究若能擴展研究對象與教育層級，將有助於檢驗提示詞素養作為「跨情境素養」的適用性與限制。

3. 倫理、權力與責任面向的探討

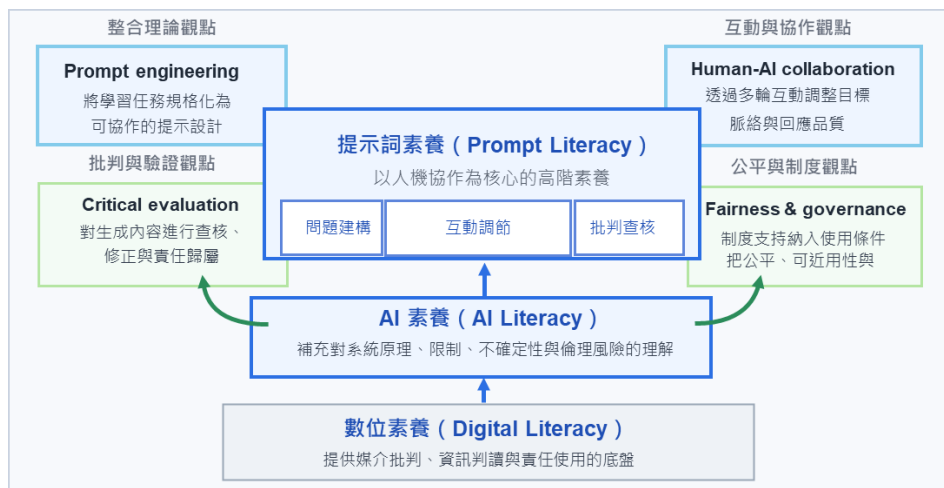
雖部分文獻已開始關注生成式 AI 的倫理與公平議題（Bispo Jr. et al., 2024），但在提示詞素養研究中，倫理往往仍被視為附加考量，而非能力構念的核心部分。然而，提示設計與互動策略本身即隱含價值選擇與權力關係，例如哪些觀點被納入、哪些被排除，以及使用者如何回應模型可能產生的偏誤與不確定性。若忽略此一面向，提示詞素養容易被窄化為效率導向能力，而未能回應教育對公民責任與批判思考的期待。未來研究可進一步探討倫理反思如何被內嵌於提示互動歷程之中，並成為可培育與可評估的能力要素。

4. 提示詞素養的理論定位與整合意義

就圖 3 的定位而言，素養概念提供理解、判斷與責任實作的基本立場；數位素養提供媒介批判與資訊責任使用的底盤；AI 素養補充對系統原理、限制與倫理風險的理解；而提示詞素養則把問題建構、互動引導、驗證修正與責任歸屬推到人機協作歷程中心。透過此一脈絡圖，本文進一步展現提示詞素養的概念意義：它並非孤立的技巧集合，而是建立在數位素養與

AI 素養之上的跨理論整合節點。

圖 3
提示詞素養的理論脈絡圖



註：出自作者依據 Street (1984)、Long & Magerko (2020)、Ng et al. (2021)、Kasneci et al. (2023)、D'Ignazio & Klein (2020) 與 Tour & Zadorozhnyy (2025) 之相關觀點繪製。

參、結論

本研究透過回顧式概念分析，旨在釐清提示詞素養的理論來源、核心能力構面與層次化理解，並回應當前研究中概念分散、操作不一與主張邊界未明的問題。整體而言，本文將提示詞素養定位為建立在數位素養與 AI 素養之上的人機協作導向高階素養，並以此作為後續研究與教育實務對話的基礎。

一、提示詞素養的概念定位

本研究將提示詞素養定位為一項人機協作導向的高階素養。此一定義使提示詞素養得以嵌入既有素養理論的延續脈絡，而非被視為技術浪潮下

的短期應用概念。透過回顧數位素養與 AI 素養的理論演進，本文指出其核心在於學習者是否能於互動歷程中展現目標界定、策略調節與批判評估的能力。

二、三構面與三層次框架的整合意義

本研究進一步以「設計、調節、評估」三個核心構面整合既有研究，並提出三層次提示詞素養框架，用以說明基礎素養、互動操作與高階反思之間的動態關係。此一框架的意義，在於將分散於不同研究中的提示型態、互動行為與修訂策略，統整為較可比較的概念結構。

值得注意的是，現有提示詞素養實證研究之所以呈現「部分顯著、部分有限」的結果，並不只是研究品質高低不同，更與研究所處理的任務類型、學習者前置素養、互動深度與成效指標有關。以 ChatGPT 輔助英語寫作修訂研究為例，Hwang et al. (2024) 發現，當學習者多以一般化、未對準目標的提示與系統互動時，成效多集中於表層修訂，對內容、組織與連貫等高階面向的改善相對有限。相對地，Kim et al. (2025) 在學術寫作任務中指出，高 AI 素養學生較常採用具脈絡、描述性且可支持協作的提示模式，且在內容、結構與表達表現上皆優於低 AI 素養學生。

此外，不同任務導向的研究亦顯示，學習者在與大型語言模型互動時，往往需要經歷任務辨識、反覆提示互動與結果修整等歷程；而在數位內容創作情境中，提示詞素養較高的學生通常較能把 AI 視為可修改、可精煉的協作夥伴，而非直接提交的答案來源。這些研究結果顯示，提示詞素養的表現差異與任務情境及互動歷程密切相關，亦支持其作為可培育互動能力的研究價值。

三、制度化培育的教育意涵

從教育實務角度來看，提示詞素養若要從概念主張走向制度化培育，關鍵不在於額外增加單次提示技巧訓練，而在於將提示設計、互動調節與

批判評估嵌入學科任務與學習目標之中。這意味著教師角色需從「防弊者」轉向「互動引導者」，並在課程、評量與治理支持上形成一致的培育參照。因此，本研究所稱的制度化培育意涵，重點不在宣稱既有制度模型已然完成，而在指出：若欲使提示詞素養成為可持續發展的教育能力，便必須在課程設計、教師專業發展、歷程性評量與公平治理之間建立清楚連結。

四、研究限制與後續建議

本研究仍有若干限制。首先，本研究以概念整合為主，雖已納入近期實證研究，但對不同教育層級、在地脈絡與非英語情境的涵蓋仍有限。其次，提示詞素養作為新興構念，其操作化與評量方式仍在發展中，本研究提出的三層次框架與制度化培育意涵，仍須透過後續實證研究持續檢驗與修正。正因如此，本研究更適合被理解為一個概念整合與研究議程建構的起點，而非最終定論。

後續研究可進一步把提示詞素養操作化為可評量的多維指標，並在語文、跨學科與不同教育階段情境中，以設計型研究、準實驗或歷程分析方式檢驗其可遷移性與學習成效；同時也應把公平、可近用性與治理支持納入研究設計，以回應生成式 AI 教育應用中的制度條件與差異處境。

總結而言，本研究透過理論整合與文獻回顧，為提示詞素養提供較清楚的概念邊界與能力定位，並指出其作為人機協作導向高階素養的研究與實務潛力。未來若能在不同教育情境中持續檢驗與深化此一框架，提示詞素養將有機會成為理解與引導生成式 AI 教育應用的重要理論支點。

參考文獻

- Bispo Jr., E. L., dos Santos, S. C., & De Matos, M. V. A. B. (2024). Equity issues derived from use of large language models in education. In *New Media Pedagogy: Research Trends, Methodological Challenges, and Successful Implementations* (pp. 425-440). Springer.

- https://doi.org/10.1007/978-3-031-63235-8_28
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
<https://doi.org/10.48550/arXiv.2005.14165>
- Chen, C., Lee, S., Jang, E., & Sundar, S. S. (2024). Is your prompt detailed enough? Exploring the effects of prompt coaching on users' perceptions, engagement, and trust in text-to-image generative AI tools. In *TAS '24: Proceedings of the Second International Symposium on Trustworthy Autonomous Systems* (Article 9, pp. 1-12). Association for Computing Machinery.
<https://doi.org/10.1145/3686038.3686060>
- Chu, K. Y., & Chiu, P. M. H. (2025). Prompt design and AI response quality in university students' academic learning tasks. *Proceedings of the Association for Information Science and Technology*, 62(1), 1402-1404.
<https://doi.org/10.1002/pra2.1417>
- D'Ignazio, C., & Klein, L. F. (2020). *Data feminism*. MIT Press.
<https://doi.org/10.7551/mitpress/11805.001.0001>
- European Union. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. EUR-Lex.
<https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Federiakin, D., Molerov, D., Zlatkin-Troitschanskaia, O., & Maur, A. (2024). Prompt engineering as a new 21st century skill. *Frontiers in Education*, 9, 1366434.
<https://doi.org/10.3389/feduc.2024.1366434>
- Hwang, Y., & Lee, J. H. (2025). Exploring students' experiences and perceptions of human-AI collaboration in digital content making. *International Journal of Educational Technology in Higher Education*, 22, 44.
<https://doi.org/10.1186/s41239-025-00542-0>
- Hwang, M., Jeens, R., & Lee, H. K. (2024). Exploring learner prompting behavior and its effect on ChatGPT-assisted English writing revision. *The Asia-Pacific Education Researcher*, 34(3), 1157-1167.
<https://doi.org/10.1007/s40299-024-00930-6>
- Kang, J., & Yi, Y. (2023). Beyond ChatGPT: Multimodal generative AI for L2 writers. *Journal of Second Language Writing*, 62, 101070.
<https://doi.org/10.1016/j.jslw.2023.101070>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Geyer, C., Gündogdu, E., Ischen, C., Kasneci, G., Löser, A., Roessler, M., & Stefes, S. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103,

102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kim, J., Yu, S., Lee, S. S., & Detrick, R. (2025). Students' prompt patterns and its effects in AI-assisted academic writing: Focusing on students' level of AI literacy. *Journal of Research on Technology in Education*, 1-18.
<https://doi.org/10.1080/15391523.2025.2456043>
- Koong Lin, H. C., Lu, R. S., & Wang, T. H. (2025). Writing is coding for sustainable futures: Reimagining poetic expression through human-AI dialogues in environmental storytelling and digital cultural heritage. *Sustainability*, 17(15), 7020.
<https://doi.org/10.3390/su17157020>
- Lankshear, C., & Knobel, M. (2011). *New literacies: Everyday practices and social learning* (3rd ed.). Open University Press.
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1-16). Association for Computing Machinery.
<https://doi.org/10.1145/3313831.3376727>
- Moorhouse, B. L., Ho, T. Y., Wu, C., & Wan, Y. (2025). Pre-service language teachers' task-specific large language model prompting practices. *RELC Journal*.
<https://doi.org/10.1177/00336882251313701>
- New London Group. (1996). A pedagogy of multiliteracies: Designing social futures. *Harvard Educational Review*, 66(1), 60-92.
<https://doi.org/10.17763/haer.66.1.17370n67v22j160u>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). AI literacy: Definition, teaching, evaluation and ethical issues. *Proceedings of the Association for Information Science and Technology*, 58(1), 504-509.
<https://doi.org/10.1002/pra2.487>
- Organisation for Economic Co-operation and Development. (2025). *AI literacy framework for primary and secondary education* (Review draft). OECD.
<https://oecd.ai/en/resources/ailit-framework>
- Plaatjies, B. O., & van Wyk, M. M. (2025). Prompt literacy as an enhancer of students' academic writing in higher education institutions: A systematic literature review. *Journal of Teaching and Learning*, 19(4), 114-134.
<https://doi.org/10.22329/jtl.v19i4.10017>
- Qian, Y. (2025). Pedagogical applications of generative AI in higher education: A systematic review of the field. *TechTrends*, 69(5), 1105-1120.
<https://doi.org/10.1007/s11528-025-01100-1>
- Street, B. V. (1984). *Literacy in theory and practice*. Cambridge University Press.
- Tour, E., & Zadorozhnyy, A. (2025). Conceptualizing and operationalizing prompt literacy for English language learners. *Journal of Adolescent & Adult Literacy*, 69(3), e70020. <https://doi.org/10.1002/jaal.70020>
- UNESCO. (2023). Guidance for generative AI in education and research. UNESCO.

<https://doi.org/10.54675/EWZM9535>

UNESCO. (2024). *AI competency framework for teachers*. UNESCO.

<https://doi.org/10.54675/ZJTE2084>

Yang, H. (2025). EFL student engagement with ChatGPT in college reading classes via prompts and perceptions. *Korean Journal of English Language and Linguistics*, 25, 686-706. <https://doi.org/10.15738/kjell.25.202505.686>

Zamfirescu-Pereira, J. D., Wong, R. Y., Hartmann, B., & Yang, Q. (2023). Why Johnny can't prompt: How non-AI experts try (and fail) to design LLM prompts. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems (CHI '23)*. Association for Computing Machinery.

<https://doi.org/10.1145/3544548.3581388>

Conceptualizing Prompt Literacy: A Layered Framework and Implications for Institutional Cultivation—A Conceptual Review Approach

Ming-Wen Chang¹ Kuo-Sheng Chen² Chien-Chih Chen³

¹Commissioner, Education Department, New Taipei City Government

²Headquarter Counselor, Education Department, New Taipei City Government

³Associate Professor, Department of Educational Management, National Taipei University of Education

Abstract

This study examined prompt literacy in educational contexts shaped by generative artificial intelligence and large language models. It adopted a conceptual review approach and integrated scholarship on literacy, digital literacy, AI literacy, and human-AI interaction to clarify the concept's theoretical position, core competence dimensions, and layered development. The review found that prompt literacy was not merely a technique for composing longer or more detailed prompts. Rather, it functioned as an advanced, human-AI-collaboration-oriented literacy built upon digital and AI literacy. The analysis identified three interrelated competence dimensions: prompt design, interaction regulation, and critical evaluation. Prompt design referred to translating learning goals into workable task specifications; interaction regulation referred to strategic adjustment across multiple rounds of dialogue; and critical evaluation referred to examining, revising, and taking responsibility for AI-generated outputs. Based on these findings, the study proposed a

three-level framework that connected foundational digital and AI literacy, generative-AI interaction and prompt operation, and higher-order reflection. The framework showed that the levels were not a linear hierarchy but a dynamic relationship in which foundational understanding supported interaction, interaction generated evidence for evaluation, and higher-order reflection fed back into subsequent prompt design and AI understanding. The review also distinguished horizontal competence dimensions from vertical developmental levels, thereby reducing conceptual overlap among prompt design, AI literacy, and broader digital literacy. This distinction was intended to strengthen its use as a shared reference for educational research and practice. The study further discussed implications for institutional cultivation in curriculum design, teacher professional development, assessment, and governance support. It suggested that prompt literacy should be cultivated through subject-based tasks, teacher-guided reflection, process-oriented assessment, and equitable governance conditions rather than through isolated prompt-writing tips. The contribution of this study lay in offering an integrated conceptual framework that could support future empirical research, curriculum development, and educational policy discussion on responsible generative AI use.

Keywords: AI literacy; human-AI collaboration; large language models; generative artificial intelligence; prompt literacy

