

國家科學及技術委員會

研究誠信電子報

第 64 期

2026 年 3 月

► 案例介紹

研究計畫重複向不同學門提出申請，涉及違反學術倫理

甲君於 113 年度專題研究計畫申請期間，同時向二個不同學門提交 A、B 二件研究計畫，經職權發現，二件計畫申請書內容高度相似，爰依規定啟動調查。本案經審查認定，甲君所提二件計畫，除計畫名稱不同外，申請書內容完全雷同，甲君有以「同一研究計畫同時重複向本會提出申請」之情事，構成本會學術倫理案件處理及審議要點第3點第8款「其他違反學術倫理行為」，決議予以書面告誡之處分。

根據本會補助專題研究計畫作業要點第26點(五)，「同一研究計畫不得同時重複向本會提出申請，違反規定者，依本會學術倫理案件處理及審議要點規定處理。」因此，研究人員若涉及「一案多投」（即同一研究計畫重複申請補助），依規定將啟動學術倫理案件的調查，視個案事實及情節輕重，除予以書面告誡或停權之處分外，涉及重複申請的計畫也可能面臨經費追繳，例如在第36期電子報案例中，若計畫已獲得補助並執行，本會得追回該研究計畫之補助款。

關於「一案多投」，本會過去已多次宣導，在第47期電子報中，也整理過去常見的一案多投錯誤態樣，就是希望研究人員可以重視「一案多投」涉及違反學術倫理的議題。至於如何界定研究計畫內容是否構成「高度相似」，通常不只是從文字重複的量化程度來看，更關注在二件計畫實質內容的差異，提醒研究人員在撰寫及提交計畫申請書時，應審慎審閱、確認申請書內容及版本，並遵守相關學術倫理規範，以維護學術誠信及資源分配之公平性。

此外，學研機構亦應落實對申請、執行研究計畫人員的監督管理責任，尤其是對新進人員，應針對加強輔導其學術倫理知能，並予以學術倫理教育訓練，避免因不諳規定而衍生學術倫理爭議。

► 專欄文章

生成式人工智慧於學術寫作之應用原則與研究責任探討

關於本期：

生成式人工智慧 (Generative Artificial Intelligence, GenAI) 近年快速進入高等教育與學術研究場域，對研究計畫撰寫與學術論文產出產生實質影響。然而，學術內容的原創性、正確性、有效性與完整性，仍須由人類作者負起完全責任。GenAI 不得列為作者或合著者，研究責任亦無法轉移至工具本身，否則將動搖研究誠信的核心基礎。

在缺乏清楚使用界線的情況下，GenAI 的引入可能存在許多學術風險。這些風險不再僅限於傳統意義上的抄襲，而是進一步延伸至研究造假、研究變造、錯誤歸因、不當引用，以及未揭露人工智慧 (Artificial Intelligence, AI) 使用所造成的不透明問題。尤其在學術寫作過程中，GenAI 能以高度流暢且極具學術語感的語言生成內容，使研究人員在不自覺的情況下，將原本應由人類承擔的學術判斷外包給工具，形成新的研究誠信盲點。

一、生成式人工智慧進入學術寫作場域的倫理挑戰

多數國際學術出版社、研究倫理委員會與高等教育機構已逐步形成共識 (Luo, 2024)，認為 GenAI 工具可作為研究與寫作歷程中的輔助工具，特別是在語言潤稿、翻譯與結構整理等層面，確實有助於降低語言門檻並提升表達清晰度。在論文或研究計畫撰寫過程中，GenAI 工具的實際使用方式，可能引發不同層次的學術誠信風險。本文從實務角度整理研究人員於撰寫論文或計畫書時使用 GenAI 的三種常見型態，分別說明其潛在的學術倫理爭議，並對應提出相應的實務原則。

(一) 常見風險作法一：使用 GenAI 直接生成文獻探討與引用清單

最具風險、亦最不建議的使用方式，是直接要求 GenAI 平臺撰寫特定主題的文獻探討段落，並同時生成文獻清單。即便部分平臺（如 Perplexity AI）會提供可點擊的文章連結，讓使用者進行來源查證，實務上仍可觀察到，其所推薦的文獻多半集中於 Open Access 文章，引用範圍有限，且仍需由人類作者自行判斷文獻是否真正支撐其論述。若研究人員僅因工具附上連結便降低查證標準，相關風險並未因此消失。

GenAI 過去最為嚴重的「幻覺（Hallucination）」問題，常表現在生成看似合理、實際卻不存在的學術資訊，例如假作者搭配假標題、真作者搭配假文章，或虛構的論文標題、數位物件識別碼（Digital Object Identifier, DOI）與統一資源定位符（Uniform Resource Locator, URL）。例如：當系統依美國心理學會（American Psychological Association, APA）發布的第 7 版論文寫作格式生成引用¹時，往往在格式上看似相當完整，文獻探討內容亦看起來合理可信，但其風險已不僅是抄襲，而可能進一步構成研究造假、研究變造、錯誤歸因、假引用，或未揭露使用 AI 的不透明行為。

隨著 GenAI 持續更新，其錯誤率已有所下降，部分平臺亦開始主動提醒使用者相關風險。然而，爭議最為集中的時期，正落在 GenAI 已廣泛被使用、但研究人員尚未建立正確使用觀念之階段。即便目前多數工具已能提供可供查證的連結，研究責任仍無法轉移，相關風險依然存在。

要回應上述操作方式所涉及的學術倫理風險，研究人員必須針對 GenAI 所生成的內容與引用，透過其他文獻搜尋工具和學術資料庫進行「反查證」。AI 不對內容真實性負責，所有學術歸屬與引用正確性，仍須由作者承擔。然而，逐一反

¹ 以引用作者 A 在 2026 年出版的英語國際期刊論文為例，其 APA 第 7 版寫作格式為：
Author, A. A. (2026). Title of the article. *Title of the Journal*, volume number (issue number), page range.
<https://doi.org/xxxxx>

查驗生成內容之出處的正確性不僅勞心費力，其效率未必高於研究者親自進行文獻探討；更重要的是，直接複製使用 GenAI 所生成的文獻內容，本身也並非值得鼓勵的研究態度。文獻探討本是研究動機與理論背景的重要基礎，若在此階段出現偏誤，等同於後續研究建立在不穩固甚至錯誤的基礎之上。

那麼，研究人員最基本應具備的「研究態度」究竟為何？舉例而言，若研究者以遊戲式學習為研究主題，是否理應至少掌握該領域中最新且關鍵的十篇核心文獻，並實際持有其完整論文檔案？若連這樣的基本準備都未能做到，所面臨的恐怕已不僅是學術誠信問題，而是更根本的研究素養挑戰。

（二）常見風險作法二：以 GenAI 摘述既有論文卻忽略引用溯源

相較於直接生成文獻探討，另一種較為常見、看似安全的作法，是研究人員先蒐集並上傳論文全文，請 GenAI 協助摘要其研究貢獻，再由作者重組撰寫文獻探討並引用該論文。此方式雖降低假文獻風險，卻仍潛藏錯誤歸因與二次引用的學術倫理問題。

即便是使用具備文件封閉性的大型語言模型(Large Language Model, LLM)，例如 Google NotebookLM，當研究人員將作者 A 的論文全文電子檔案上傳至平臺後，透過擷取增強生成 (Retrieval Augmented Generation, RAG)，提高了資訊檢索的準確性、效率和客製化，透過對話要求可以讓研究人員從這樣的 GenAI 工具中取得看似合理的客製化摘要內容，例如：「A 的主要貢獻在於建構一個新的理論框架，用以解釋 XXX。」然而，NotebookLM 的回答內容中，可能已包含作者 A 對作者 B 論文的轉述、延伸、引用或整合。雖然作者 B 的論文確實曾被 A 所引用，但研究人員卻僅依據 GenAI 的摘要生成來源進行引用，最終只引用了作者 A 的研究，可能會涉及學術不當引用。

要回應上述操作方式所涉及的學術倫理風險，建議研究人員可與 GenAI 平臺進行多次來回對話以釐清內容，但更值得鼓勵的做法，仍是親自回到原始所提供

的論文全文(例如:PDF)檔案內容中進行比對,由人類作者自行判斷究竟是誰最早提出該概念或方法,並確認是否需要補充引用。因為像 Google NotebookLM 這種工具,既然資料來源是人類提供,比起開放式大型語言模型而言較不用憂心幻覺,而且引用的資訊來源都是使用者可以直接查核的,除了留意抄襲風險,也要留意文獻溯源。

將文獻全文檔案提供給 GenAI 平臺,其中一項重要目的,是協助研究人員快速定位研究重點,例如研究貢獻、關鍵段落、理論與研究方法,或問卷與研究工具在文中出現的位置。然而,在此過程中,研究人員仍應主動回到原始文獻,親自檢視其參考文獻與文內引註(in-text citation),以作出審慎且負責任的判斷,而非因依賴 GenAI 而產生「後設認知懶惰」(Metacognitive laziness)(Fan et al., 2025),將原本應由人類完成的學術判斷全然外包給 GenAI。

特別是在 GenAI 摘要中出現「framework」、「model」、「theory」、「scale」、「algorithm」,或「based on」、「adapted from」、「following previous studies」等語句時,更需要研究人員主動進行引文追溯(citation tracing),人工確認該概念究竟為作者 A 的原創,或是 A 引用自作者 B。若屬後者,則應進一步找出作者 B 的原始論文加以引用,同時引用 A 與 B 的研究,才是較為完整且妥適的作法。依此例而言,若作者同時漏引 A 與 B 的論文,便存在明確的抄襲風險。這並非因為 GenAI 以換句話說的方式重述同一概念便可免責,因為其背後仍涉及具體且可識別的人類作者 A 與 B 之產出。即便通過論文比對系統,也僅代表研究人員未「複製文字」,並不能保證未「挪用思想」。

根據前述情況,若研究人員僅漏引 B 的論文,卻將 GenAI 對 A 論文的摘要結果誤視為原創歸屬,顯示該研究人員所欠缺的能力,已非 GenAI 的操作技巧,而是沒做到引文追溯,或為了求快直接全然採納這類 GenAI 的整理結果,意味著其未能妥善履行作為作者所應承擔的學術判斷責任。

此類錯誤往往反映的不是 GenAI 工具本身的問題，而是研究者未能履行應有的學術判斷責任。若研究人員具備引文追溯能力卻選擇完全信任 AI 輸出，等同於主動放棄學術判斷權，將風險轉嫁給工具，這正是 GenAI 時代研究誠信最核心的挑戰之一。

(三) 常見風險作法三：生成不可查證的情境與案例敘事

最後要提的第三種 GenAI 使用風險，在於其所提供之案例與情境的真實性。未加標示的模擬情境，或看似真實卻不可查證的敘事內容，其所構成的問題已不僅是引用不當，而是涉及研究造假與誤導。所謂情境幻覺(Scenario hallucination) 與捏造式範例(Fabricated exemplification)，即是另一類常見的 GenAI 幻覺現象。

這類問題通常並非抄襲，因為抄襲的前提在於指向特定且可識別的來源；然而，若將 AI 捏造的情境描述為「真實發生的案例」，這類內容往往語言自然、情節合理，卻缺乏可驗證的真實來源。若研究人員未清楚標示其為模擬情境或假設案例，而將其寫成「實際發生的案例」，所涉及的已不只是引用不當，而是錯誤陳述甚至研究造假。

在部分教學或概念性論述中，若能明確標示為模擬案例或假設情境，甚至說明係於何種版本的 GenAI、在何種提示條件下所生成，或仍可被接受。然而，在嚴謹的學術論文中，未經驗證的敘事內容極易誤導讀者，破壞研究可信度。部分國際出版社亦明確限制使用 AI 生成的圖像 (Elsevier, n.d.) 與情境模擬內容，除非該研究本身即以 AI 為研究方法，且能清楚交代生成流程與技術參數。

二、從防抄襲走向責任歸屬的研究誠信觀

傳統研究學術倫理訓練多半將焦點放在「抄襲防範」上，但在 GenAI 時代，最具破壞性的問題往往不是文字重複，而是錯誤歸因、假引用，以及學術判斷被無意識地外包。GenAI 能以高度流暢、且極為學術化的語言生成此類判斷性語句，

使研究人員在不自覺的情況下，將原本應由人類承擔的學術判斷，悄然轉交給 AI。這類錯誤通常不會被相似度比對系統偵測到，卻極容易在人工審稿階段被識破，進而對研究信譽造成實質傷害。此類論文寫作錯誤往往在數月後的審稿階段甚至已經接受刊登後才被發現，這樣延遲後果的效應反而使研究人員對風險的覺察太低。

因此，當前研究誠信的核心，除了是否抄襲，在 GenAI 時代更應該強調「誰做出判斷、誰為結果負責」。研究人員可以合理使用 GenAI 改善語言表達或進行腦力激盪，但不能讓 AI 決定研究貢獻的價值、理論的重要性或結論的合理性，所有研究人員都應該有「判斷責任歸屬於作者我自己」的自知，而不是直接讓 GenAI 取代主要的撰寫工作。

三、結語：以人類主導、AI 輔助為原則

從國際學術出版實務觀之，主流共識已相當明確：GenAI 僅能作為輔助工具，用於語言潤稿、格式修正與初步整理，不得取代研究構想、資料詮釋主責、學術判斷與結論形成等核心學術工作。多數期刊亦要求研究人員於致謝或揭露聲明中，清楚說明 GenAI 的使用工具、版本、目的與範圍，以確保研究歷程的透明性與可溯源性。

人機協作的發展方向，並非追求效率導向的捷徑，而是在人類主導的前提下，善用 AI 強化研究者的批判思考、論證與寫作能力。各國頂尖大學與學術機構的政策趨勢亦顯示，「透明揭露」、「人類責任」、「維持原創性」與「培養 AI 素養」已成為共同原則；對研究者與教師而言，清楚界定人與工具的角色分工，並持續強化學術倫理意識，是 AI 時代維繫學術誠信與研究品質的關鍵基礎。

四、參考文獻

Elsevier. (n.d.). *Generative AI policies for journals*. Elsevier. Retrieved January 09,

2026, from <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals>

Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y., ... & Gašević, D. (2025). Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2), 489-530. <https://doi.org/10.1111/bjet.13544>

Luo, J. (2024). A critical review of GenAI policies in higher education assessment: A call to reconsider the “originality” of students’ work. *Assessment & Evaluation in Higher Education*, 49(5), 651-664.

致謝

本文作者：許庭嘉特聘教授 / 國立臺灣師範大學科技應用與人力資源發展學系
(本文內容僅代表作者個人觀點，不代表主管機關立場)

► 資訊補給站

111 年 1 月至 115 年 1 月學術倫理案件統計

整理近5年本會處理學術倫理案件相關統計資料，提供各界參考。

學術倫理案收件與處理情形（統計自 111 年 1 月至 115 年 1 月）

單位：案件數

檢舉方式	具名	121
	未具真實姓名或聯絡方式	20
	職權發現	68
受理結果	不成案	56
	無違反學倫	67
	審查中	40
	有違反學倫	46
合計		209

備註：

1. 統計期間為 111/1/1~115/1/31。
2. 依「國家科學及技術委員會學術倫理案件處理及審議要點」第 2 點規定，本要點適用於申請或取得本會學術獎勵、專題研究計畫或其他相關補助之研究人員，爰申請或取得本會獎補助，疑有違反學術倫理行為者，為本會審議之範圍。
3. 不成案原因包括：事證不足、非本會業管範圍、前案事證已處理。
4. 「有違反學倫」之案件數以收件年度統計，非以處分年度統計。同一案件可能涉及多人。

「違反學術倫理案件」的行為態樣及處分情形

(一) 違反之行為態樣 (統計自 111 年 1 月至 115 年 1 月)

單位：人次

違反之行為態樣	造假	4
	變造	6
	抄襲	11
	自我抄襲 (含隱匿及未適當引註)	7
	重複發表	2
	代寫	2
	影響論文審查	0
	其他	27
	合計	59

備註：

- 統計期間為 111/1/1~115/1/31。
- 違反態樣請參照「國家科學及技術委員會學術倫理案件處理及審議要點」第 3 點；同一人有多種違反態樣，以款次在前計算。
- 108.11.25 修正本會學術倫理案件處理及審議要點，將「隱匿其部分內容為已發表之成果或著作」、「研究計畫或論文大幅引用自己已發表之著作，未適當引註」兩款，整併為「自我抄襲」，並新增「代寫」之態樣；依現行規定，共有 8 款違反學術倫理之行為類型：
 - 造假：虛構不存在之申請資料、研究資料或研究成果。
 - 變造：不實變更申請資料、研究資料或研究成果。
 - 抄襲：援用他人之申請資料、研究資料或研究成果未註明出處。註明出處不當情節重大者，以抄襲論。
 - 自我抄襲：研究計畫或論文未適當引註自己已發表之著作。
 - 重複發表：重複發表而未經註明。
 - 代寫：由計畫不相關之他人代寫論文、計畫申請書或研究成果報告。
 - 以違法或不當手段影響論文審查。
 - 其他違反學術倫理行為，經本會學術倫理審議會議決通過。

(二) 處分情形 (統計自 111 年 1 月至 115 年 1 月)

單位：人次

處分情形	書面告誡	17
	停權 1-2 年	27
	停權 3-10 年以上	4
	追回補助費用、獎勵 (費)、獎金或獎勵金	4
	撤銷獎項	0

備註：

- 統計期間為 111/1/1~115/1/31。
- 處分方式請參照「國家科學及技術委員會學術倫理案件處理及審議要點」第 13 點：學術倫理審議會就違反學術倫理行為證據確切者，得按其情節輕重對當事人作成下列一款或數款之處分建議：(一) 書面告誡。(二) 停止申請及執行補助計畫、申請及領取獎勵 (費) 一年至十年，或終身停權。(三) 追回部分或全部補助費用、獎勵 (費)、獎金或獎勵金。(四) 撤銷所獲相關獎項。
- 受「停權」處分者，共有 4 人同時追回獎補助費用。
- 依 113 年 5 月 2 日修正前之「國家科學及技術委員會學術倫理案件處理及審議要點」第 9 點第 1 項第 1 款第 3 目規定，審查小組審查結果認定違反學術倫理行為，未嚴重違反該學術社群共同接受之行為準則，或未嚴重影響本會審查判斷或資源分配公正之虞者，無須提交學術倫理審議會複審，應視情形為適當之處理。