

國家科學及技術委員會

研究誠信電子報

第 61 期

2025 年 6 月

► 案例介紹

二件研究計畫執行內容雷同，未清楚揭露說明

甲君向本會所提A研究計畫，經發現，與乙君向其他單位所提B計畫，二件計畫之成果報告內容有相當程度雷同，且二件計畫期程亦有重疊，為釐清二件計畫彼此間的關係，爰啟動調查。經查，甲乙二人曾合作發表研究成果，依既有的研究成果，再各自提出研究計畫申請。本案中，乙君同時為A計畫之共同主持人。經審查認定，A計畫之執行內容，與B計畫於核心部分，內容有相當之重複，甲君於申請及繳交成果報告時，未適當註明另一研究計畫之內容，其行為有本會學術倫理案件處理及審議要點第3點第8款「其他」情事，決議予以停權1年等處分。

研究團隊之成員各自提出計畫申請時，應審慎處理計畫申請、成果報告撰寫等事宜，避免衍生學術倫理爭議：實務上經常有研究人員合作研究，再各自以既有的研究成果為基礎展開後續研究。從事延伸性的研究本身當然未違反學術倫理，但應於後續的研究計畫中說明已完成之研究成果及貢獻，避免因未經註明而衍生學術倫理爭議，提醒研究人員注意。

此外，依本會補助專題研究計畫作業要點第26(六)點規定，以同一研究計畫向本會及其他機構(含國內外、大陸地區及港澳)申請補助時，應於計畫申請書內詳列申請本會及其他機構補助之項目及金額，同一項目及金額不得重複申請補助。本會專題研究計畫申請書表格亦要求計畫主持人揭露是否同時有其他單位提供補助項目。若為向本會生科處申請之研究計畫，亦須填寫生科處專屬表格(NSCB09)-「揭露本次計畫申請與政府補助計畫之相關性」，由申請人揭露本次計畫申請與政府補助計畫(如:本會、經濟部、衛福部、教育部、農業部、中研院、國家衛生研究院；不含學校內/醫院院內計畫)之相關性，併予說明。

► 專欄文章

您的作品或是純AI的輸出？—從著作權與CC授權學習AI輔助的標示要素

關於本期：

生成式 AI 工具從 2022 年末開始普及，並逐步擴大被應用到教育、研究領域。然而在採用 AI 輔助創作與發表時，有哪些是工具性的使用？有哪些是輔具性的使用？這牽涉到法律上著作權利地位的主張，也連帶影響到學術誠信的揭露責任。本文導引讀者建立著作權上基礎認知，進而梳理到實例的應用，其後探討 CC 授權條款與主流期刊組織，就 AI 輔助創作的標示要素，以促進讀者對 AI 技術使用的學術課責，有更紮實的理解和實踐。

您使用生成式 AI 輔助而產出的作品，是屬於 AI 的？您的？還是各分配一半？或者說，人類就創作投入程度的高或低，會不會影響相關的判斷？而若產出的作品確實能被著作權利保護，那在公開發表或研究投稿時，是不是應該將 AI 輔助的過程，向公眾或審核單位進行標示與披露？若不標示，會不會有學術倫理上涉及抄襲他人著作的課責爭議？但倘若會有相關課責問題，難道生成式 AI 就是不能用在研究輔助上嗎？

一、生成式 AI 產出作品的著作權適格

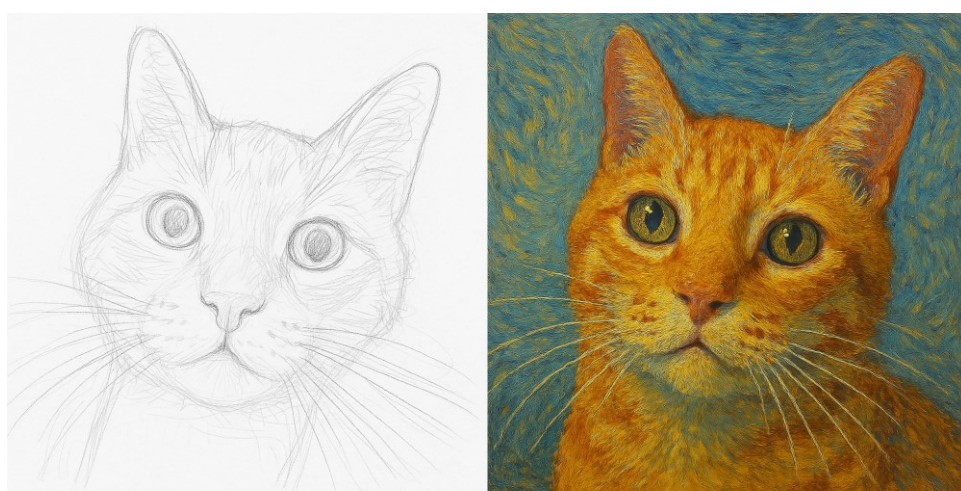
參酌美國著作權辦公室，就 AI 涉及著作權爭議所作的研究整理報告 (U.S. Copyright Office, 2024-2025)，其指出依據國際著作權法通則，以及各國司法判決的例示，任何作品要得到著作權利的保護，必須要有人為的創作性被參採其中，故純粹 AI 演算的成果，僅能視為亂數或純粹數理演算的輸出 (output)，此等事實性的輸出縱然外觀華美，尚不能直接取得著作保護。以我國為例，《著作權法》第 3 條亦明定「著作是指屬於文學、科學、藝術或其他學術範圍的創作。」也就

是說著作必須有創作性，而創作性的認定，雖然依循「美學不歧視 (minimal requirement of creativity) 」法則，毋須申論作品美醜等表達層次的「高或低」，但仍必須考究其參採人為創作性的「有或無」。

1. 人為創意導引在前的狀況

然同篇報告裡也明確指出 (U.S. Copyright Office, 2024-2025)，人類透過操作生成式 AI 工具所取得的成果，得否主張著作權，核心關鍵點取決於在創作過程中，是否有被注入人類創意，那麼首先操作者必須是人類，不能完全依賴人工智慧純自動化的行為；再者，需判斷該注入是否為創意的具體引導，並進一步評估操作者對輸出表達是否具實質控制能力。

如圖 1 所示，當人類依一般創作方式，採硬筆筆描或電腦繪圖的方式完成草稿後，此草稿即為受著作權保護之美術或圖形著作，其後再透過提示詞(prompt)輸入，委由生成式 AI 工具將其轉為梵谷風格的油畫圖，由於在素描圖的階段，已確立了人類創意的注入是存在的，後續雖然轉換繪畫風格時，委由 AI 工具自動生成，然整體觀之並不排除該成果的最後輸出有參採人類創意，並且亦是依循人類創意的導引，故最後油畫成品操作者，得以主張相關著作權利。



"Illustration of AI Colorization Following a Human Sketch" made in 2025 by Lucien C.H. Lin.
This image was generated through ChatGPT 4o under the direction and creative concept of the author, and is dedicated to the public domain under the CC0-1.0 license.

圖 1：人為創意導引在前之後指示生成式 AI 填圖之流程示意 (此示意圖左幅由本文作者手繪後掃描為電子檔，再透過 ChatGPT 4o 進行生成式 AI 轉梵谷繪風輸出產出右幅。)

2. 人為創意修潤在後的狀況

再依圖 2 所示，AI 生成工具的操作者，透過指示輸入的方式，要求 AI 工具產出貓掌的漫畫風格圖片，然而產出的圖片顯示出 7 根手指，與一般常情不符，其後繪師再透過人工手繪或繪圖筆修的方式，將 7 根手指調整回 5 根手指。一般而言，左圖屬於純輸入詠唱的 AI 輸出成果，預設不具著作權保護資格，因操作者對成果表達通常缺乏實質掌控地位，然右圖在原輸出的基礎上，若能明確加以人工修潤，此時修潤階段人類創意的注入確實存在，實務上經此修潤流程，最後修潤成品的操作者，多得以主張其相關著作權利。



"Illustration of Human Refinement and Adaptation of an AI Generated Work" made in 2025 by Lucien C.H. Lin.
This image was generated through ChatGPT 4o under the direction and creative concept of the author, and is dedicated to the public domain under the CCB-1.0 dedication.

圖 2：取用 AI 生成素材其後修潤注入人為創意之流程示意（此示意圖左幅由本文作者指示 ChatGPT 4o 進行生成式 AI 輸出，其後再透過繪圖軟體手動修正相關內容而產出右幅。）

二、生成式 AI 輔助創作的爭議風險與正確認知

承前所述，AI 協力產出的成果得否成立著作權利，關鍵要素在於人類創意能否被證實存在。至於產出後著作權利應如何分配，在多人創作的狀況下，著作權法的通則是依各人參與創作的程度來進行分配，然而 AI 工具畢竟不是人，所以除非 AI 工具的提供者另訂額外契約，不然一般實務都是交由 AI 工具的操作者，決定其是否就相關輸出成果主張著作權。然而，不論操作者是否主張著作權利，以下二則關鍵影響要素，應該要被留意並適當處理。

1. 生成式 AI 工具的使用風險

生成式 AI 工具的使用，可能涉及隱私、個人資料、業務機密等等面向的使用風險，例如歐洲的學術出版公司 Elsevier 即透過政策專頁 (Elsevier, 2024a)，規範文章評鑑的審查者，不得使用生成式 AI 工具來輔助審查，我國中央政府 2023 年發布的《行政院及所屬機關(構)使用生成式 AI 參考指引》，以及臺北市政府 2024 年發布的《臺北市政府使用人工智慧作業指引》，亦作了相類的行政指導。主要是，涉及雲端式外網方案的使用時，使用者將無法查察 AI 服務收錄了哪些隱私或機密資訊，然除此之外，生成式 AI 工具當前最為爭議的風險之處，在於對他人著作未經授權進行了重製、改作，或簡稱為 AI 抄襲的爭議。

究其要理是著作權的核心保護標的，在於著作獨特的創作表達形式，而目前市面主流的生成式 AI 工具，在訓練過程，皆大量使用他人著作，尤其是涉及影音互動的多模態 (Multimodal) 模型更是如此。AI 服務公司的立場是將他人作品的內容，視為事實資訊之邏輯梳理，並就這些內容的表達邏輯或事實性概念來進行學習，進而依著作權法主張利用標的僅為著作背後的概念、思想或發現，而不及於原作的表達，然後聲明相關 AI 的訓練與應用，不受到著作權相關限制。例如我國《著作權法》第 10-1 條指示：「依本法取得之著作權，其保護僅及於該著作之表達，而不及於其所表達之思想、程序、製程、系統、操作方法、概念、原理、發現。」

然而這樣的主張涉及灰色地帶，畢竟何謂著作的表達及其背後的概念，常有區分界限上的模糊與困難，以發展趨勢來觀察，現階段全球各地，皆開始肇生生成式 AI 訓練相關的司法爭訟 (ChatGPTiseatingtheworld.com, 2024)。主要是相關著作的原始權利人，與生成式 AI 訓練公司間在司法審判上的拉扯。所以說，利用生成式 AI 工具在創作或研究時，操作者仍應避免生成與他人原有作品，在創作表達與表達特徵上過於近似的 AI 輸出物，從而降低涉入侵權爭議的風險。

2. 工具性利用與輔具性分工的差異與影響

另一個影響要素，則是將 AI 工具與人的協作與互動，釐清為工具性的利用，或是輔具性的分工。所謂工具性的利用，意指生成式 AI 工具，純然在創作過程，僅單純為人操作指揮的工具，這種狀況下，工具的操作者獨得相關產出的所有創意注入，但倘若這般機械性、工具性的利用過程，操作者無法成功主張其人類的創意注入，或甚至操作者無法預見 AI 輸出成果的表達形式的話，舉美國 *Thaler v. Perlmutter* 一案為例 (*Thaler v. Perlmutter*, 2023)，承審法官直接判定無人類創意注入的 AI 輸出成果，無法得到任何著作權利的保護。所以工具性利用 AI 的結果，就著作權受保護與否的立場，大致是全有或全無，若有人類創意注入，則創意的注入者，即為該成果的作者，而若無人類創意的注入，則相關輸出成果即使外觀如何華美，依當前著作權法亦難以主張受到保護。

不過在輔具性分工的立場上，情況稍為複雜，此處輔具性分工，指的是產出作品的過程，確有人類創意的注入，但該創意注入卻也帶有生成式 AI 工具的分工，亦即 AI 工具參與相關創意的表達進行調整或修改。就現行著作權法的立場，只要事前或事後有人類創意的導引或修潤，生成式 AI 工具的輔助成果，亦得受到著作權利保護。然而所謂輔具性的分工，應該要被進一步辨識並釐清的主要緣由，在於此類模式，有其 AI 自動化或半自動化沉潛運作的事實，也就是說，操作者對於生成式 AI 工具調整或增補其著作表達，處於有所預測、但又無法完全掌握的狀態，這樣的狀態在公開發表、文章投稿，甚至學術研究等情境，就有必須被適當揭露的責任基礎。因為涉及公開行為時，將連結到溯源課責的要求，例如文章經過生成式 AI 輔具性分工的調修後，或會被嵌入非真實的編撰資訊，那該等文章的發表，應該將 AI 工具修潤的狀態予以具體披露，以提示讀者留意資訊呈現上的真實性，另外就文章投稿、學術研究，更是必須建立在公平公開的競賽基礎上，若任一造使用了生成式 AI 輔具分工，卻未揭露，而讓成果視為全然原創，最後導致評比機制失衡，並損害其信賴基礎。

三、生成式 AI 輔助下的標示義務

前述生成式 AI 工具應用上的二則關鍵影響要素，包括使用風險的控管，以及工具輔具的區隔，都可以透過善盡標示義務來進行相對應的處理。因此，多數主流生成式 AI 平台，皆於其授權條款或使用政策，規範使用者於後續發布或傳輸 AI 輸出成果時，應明確標示相關成果為 AI 生成，並且要求揭露工具或平台的名稱。

1. 適當標示將得以緩解著作侵權爭議

寫作風格、繪畫風格、程式碼撰寫風格，歷來不受著作權的直接保護。然而，創作風格 (style) 雖然不受保護，應用 AI 工具在作品風格的轉換和套用上，仍應考究被轉換原著作的權利保護狀態。舉 2025 年 3 月 OpenAI ChatGPT 造成的「吉卜力風格」熱潮為例，如下圖所示，操作者取用自行拍攝的照片進行吉卜力風格轉檔，原則上 ChatGPT 皆可完成圖像的風格轉換，然若是取用其他著名的動漫角色或名人肖像，則 ChatGPT 多會依其內部政策規範 (OpenAI, 2025)，拒絕生成。因此，若操作者轉換的是自行拍攝或繪製的圖像，則侵權風險相對較低。然而，若使用未經授權的他人作品，即便風格本身不受保護，AI 生成的輸出結果仍可能與原作構成實質近似 (substantial similarity)，從而引發侵權爭議。



"Consideration in AI-Based Transformation of preexisting Works", made in 2025 by the prompt author Lucien C.H. Lin. The image on the left was taken by the prompt author, and the image on the right was generated by ChatGPT-4o in the style of Studio Ghibli, based on instructions provided by the prompt author as well, and is likewise contributed to the public domain under CC0-1.0.

圖 3：AI 轉換既成著作之具體考量示意圖 (此示意圖左幅由本文作者拍攝，再透過 ChatGPT 4o 進行生成式 AI 轉吉卜力繪風輸出產出右幅。)

若生成式 AI 應用之結果，確實與他人原作呈現實質近似，依著作權法，仍有機會援引內設的「合理使用 (fair use) 」機制，作為抗辯。依我國《著作權法》第 65 條之規定，著作之合理使用，應審酌一切情狀，包括利用之目的及性質，被利用著作本身之性質，所利用的質量比例，以及利用結果對原著作潛在市場與現在價值的影響。然常被人忽略的是，《著作權法》第 64 條亦要求，當主張合理使用时，除非原作者不明、或不具名，不然仍應依合理方式明示該著作的出處，包括原著作人之姓名或名稱。進一步說，並不是說有註引出處，就必然能夠成功主張合理使用，但若能註引出處，則依《著作權法》第 65 條所訂應予審酌之各項評判，在個案上當更易建立清楚的分析與抗辯立場，導引承審法官認同原著作價值未被負面影響，來免除侵權責任。

2. 以 Elsevier 投稿政策與 CC 授權為引說明符合倫理課責的標示要素

由於生成式 AI 工具使用上的標示方式，嚴格來說目前尚未有單一標準，多數期刊是建議在文章的研究方法、致謝，或相應的聲明區塊，就所使用的生成式 AI 工具、使用方式，來進行說明與揭露。本文即以 Elsevier 的投稿政策，與 CC 授權領域兩個指標性組織，進行分析與經驗分享。

首先，直接使用生成式 AI 工具來輔助寫作，基本上並不受到 Elsevier 投稿政策的允許 (Elsevier, 2024b)。Elsevier 當前立場是學術求真求實，相關寫作經過生成式 AI 的輔具性分工，往往增生作者原所未有的創意，而混淆了研究資料在原創性和真實性的呈現。然而，Elsevier 並沒有完全禁絕生成式 AI 工具的使用，其對生成式 AI 工具的規範和限制，僅適用於寫作過程，而不及於研究過程，也就是作者可利用 AI 工具來就資料進行分析或要點整理，然若要用在寫作上，必須符合額外的例外條件。進一步說，Elsevier 的投稿政策，並不建議將生成式 AI 工具用於文章裡的圖表、影像、或相關藝術作品的美化，然而當生成式 AI 工具的使用，本就是研究設計或研究方法要探討和分析之標的，或僅是就文章可讀性或語言進行潤飾，便可被例外使用，例如本文的部份示意圖便為此種類型的應用。而當利用

到生成式 AI 工具時，Elsevier 要求文章投稿與發表時，必須達到右列三項要點的揭露與備存：(1)生成式 AI 或 AI 輔助工具的使用方式，(2)AI 模型或工具的名稱及製造商，以及(3)相關標的未經 AI 調整前的原始版本，以在需要查對比較原創性時，得據以查察。

另外一個在標示實務得被學習、師法對象，是 CC 組織(Creative Commons)。這是因為在 CC 授權的領域，本推廣多元共創，透過「以善意換取善意」理念的實踐，導引全球各美術館、圖書館、檔案館、博物館，以及學術期刊的內容，得採 CC 授權條款加速流通。故在過去二十年，已累積不少取材後相關改作或集合作品如何被明確標示尊重的實務經驗，並編撰於「建議的著作權標註實務」專頁(Creative Commons, 2024)裡。基本上 CC 的標註實務強調「TASL 四要素」，包括：(1) T = Title (作品名稱)、(2) A = Author (作者/創作者)、(3) S = Source (作品來源)，以及(4) L = License (授權條款)。這些要素可以採彈性方式呈現，只要相關資訊完整且被具體述明即可，另外，CC 的標註實務特別強調，原作若已被修改或調整，就應於後續發布時，明確標示該被調整的狀態，並簡要述明採何種方式調整。

綜合前述的著作權法律邏輯，以及 Elsevier 與 CC 全球組織的建議要點，當研究或寫作涉及生成式 AI 輔具性的分工與利用時，下列的標示模式，或為現時得具參考價值的標示作法：

[物件名稱]，[操作者真名、筆名、或可被辨識之別名]，[使用的生成式 AI 方案]，[AI 協力的標註]，[物件的授權狀態]，[生成日期]，以及[作品來源或生成式 AI 工具的使用方式]。

依此原則以下列圖 4 為示範例，[物件名稱]為「測試 ChatGPT 4o 對中英字顯示的正確性與可調整性」，[使用的生成式 AI 方案]為「**ChatGPT 4o**」，[AI 協力的標註]為其打上「**AI 生成圖片**」，[物件的授權狀態]若該 AI 服務並未有干涉契約，且被轉檔作品的著作權利已由操作者所有，或操作者已被充份授權，則可由操作者決定，此處範例採「CC0-1.0 公眾領域貢獻宣告釋出」，若生成時間具備參考價值，則可註明生成日期「2025.4.8」，最後「**作品來源或生成式 AI 工具的使用方式**」，指所有與該照片生成或研究理絡有關的描述，皆可在這個區塊進行補充與表達。



"ChatGPT-4o testing for the accuracy and adjustability of text rendering" by the prompt author/operator: Lucien C.H. Lin, AI-generated image, provided under CC0-1.0, Date of operation: April 8, 2025.
Two photos of the operator were used as input, with the prompt: "Transform into Ghibli style, and pay attention to the correct presentation of language text in the original image. For the image with the silver-gray cat, please mirror the image left to right, but do not mirror the language or numerical characters visible in the original photo."

圖 4：「測試 **ChatGPT 4o** 對中英字顯示的正確性與可調整性」，操作者：林誠夏(Lucien C.H. Lin)，使用的生成式 AI 方案：ChatGPT 4o (OpenAI)，**AI 生成圖片**，採 CC0-1.0 發布，操作日期：2025.4.8，使用 2 張操作者自拍照片，輸入詠語「轉換為吉卜力風格，並留意原圖裡中英文字的正確呈現，銀灰毛色貓那張圖，請轉換時將影像的左右反轉，然原照片顯示之中英數字，不被反轉。」

四、 結語

總結本文，就尊重原創和知識溯源的立場來看，學術引用的核心爭議點不在於引用他人著作，而往往是引用了沒有講，讓人誤會引用者即為原創，而在生成式 AI 工具的應用上也相仿，重點不是使用了 AI 工具，而在於使用 AI 工具卻去遮蔽掩飾，其後肇生了更為難的侵權爭議和學術課責的問題。故若能善盡相關標示義務，當得相當程度降減產生爭議與磨擦的風險。

進一步說，相關的標示要素，也並非要窮盡相關資訊，課予利用人過高的負荷。在 CC 授權的領域裡，姓名標示一向主張的是「合理方式之應用 (in any reasonable manner) 」，所以若因版面、載體或篇幅所限，亦可採前述粗體項目重點列示，佐以隨頁註或說明專頁集中補充的方式來處理。也就是說，涉及生成式 AI 輔具式的分工情境時，操作者若能具體述明他使用 AI 工具的方式和歷程，於相關成果公開發表或公開傳輸時，善盡標示義務，當能有效的降減爭議，並合理適份的將 AI 工具，用於研究輔助或作品發表上。

五、參考文獻

ChatGPTiseatingtheworld.com. (2024, August 27). *Master list of lawsuits v. AI, ChatGPT, OpenAI, Microsoft, Meta, Midjourney & other AI cos.*
<https://chatgptiseatingtheworld.com/2024/08/27/master-list-of-lawsuits-v-ai-chatgpt-openai-microsoft-meta-midjourney-other-ai-cos/>

Creative Commons. (2024, December 4). *Recommended practices for attribution.*
<https://creativecommons.org/licenses/by/4.0/#:~:text=Recommended%20practices%20for%20attribution>

Elsevier. (2024a). *The use of generative AI and AI-assisted technologies in the review process for Elsevier.* <https://www.elsevier.com/about/policies-and-standards/the-use-of-generative-ai-and-ai-assisted-technologies-in-the-review-process>

Elsevier. (2024b). *The use of generative AI and AI-assisted technologies in writing for Elsevier.* <https://www.elsevier.com/about/policies-and-standards/the-use-of-generative-ai-and-ai-assisted-technologies-in-writing-for-elsevier>

OpenAI. (2025, January 29). *Usage policies.* <https://openai.com/policies/usage-policies/>

Thaler v. Perlmutter, No. 1:22-cv-01564 (D.D.C. 2023).

<https://www.copyright.gov/ai/docs/district-court-decision-affirming-refusal-of-registration.pdf>

U.S. Copyright Office. (2024–2025). *Copyright and Artificial Intelligence Report*.

<https://www.copyright.gov/policy/artificial-intelligence/>

備註：本文所示圖片，已於 2025 年 4 月 14 日由教育部臺灣學術倫理教育資源中心、國立陽明交通大學、社團法人臺灣開放式課程暨教育聯盟，以及臺灣學術倫理教育學會共同主辦之「生成式 AI 的開源應用與倫理標示」講座中公開發表。

致謝

感謝周倩教授（國立陽明交通大學副校長）針對本文給予寶貴的修正建議。

本文作者：林誠夏 / CC Taiwan 計畫主持人、鈞理知識產權事務所法制顧問

文字編修：詹韻蓉計畫助理管理師 / 國立陽明交通大學人文與社會科學研究中心

（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

110 年 1 月至 114 年 4 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 110 年 1 月至 114 年 4 月）

單位：案件數

檢舉方式	具名	105
	未具真實姓名或聯絡方式	15
	職權發現	52
受理結果	不成案	46
	無違反學倫	53
	審查中	35
	有違反學倫	38
合計		172

備註：

1. 統計期間為 110/1/1~114/4/28。
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 110 年 1 月至 114 年 4 月)

單位：人次

違反之行為態樣	造假	1
	變造	5
	抄襲	11
	自我抄襲 (含隱匿及未適當引註)	4
	重複發表	2
	代寫	2
	影響論文審查	0
	其他	22
	合計	47

備註：

- 統計期間為 110/1/1~114/4/28。
- 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 110 年 1 月至 114 年 4 月)

單位：人次

處分情形	書面告誡	12
	停權 1-2 年	23
	停權 3-10 年以上	1
	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 110/1/1~114/4/28。
- 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。