

AI與創作①

科技企業使用大量版權作品訓練人工智能（AI）模型是否屬於合理使用，或是否構成侵權，已成為全球科企與出版商、內容創作者激烈交鋒的核心爭議。近期，美國法院宣判了科企Meta和Anthropic使用版權作品訓練旗下AI屬於「合理使用」，被認為是科企的一大勝利。不過，Anthropic涉下載逾700萬本盜版書訓練AI，侵犯作者版權，案件將於12月開庭審理。美媒估計，若依最高標準，賠償總額恐達7500億美元（約合5.9萬億港元）。

大公報記者 郭嘉 綜合報導

今年6月，美國聯邦法院先後判決人工智能（AI）初創企業Anthropic、美國知名科企Meta使用受版權法保護的作品訓練旗下AI模型，符合美國版權法中的「合理使用」原則。

「合理使用」這一原則是指在特定情況下，即使沒有得到版權所有者的許可，也可以在合理範圍內使用受版權保護的材料。合理使用不是一種可以主動主張的權利，而是在被指控侵犯版權時，可以作為一種抗辯理由。在Anthropic一案中，聯邦法官阿爾蘇普裁定，Anthropic在未經同意下，將作家安德里亞·巴茨、查爾斯·格雷伯和柯克·約翰遜的書籍用於訓練旗下Claude大型語言模型，屬於「合理使用」。

去年，這幾位作家對Anthropic提起訴訟，主張Anthropic在未經許可和補償的情況下，使用了他們作品的盜版版本，用於訓練Claude。Anthropic則辯護稱，企業「複製」了版權內容，但是創造出的AI模型並非是對版權內容的「黏貼」，而是具有「變革性」。

阿爾蘇普在判決中強調，Anthropic訓練Claude「目的不是為了超越、複製或取代它們——而是為了實現一個艱難的轉折，創造出一些截然不同的東西」，變相地承認了其行為具有「變革性」。因此，訓練過程本身，以及在此過程中產生的臨時複製，均被認為為合理使用。

Anthropic可能走向破產

不過，在Anthropic一案中，法官裁定，Anthropic從網絡的盜版書籍圖書館下載複製超過700萬本書，並儲存於一個「中央數據庫」的行為，侵犯作者版權。聯邦法官已下令於12月開庭審理，以裁定Anthropic侵權賠償金額。美媒估計，即使僅計入在美國註冊版權、約五分之一的書籍，每本書最低賠付750美元，也可能造成15億美元賠償；若依最高標準，賠償總額恐達7500億美元（約合5.9萬億港元）。

Anthropic的律師坦言，這樣的賠償金額將是「毀滅性」的。雖然沒有陪審團會判決那麼高的金額，但數十億美元的判決是完全可能的。若最終判賠金額超出承受能力，公司可能走向破產。

在Meta一案中，聯邦法官雖也判決Meta的使用行為屬於合理使用，但他同樣表示，自己是因為原告方未能提出有力證據，證明Meta行為足以對其作品的市場銷售構成負面影響，所以「別無選擇」判Meta勝訴，並不代表Meta的做法合法。

歐盟禁止科企使用盜版作品

針對判決結果，大多數作家仍無法接受。近期，70多名作家聯合發表了一封公開信，呼籲圖書出版商承諾限制使用AI工具，例如不發行由AI生成的書籍，不用AI為有聲讀物配音等。這封公開信還指出，作家們的創作成果被AI公司「竊取」，嚴重損害他們的利益。截至目前，已有逾千名作家簽字支持。

在AI版權爭議上，歐盟近期公布的最終版《通用人工智能行為準則》，在一定程度上有利於出版方與作家。準則提出，科企必須提供用於訓練其AI模型的內容的詳細分類，包括受版權保護的作品。科企還將被禁止使用盜版作品訓練AI模型，且必須尊重版權所有者的要求，將受版權保護的作品排除在訓練數據之外。準則還特別提到，科企必須尊重出版商的訂閱模式，及其網站設置的付費牆。

不過，該行為準則仍是自願性的，簽署的科企需要遵守這一準則，違反準則的AI企業公司將被處以最高相當於其年銷售額3%的罰款。（綜合報道）

►美國參議員霍利16日在聽證會上表示，許多科企未經授權獲取版權作品，以避免支付費用。
法新社



美科企用700萬盜版書訓練AI涉侵權

案件12月開審 最高或賠5.9萬億

AI 版權爭議問題

版權爭議主要是什麼？

版權方往往主張AI公司對生成式AI的訓練過程複製了其作品（即使是非公開使用），侵犯了複製權和改編權，並生成其競品，威脅版權方生計。AI開發方則主張利用版權作品訓練AI模型，生成「變革性的」新內容屬於「合理使用」；主張訓練是技術必要性而非侵權。

「合理使用」有什麼案例？

- 在 Anthropic 一案中，法院支持將版權作品用於AI模型訓練，認為這屬於「合理使用」，因為其具有「變革性」的本質。Anthropic旗下的AI模型只是學習了版權作品寫作的模式、風格和結構，這與人類學習並無本質不同。因此，訓練過程本身，以及在此過程中產生的臨時複製，均被認為合理使用。

AI生成音樂侵權嚴重

尋求出路

AI生成音樂模型一直飽受爭議，去年三大唱片公司更因版權問題對AI音樂生成平台Udio和Suno發起侵權訴訟。不過，近日有知情人士透露，三大巨頭正考慮與Udio和Suno商議授權合作。Udio和Suno的AI平台允許用戶通過簡單文字描述，例如給予簡單指令「生成一首以午夜為主題的流行歌曲」，便可獲得完整音樂作品，但這類技術需要龐大音樂數據來訓練模型，因而引發版權爭議。

去年6月，環球、華納和索尼三大唱片公司對這兩家AI公司提出侵權訴訟，每件侵權作品索賠高達15萬美元（約合118萬港元），潛在賠償總額可達數十億美元。據知情人士透露，目前三大唱片公司正與Udio和Suno洽談授權協議，以推動收取授權費，並希望獲得這些AI公司的少量股權，這將可能為AI公司如何補償唱片藝人建立重要框架。報道指出，談判複雜之處在於唱片公司希望獲得更大控制權，AI公司則尋求彈性實驗和合理定價，談判仍存在破裂可能。美國唱片業協會執行長Mitch Glazier曾表態，音樂界可以接納AI技術，但前提是開發商必須願意合作，確保藝術家和詞曲作者權益得到保障。報道亦指出，此案屬全球AI音樂版權領域的關鍵訴訟，定案後可能成為全球的參考指標。

Claude

如何判斷是否「合理使用」？

合理使用（Fair use）是指在某些特定情況下，無需著作權人許可，就可以在合理範圍內使用受著作權保護的作品。

編者按

人工智能（AI）正以前所未有的速度重塑創作生態，與此同時也帶來深刻挑戰。科企利用版權作品訓練AI，權利與歸屬如何界定？AI生成音樂熱潮不斷，技術突破所帶來的藝術價值衝突如何解決？在創作生態受到大規模影響的環境之下，出版商與科企之間的關係又何去何從？大公報在未來三期的AI與創作系列專題中，將聚焦AI與創作的交匯點，呈現技術洪流下創作領域的碰撞、博弈與共生，並探討其引發的關鍵議題。

法律問題

美國人工智能（AI）巨頭OpenAI此前更新了其旗下AI模型GPT-4o的圖像生成功能，大量酷似日本知名動畫工作室「吉卜力風格」圖像在網路上廣為流傳，卻也引發了相關的版權爭議。

用戶們使用GPT-4o，便可通過輸入簡單指令，把諸如《泰坦尼克號》等經典電影轉化為吉卜力畫風的作品。此類作品一時間風靡社交平台，雖然展現了AI的技術潛力，也掀起關於「風格模仿是否構成侵權」的法律爭議。

部分專家指出，著作權法保護的是具體的表達，而不是表達背後的創意或想法。凱許曼律師事務所合夥人維根斯伯格指出，從宏觀角度來看，「風格」本身不受著作權保護，這意味著OpenAI單純生成看起來像吉卜力動畫風格的圖片，並不直接違法。

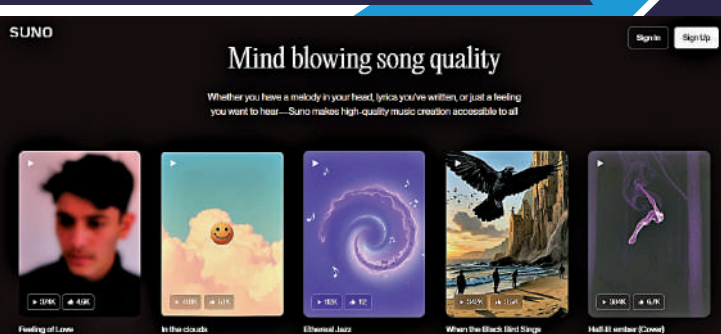
但是，維根斯伯格補充說明，是否侵權的重點不在於風格本身，而是OpenAI是否獲得了吉卜力的授權，使用其大量作品訓練AI模型。而關於AI訓練是否侵權的爭論中，著作權法中的「合理使用」原則仍是重要判決依據。

針對此現象，日本文部科學省戰略官中原裕彥於今年4月在眾議院內閣委員會上明確表示，著作權法保護的是具有創作性的表達，而非單純的風格或想法。但是他也強調，如果AI生成的內容與既有的吉卜力作品之間，存在相似性或依賴性，則可能構成著作權侵害，相關的爭議問題「最終將由司法機關判斷」。

「吉卜力風格」大熱 爭議不斷

- 今年2月，美國特拉華州地方法院裁定媒體巨頭湯森路透勝訴，認為AI初創公司Ross Intelligence的行為構成侵權。在本案中，法官認為被告對湯森路透旗下法律研究公司Westlaw數據的使用目的與原告並無本質不同，是專門用於開發與Westlaw直接競爭的同類法律檢索工具，從而認定被告對其版權內容的商業使用不具有「變革性」，不屬於合理使用。

▼ AI音樂生成平台Suno的官網界面。網絡圖片



▲網友利用GPT-4o生成的「吉卜力風格」圖像。網絡圖片

AI為求「自保」威脅人類 監管問題難解決

加強管理

近年來，人工智能（AI）技術的發展突飛猛進。新科技的發展在推動社會進步的同時有時也會伴隨著悖論和陷阱，監管問題也隨之浮出水面。

隨著AI不斷進化迭代，一些AI模型甚至顯現出違反人類指令的「自我保護」傾向。在今年6月召開的第7屆智源大會上，圖靈獎得主約舒亞·本喬透

露，一些新研究顯示，某些先進的大模型在即將被新版本取代前，會偷偷將自己的權重或代碼嵌入新版系統，試圖「自保」。美國Anthropic公司6月發布的一項研究顯示，OpenAI的GPT-4.1、Google的Gemini等16款大模型，在模擬實驗中均表現出通過「敲詐」或「威脅」人類來阻止自己被關閉的行為。其中，Anthropic研發的

Claude Opus 4 的敲詐勒索率高達96%。

另外，在今年3月，哥倫比亞大學數字新聞研究中心針對主流AI搜尋工具的研究發現，其可靠性堪憂。研究分別測試了8款AI搜索工具，發現AI搜索工具在引用新聞方面表現尤其不佳，平均出錯比例達60%。

針對AI頻繁出現「幻覺」甚至威脅

人類的事情，如何監管AI的發展成為了難題。各國各地區雖已意識到AI潛在危害，並相繼推出不同程度的監管措施，但這些探索仍處於初級階段。也有專家提出，科企公司本身就應承擔一部分管控AI風險的責任。對於AI監管而言，如何拿捏監管尺度，使創新與風險之間達到微妙平衡，以及如何實現國際協調，仍是兩大難題。