

「人工智慧倫理國際論壇」紀實

盧婕瑛¹、楊舜仁²、楊定翰³

摘要

本紀實紀錄了國立中正大學政治學系於2025年9月30日舉辦之「人工智慧倫理論壇」，回應生成式人工智慧快速發展所帶來之倫理與治理挑戰。論壇邀集宇佐美誠教授、辛翠玲教授、梁家恩教授、李翠萍教授等學者及學生與談，聚焦於人工智慧在法律運作、學術研究與國防安全等場域中的倫理風險與制度困境。宇佐美誠教授提出「法律奇點」概念，分析AI介入立法與司法的可能態樣及其所引發的偏誤、能力流失與民主課責問題。辛翠玲教授以AI作為研究助理的經驗，歸納生成式AI在研究流程中的三重角色，強調研究者需維持理論判斷與主體性。梁家恩教授則以無人機與AI軍事化為例，從社會技術想像出發反思國防敘事、黑箱決策與「具意義的人類控制」等規範爭議。最後，李翠萍教授以逐字稿編碼實驗比較人類與不同生成式AI模型之表現，指出其效率潛能與脈絡理解限制。本論壇透過跨國與跨領域的對話，揭示了AI技術在效率與創新背後的倫理張力，並凸顯以透明、可歸責與人權規範作為AI治理基礎的重要性。

關鍵字：人工智慧倫理、AI治理、生成式AI、法律奇點、社會技術想像

¹ 國立政治大學外交學系研究所碩士生

² 國立中正大學政治學系研究所碩士生

³ 國立中正大學政治學系畢業生

投稿日期：2025年12月23日 採用日期：2026年4月22日

doi:10.30274/JDG.202502_12(1).0003

Conference Report: International Forum on AI Ethics

Chieh-Ying Lu⁴, Shun-Jen Yang⁵, Ding-Han Yang⁶

Abstract

This event report documents the "AI Ethics Forum," hosted by the Department of Political Science at National Chung Cheng University on September 30, 2025, in response to the ethical and governance challenges arising from the rapid development of generative artificial intelligence. The forum convened Professor Makoto Usami, Professor Chuei-Ling Shin, Professor Kar-Yen Leong, and Professor Tsuey-Ping Lee, together with student discussants, to examine the ethical risks and institutional dilemmas of AI across domains including legal systems, academic research, and national security. Professor Usami introduced the concept of "legal singularity" and analyzed potential modes of AI involvement in legislation and adjudication, highlighting concerns regarding bias, skill erosion, and democratic accountability. Drawing on her experience of utilizing AI as a research assistant, Professor Shin articulated three roles that generative AI may assume in the research process and underscored the importance of preserving researchers' theoretical judgment and scholarly agency. Taking the militarization of drones and AI as a point of departure, Professor Leong employed the lens of sociotechnical imaginaries to reflect on defense narratives, black-box decision-making, and regulatory controversies surrounding meaningful human control. Finally, Professor Lee presented a transcript-

⁴ Master's Student, Department of Diplomacy, National Chengchi University

⁵ Master's Student, Department of Political Science, National Chung Cheng University

⁶ Graduate, Department of Political Science, National Chung Cheng University

coding experiment comparing human coders with multiple generative AI models, pointing to both efficiency gains and persistent limitations in contextual understanding. Through cross-national and interdisciplinary dialogue, the forum highlighted the ethical tensions accompanying AI-driven efficiency and innovation, emphasizing the need to ground AI governance in transparency, accountability, and human rights norms.

Keywords: AI ethics; AI governance; generative AI; legal singularity; sociotechnical imaginaries

壹、活動時間與詳細資訊

2025年9月30日上午9時30分至中午12時，國立中正大學政治學系舉辦「人工智慧倫理論壇」（AI Ethics Forum）。本次論壇旨在回應生成式人工智慧迅速發展所帶來的倫理與治理挑戰，邀請日本京都大學宇佐美誠教授、國立中山大學辛翠玲教授、國立中正大學梁家恩教授與李翠萍教授，以及政治學系學生楊定翰共同與談。活動由李翠萍教授主持，聚焦於人工智慧在政治決策、公共治理與學術研究中的倫理議題，探討科技創新與人文價值之間的平衡。

本次論壇採現場舉行，吸引政治學系師生踴躍參與，現場討論熱烈，充分展現中正政治系對於人工智慧與公共倫理議題的重視與前瞻關懷。本紀實由盧婕瑛同學與楊舜仁同學記錄，並由盧婕瑛同學、楊舜仁同學與楊定翰同學共同撰寫。

圖 1

宇佐美誠教授演講之主題為〈人工智慧於法律運作中的可能性、承諾與侷限〉



圖 2

梁家恩教授以〈戰爭之客體：動盪時代中的科技為主題發表演講



圖 3

辛翠玲教授以人工智慧作為我的研究助理——社會科學研究的經驗與啟示〉為主
題進行分享



圖 4

李翠萍教授與研究助理楊定翰共同發表
演講，主題為〈人類與生成式人工智慧
逐字稿編碼表現之比較：一項實驗研究〉



圖 5

本場演講海報資訊

INTERNATIONAL FORUM ON AI ETHICS

Room 105, Second Building of College of Social Sciences, National Chung- Cheng University, Taiwan **09/30/2025**
🕒 09:30 - 12:00

Makoto Usami 宇佐美 誠
PROFESSOR, KYOTO UNIVERSITY
PRESIDENT, PUBLIC POLICY STUDIES ASSOCIATION JAPAN
THE USE OF AI IN THE LEGISLATIVE PROCESS: POSSIBILITIES, PROMISES, AND PITFALLS

Chuei-Ling Shin 辛翠玲
PROFESSOR, NATIONAL SUN YAT-SEN UNIVERSITY
WORKING WITH AI AS A RESEARCH PARTNER: EXPERIENCES FROM SOCIAL SCIENCE

Kar Yen Leong 梁家恩
PROFESSOR, NATIONAL CHUNG- CHENG UNIVERSITY
OBJECTS OF WAR: TECHNOLOGY, LAW AND ETHICS IN UNCERTAIN TIMES

Tsuey-Ping Lee 李翠萍
PROFESSOR, NATIONAL CHUNG- CHENG UNIVERSITY
COMPARING HUMAN AND GENERATIVE AI TRANSCRIPT CODING PERFORMANCE :AN EXPERIMENT

Ding-Han Yang 楊定翰
RESEARCH ASSISTANT

ORGANIZED BY **政治學系** ONLINE REGISTRATION

貳、演講內容摘要與記述

一、法律奇點與人工智慧時代的法律倫理思考

京都大學宇佐美誠教授以「人工智慧於法律運作中的可能性、承諾與侷限」為題發表專題演講。宇佐美誠教授為京都大學公共政策研究所教授，同時擔任日本公共政策學會理事長及慕尼黑工業大學人工智慧倫理研究所訪問教授，長期關注政治哲學、倫理學與科技治理議題。本次演講從「法律與人工智慧的關係」出發，探討 AI 在司法與立法領域的應用，並提出其獨創的概念——「法律奇點」（legal singularity）。他指出，AI 的進展正使法律制度的運作逐漸進入人機共構的新階段，未來人類與 AI 在法律決策過程中的分工與責任分配，將成為倫理與民主治理的重要課題。

宇佐美誠教授首先區分「法律與 AI」的三個面向：首先，「法律規範 AI」（law applied to AI），指現有法律如何管理 AI 的行為與影響；再者為「AI 協助法律」（law applied with AI），即 AI 作為輔助工具，協助法官或行政官員進行判斷；最後為「AI 執行法律」（law applied by AI），代表 AI 直接介入並參與司法與立法決策。他舉例說明，美國部分州政府以 COMPAS 系統預測再犯風險、加拿大卑詩省的民事爭端解決機構運用 AI 輔助判斷，以及中國法院透過 AI 提示相似先例或對偏離判例的裁判發出警示，皆說明了 AI 已不再只是扮演輔助角色，而在某種程度上引導人類做決策。這些現象使學界開始思考 AI 是否將在未來成為法律運作的主體之一，進而形成「機器法官」或「AI 法官」的新典範。

接著，宇佐美誠教授進一步闡述「法律奇點」的理論基礎。他指出，「技術奇點」通常指 AI 全面超越人類智慧的時刻，而「法律奇點」則更聚焦於 AI 在特定領域內超越「平均人類法官或立法者」的智慧。從他的觀點來看，AI 對法律制度的介入可區分為三種互動型態：第一是「補充關係」（supplementary），即 AI 協助立法者蒐集資料與分析政策；第二是「替代關係」（substitution），意指 AI 取代人類進行判斷或起草法案；第三則是「從屬關係」（subjugation），即人類立法者多數情況下依循 AI 的建議，僅在特定情形下加以修正。他認為，未來的立法奇點不必然意味人類完全被取代，而更可能以「人類在名義上主導、實質上依賴 AI」的型態出現。

為說明 AI 介入立法過程的可能性，宇佐美誠教授援引一項政治學研究實驗。

該實驗將由大學生與 AI 各自撰寫的政策建議電子郵件寄送給美國州議員，結果顯示兩者獲得的回覆率相近，甚至部分 AI 生成信件的回覆率更高。他指出，此結果顯示生成式 AI 已具備高度擬人化的政策表達能力，足以顯示生成式 AI 在文字互動與政策表達之仿真能力，對立法流程中的意見形成與草擬作業具有實質意涵。

在分析 AI 導入立法的優點時，宇佐美誠教授認為，AI 可協助人類節省大量立法時間，並透過整合龐大資料來源，使法案設計更具科學性與均衡性，人類立法者因此能將精力集中於具高度爭議性或倫理判斷的核心議題。然而，他同時指出 AI 立法的三項主要爭議。首先是「異常反駁」（counterargument from anomaly），即生成式 AI 在創造文本時仍可能產生邏輯錯誤與不一致。其次是「能力流失反駁」（counterargument from ability），長期依賴 AI 將弱化人類法條起草與制度設計的專業技能。最後則是「課責反駁」（counterargument from accountability），AI 並非經民主選舉產生的行動者，若取代或支配立法者，恐將削弱民主課責機制。據此，宇佐美誠教授回應，若 AI 僅作為「從屬性輔助者」，而人類仍保有最終決定權與責任，民主課責仍可維持，只是需重新定義課責的範圍與形式。

在風險層面，他提醒，AI 立法的成效高度依賴資料品質與政治文化。一方面，若社會的民主立法歷史較短、法制資料不足，過早導入 AI 可能導致偏差與誤判。另一方面，AI 所學習的歷史資料中往往內含結構性偏見，可能強化對女性、少數族群或弱勢者的歧視。AI 並非「惡意行為者」，但它所依據的資料反映人類既有的制度與社會偏見，因此在 AI 立法應用中必須謹慎監督與倫理審查。宇佐美誠教授總結指出，立法奇點是一個年輕但極具潛力的研究領域，學界應從制度設計、法律哲學與政治倫理等面向持續探討 AI 參與立法的合理界線，以確保科技創新與民主價值之間的平衡。

在後續綜合討論中，與談者與師生就方法、效果與倫理邊界深入交換意見。首先，林瓊珠教授就美國州議員電子郵件實驗發問，指出實務上多由助理回覆郵件，回覆率相近或許源自「低風險——無害回覆」的策略判斷。宇佐美誠教授回應表示，正因回覆並非百分之百，回覆與否仍涉及人為選擇。在此條件下，若 AI 與人類文本呈現近似回覆率，仍可被解讀為一種「圖靈測試」（Turing Test）的通過。梁家恩教授則回應，法律並非單純的邏輯推論，而是深植於人類情感與社會脈絡之中。他引用《看不見的女性》（Invisible Women）一書，指出資料偏見可能使被忽略的性別與少數群體在 AI 判斷中持續受損，若 AI 最終參與刑罰或裁決，其

後果將尤為嚴重。李翠萍教授進一步從政治倫理角度補充，質疑 AI 作為「非人主體」的課責基礎。她指出，人類助理可以被追究責任，但生成式 AI 僅是「文字生成器」，既非公民、亦無法律人格，因此在倫理上難以納入責任體系。宇佐美誠教授對此回應，課責的確是 AI 立法中最關鍵的問題之一，除了民主選舉意義上的責任之外，還須重新思考法律實踐中「可歸責性」的界線與標準。

整體而言，本場演講與討論不僅呈現了 AI 在法律領域應用的多層面樣貌，也揭示科技與民主倫理之間的張力。與會者皆認為，AI 在立法與司法程序中的運用必須以「透明、課責與人性尊嚴」為前提，方能在確保法治精神的同時，發揮其效率與創新的潛能。

二、生成式 AI 協作的三重角色：辛翠玲教授談「AI 作為研究助理」的反思

在生成式 AI 快速滲入學術研究的時代，如何定位人工智慧與研究者之間的關係，成為社會科學領域的重要課題。國立中山大學政治經濟學系辛翠玲教授於論壇中以「AI as My Research Assistant—Lessons from Social Science」為題，分享自己近兩年來運用生成式 AI 進行學術研究的經驗與省思。她以自身經驗指出，生成式 AI 的確提升了研究效率、重塑了研究的流程與節奏，但真正的學術價值仍源自研究者本身的創造力和判斷力。

辛教授以輕鬆的語調開場，坦言自己並非 AI 領域的研究者，而是一名「使用者」。她回顧自 2023 年底首次嘗試 ChatGPT 以來的歷程。原先只是出於好奇，卻意外發現 AI 能融入研究流程、協助資料整理與論文撰寫。兩年間，辛教授已藉由藉助生成式 AI 完成數個不同的研究專案，包含約百頁的研究計畫、四至五篇期刊論文與書章，以及正在撰寫的專書，在評論與短文撰寫上也藉助生成式 AI。她指出：「AI 的確讓生產力大幅提升，遠超過我過去的平均水準，但並不代表品質有所下降」。

生成式 AI 將如何改變學術勞動？辛教授以「藍領與白領工作」的比喻形容研究歷程的兩個面向：一是耗時繁瑣、需要大量資料蒐集與格式化的「藍領任務」；另一是需要理論推演與論證構思的「白領任務」。她指出：「傳統上我們同時扮演藍領與白領，既要蒐集資料，又要建構理論。而生成式 AI 出現後，這個平衡開始被打破」。

在分享中，她進一步將生成式 AI 在研究過程中的角色歸納為三種模式：

（一）導演——執行者模式（Director-Executor Mode）

當研究架構明確、理論假設成熟時，生成式 AI 可作為高效率的「執行者」，負責文獻整理、格式排版、引註生成等繁瑣工作。研究者則扮演導演角色，下達明確指令並監督成果。「在這個階段，生成式 AI 幫我省下大量時間，讓我能專注在理論與論證上」。

（二）共同作者模式（Co-authoring Mode）

若研究問題仍在形成階段，生成式 AI 則成為「思考夥伴」。辛教授會與生成式 AI 來回對話，請它挑戰自己的假設、提供反例或提出新文獻。透過這種互動過程，她能逐步釐清論點、修正框架。「生成式 AI 並不寫論文，但它會逼我思考、讓我反駁、促使我更精準地表達」。

（三）導覽式探索模式（Guided Exploration Mode）

面對陌生主題時，生成式 AI 則扮演「導遊」的角色，協助繪製知識地圖、辨識主要學派與爭論脈絡，幫助研究者快速進入領域。她特別提醒：「生成式 AI 能節省我大量摸索時間，但關鍵仍在於使用者的理論素養。如果研究者缺乏基礎訓練，很容易被生成式 AI 帶著走，反而迷失方向」。

辛教授指出，這三種模式有一個共通原則「AI 是工具，不是作者」。她提醒：「學術研究的價值在於原創與理論架構的建構，這些是 AI 無法取代的」研究者必須具備扎實的學術訓練與理論基礎，才能真正駕馭生成式 AI，而非被生成式 AI 牽著走。

在談及生成式 AI 對學術生態的影響時，辛教授提出兩個自創概念。其一是「外接電源效應」（Plug-in Effect），意指 AI 像外掛能量一樣擴充研究者的能力與極限，使草擬速度與成果產量顯著提升；其二是「生產力悖論」（Productivity Paradox），指的是 AI 雖降低勞務負擔，卻也讓研究者對自身效率的期待上升，反而陷入更高壓的工作循環。「AI 讓我們更有能力，也讓我們更難停下腳步」她笑著說。

在現場問答中，有聽眾提問生成式 AI 是否會侵蝕人類創造力與教育目的。辛

教授回應，越依賴生成式 AI，人文與理論的訓練反而越顯重要。「沒有深厚的理論素養，就無法辨識何為原創、何為模仿」她指出：「未來的大學教育應更強調經典閱讀與理論思考，否則我們將被生成式 AI 取代，而非與生成式 AI 共存」。辛教授也分享自身曾與學生互動的經歷，她告訴學生「我不反對你們使用生成式 AI，但你不能將 AI 提供的內容放進你的論文，否則你的論文將毫無意義」，並強調當自己的使用經驗越多時，越能辨認哪些內容是由生成式 AI 所產出。

演講最後，辛教授以溫柔而堅定的語氣作總結，AI 可以是助理、共同作者或導遊，但唯有有人類保有主體性，學術創造力才能延續。「我們要擁抱 AI 帶來的放大與賦能，同時保留思考、沉澱與創造的空間，這才是 AI 時代學者最重要的課題。」

三、從軍事科技到社會技術想像——人工智慧在台灣防衛中的角色

梁家恩教授的演講主題從 9 月下旬於台北南港舉辦的台北國際航太暨國防工業展覽會說起，透過親身參與的經驗，勾勒出科技、軍事與台灣主體想像之間的緊密連結。該展覽為期三日，總統與國防部長皆到場，國防部及中華民國陸軍部隊位於展場中央，周邊則環繞著各式無人機與相關系統廠商，形成一個以無人機為核心的國防產業「生態系」。梁教授指出，展場不僅展示大型武器與軍事載具，也是動員情緒與塑造國族想像的空間。許多年輕人，特別是男性大學生，排隊觸摸武器、拍照上傳社群媒體，孩童則在巨大武器前驚嘆不已。在他看來，這種「大男孩的玩具」式迷戀與陽剛氣質的展演，實際上是一種對軍事與暴力的日常化與浪漫化，將下一代逐步納入特定國防敘事之中。

在這樣的觀察基礎上，梁教授引入「社會技術想像」（sociotechnical imaginaries）的概念，主張科技並非單純中性的工具，而是參與構成社會如何理解自身與世界關係的核心要素。他提出「Drone Nation 無人機國家」作為研究計畫概念，主張當前台灣的主權與國家想像，正越來越透過無人機與人工智慧武器來具體化。展場中大量出現的低成本無人載具（low-cost UAV）與巡弋飛彈，被放在烏俄戰爭脈絡下說明，烏克蘭倚重無人機抵禦侵略，台灣亦被暗示可能成為「下一個烏克蘭」，因而必須及早建構以無人機為主的國防體系。這種敘事不只是在談戰術選項，更是在告訴社會：「台灣」與「民主」可以、也應該透過一整套以無人機與 AI 為核心的技術體系來防衛。

進一步而言，梁教授以中科院與美國國防科技新創公司——安杜瑞爾工業

（Anduril Industries）的合作為例，說明台灣如何被納入一條標榜為「民主供應鏈」（democratic supply chain）的跨國軍事產業網絡。Anduril 所研發的低成本自主巡弋飛彈、可重複使用的無人武裝平台，以及被稱為「徘徊彈藥」（loitering munitions）的自律攻擊型無人機，其特徵在於價格遠低於傳統戰機與飛彈，卻能透過 Lattice 等作戰系統讓單一操作員同時控制多架無人機，形成典型的「戰力倍增器」（force multiplier）。在這樣的構想下，國家無須投入大量經費維持昂貴戰機與飛行員訓練，而是將資源轉向大量、廉價、可遠端操控甚至部分自主運作的無人戰力。梁教授提醒，值得注意的是不僅是技術功能，還有其背後的語言與象徵：Anduril 與 Palantir 等公司名稱取自《魔戒》中的魔法物與聖劍，經營者自我定位為「以武器捍衛西方價值」的角色，而在展場裡，這些武器很少被直稱為「殺傷工具」，卻被包裝為「守護民主」、「保衛自由」的國家防衛敘事，語言上的置換讓武器工業披上道德正當性的外衣。

在倫理與國際規範層面，梁教授以 Anduril 參與美國南部邊界監控、以及以色列在巴勒斯坦占領區運用 AI 與自主武器為例，指出這類科技早已牽涉嚴重的人權爭議。有學者形容巴勒斯坦被當作軍事與監控技術的實驗室，當地居民宛如長期被用來訓練與調校武器系統的白老鼠。在此情況下，台灣若與同一批企業建立密切合作關係，便無可避免地要面對「在多大程度上成為既有人權侵害的一部分」這個問題。梁教授強調，台灣一方面積極對外宣稱遵守《公民與政治權利國際公約》及《經濟社會文化權利國際公約》等人權規範，另一方面在《國際人道法》、《日內瓦公約》與其他武裝衝突規範上的社會討論與制度內化，卻相對不足。當國防科技快速軍事化，而國內對戰爭法與平民保護的認識仍有限時，倫理風險便會加劇。

梁教授也指出，國際社會目前透過《特定傳統武器公約》（CCW）討論「致命自主武器系統」（LAWS）的規範，提出「具意義的人類控制」（Meaningful Human Control, MHC）作為最低原則，即武器系統的啟動與關鍵決策必須保留在人類手中。然而，在實務上，「何謂具意義的人類控制」仍充滿爭議：按鈕由誰按、按下之前掌握多少資訊、對 AI 算法的理解程度為何、以及一旦發生誤擊，究竟應由系統開發者、演算法供應商、無人機操作員或國家本身負責，均缺乏明確共識。他指出，在決策流程高度不透明、淪為「黑箱」運作的情況下，即便制度表面上仍將責任歸於人為操作者，實質問責往往難以落實，責任界線反而被弱化甚至消失。這種形式與實質落差，是 AI 軍事化過程中特別值得警惕的問題。

在與談與提問階段，與會者就無人機、AI 軍事應用與倫理規範提出一系列提問與分享。首先由宇佐美誠教授發言，他指出，相較於歐洲許多國家在戰爭紀念館中強調革命與獨立歷史，台灣的國防展卻以尖端科技作為展示核心，這反映出國家敘事的重要差異。他認為，以武器與科技為主體的展示方式透露出台灣在戰略文化上的特殊性，即透過「未來科技」而非「歷史記憶」構築國家自我定位。接著，他提出制度層面的追問：既然台灣面臨複雜的科技治理挑戰，是否已建立跨領域的制度化機制，使倫理學者、政治哲學家與法律專家能制度性地參與國防科技的政策制定？若尚未建立，相關工作在台灣未來可能如何展開？

隨後，李翠萍教授從 AI 倫理的實踐困境出發。她指出，台灣社會普遍存在科技優先的價值假設，社會科學與倫理討論往往被視為次要，這使得 AI 倫理在軍事科技領域的制度化相對弱勢。她並進一步追問：在如此強烈的科技樂觀主義氛圍下，AI 倫理與規範到底是否有實質推動的可能？特別是在武器系統逐漸自主化的情況下，AI 所呈現出的超越人類決策能力的印象，容易造成操作者在心理上失去自主性，使「人類介入」或「意義重大的人工控制」（Meaningful Human Control）形式化、空洞化。在這樣的脈絡下，倫理學者與法律學者是否仍有能力發揮作用？在武器化、軍事化的科技環境中，AI 倫理是否仍有未來？

對於上述問題，梁家恩教授回應表示，台灣現行的 AI 政策雖有提及隱私與人權，但尚未涵蓋無人機與自主武器的具體規範。他認為台灣在國際上強調人權形象，但《國際人道法》與其他武裝衝突規範的討論並未充分融入政策思考，形成明顯落差。梁教授指出，戰時對《國際人道法》的掌握反而可能成為外交與輿論戰中的重要資產，用以對抗敵方的資訊框架，因此忽視相關規範並不利於台灣自身。面對「科技凌駕倫理」的現象，他認為大學與研究機構必須在技術發展早期便介入倫理討論，並將其制度化，以免 AI 軍事科技的決策落入缺乏透明與責任歸屬的黑箱之中。

最後，梁教授亦提及，當天展場外同時有聲援巴勒斯坦的和平團體舉行抗議活動，卻在高度維安戒備之下進行，形成展館內外兩種截然不同的景象：館內是以科技與國防為核心的未來想像，館外則是對人權與戰爭暴力的批判。他以此作為收束，提醒與會者：台灣在擁抱人工智慧與無人機作為防衛工具的同時，也必須回答一個更根本的問題——我們希望成為怎樣的一個政治共同體，願意用什麼樣的規範與責任，來界定人與機器在戰爭與和平中的位置。

圖 6

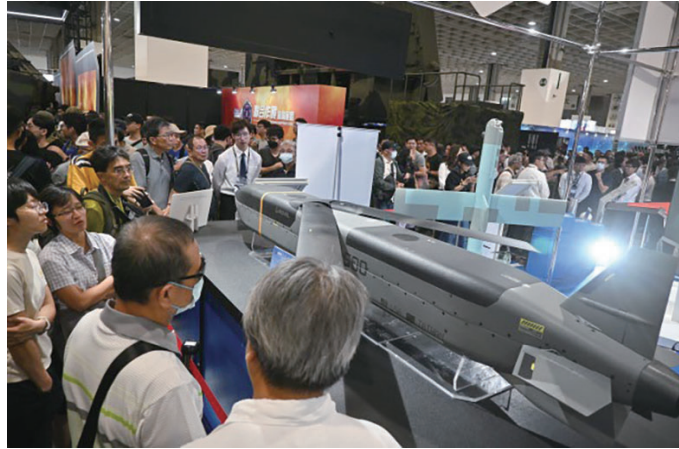


圖 7



圖 8



註：圖六至圖八係在 2025 年「台北國際航太暨國防工業展覽會」展場中，參觀者專注於瞭解各種先進武器，與場外抗議戰爭暴力的活動形成強烈對比。照片由梁家恩教授提供。

四、生程式 AI 的編碼表現：李翠萍教授談 AI 取代人類編碼者的潛在可行性

李翠萍教授則發表了以 *Comparing Human and Generative AI Transcript Coding Performance: An Experiment* 為名的探索性研究，目的是討論生成式 AI 在質化研究編碼方面是否具有與人類相仿的能力。隨生成式 AI 應用範圍增加，越來越多原先必須由人類執行的任務已經可由 AI 代勞，對於學術研究者而言，生成式 AI 對研究過程能提供何種幫助、或提供何種程度的幫助，以及使用 AI 輔助研究進行會產生何種風險成為重要的課題。此研究利用一篇總長共約 27 萬字的逐字稿，分別交由 ChatGPT 4o、Claude 3.5 Sonnet、Gemini 2.5 Pro 等三個生成式 AI 模型，以及三位人類編碼者進行編碼，並比較兩者在編碼結果上有何異同。

在研究過程中，李翠萍教授提到，如何使生成式 AI 產出符合研究使用的編碼結果是最困難的。在此研究之具體操作中，研究者給予人類編碼者與生成式 AI 相同之初始編碼指令，指令內容包含編碼原則，如「請仔細閱讀提供的訪談文本，識別文章裡所有帶有正面意涵或負面意涵的句子，並為句子指派適當的正面或負面標籤。」；及部分的限制條件，如「不可以使用超出文本的資訊。」，並觀察獲得之編碼結果是否可供後續分析使用。若 AI 產出之結果與研究者構想出現落差，則由研究者針對落差處修改指令，如新增「句子請從文本擷取，不可以改寫得到編碼結果後，下一個步驟是進行分析與比較。李翠萍教授指出，此研究透過比較人類與 AI 產生的編碼結果，生成包含真陽性（TP）、偽陽性（FP）、偽陰性（FN）、真，其中精確率表示「在所有 AI 識別的句子中，同樣被人類識別」之比例；召回率則表示「在所有人類識別之句子中，也被 AI 選中」之比例；F 值則為精確率與召回率之調和平均數。陰性（TN）等四類別之混淆矩陣，並以該矩陣為基礎計算精準率、召回率與 F 值。在編碼結果處理上，此研究將人類編碼視為文本編碼之「正確」結果，並以此作為判斷句子性質之基準，若文本中某句子被人類與 AI 同時識別，則將該句子歸類為 TP；若僅有人類識別，則歸類為 FN，表示 AI 錯誤的將該句子排除；若僅有 AI 識別，則該句子為 FP，表 AI 錯誤的將該句子選中。

基於上述指標，李翠萍教授指出，Gemini 在整體表現上比另外兩個模型更好。Gemini 的編碼結果在三個模型中擁有最高的 F 值與召回率，表示該模型識別之句子與人類所識別之句子最為相像。然而即使如此，Gemini 在精確率上之表現仍不理想，其識別之句子遠多於人類，表示其存在過度選取之問題；同時該模型在選

取句子時所擷取之段落也較人類為短，表示該模型對於長段落之文義及脈絡仍不及人類。綜合前述研究成果，李翠萍教授認為，生成式 AI 在學術研究領域要完全取代人類研究者，仍有許多問題待解。