

初探生成式 AI 風險感知的成因與第三人效果：以 ChatGPT 為例

曾懷寬、李芷容、吳泰毅

摘要

隨著 ChatGPT 等生成式 AI 愈來愈融入人類生活，人工智慧科技對於人類生活的負面影響也值得關注。本研究以第三人效果 (the third-person effect) 作為理論基礎，一方面探討民眾的 ChatGPT 使用經驗、AI 素養對生成式 AI 風險第三人感知的影響，另一方面也檢驗上述因素是否影響其支持政府規範 AI 科技，並將樣本群體劃分為「X世代」、「千禧世代」及「Z 世代」，剖析不同世代對生成式 AI 風險的認知與反應。研究以《2024 台灣網路報告》的數據為分析資料 ($N = 979$)，結果發現，三個世代的民眾普遍展現出第三人效果傾向，認為 AI 科技風險對他人的影響遠大於對個人的影響；不過，ChatGPT 使用經驗在三個世代中均無法顯著預測第三人感知，而 AI 素養也僅對 X 世代產生顯著負向影響。此外，ChatGPT 使用經驗對千禧世代與 X 世代的法律規範態度有顯著負向影響，而 AI 素養對 Z 世代法律規範態度有顯著負向影響，至於三個世代的第三人效果則無法顯著預測對 AI 科技法律規範的態度。研究結果有助於增進了解當前台灣民眾對於生成式 AI 的風險認知，並可作為政府制訂相關政策時參考。

- ◎ 關鍵字：AI 素養、AI 使用經驗、ChatGPT、生成式 AI、生成式 AI 科技風險感知、第三人效果、世代差異
- ◎ 本文第一作者曾懷寬為國立陽明交通大學應用藝術研究所博士候選人；第二作者李芷容為國立陽明交通大學傳播研究所碩士生；第三作者吳泰毅為國立陽明交通大學傳播研究所副教授。
- ◎ 通訊作者為曾懷寬，聯絡方式：Email：miazeng.hs07@nycu.edu.tw；通訊處：300 新竹市東區大學路1001號人社二館1樓。
- ◎ 收稿日期：2024/02/10 接受日期：2025/07/24

Exploring the Antecedents of Risk Perception toward Generative AI and Its Third-Person Effect: A Case Study of ChatGPT

Huai-Kuan Zeng, Zhi-Rong Li, Tai-Yee Wu

Abstract

With the increasing integration of generative AI, such as ChatGPT, into daily life, the potential negative influence of this technology on human lives warrants attention. The current study employs the third-person effect (TPE) as a theoretical framework to investigate the impact of individuals' ChatGPT use experience and AI literacy on their third-person perception of generative AI risks. Furthermore, it investigates whether these factors affect the support for government regulation of AI technology. The sample was divided into Generation X, Millennials, and Generation Z to explore the variations in AI risk perceptions and responses across these generations. This study analyzed data from the 2024 Taiwan Internet Survey ($N = 974$). The results revealed that all three generations exhibited a tendency toward the third-person perception, perceiving that the risks of AI technology as having a greater impact on others than on themselves. However, ChatGPT usage did not significantly predict the third-person effect across any generation. AI literacy had a significant negative impact on the third-person perception only within Generation X. Furthermore, ChatGPT usage exerted a significant negative influence on attitudes toward AI regulation among Millennials and Generation X. In contrast, AI literacy had a significant negative impact on such attitudes among Generation Z. Notably, the third-person perception failed to exert a significant impact on attitudes toward AI regulation across three generations. These findings contribute to a better understanding of Taiwanese citizens' risk perceptions regarding generative AI and can serve as a reference for policymakers in formulating relevant policies.

- ⊙ Keywords: AI literacy, AI usage experience, cohort, generative artificial intelligence, generative AI risk, ChatGPT, the third-person effect
- ⊙ The first author, Huai-Kuan, Zeng , is a Ph.D. candidate in the Institute of Applied Arts at National Yang Ming Chiao Tung University. The second author, Zhi-Rong Li, is a master's student in the Institute of Communication Studies at National Yang Ming Chiao Tung University. The third author, Tai-Yee Wu, is an Associate Professor in the Institute of Communication Studies at National Yang Ming Chiao Tung University.
- ⊙ Corresponding author: Huai-Kuan, Zeng, email: miazeng.hs07@nycu.edu.tw. address: 1F, HSS Building 2, No. 1001, Daxue Rd. East Dist., Hsinchu City 300, Taiwan.
- ⊙ Received: 2024/02/10 Accepted: 2025/07/24

壹、緒論

近年來，生成式 AI（generative AI，如：ChatGPT、Gemini 與 Claude）的發展日漸成熟與多元，對於民眾的工作與日常生活扮演了重要的角色，卻也逐漸浮現負面影響。例如 ChatGPT 生成帶有偏見的內容（Vartak, 2023），甚至產生幻覺（hallucinations），以貌似事實的誤導性資訊來回應使用者（Alkaissi & McFarlane, 2023）；也有研究指出，愈來愈多青少年受到 AI 生成的虛假內容所誤導（Common Sense Media, 2025）。有鑑於此，人們開始從法律層面來制定 AI 的問責機制，以保障 AI 科技的合理使用，並防制 AI 科技的不當濫用（e.g. Bensinger, 2024; European Parliamentary Research Service, 2023）。

儘管 AI 研究目前還是熱門的研究議題，但既有文獻多著重於 AI 科技的採用（例如：C. Wang, Boerman, Kroon, Möller & Vreese, 2024；吳泰毅、鄧玉玲，2023；許馥嘉、吳泰毅，2025），關於民眾對生成式 AI 的風險感知為何，受到哪些因素影響，以及對於 AI 科技進行規範的態度等，仍有待探究。據此，本研究以第三人效果（the third-person effect, TPE; Davison, 1983）作為理論基礎，檢驗引發台灣民眾 AI 科技風險的第三人感知（the third-person perception）之影響因子（antecedents），以及後續的可能結果（outcomes）。同時，由於不同世代的數位科技掌握能力存在差異（e.g. Karampelas, 2023；Prensky, 2001），數位落差更是影響 AI 素養的關鍵因素（e.g. Celik, 2023a），對於各個世代如何理解與使用 AI，現有研究仍關注不足，因此本研究也將進一步比較「X 世代」（1965—1980 年出生）、「千禧世代」（1981—1996 年出生）及「Z 世代」（1997—2012 年出生）對於生成式 AI 科技在使用、素養、風險感知及第三人效果以及法律規範態度上的異同。

第三人效果係指人對於媒體所造成的影響存有主觀的偏見，就媒體的負面影響而言，人們傾向高估媒體對「他人」（即：非你、非我的第三人）的影響，卻低估媒體對「自身」的影響，形成一種名為第三人感知的錯估落差，進而影響人們的態度與行為（Davison, 1983），例如支持色情內容的審查（Lo & Wei, 2002；Zhao & Cai, 2008；Zhou & Zhang, 2023）、以及推廣和宣傳正向訊息及行為（Sun, Shen & Pan, 2008）等。

本研究聚焦於台灣民眾對生成式 AI 科技風險的第三人感知，在影響此感知的因素方面，本研究著眼於 ChatGPT 使用經驗以及 AI 素養（AI literacy）。過去文獻指出，民眾自身的媒體使用經驗會影響他們評估媒體內容對他人的影響程度（e.g. Chen & Atkin, 2021；Cho, Shen & Peng, 2021；Hong, 2023）。本研究著眼於 ChatGPT 的使用，當民眾 ChatGPT 的使用經驗愈多，愈可能經常接觸到 ChatGPT 生成錯誤、帶有偏見甚至 AI 幻覺的資訊，因此認為 AI 生成內容可能會為他人帶來較大的負面影響。此外，AI 素養是指使用者正確辨別、正當使用、清楚評估 AI 科技的能力（B. Wang, Rau & Yuan., 2022），AI 素養愈高的使用者，會自認個人具備足夠的知識來辨析資訊的正確性，而他人應不具備與他們等量的之時，因而容易高估他人受到 AI 科技的負面影響。

此外，台灣民眾對於政府規範 AI 科技的態度，也是本研究關心的議題。過去文獻指出，第三人感知會導致人們支持對於負面媒介或其內容進行審查或規範，像是支持社群平台規範錯假訊息（Chung & Kim, 2021；Chung & Wihbey, 2024；Jang & Kim, 2018）、採取保護網路隱私的措施（Chen & Atkin, 2021）以及贊成色情內容審查（Lo & Wei, 2002; Zhao & Cai, 2008）。由於政府在制定內容分級系統和執行審查政策扮演了重要的角色，本研究也將檢驗民眾對於生成式 AI 風險所引發的第三人感知，是否影響其對於政府規範 AI 科技的態度。

綜上所述，本研究以第三人效果為理論基礎，試圖了解台灣民眾對生成式 AI 的風險是否存在第三人感知，同時檢驗民眾的 ChatGPT 使用經驗和 AI 素養對此第三人感知的影響，進而探討此第三人感知與支持政府規範 AI 科技之間的關聯性，並嘗試剖析可能的世代差異。研究結果期望能有助於擴充第三人效果應用於 AI 風險議題的實證文獻，探索不同世代的 AI 使用樣貌與特性，並可提供政府相關部門在制定政策、平衡 AI 風險與推廣 AI 教育的應用指引。

貳、文獻回顧

一、 人工智慧風險與第三人效果

AI 為近年新興的科技工具，由於功能多元且強大，對民眾的日常生活扮演了重要的角色。然而，由於入門操作門檻低，一般民眾容易忽略 AI 科技隱藏的風險，例如生成不正確或帶有偏見的資訊（Vartak, 2023），以及使用的過程中對於個資的保障不明（Gupta, Akiri, Aryal, Parker & Praharaj, 2023）。為了解決 AI 衍生的負面影響，國際組織（例如：European Parliament, 2023）與學術文獻（Toreini et al., 2020）皆已指出，AI 科技的問責（accountability）、監督（auditability）機制是提升 AI 可信賴性（trustworthy AI）的重要因素之一。這些觀點顯示，政府應積極介入，以保障 AI 的合理與正當使用。在此背景下，第三人效果提供了重要的視角，有助於理解民眾支持政府制定 AI 科技法律規範的心理動機。

自我強化（self-enhancement）為解釋第三人感知的重要心理機制，根據自我強化理論，承認個人受到媒介影響是個人不想要的特質，但藉由假定他人更為脆弱而受到媒介效果的影響，個人可以維護個人正面的形象，並強化個人比他人優越的想法的樂觀偏誤（Perloff, 1999；2009）。因此，若媒介訊息被認為是不良的特質，人們傾向認為個人能夠抵禦負面媒介訊息的影響，而他人較容易受到媒介的負面影響，例如：暴力內容（Boyle, McLeod & Rojas, 2008）、色情（Lo & Wei, 2002；Rosenthal, Detenber & Rojas, 2018）、假新聞（Chung, & Kim, 2021；Lyons, 2022）、網路隱私（Chen & Atkin, 2021）。這些研究主題涵蓋暴力內容、色情內容、假新聞、網路隱私等多元主題，因此，第三人效果理論應也能應用於探討生成式 AI 科技風險。由於生成式 AI 產生的內容可能存有偏見或是產出不正確的資訊等負面特質，根據自我強化觀點，藉由假定他人更容易受到媒介的影響，有助於維護個人的正面形象並強化個人比他人優越的樂觀偏見。因此，台灣民眾在評估生成式 AI 科技所帶來的風險時，應認為他人更容易受到 AI 科技的負面影響，而自己則較不易受到 AI 科技的影響，因此，本研究提出研究假設：

H1：台灣民眾認為生成式 AI 風險對自身的影響較小，對他人的影響較大。

二、ChatGPT 使用經驗與生成式 AI 風險的第三人感知

根據自我強化的觀點，為了維護個人的正面形象，人們傾向認為自身的能力優於他人，而產生他人更容易受到媒介影響的樂觀偏誤（Perloff, 1999；2009）。過去實證研究顯示，個人特定的媒體使用經驗，是預測第三人感知的重要變數。例如：Lev-On（2017）的研究指出，對 Facebook 功能愈熟悉的使用者，會愈認為個人更擅長處理 Facebook 的風險，並認為其他對 Facebook 不熟悉的使用者則更容易受到 Facebook 的風險影響。在網路隱私的議題上（e.g. Chen & Atkin, 2021）當民眾愈線上搜尋資訊的頻愈高，愈傾向認為個人比他人更注重網路隱私的風險，而他人則不具備這樣的能力，因此會推薦他人採取保護措施以避免隱私外洩。

由於 ChatGPT 等生成式 AI 可能產生不盡正確或帶有偏見的內容（Vartak, 2023），或是因為不當使用而生成造假影像而散播不實資訊（Broussard et al., 2019；King & Meinhardt, 2024），因而被認為是具有負面特質的媒體。根據自我強化的觀點，當人們對 ChatGPT 的使用經驗愈多，他們會對於生成式 AI 科技更為熟悉，因而並認為個人具備足夠的知識與自信來妥善運用這些生成式 AI。因此，他們傾向認為個人更有本事抵擋這些 AI 科技的負面風險，據此，本研究提出假設 2：

H2：ChatGPT 使用經驗愈高，生成式 AI 風險的第三人感知愈高。

三、AI 素養與生成式 AI 風險的第三人感知

「素養」源自於個人具備運用書寫語言來表達個人想法與溝通的能力（Long & Magerko, 2020）。過去文獻（e.g. Long & Magerko, 2020；Ng, Leung, Chu & Qiao, 2021；B. Wang et al., 2022）將 AI 素養擴及至多面向的能力，包含：（1）了解與理解 AI 技術：知道 AI 的基本功能並運用於日常生活。（2）運用 AI：在不同情境中有效應用 AI 技術協作。（3）評估與創建 AI：運用高階思考能力來批判性評估 AI 提供的訊息以及設計內容。（4）AI 倫理：考慮其倫理問題的能力，例如：公平、問責、透明性與倫理。

因此，除了自我強化的觀點以外，控制需求（need of control）為另一個解釋第三人感知的心理機制（Perloff, 2009）。若個人相信每一則媒體訊息或刺激都會對他們造成劇烈的影響，可能使個人陷入精神崩潰，因此，藉由假定個人不會受到大眾媒介的影響，人們才得以在媒體所主導的世界中正常生活，並從生活中獲得滿足，並理性地將媒體納入人們的日常生活中。同理，由於 AI 素養是指個人使用 AI 科技的知識與能力，擁有較高 AI 素養的人，可能會認為自己更有能力掌握 AI 科技，而他人的 AI 掌握能力不及他們，因而影響他們對於生成式 AI 風險的第三人感知。

Perloff (1999) 指出，個人的知覺知識會影響到他們對媒介影響力的看法，甚至比個人的實際知識更能預測閱聽眾對他人所受到媒介影響力的評估（e.g. Li, Shi, Zhao, Zhang & Zhong, 2024）。Salwen 與 Dupagne (2001) 在調查電視暴力內容時更發現，相較於人口變項（如：性別、年齡、意識形態及媒體使用經驗），個人的知覺知識更能預測第三人感知。在網路隱私（Chen & Atkin, 2021）及 Covid-19 相關的假新聞（J. Yang & Tian, 2021）的研究主題中，也發現當使用者認為自身擁有更高的知識，會強化第三人效果。根據上述文獻，當個人認為自身具備足夠的 AI 知識（即 AI 素養），應會認為個人相較他人具備足夠的知識來辨識 AI 科技生成內容的正確性以及隱藏的偏見；反之，因他人不具備與他們同等知識，因而不易發覺到 AI 隱藏的風險而受到影響。因此，本研究提出以下假設：

H3：AI 素養中的（a）覺察；（b）使用；（c）評估能力愈高，生成式 AI 風險的第三人感知愈高。

四、生成式 AI 風險的第三人感知與法律規範 AI 的態度

過去研究指出，第三人感知可顯著預測人們的行為意圖（Perloff, 2009；Sun et al., 2008）。Sun 等人（2008）提出三種受到第三人感知所引發的行為：限制（restriction）、更正（correction）以及推廣（promotion）。首先基於保護他人不受到負面媒體內容的影響，人們會傾向於支持政府或相關機構採取審查和限制，例如：Rosenthal et al. (2018) 的研究指出，新加坡民眾面對電影色情內容審查時，傾向認為他人不具備足夠判斷與理解影片內容的能力而容易受到色情內容的影響，因此認為媒

介審查（censorship）為限制他人接觸負面媒介內容的有效措施。第二，當人們認為特定媒體內容會對他人造成負面影響，人們會採取更正行動，例如：路森、羅文輝、魏然（2023）調查 Covid-19 疫情的虛假資訊發現，當人們認為虛假資訊會影響他人，他們會透過分享 Covid-19 查核訊息，來避免他人受到虛假訊息所誤導。最後，有些人會透過分享公益廣告、或是具體的保護措施，以推廣法規或是糾正負面的媒體內容（Chen & Atkin, 2021; Sun et al., 2008）。

由於生成式 AI 可能產生不實資訊或具偏見的內容，國際機構開始透過公權力的介入，期望透過法律來規範 AI 科技的使用。例如，歐盟提出世界上第一個 AI 法律，著眼於 AI 風險分級，與規範發行者與使用者應遵循的義務（European Parliament, 2023）。在亞洲方面，台灣與韓國政府近年來也個別開始起草 AI 基本法，明確界定風險及責任分級以及降低人工智慧風險（中央社，2025；徐子苓，2024）；而日本政府有感於生成式 AI 產出的假訊息和錯誤資訊，也正式考量導入 AI 的法律規範，以致力提升 AI 的內容品質以及社群媒體上詐騙廣告的防範（經濟部台日產業合作推動辦公室，2024）。不過，並非各國都尋求更嚴格的 AI 法律規範，例如美國對於 AI 的風險管理政策因政權移轉而由緊轉鬆，前任拜登政府簽署行政命令以消弭 AI 帶來的風險，現任政府則致力於取消相關政策和法規，以維持和增強美國在全球人工智慧領域的主導地位（賴又豪，2024）。

至於民眾對於 AI 科技是否由政府、法律規範，亦抱持著歧見。例如，一個含括 30 國的大型跨國調查（ $N=22,816$ ）指出，民眾對於 AI 科技的態度看法兩極（Ipsos, 2023）。美國民眾一方面認為現有的 AI 科技規範政策仍不足（Faverio & Tyson, 2023），一方面也擔憂 AI 科技的規範會限制創新，削弱美國在 AI 領域的競爭力（Bensinger, 2024）。日本民眾對於 AI 態度並未擁有明確的立場，一方面對於 AI 帶來的生活便利化有所期待，一方面也擔憂 AI 的使用對個人隱私權的侵害（關鍵評論網，2023）。

綜合上述文獻回顧，台灣民眾如何看待 AI 科技的規範是值得關注的議題。許多實證研究結果顯示（e.g. Hoffner et al, 2001；Jang & Kim, 2018；J. Yang & Tian, 2021；Zhao & Cai, 2008），針對爭議性內容，人們會基於保護他人不受到媒體內容的影響，而更傾向於支持政府或相關機構採取審查和限制。由於生成式 AI 可能產生不實資訊

或具偏見的內容，使用者會擔心他人受到這些內容的誤導，因而更傾向藉由公權力的介入，透過法律來規範 AI 科技的使用。因此，本研究提出以下假設 4：

H4：生成式 AI 風險的第三人感知與支持政府規範 AI 科技呈正相關。

五、ChatGPT 使用經驗、AI 素養與法律規範 AI 的態度

最後，基於自我強化理論，面對 AI 科技個人可能會存在樂觀偏誤。例如，對於 AI 科技抱持最正向態度的先驅者，不僅對 AI 科技抱持更高的信任感，並認為 AI 產品與服務相較傳統人力更能滿足個人需求，讓生活過得更好（吳泰毅、鄧玉羚，2023）。因此，可預期當個人使用 ChatGPT 的經驗愈多，應自覺個人有足夠的能力掌握這些生成式 AI，而他人未必具備同樣的能力，而有機會受到生成式 AI 科技的負面影響。同理，AI 素養反映了使用者對於 AI 的基本認知、辨識與使用能力，根據計畫行為理論（Ajzen, 1985），知覺行為是指個人主觀知覺能行駛特定行為所具備的能力，當個人的知覺行為控制愈高，會認為自己更有能力來從事特定行為。過去學者（e.g. C. Wang et al., 2024）也指出，當使用者的 AI 素養愈高，愈能提升他們對生成式 AI 的態度、知覺行為控制。因此，從控制需求的觀點來看，由於人工智慧屬於創新科技，相較於傳統媒體，使用者需具備更豐富且多元的知識與技術方能駕馭它，這意味著擁有較高 AI 素養的人，可能會認為自己更有能力掌握 AI 科技。因此，本研究推測，當個人對 ChatGPT 的使用經驗愈多、AI 素養愈高，他們愈對自己控制 AI 風險有格外的信心，因此不希望透公權力過度干預他們使用 AI 科技。因此，本研究提出以下兩個假設：

H5：ChatGPT 使用經驗與支持政府規範 AI 科技呈負相關。

H6：AI 素養中的（a）覺察；（b）使用；（c）評估能力與支持政府規範 AI 科技呈負相關。

六、ChatGPT 使用經驗、AI 素養、生成式 AI 風險第三人感知與法律規範態度之世代差異

數位落差 (Digital divide) 廣義上是指使用者獲取與使用新興資訊科技所存在的差距，除了電腦與網路，也包含手機及數位電視等數位產品 (Y. Chang, Wong & Park, 2016; van Dijk, 2006)。它不僅指個人或是群體在數位產品的近用差異，更是一個關鍵、複雜且動態的全球現象 (van Dijk & Hacker, 2003)，例如：已開發國家和開發中國家資訊與通訊科技 (information and communication technology, ICT) 的發展、近用和使用程度存在顯著差異 (Y. Chang, Shahzeii, Kim & Park, 2012)。

然而，van Dijk (2002) 指出，實體使用科技不足理解數位落差，因此，學者 (Y. Chang et al., 2016; van Dijk, 2002; 2006) 提出更廣泛的理論來解釋數位落差的概念，包含：(1) 動機近用 (motivational access)：個人學習使用與採用數位科技的動機與意願；(2) 物理近用 (physical access)：個人有機會接觸與使用科技，包含採用該科技的成本與基本設施；(3) 技能近用 (skill access)：不只強調個人操作科技的能力，更要能理解該科技的優點與缺點；(4) 使用近用 (usage access)：個人實際使用科技，並運用工作任務的能力。儘管 AI 技術能協助使用者解決日常許多任務，但其生成的內容仍可能存在不實或帶有偏見的內容。因此，個人若要實際接觸、妥善運用並批判性思考 AI 生成的內容，就必須具備足夠的 AI 科技的動機及技能。Celik (2023a) 的研究指出，數位落差可正向影響 AI 素養。因此，可預期數位落差的各個面向對個人 AI 素養以及 AI 風險評估扮演關鍵角色。

然而，不同世代民眾近用 AI 的動機、成本、技能及使用可能存在顯著差異，因而影響到其 AI 素養的發展以及對於 AI 風險的認知。例如，學者 Prensky (2001) 提出「數位原民」(Digital-Native) 及「數位移民」(Digital-Immigrant) 的概念，描述在不同人生階段接觸、學習及運用數位科技 (例如：電腦、手機等電子產品) 的使用者。數位原民是指在年輕時代即接觸數位科技、並隨著新興科技成長的族群；相對地，數位移民則是指在傳播科技出現在他們中老年時期的族群。相較於較晚接觸新興數位科技的數位移民，數位原民對於這些數位科技有較高的熟悉程度及掌握能力，並在新科技的接受程度、使用能力與應用能力與數位移民有所區別。而這項世代差異，

同樣也預期會延伸至使用者對 AI 科技的近用、動機與技能，因而影響 AI 科技風險評估以及 AI 素養的建立與發展。

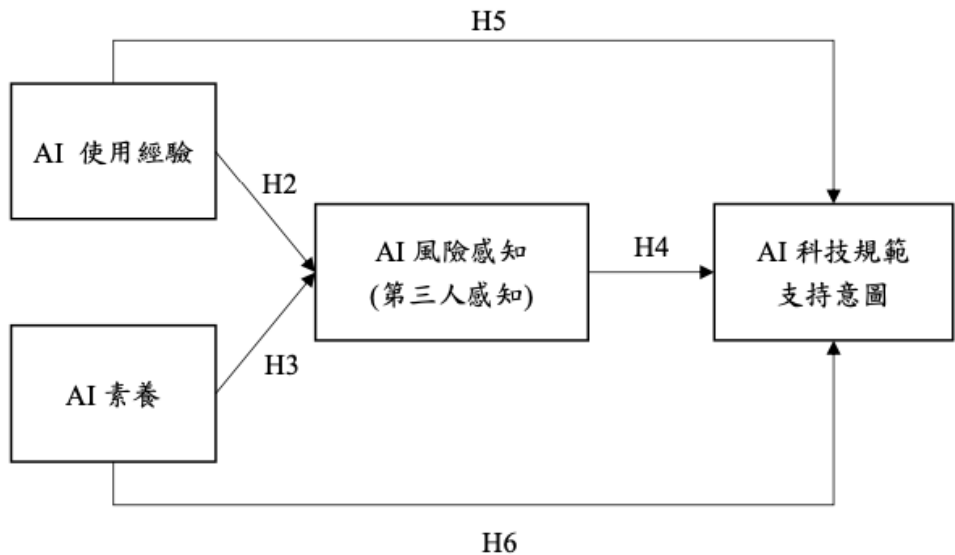
除此之外，隨著 ChatGPT 等生成式 AI 科技的問世，「AI 原民」(AI-natives) 也隨之出現。他們不僅擅長運用 AI 工具，更發展出「AI 直覺」(AI intuition) 以運用於學習及未來的職場中 (Karampelas, 2023)。然而，由於人工智慧屬於創新科技，相較於傳統媒體，使用者需具備更豐富且多元的知識與技術方能駕馭它，不同世代的民眾對於 AI 科技使用的觀點與法律規範態度可能存在不同的世代偏好。例如，許多大型研究調查顯示 (e.g. Ipsos, 2023; Pew Research Center, 2023; 財團法人網路資訊中心, 2024)，不同世代對於人工智慧抱持兩極的態度，相較於熟齡世代，年輕世代不僅更願意搶先體驗創新科技，對 AI 科技亦抱持較為樂觀的看法。相對地，由於熟齡世代較晚接觸生成式 AI 科技，相較熟悉數位工具並善用 AI 科技進行協作的年輕世代，他們可能因為缺乏相關的知識與經驗，較難察覺到生成式 AI 可能生成不正確的資訊。例如，C. Wang et al. (2024) 針對荷蘭民眾進行的研究發現，年紀較長、教育程度較低，且隱私保護技術較為薄弱的民眾，是最容易受到 AI 科技負面影響的族群。因此，不同世代在 AI 科技使用經驗與 AI 素養可能存在差異，進而在感受到 AI 風險時的行為反應，展現不同的世代傾向。

本研究聚焦於 X 世代、千禧世代與 Z 世代，主要是依循社會科學領域 (e.g. BBC, 2017; Pew Research, 2019) 對於世代劃分主要方式，探討 X 世代、千禧世代與 Z 世代形成連貫的世代脈絡。另一方面，這三個世代長期處於數位環境，更是生成式 AI 技術潛在的主要使用者與內容接收者。他們如何看待 AI 可能產生的不實資訊、偏誤內容與真偽模糊的現象，將有助於理解不同世代對 AI 風險的認知與反應。因此，本研究提出以下研究問題 1、2：

RQ1：不同世代在本研究所提出的各項假設關係中，是否呈現出世代差異？

RQ2：不同世代的 (a) AI 素養能力 (包含：覺察、使用、評估)、(b) ChatGPT 使用經驗、(c) 生成式 AI 科技風險第三人感知是否存在差異？

圖一：本研究研究架構



註：H1 及 RQ1、RQ2 未列於研究架構中

參、研究方法

一、研究樣本

本研究採用財團法人網路資訊中心《2024 台灣網路報告》的調查資料進行研究分析。該研究對象為居住在全台灣 22 縣市、18 歲以上人口，並再分為住宅電話抽樣（以全台灣的都、縣、市為分層依據，再依各層的家戶比例抽取市話號碼）和手機電話抽樣（隨機撥打號碼）兩種；調查完成後再將住宅、手機的抽樣資料合併為雙底冊，並依事後分層組合估計法進行數據加權。研究調查期間為 2024 年 6 月 19 日至 29 日，合計有效樣本數為 $N = 1,898$ 。由於本研究目的在理解以「Z 世代」、「千禧世代」、「X 世代」這三個不同世代的 AI 使用情形與觀點，並以 AI 相關題目作為分析題組，刪除掉不屬於這三個世代的樣本（1964 年以前出生， $N = 95$ ），填答不完整、部分題目回答沒有相關經驗、不知道、或拒答的參與者後，以樣本數 $N = 979$ 進行分析。

二、測量變項

（一）ChatGPT 使用經驗：

ChatGPT 使用經驗定義為台灣民眾過去三個月使用 ChatGPT 的頻率。研究結果顯示，整體台灣民眾過去三個月使用 ChatGPT 的頻率為：從來沒有（57.2%）、很少（19.9%）、有時（11.8%）、經常（7.4%）、總是（3.7%）；選項採 Likert 五點量表，從 1「非常不同意」至 5「非常同意」（ $M = 1.81$ ， $SD = 1.13$ ）。

（二）AI 素養：

此題組依據 B. Wang et al. (2022) 所提出的 AI 素養面向——覺察、使用、評估等三個指標及定義，改寫 AI 素養相關之題項，各個面向能力的分別定義為（1）覺察：使用者了解與理解 AI 基本功能並運用於日常生活的能力：「我有能力分得清楚哪些產品與服務是運用 AI 人工智慧，哪些不是（例如：網路上的 AI 客服與真人客服）」（ $M = 3.83$ ， $SD = 1.08$ ）。（2）使用者在不同情境中有效應用 AI 技術協作的的能力：「我有能力運用 AI 產品或服務完成我想要做的事情」（ $M = 3.49$ ， $SD = 1.24$ ）

（3）使用者能批判性評估 AI 提供的訊息以及設計內容之能力：「在使用過 AI 產品或服務後，我很清楚地知道它們有哪些優點和缺點」（ $M = 3.81$ ， $SD = 1.10$ ）；選項採 Likert 五點量表，從 1「非常不同意」至 5「非常同意」。

（三）AI 風險感知：

此題改編自 Chung 與 Kim (2021) 所發展的測量方式。題目先以「目前 AI 人工智慧所生成的內容，可能存有偏見或是產出不正確的資訊。」為前述句，再進一步詢問研究參與者 AI 生成內容分別對「個人」以及「一般大眾」的影響程度，包含 2 題：（1）對「個人」的風險感知：請問您認為，這對您個人有沒有影響？（ $M = 2.17$ ， $SD = 1.09$ ）（2）對「一般大眾」的風險感知：請問您認為，這對一般大眾有沒有影響？（ $M = 3.53$ ， $SD = 1.12$ ）。平均數愈高，代表民眾認為 AI 生成內容的風險愈高。此外，所對第三人效果感知係指對個人的風險感知與對他人的風險感知差距（ $M = 1.36$ ， $SD = 1.27$ ）。選項採 Likert 五點量表（從 1「非常不同意」至 5「非常同意」）。

（四）法律規範（民眾支持政府規範 AI 科技的態度）：

此題目改編自 Salwen 與 Dupagne (1999) 之題組，以了解台灣民眾支持政府透過法律規範 AI 科技的態度。調查時電訪員首先向受訪者詢問：「請問您認為 AI 人工智慧產品或服務是否應受到政府的規範？」，選項包含：支持由政府規範則編碼為 1，未選擇支持由政府規範則編碼為 0；結果顯示，支持政府機關規範（支持， $N = 508$ ，51.9%；不支持， $N = 471$ ，48.1%）。

（五）控制變項：

本研究的控制變項包括性別、年齡、教育程度、上網經驗和使用其它類型的 AI 經驗。在性別方面（49.2% 為男性），將變項事後重新編碼為：男 = 1，女 = 0；年齡方面，平均數 $M = 37.13$ ， $SD = 11.36$ ；教育程度方面，小學及以下：0.4%；國、初中／初職：2.5%；高中、職：22.7%；專科：9.6%；大學：49.5%；碩士及以上：15.3%；事後重新編碼為：小學及以下 = 1；國、初中／初職 = 2；高中、職 = 3；專科 = 4；大學 = 5；碩士及以上 = 6）。

此外，台灣民眾過去三個月的上網經驗次數分配為：很少（一週不到一次）：0.1%；一週數次（一週一至六次）：1.5%；一天一次：2.1%；一天數次：36.8%；幾乎一直上網（次數多到無法計算）：59.6%；事後重新編碼：很少（一週不到一次）= 1；一週數次（一週一至六次）= 2；一天一次 = 3；一天數次 = 4；幾乎一直上網（次數多到無法計算）= 5）。最後，台灣民眾過去三個月來使用另一種 AI 工具—數位語音助理的次數分配為：從來沒有（51.5%）、很少（24.1%）、有時（10.4%）、經常（9.4%）、總是（4.6%）；選項採 Likert 五點量表，從 1「從來沒有」至 5「總是」（ $M = 1.91$ ， $SD = 1.18$ ）。

三、共線性檢定

為了確認變項間是否有共線性問題，本研究透過 SPSS 24 版，將 ChatGPT 使用經驗、AI 素養、第三人感知、支持政府規範等變項逐一放入多元迴歸，研究結果顯示，整體民眾的變異數膨脹因數（Variance Inflation Factor, VIF）為（1.01~1.61）；各世代的變異數膨脹因數分別為：Z 世代（1.02 ~ 1.66）、千禧世代（1.01~1.45）、X 世代（1.01~1.63），顯示變異數膨脹有限（Field, 2018）。此外，各變項相關分析結果請見表一。

表一：預測變數之相關分析

整體民眾	1	2	3	4	5	6
1. ChatGPT 使用經驗	--					
2. AI 素養：覺察	.16**	--				
3. AI 素養：使用	.35**	.37**	--			
4. AI 素養：評估	.30**	.42**	.56**	--		
5. 第三人感知	-.08*	-.04	-.07*	-.05	--	
6. 支持政府規範	-.20**	-.01	-.12**	-.11**	.03	--
Z 世代	1	2	3	4	5	6
1. ChatGPT 使用經驗	--					
2. AI 素養：覺察	.09	--				
3. AI 素養：使用	.33**	.35**	--			
4. AI 素養：評估	.31**	.32**	.58**	--		
5. 第三人感知	-.14*	-.13*	-.04	-.12	--	
6. 支持政府規範	-.07	.002	-.08	-.15*	.09	--
千禧世代	1	2	3	4	5	6
1. ChatGPT 使用經驗	--					
2. AI 素養：覺察	.11*	--				
3. AI 素養：使用	.31**	.32**	--			
4. AI 素養：評估	.25**	.38**	.50*	--		
5. 第三人感知	-.06	.06	.01	.003	--	
6. 支持政府規範	-.19**	.01	-.12*	-.07	.001	--

X 世代	1	2	3	4	5	6
1. ChatGPT 使用經驗	--					
2. AI 素養：覺察	.17**	--				
3. AI 素養：使用	.29**	.37**	--			
4. AI 素養：評估	.27**	.44**	.56**	--		
5. 第三人感知	-.002	-.11	-.17**	-.09	--	
6. 支持政府規範	-.20**	.07	-.01	-.03	-.01	--

註：* = $p < .05$, ** = $p < .01$

肆、研究結果

本研究以 SPSS 作為研究分析工具，在進行假設驗證前，為了解不同年齡世代的民眾對於 AI 使用經驗、AI 素養、AI 風險感知以及支持規範 AI 科技的差異。本研究將資料分割為三個不同世代，分別以成對樣本 T 檢定、多元迴歸與羅吉斯迴歸 (Logistic regression) 與單因子變異數分析進行假設檢定。

一、台灣民眾的 AI 風險第三人感知

研究假設 1 在探究台灣民眾評估對生成式 AI 科技風險其自身的影響、對他人的影響之間是否存在差異。成對樣本 T 檢定結果顯示，台灣整體民眾認為，AI 科技風險對他人的影響 ($M = 3.53, SD = 1.12$) 顯著高於對個人的影響 ($M = 2.17, SD = 1.09$)，顯示民眾評估 AI 科技風險對個人及他人的影響確實存在顯著差異： $t(978) = -33.64, p < .001$ 。進一步分析不同世代評估 AI 科技風險對個人及對他人的影響是否存在差異，研究結果顯示，Z 世代 ($t(246) = -15.65, p < .001$)、千禧世代 ($t(431) = -24.35, p < .001$) 以及 X 世代 ($t(299) = -17.55, p < .001$) 的民眾均認為 AI 科技的風險對他人的影響顯著高於個人，證實第三人感知存在於這三個世代。基此，假設 1 獲得成立。研究結果詳見表二。

表二：台灣民眾評估 AI 科技風險對個人及對自身的影響

	三個世代	Z 世代 (1997-2012 年出生)	千禧世代 (1981-1996 年出生)	X 世代 (1965-1980 年出生)
<i>N</i>	979	247	432	300
對個人的影響 <i>M(SD)</i>	2.17(1.09)	2.17(1.00)	2.15(1.10)	2.18(1.13)
對他人的影響 <i>M(SD)</i>	3.53(1.12)	3.36(1.08)	3.65(1.09)	3.49(1.17)
<i>T</i> 值	<i>t</i> (978)= -33.64***	<i>t</i> (246)= -15.65***	<i>t</i> (431)= -24.35***	<i>t</i> (299)= -17.55***

註：***= *p* <.001

二、依變項：第三人效果感知

研究假設 2、3 分別推測 ChatGPT 使用經驗以及 AI 素養對生成式 AI 科技風險的第三人感知具正向影響。本研究以階層迴歸分析，以第三人感知作為依變項，在模型第一層置入性別、年齡、教育程度、近三個月的上網頻率以及數位語音助理等控制變項，第二層置入 ChatGPT 使用經驗，以及 AI 素養中的覺察、使用與評估等三個能力面向。

(一) 假設 2 檢定

研究結果顯示，台灣整體民眾認為，ChatGPT 使用經驗 ($\beta = -.06, p = .084$) 無法預測第三人感知，研究結果詳如表三。進一步分析不同世代 ChatGPT 使用經驗與第三人感知之間的關係，研究結果顯示 Z 世代 ($\beta = -.12, p = .102$)、千禧世代 ($\beta = -.06, p = .240$) 以及 X 世代 ($\beta = .07, p = .300$) 在 ChatGPT 使用經驗上，皆無法顯著預測第三人感知，故假設 2 未獲得證實，研究結果詳如表四。

(二) 假設 3 檢定

1、整體民眾

研究結果顯示，AI 素養中的覺察 ($\beta = -.01$, $p = .793$)、使用 ($\beta = -.04$, $p = .304$) 及評估 ($\beta = -.02$, $p = .682$) 等三面向的能力皆無法預測第三人感知，故假設 3(a)、3(b)、3(c) 皆未獲得證實。研究結果詳如表三。

2、Z 世代

再以三個世代的使用者進行區分，就 Z 世代而言，研究結果亦顯示 AI 素養中的覺察 ($\beta = -.13$, $p = .058$)、使用 ($\beta = .15$, $p = .069$) 及評估 ($\beta = -.13$, $p = .101$) 皆無法預測第三人感知，再次證實假設 3(a)、3(b)、3(c) 不成立。研究結果詳如表四。

3、千禧世代

同樣地，千禧世代的使用者在覺察 ($\beta = .07$, $p = .181$)、使用 ($\beta = .01$, $p = .916$) 或評估 ($\beta = .002$, $p = .972$) 等 AI 素養面向方面，對第三人感知皆不具顯著預測力，因此假設 3(a)、3(b)、3(c) 仍不成立。研究結果詳如表四。

4、X 世代

另外，X 世代的資料顯示，AI 素養中的覺察 ($\beta = -.04$, $p = .544$) 與評估 ($\beta = .02$, $p = .812$) 皆無法顯著預測第三人感知；然而，使用能力 ($\beta = -.17$, $p = .016$) 則可顯著負向預測第三人感知，但是與假設的預測方向相反，故假設 3(a)、3(b)、3(c) 皆未獲得證實。研究結果詳如表四。

總結而言，台灣整體民眾的樣本顯示，ChatGPT 使用經驗以及 AI 素養中的覺察、使用、評估等三項能力均無法顯著預測第三人感知；進一步分析不同世代樣本的認知差異，三個世代民眾的 ChatGPT 使用經驗均無法顯著預測第三人感知。關於 AI 素養對第三人效果的影響，Z 世代、千禧世代的 AI 素養中的三項能力（覺察、使用、評估）均無法顯著預測第三人感知；然而，X 世代 AI 素養的使用能力與第三人感知呈負相關，但 AI 素養的覺察與評估能力則無此效果。

表三：整體民眾階層迴歸分析結果

預測變數	第三人感知			
	R^2	β	t	p
階層 1	.01			
性別 (男性= 1)		.09**	2.66	.008
年齡		.04	1.30	.195
教育程度		.02	.54	.588
網路使用		-.02	-.74	.460
數位語音助理		-.05	-1.59	.112
階層 2	.01			
ChatGPT		-.06	-1.73	.084
AI 素養：覺察		-.01	-.26	.793
AI 素養：使用		-.04	-1.03	.304
AI 素養：評估		-.02	-.41	.682

註： **= $p<.01$ 。依變項為第三人感知

表四：各世代民眾階層迴歸分析結果

預測變數	Z 世代 (1997-2012 年出生)				千禧世代 (1981-1996 年出生)				X 世代 (1965-1980 年出生)			
	R ²	β	t	p	R ²	β	t	p	R ²	β	t	p
階層 1	.04				.01				.01			
性別 (男性=1)		.04	.68	.498		.09	1.80	.073		.08	1.29	.198
年齡		.16	2.32	.021		-.08	-1.71	.088		.14*	2.24	.026
教育程度		-.08	-1.14	.255		<.001	-.01	.995		.02	.26	.798
網路使用		-.17**	-2.64	.009		.06	1.29	.197		-.03	-.51	.609
數位語音助理		-.09	-1.34	.181		-.07	-1.48	.139		-.02	-.36	.720
階層 2	.06				.01				.03			
ChatGPT		-.12	-1.64	.102		-.06	-1.18	.240		.07	1.04	.300
AI 素養：覺察		-.13	-1.91	.058		.07	1.34	.181		-.04	-.61	.544
AI 素養：使用		.15	1.82	.069		.01	.11	.916		-.17*	-2.42	.016
AI 素養：評估		-.13	-1.65	.101		.002	.04	.972		.02	.23	.815

註：* = $p < .05$, ** = $p < .01$ 。依變項為第三人感知

三、依變項：法律規範（民眾支持政府規範 AI 科技）

假設 4、5、6 分別檢驗第三人感知（H4）、ChatGPT 使用經驗（H5）以及 AI 素養（H6）與民眾支持政府規範 AI 科技的關係。本研究以階層羅吉斯迴歸分析，依變項為支持政府規範（支持 = 1，不支持 = 0），並於第一層置入性別、年齡、教育程度、最近三個月的上網頻率以及數位語音助理使用經驗等控制變項，第二層置入 ChatGPT 使用經驗與 AI 素養的覺察、使用與評估等三項能力，第三層置入第三人感知。本研究檢視了標準化殘差與 Cook's Distance，以評估模型的極端值與影響值。結果顯示，各世代的標準化殘差範圍如下：整體樣本（.001~.12）

Z 世代（-1.53 ~ 2.43）、千禧世代（-2.00 ~ 2.03）與 X 世代（-2.72 ~ 1.91）之間，超過 2.58 的觀察值未超過 1%，符合過去文獻（Field, 2018）不得有超過 1% 的觀察值標準化殘差的絕對值大於 2.58。此外，各世代的 Cook's Distance 均小於 1，顯示個別資料點對整體模型無過度影響，模型推論具備穩定性（Cook & Weisberg, 1982）。

（一）假設 4 檢定

首先檢驗第三人感知對支持政府規範 AI 科技的正向影響，整體民眾與各世代研究結果如下：

1、整體民眾

台灣總體民眾的資料顯示，整體模式的預測力達顯著： $\chi^2(10) = 89.85, p < .001$ ；Hosmer-Lemeshow 適配度檢定，結果顯示 $\chi^2(8) = 21.31, p = .006$ ，模型與資料存在顯著差異，推論時需審慎解讀。然而，第三人感知對支持政府規範雖有正向影響，但未達統計上的顯著水準（ $B = .02, SE = .05, \text{Wald's } \chi^2(1) = .14, p = .708$ ），故假設 4 未獲得支持，研究結果詳如表五。

2、Z 世代

進一步檢驗不同世代的樣本，Z 世代使用者的研究結果顯示，整體模式的預測力未達顯著： $\chi^2(10) = 13.76, p = .184$ ；Hosmer-Lemeshow 適配度檢定結果顯示 $\chi^2(8) = 18.43, p = .018$ ，模型與資料存在顯著差異，推論時需審慎解讀。然而，Z 世代的第三人感知對支持政府規範雖有正向影響，但未達統計上的顯著水準（ $B = .17, SE = .12,$

Wald's $\chi^2(1) = 1.91, p = .167$ ），故假設 4 未成立，研究結果詳如表六。

3、千禧世代

同樣地，千禧世代民眾的分析結果顯示，整體模式的預測力達顯著： $\chi^2(10) = 29.42, p = .001$ ；Hosmer-Lemeshow 適配度檢定結果顯示 $\chi^2(8) = 22.21, p = .005$ ，模型與資料存在顯著差異，不宜過度推論研究結果。研究結果發現，千禧世代的第三人感知對支持政府規範具正向關係，但未具顯著性（ $B = .001, SE = .08, \text{Wald's } \chi^2(1) < .001, p = .987$ ），故假設 4 未成立，研究結果詳如表六。

4、X 世代

進一步確認 X 世代的使用者，整體模式的預測力達顯著： $\chi^2(10) = 25.38, p = .005$ ；Hosmer-Lemeshow 適配度檢定結果顯示 $\chi^2(8) = 15.17, p = .056$ ，模型與資料未存在顯著差異。X 世代民眾的第三人感知雖負向預測對政府規範的支持，但未達統計上的顯著水準（ $B = -.02, SE = .10, \text{Wald's } \chi^2(1) = .04, p = .845$ ），故假設 4 未成立。

整體而言，在三個世代的樣本中，第三人感知均無法顯著預測 AI 科技法律規範的態度，研究結果詳如表六。

（二）假設 5 檢定

假設 5 檢驗 ChatGPT 使用經驗對支持政府規範 AI 科技的負向影響，整體民眾與各世代研究結果如下：

1、整體民眾

台灣整體民眾羅吉斯迴歸分析整體模式的預測力達顯著： $\chi^2(9) = 89.71, p < .001$ ；Hosmer-Lemeshow 適配度檢定結果顯示 $\chi^2(8) = 27.82, p = .001$ ，模型與資料存在顯著差異，推論時需審慎解讀。研究結果顯示，ChatGPT 使用經驗對支持政府規範有顯著負向影響（ $B = -.30, SE = .07, \text{Wald's } \chi^2(1) = 19.94, p < .001$ ）。換言之，整體民眾，ChatGPT 使用經驗愈高，愈不支持政府規範 AI 科技，故假設 5 成立，研究結果詳如表五。

2、Z 世代

進一步區別為不同世代，Z 世代樣本整體模式的預測力未達顯著： $\chi^2(9) = 11.83, p = .223$ 。Hosmer-Lemeshow 適配度檢定結果顯示 $\chi^2(8) = 9.54, p = .299$ ，模型與資料未存在顯著差異。ChatGPT 使用經驗對支持政府規範為負向影響，但未具顯著性（ $B =$

$-.04, SE = .13, \text{Wald's } \chi^2(1) = .11, p = .745$), 故假設 5 未成立, 研究結果詳如表六。

3、千禧世代

千禧世代參與者的整體模式的預測力達顯著： $\chi^2(9) = 29.42, p = .001$ 。Hosmer-Lemeshow 適配度檢定結果顯示模型與資料存在顯著差異，推論時需審慎解讀 ($\chi^2(8) = 20.08, p = .010$)。千禧世代的 ChatGPT 使用經驗愈多，愈不支持政府規範，且達顯著 ($B = -.35, SE = .11, \text{Wald's } \chi^2(1) = 10.66, p = .001$)，因此假設 5 獲得證實，研究結果詳如表六。

4、X 世代

同樣地，本研究檢驗 X 世代 ChatGPT 使用經驗與支持政府規範兩者間的關係。研究結果顯示，整體模式的預測力達顯著： $\chi^2(9) = 25.34, p = .003$ 。Hosmer-Lemeshow 適配度檢定亦顯示模型與資料未存在顯著差異 ($\chi^2(8) = 10.02, p = .264$)。X 世代民眾的資料顯示，他們的 ChatGPT 使用經驗與會支持政府規範 AI 科技呈顯著負相關 ($B = -.52, SE = .15, \text{Wald's } \chi^2(1) = 12.14, p < .001$)，故假設 5 獲得支持，研究結果詳如表六。

整體而言，台灣整體民眾的 ChatGPT 使用經驗可負向顯著預測 AI 科技法律規範的態度。進一步剖析不同世代的差異，研究結果顯示，千禧世代與 X 世代的 ChatGPT 使用經驗可負向顯著預測 AI 科技法律規範的態度，但 Z 世代則無此效果。

(三) 假設 6 檢定

假設 6 預測 AI 素養是否可以負向預測支持政府規範 AI 科技，整體民眾與各世代研究結果如下：

1、整體民眾

假設 6 預測 AI 素養是否可以負向預測支持政府規範 AI 科技。研究結果顯示，整體模式的預測力達顯著： $\chi^2(9) = 89.71, p < .001$ ；Hosmer-Lemeshow 適配度檢定，結果顯示模型與資料存在顯著差異，推論時需審慎解讀 ($\chi^2(8) = 27.82, p = .001$)。進一步檢驗台灣整體民眾 AI 素養三項能力對支持政府規範 AI 科技的影響，研究結果顯示：AI 素養中的覺察能力雖然對支持政府規範有正向影響，但未達統計上的顯著水準 ($B = .14, SE = .07, \text{Wald's } \chi^2(1) = 3.76, p = .053$)；使用能力則呈負向關係 ($B = -.06, SE = .07, \text{Wald's } \chi^2(1) = .72, p = .397$)，同樣未具顯著性；評估能力對支持政府規

範亦為負向且不顯著的影響 ($B = -.09$, $SE = .08$, Wald's $\chi^2(1) = 1.40$, $p = .237$)，故假設 6(a)、6(b)、6(c) 未成立，研究結果顯見表五。

2、Z 世代

進一步檢驗不同世代的差異，Z 世代整體模式的預測力未達顯著： $\chi^2(9) = 11.83$, $p = .223$ 。Hosmer-Lemeshow 適配度檢定，結果顯示 $\chi^2(8) = 9.54$, $p = .299$ ，模型與資料未存在顯著差異。AI 素養的三項能力對支持政府規範的影響，分析結果如下：覺察能力 ($B = .14$, $SE = .17$, Wald's $\chi^2(1) = .63$, $p = .427$) 與使用能力 ($B = .02$, $SE = .18$, Wald's $\chi^2(1) = .02$, $p = .901$) 對支持政府規範有正向影響，但皆未具顯著性；然而，評估能力對支持政府規範則有負向顯著影響 ($B = -.43$, $SE = .21$, Wald's $\chi^2(1) = 4.12$, $p = .042$)。故假設 6(a)、6(b) 未獲得證實，假設 6(c) 獲得證實，研究結果詳見表六。

3、千禧世代

其次檢驗千禧世代的閱聽眾，整體模式的預測力達顯著： $\chi^2(9) = 29.42$, $p = .001$ 。Hosmer-Lemeshow 適配度檢定，結果顯示 $\chi^2(8) = 20.08$, $p = .010$ ，模型與資料存在顯著差異，推論時需審慎解讀。AI 素養的三個面向對政府規範的影響，研究結果如下：AI 素養中的覺察能力雖對支持政府規範有正向影響，但未達統計上的顯著水準 ($B = .07$, $SE = .12$, Wald's $\chi^2(1) = .38$, $p = .537$)；AI 素養的使用能力 ($B = -.14$, $SE = .10$, Wald's $\chi^2(1) = 1.93$, $p = .165$) 與評估能力 ($B = -.05$, $SE = .11$, Wald's $\chi^2(1) = .19$, $p = .664$) 對支持政府規範雖有負相關，但同樣未具顯著性。故假設 6(a)、6(b)、6(c) 均未獲得證實。研究結果詳見表六。

4、X 世代

本研究延伸分析 X 世代的效果，整體模式的預測力達顯著： $\chi^2(9) = 25.34$, $p = .003$ 。Hosmer-Lemeshow 適配度檢定，結果顯示 $\chi^2(8) = 10.02$, $p = .264$ ，模型與資料未存在顯著差異。進一步檢視 AI 素養三項能力與支持政府規範 AI 科技之關聯，AI 素養的覺察能力 ($B = .22$, $SE = .11$, Wald's $\chi^2(1) = 3.62$, $p = .057$) 與使用能力 ($B = .08$, $SE = .12$, Wald's $\chi^2(1) = .46$, $p = .500$) 對支持政府規範雖有正向影響，但同樣未達顯著；評估能力則呈現負向關係，但未達顯著 ($B = -.04$, $SE = .13$, Wald's $\chi^2(1) = .12$, $p = .734$)。故假設 6(a)、6(b)、6(c) 均未獲得證實。研究結果詳見表六。

整體而言，台灣整體民眾的三項 AI 素養能力均無法顯著預測 AI 科技法律規範的

態度，千禧世代與 X 世代民眾 AI 素養中的三項能力亦無法顯著預測 AI 科技法律規範的態度；然而，Z 世代的評估能力則對 AI 科技法律規範的態度呈現負向關係。

表五：台灣整體民眾羅吉斯迴歸分析結果

	支持政府規範 AI 科技		
	B(SE)	Wald(df)	Exp(B)
階層 1			
性別 (男性= 1)	-.14(.13)	1.05	.87
年齡	.04(.01)***	44.84	1.04
教育程度	.18(.06)**	8.35	1.20
網路使用	.24(.11)*	4.69	1.27
數位語音助理	-.02(.06)	.17	.98
Pseudo R ²	.06		
階層 2			
		-.06	-1.73
ChatGPT使用經驗	-.30(.07)	19.94	.74
AI 素養：覺察	.14(.07)	3.76	1.15
AI 素養：使用	-.06(.07)	.72	.94
AI 素養：評估	-.09(.08)	1.40	.91
Pseudo R ²	.09		
Block 3			
第三人感知	.02(.05)	.14	1.02
Pseudo R ²	.09		

- 註：1. * $p < .05$, ** $p < .01$, *** $p < .001$ ；
 2. 依變項為支持政府規範
 3. 整體民眾： $\chi^2(10) = 89.85, p < .001$

表六：各世代民眾階層羅吉斯迴歸

	Z 世代 (1997-2012 年出生)			千禧世代 (1981-1996 年出生)			X 世代 (1965-1980 年出生)		
	B(SE)	Wald(df)	Exp(B)	B(SE)	Wald(df)	Exp(B)	B(SE)	Wald(df)	Exp(B)
階層 1									
性別 (男性=1)	-.12(.27)	.18	.89	-.31(.20)	2.42	.74	-.02(.25)	.01	.98
年齡	.06(.05)	1.34	1.06	.04(.02)*	4.32	1.05	.06(.03)*	4.88	1.06
教育程度	.13(.15)	.77	1.14	.13(.10)	1.57	1.14	.18(.11)	2.89	1.20
網路使用	.32(.22)	2.09	1.37	.21(.17)	1.45	1.23	.23(.19)	1.49	1.26
數位語音助理	.08(.13)	.32	1.08	-.06(.08)	.61	.94	-.07(.10)	.47	.93
Pseudo R ²	.02			.02			.03		
階層 2									
ChatGPT	-.04(.13)	.11	.96	-.35(.11)**	10.66	.70	-.52(.15)***	12.14	.59
AI 素養：覺察	.14(.17)	.63	1.15	.07(.12)	.38	1.08	.22(.11)	3.62	1.24
AI 素養：使用	.02(.18)	.02	1.02	-.14(.10)	1.93	.87	.08(.12)	.46	1.08
AI 素養：評估	-.43(.21)*	4.12	.65	-.05(.11)	.19	.95	-.04(.13)	.12	.96
Pseudo R ²	.03			.06			.07		
階層 3									
第三人效果感知	.17(.12)	1.91	1.18	.001(.08)	<.001	1.00	-.02(.10)	.04	.98
Pseudo R ²	.05			.07			.08		

註：* = $p < .05$, ** = $p < .01$, *** = $p < .001$ ；依變項：支持政府規範

Z 世代： $\chi^2(10) = 13.76, p = .184$

千禧世代： $\chi^2(10) = 29.42, p = .001$

X 世代： $\chi^2(10) = 25.38, p = .005$

四、不同年齡世代之 ChatGPT 使用經驗、AI 素養、第三人感知之差異

為了檢驗研究問題 2，本研究以單因子變異數分析來檢驗不同世代對於 ChatGPT 使用經驗、AI 素養、第三人感知是否存在差異。研究結果顯示如下：

（一）ChatGPT 使用經驗：

在 ChatGPT 的使用經驗方面，不同世代的使用經驗有顯著差異： $F(2, 975) = 37.48$, $p < .001$ ，Z 世代 ChatGPT 的使用經驗（ $M = 2.31$, $SD = 1.71$ ）顯著高於千禧世代（ $M = 1.71$, $SD = 1.05$ ）及 X 世代（ $M = 1.53$, $SD = .96$ ），同時千禧世代的使用經驗也顯著高於 X 世代。

（二）AI 素養：

不同世代的 AI 素在覺察、使用及評估等三個面向均存在顯著差異。在 AI 覺察能力方面（ $F(2, 975) = 31.41$, $p < .001$ ），Z 世代（ $M = 4.05$, $SD = .88$ ）顯著高於千禧世代（ $M = 3.98$, $SD = .95$ ）與 X 世代（ $M = 3.43$, $SD = 1.28$ ）；同時，千禧世代的 AI 覺察能力亦高於 X 世代。在 AI 使用能力方面（ $F(2, 975) = 32.18$, $p < .001$ ），事後比較結果顯示，Z 世代（ $M = 3.98$, $SD = .96$ ）顯著高於千禧世代（ $M = 3.43$, $SD = 1.25$ ）與 X 世代（ $M = 3.17$, $SD = 1.29$ ），千禧世代的 AI 使用能力亦高於 X 世代；在 AI 評估能力方面（ $F(2, 975) = 26.43$, $p < .001$ ），事後比較結果顯示，Z 世代（ $M = 4.13$, $SD = .82$ ）顯著高於千禧世代（ $M = 3.86$, $SD = 1.10$ ）與 X 世代（ $M = 3.47$, $SD = 1.21$ ），千禧世代的 AI 評估能力亦高於 X 世代。

（三）第三人感知：

不同世代的生成式 AI 風險的第三人感知存在顯著差異（ $F(2, 975)$, $p < .05$ ），事後比較結果顯示，千禧世代的第三人感知（ $M = 1.49$, $SD = 1.27$ ）顯著高於 Z 世代（ $M = 1.20$, $SD = 1.20$ ），但與 X 世代（ $M = 1.31$, $SD = 1.29$ ）並無顯著差別。

綜合上述結果，在 ChatGPT 的使用經驗上，Z 世代的使用經驗最高，千禧世代次之，X 世代的 ChatGPT 使用經驗最低；在 AI 素養的覺察、使用及評估等三大面向能力，Z 世代擁有最高的 AI 素養，千禧世代僅次於 Z 世代，X 世代的 AI 素養最低；第三人感知中，千禧世代的第三人感知顯著高於 Z 世代，Z 世代與 X 世代則無存在顯著

表七：台灣不同世代在 ChatGPT 使用經驗、AI 素養、生成式 AI 風險感知之差異

	<i>M</i>	<i>SD</i>	<i>F</i> (df)	Games-Howell Test
ChatGPT				
1. Z 世代	2.31	1.30		
2. 千禧世代	1.71	1.05	F(2, 975) =37.48 ***	1>2,3; 2>3
3. X 世代	1.53	.96		
AI 素養：覺察				
1. Z 世代	4.05	.88		
2. 千禧世代	3.98	.95	F(2, 975) = 31.41***	1>3; 2>3
3. X 世代	3.43	1.28		
AI 素養：使用				
1. Z 世代	3.98	.96		
2. 千禧世代	3.43	1.25	F(2, 975) =32.18 ***	1>2, 3; 2>3
3. X 世代	3.17	1.29		
AI 素養：評估				
1. Z 世代	4.13	.82		
2. 千禧世代	3.86	1.10	F(2, 975) = 26.43 ***	1>2, 3; 2>3
3. X 世代	3.47	1.21		
第三人感知				
1. Z 世代	1.20	1.20		
2. 千禧世代	1.49	1.27	F(2, 975) = 4.64*	2>1
3. X 世代	1.31	1.29		

註：* = $p < .05$, ** = $p < .01$, *** $p < .001$

差異。詳細數據見表七。

伍、討論

由於生成式 AI 是透過大型語言模型（Large Language Models, LLMs）來預測並生成近似於人類語言的內容（Endert, 2024），當中可能存在產生虛假資訊以及帶有偏見的內容（Vartak, 2023），進而對民眾帶來負面的影響（e.g. Armstrong, 2023；Common Sense Media, 2025；Melnick, 2024），因此本研究透過《2024 台灣網路報告》的調查資料，以第三人效果作為理論基礎，以了解生成式 AI 如何影響到民眾對他人預設影響力的評估，並進一步探討這樣的第三人感知如何影響後到後續的行為結果。

一、主要研究發現與啟示

（一）台灣民眾對生成式 AI 風險存在第三人感知

本研究證實台灣民眾對於生成式 AI 存在第三人感知（H1）。考量到生成式 AI，特別是 ChatGPT 對於民眾生活與工作所扮演的關鍵角色，本研究著眼於 ChatGPT 可能生成錯誤資訊或是帶有偏見的內容如何塑造台灣民眾的認知與反應。研究結果顯示，台灣整體民眾及三個世代的台灣民眾普遍認為 ChatGPT 所引發的科技風險對他人的影響遠大於對個人的影響，這項發現能呼應過去於假新聞與網路隱私的研究結果相似（e.g. Chen & Atkin, 2021；Chung, & Kim, 2021；Lyons, 2022），反映民眾對生成式 AI 科技風險的擔憂，不僅擴充第三人效果在 AI 風險的適用性（e.g. Chung, Kim, Jones-Jang, Choi & Lee, 2025），也增進理論探討不同世代對 AI 風險的認知差異。

（二）ChatGPT 使用經驗無法預測第三人感知的影響

關於第三人感知的預測因子，本研究有了意外的發現：台灣整體民眾的 ChatGPT 使用經驗（H2）無法顯著影響第三人感知。進一步區別為 Z 世代、千禧世代與 X 世代的樣本資料顯示，Z 世代與千禧世代的 ChatGPT 使用經驗與第三感知呈負相關，但未具顯著性；換言之，當他們 ChatGPT 的使用經驗愈豐富，愈會消弭感知 AI 對個人影響與對他人影響評估之間的落差。可能的解釋原因是，由於 ChatGPT 仍在早期採用的階段，民眾仍可能還在摸索與認識 AI 的風險，因此當他們 ChatGPT 的使用經

驗愈多，愈了解 AI 可能生成不實資訊以及帶有偏見的內容，因此認為個人更有機會暴露在 AI 科技的負面影響中。儘管這樣的結果與自我強化的觀點相反，但過去亦有文獻支持類似的研究結果。例如：Huang（2023）發現在疫情期間，當個人暴露在愈多 Covid-19 相關的新聞，愈可能導致負面情緒，因而認為自己容易受到疫情的負面影響。因此，當 ChatGPT 使用經驗較豐富的人，在評估個人與他人所受到的 AI 風險影響程度時，兩者的差異可能因此減少。相反地，X 世代 ChatGPT 使用經驗與第三感知呈正向關聯性，同樣未達統計上的顯著水準，意即當他們 ChatGPT 經驗愈豐富，愈認為他人容易受到 AI 科技的負面影響，這項結果大致符合自我強化的觀點，但可能因為 X 世代對數位科技較不熟悉，未必有一定自信認為個人可以清楚辨識 AI 科技生成的資料的真偽，因此在評估個人與他人受到 AI 負面影響並無顯著差異。

（三）AI 素養對第三人感知的影響僅限於 X 世代

關於 AI 素養（H3）對第三人感知的影響，台灣整體的樣本顯示，AI 素養的三項能力與對第三人感知雖呈現負相關，但均未達統計上的顯著水準。然而，進一步區分為不同世代後可發現具體差異：X 世代的使用能力可負向顯著預測第三人感知，Z 世代與千禧世代的三項 AI 素養能力與第三人感知的關係則不顯著。此結果顯示，X 世代民眾具備愈高的 AI 素養使用能力，其第三人感知也會減少。可能的解釋原因在於，高 AI 素養的民眾與低 AI 素養的民眾，在評估 AI 科技風險對個人與他人的評估可能差異不大。儘管第三人效果是基於個人評估他人所受到的媒介影響時所引發的認知扭曲（e.g. Perloff, 2009），對於個人抵擋媒介負面影響的能力存有樂觀偏誤，但 Shen, Sun 與 Pan（2018）提出不同的觀點，認為第三人效果未必是認知扭曲，有時人們在評估個人與他人之間所受到的影響落差可能是準確的。例如，當人們評估媒介對兒童的影響，普遍會判斷媒介對兒童的影響力應大於成人，而這樣的判斷結果應為準確，原因是兒童相對於成年人不具備同等的智識、能力與媒介使用經驗來判斷與抵禦媒介風險。AI 素養的使用能力，代表個人有能力運用 AI 產品或服務以達成個人想要的目的，因此，當 X 世代民眾愈有能力去使用 AI 科技，他們對於 AI 科技的風險評估可能愈趨向準確，愈能知曉 AI 科技隱藏著會生成錯誤資訊的風險，反而降低樂觀偏誤，開始認知到自身也有機會暴露於風險中，而消弭個人與他人之間的風險評估落差。

（四）第三人感知無法解釋法律規範態度

本研究發現無論台灣總體本，或是區別為三個世代民眾的樣本，民眾生成式 AI 風險第三人感知皆無法顯著預測支持政府規範 AI 科技（H4）。本研究推測，造成民眾法律規範態度不顯著的原因，很可能在於目前台灣社會上對於法律規範人工智慧科技的討論並不常見，因此民眾的看法仍不具體尚未。事實上，由於 AI 技術變化快速，國際情勢瞬息萬變，不只是台灣政府對於「人工智慧基本法」尚未有共識（余弦妙、歐芯萌，2025），歐盟以及英國也延後 AI 監管法案的執行（Courea & Stacey, 2025；Supantha, 2025），目前的美國政府亦鬆綁前任拜登政府時期的 AI 政策及法規（賴又豪，2024），顯示 AI 相關的法律監管是個相當複雜的議題，各國政府對於制定 AI 科技監管的相關法律普遍抱持著謹慎的態度。再者，因應 AI 技術持續快速變化，不同形式的 AI 科技，包含辨識型、生成式以及未來可能發展的 AI 代理人，各種 AI 科技的監管方式不同（林敬殷，2025），實際該如何立法監管，例如：實質罰則、監管力度與風險分級（余弦妙、歐芯萌，2025），各界仍在討論與協商。因此，儘管民眾普遍對於 AI 科技的潛在危害抱持第三人感知，仍持續觀望 AI 的發展，因此無法倉促決定是否支持法律規範。

不過，本研究也有另外一個意外發現，即是儘管都未達統計上的顯著水準，Z 世代與千禧世代，第三人感知與規範行為呈現正相關，但 X 世代第三人感知與規範行為反而呈現負相關，可能的原因在於相較於另外兩個世代，X 世代具備最低的 ChatGPT 使用經驗與 AI 素養，因此自認其對 AI 的掌控感較弱，因此，相較於其他兩個世代是基於利他動機（保護他人不受到 AI 負面影響）而支持法律規範的原因，X 世代民眾可能更基於自我保護動機，而支持 AI 科技監管來保護個人不受到 AI 科技的負面影響。

此外，本研究所關注的第三人效果僅有政府以法律規範 AI 科技一項，雖然效果不如預期顯著，但也不排除對 AI 風險的第三人感知仍有可能影響其它類型的行為，比如媒體素養教育的推動（e.g. Chung et al., 2025；Jang & Kim, 2018）過往研究（Jang & Kim, 2018）即指出，在假新聞議題上，第三人感知會負向預測規範假新聞，但會正向預測推動媒體素養教育。Chung 等人（2025）也有類似的發現，面對 AI 科技（即

ChatGPT)，民眾的第三人感知無法顯著預測支持規範 AI 科技，但是可以正向預測支持政府推動 AI 素養教育。這些研究結果顯示，即使是相同的風險類型，仍可能引發閱聽眾不同的行為反應。面對生成式 AI 可能「生成不實資訊」的風險，比起透過法律規範 AI 科技，民眾或許更願意採取其他行動（例如：支持媒體素養教育），藉此提升他人的媒體素養能力而不受到媒體風險的影響。因此，本研究建議未來的研究可以調查閱聽眾不同行為面的反應，以期更周延的理解使用者因應 AI 科技風險感知所採取的行為。

（五）ChatGPT 使用經驗對法律規範態度的影響限於千禧世代、X 世代

本研究發現，總體資料中難以察覺 ChatGPT 使用經驗與法律規範態度之間的明確關係，但進一步區別為不同世代後可發現具體差異，即千禧世代與 X 世代 ChatGPT 的使用經驗愈豐富，愈不支持政府規範 AI 科技（H5）。這樣的研究結果也呼應了自我強化的觀點，當人們對於生成式 AI 科技的使用經驗愈多，他們會更有自信自己比他人更能善用工具，而視相關的規範與個人無關，而這樣的想法也促使他們不希望政府對所有人（包含自己）採取規範 AI 科技的措施。然而，Z 世代 ChatGPT 的使用經驗與支持政府規範 AI 科技呈負相關，但未達統計上的顯著水準，或許反映出 Z 世代民眾的使用經驗差異，令他們對於法律規範 AI 的態度有所分歧。本研究資料顯示，Z 世代民眾的 ChatGPT 使用經驗顯著高於其他兩個世代，這與國外的發現相似（e.g. Carroll, 2025；Lake, 2025），這些 ChatGPT 的活躍用戶也許將 ChatGPT 視為「生活顧問」來解決各種疑難雜症，在重度使用下，可能更重視使用的自由，而不希望透過法律限制。不過，另一方面，部分使用者或許也因使用經驗的累積，更常發現生成式 AI 失準或失靈的情況，因此可能對於法律規範持較為樂觀的態度。未來的研究可嘗試以其他研究方法（例如：深度訪談法），深入了解 Z 世代對法律規範 AI 科技的看法。

（六）AI 素養對法律規範態度在 Z 世代中突出

從整體樣本資料觀察，AI 素養對法律規範態度並沒有顯著影響，但是區別為不同世代後，AI 素養與法律規範態度之間的關係在特定族群中突出。本研究結果顯示，Z 世代 AI 素養中的評估能力（H6c）可負顯著預測支持政府規範 AI 科技，由於 AI 素養的評估能力是指個人在應用 AI 科技產服務時，能清楚知道其優缺點。因此對於 Z 世代的民眾而言，他們自認個人清楚了解 AI 科技的長處與短處，因此深信自己不會受

到生成式 AI 產出不正確資訊的負面影響，而不希望政府過度干預 AI 科技。然而，千禧世代與 X 世代的三項 AI 素養能力皆無法顯著預測對支持政府規範 AI 科技。其可能的原因在於，對於 AI 素養較高的民眾，透過限制性措施來規範 AI 科技並不一定能有效保護他人免受到 AI 科技風險。過去文獻（Chung et al., 2025）亦指出，對於 AI 素養較高的民眾而言，比起支持規範 AI 科技，他們反而更支持政府推動 AI 素養教育，讓其他使用者學會更安全、負責任且聰明使用 AI 科技，他們認為這不僅對於保護他人不受到 AI 科技的負面風險更有助益，同時也能維護他們使用生成式 AI 的自由。

（七）ChatGPT 使用經驗、AI 素養、第三人感知以及法律規範態度的世代差異

本研究發現不同世代在 ChatGPT 使用經驗、AI 素養、第三人感知與法律規範態度展現出各具特色的使用模式與行為反應。首先，在 ChatGPT 使用以及 AI 素養方面，Z 世代作為第一個普遍應用 AI 科技於學業、職場與娛樂的世代（e.g. Chung et al., 2025），不僅 ChatGPT 的使用頻率最高，AI 素養程度亦顯著高於另外兩個世代。千禧世代在 ChatGPT 使用以及 AI 素養僅次於 Z 世代。至於做為本研究最年長的 X 世代，在 ChatGPT 使用經驗與 AI 素養上，顯著低於 Z 世代及千禧世代，這樣的現象與過去的發現相符，熟齡人口在 AI 科技的使用存在學習門檻（S. Wang et al., 2019），同時也是最容易受到 AI 負面影響的群體（C. Wang et al., 2024）。

其次，在第三人感知方面，相較於 ChatGPT 使用經驗有限與 AI 素養最低的 X 世代，千禧世代對於生成式 AI 科技風險存有較高的風險預期，顯現高度的公共風險意識。本研究推測，可能因為千禧世代成長於網路爆炸與數位轉型的時代，他們長期處於資訊高度流通與科技快速演進的環境，因此可能對於新興科技所伴隨的風險抱持更高的警覺性，在實際使用 ChatGPT 等生成式 AI 科技時，傾向以更現實、謹慎的態度來看待 AI 科技可能帶來的負面影響。而作為 ChatGPT 使用以及 AI 素養程度最低的 X 世代，儘管其 AI 科技的掌握能力不及其他世代，一旦接觸新科技，他們會迅速連結自身價值以及社會責任的判斷，因此對於相關的風險意識抱持更不樂觀的態度，傾向認為自身與他人都有機會受到 AI 科技的負面風險，因此與千禧世代相比，第三人感知沒有更突出的態度。

最後，在法律規範方面，Z 世代與相較年長的千禧世代及 X 世代，則是基於不同的因素影響法律規範態度。具備不同 AI 素養程度的 Z 世代，對於法律規範態度抱持

迥異的看法。儘管 AI 素養程度較高的 Z 世代民眾更不傾向支持法律規範，但對於 AI 素養程度較低的 Z 世代民眾，反而更支持法律規範。換言之，低 AI 素養的 Z 世代可能更基於自我保護的心理動機而支持法律規範。該研究發現也回應了 Schmierbach 等人（2023）對於第三人效果研究的反思：第三人感知固然重要，但並非這樣的「認知落差」直接導致行動，而是這樣的「認知落差」和行動都反映了同一個深層的心理特質（例如：自我保護動機或樂觀偏誤），而促成行動；同時，也呼應了其他學者對第三人效果研究的觀點（e.g. Baek, Kang & Kim, 2019; Chung & Wihbey, 2024; Ho, Goh & Leung, 2022），不只是「認知落差」，媒介對「個人」的影響也同樣會造成他們行為面的改變。

相對地，千禧世代以及 X 世代反而是個人的 ChatGPT 使用經驗會影響到他們法律規範的態度，具備不同 AI 素養的民眾對於法律規範的態度反而不具顯著的差異。儘管千禧世代的 ChatGPT 的經驗與 AI 素養不及 Z 世代頻繁，但他們見證過許多新舊媒體的變遷，或許認為自身對於新興科技的風險仍具有一定敏感度，因此當他們 ChatGPT 使用經驗愈多，愈提升對自身掌握 AI 科技的信心程度，反而展現出更低的法律規範態度。然而，對 ChatGPT 使用經驗有限的 X 世代而言，部分觀點（e.g. Wartgow, 2024）認為若他們可以透過 AI 工具能平衡他們生活與職場工作，幫助個人持續學習與提升技能，反而會對 AI 科技抱持更開放的態度。此觀點與本研究發現相近，可以預期善用 ChatGPT 使用經驗的 X 世代，可能對 AI 科技抱持更包容態度，因而不傾向透過法律規範來限制 AI 科技。

整體來說，本研究發現不同世代在 ChatGPT 使用經驗、AI 素養、第三人感知與法律規範態度呈現明顯的差異。由於數位落差包含個人在近用科技的動機、實體存取、技能與使用等層面（Y. Chang et al., 2016；van Dijk, 2002；2006）。本研究結果與 Celik（2023a）的發現相似，當使用者對 AI 科技具備較高的近用動機、機會、技能與使用，其 AI 素養往往愈高；其他類似的研究結果（Celik, 2023b）亦顯示，在教學現場中，教師對 AI 的科技知識越多，愈能在倫理層面上評估不同 AI 工具的決策結果。進一步推論，不同世代可能存在數位落差，因而影響到使用者的 ChatGPT 使用經驗、AI 素養程度以及對 AI 科技潛在風險的感知。建議未來的學者在探究其他 AI 科技風險議題時（例如：隱私外洩、深偽技術、就業衝擊），也應納入數位落差的面向，以

更全面了解影響閱聽眾風險認知的因素。

總結而言，本研究檢驗了台灣民眾對生成式 AI 科技風險引發的第三人感知之影響因子與後續行為意圖，就理論貢獻而言，研究結果不僅再度驗證第三人效果應用在 AI 科技風險的適用性，更發現總體資料可能掩蓋不同世代的特徵與行為表現，透過區分為不同世代後，有助於理解不同世代回應生成式 AI 科技風險認知與反應的異質性。就實務面而言，本研究發現突顯了 AI 素養培育的重要性，相較於透過法規限縮民眾 AI 科技的使用，推動全民 AI 素養教育或許更有助於保護民眾不受到 AI 科技的風險，同時也能維護民眾使用 AI 科技的自主性。此外，本研究亦發現 ChatGPT 在不同世代間的採用可能存在數位落差，因此，未來若要制定提升 AI 素養的相關政策時，應聚焦於不同世代的數位落差近用面向（動機、實體存取、技能以及使用），以期能貼近不同世代的需求。

二、研究限制與建議

首先，由於台灣網路報告屬於全國性的大規模調查，整體調查題數較多，因此各題組在問卷設計時有數量限制，主要變項以單一題項而非使用多個題目（量表）來測量；儘管使用單一題目能大幅縮減問卷長度及節省作答時間，但是單一題向無法構成量表，因此無法測量信度，測量結果的信效度可能因此受限，後續研究結果的討論與推論應更為謹慎。特別是第三人效果所引發的行為面不只使用單題，且選項僅包含支持vs.不支持等兩個選項，可能不如個別詢問受訪者對於政府規範方式的同意程度來的精確，因此測量結果難免有解釋上的侷限。未來學者建議可以使用更完整的量表再次驗證本研究發現，考慮以多個題目來測量各個構面，以確保資料的品質與研究的嚴謹性。

第二，本研究只關注於單一行為面的影響（AI 科技規範），相同的媒介所衍生的不同風險類型，很有可能透過第三人感知引發民眾採取不同的行為意圖。例如：糾正性以及推廣性行為（Sun et al., 2008）。由於生成式 AI 可能生成錯誤資訊，過去關於假新聞的第三人效果研究顯示，當民眾認為假新聞會對他人產生負面後果時，他們會採取糾正行為（e.g. 路森等學者，2023；Chung, & Kim, 2021；Lyons, 2022）。因此，

未來的可納入更多元的風險主題，以近一步了解民眾面對不同風險類型時所促發的行為，這不僅可以拓展第三人效果於 AI 風險範疇中可能引發的行為效果，也能提供政府與教育單位有效的政策性建議。

第三，受限於人力、經費以及時間能力所及，而未能涵蓋所有可能影響第三人感知與民眾行為面之間關係的潛在變數。首先，民眾對於 AI 科技的態度可能會影響閱聽眾對於 AI 科技的信任感以及持續採用意圖，例如：AI 恐懼 (Molina & Sundar, 2024) 及 AI 科技態度 (van Deursen & Dijk, 2019)。特別是機器捷思 (machine heuristic; S. Lee, Oh & Moon, 2023 ; Molina & Sundar, 2024 ; Sundar & Kim, 2019) 強調使用者透過對機器的既有印象來判斷其表現，因此具備較高正向機器捷思的使用者，對於 AI 科技的信任感愈高，可預期他們對 AI 科技的風險感知可能與一般使用者有所差異。其次，人們對於議題的主觀嚴重性以及主觀效能同樣會影響其對議題的風險感知。具體而言，當閱聽眾知覺風險的主觀嚴重性高、並提供具體預防措施，人們會認為議題更為真實且切身相關，因而認為風險對於個人的影響遠大於他人的影響 (H. Lee & Park, 2016)。再者，民眾對於 AI 科技發行公司 (e.g. Open AI, Meta & Google) 的信任程度 (e.g. NL Times, 2025)，以及對政府施政的信心程度 (e.g. Jang & Kim, 2018 ; Lin, 2014) 同樣也是影響民眾風險感知的潛在影響變數。最後，為了更細緻理解不同世代對 AI 科技的認知與反應，未來的研究可納入數位落差的能力面向 (e.g. Y. Chang et al., 2016 ; van Dijk, 2002 ; 2006)，以期更完整理解促成閱聽眾行為改變的心理動機。

第四，本研究採用次級資料進行分析，主要基於《2024 台灣網路報告》為大型調查，以了解台灣民眾的網路近用以及網路應用服務的使用現況，該資料具有大樣本且較具代表性的特質，能更有效反映台灣民眾對於 AI 科技風險的感知與法律規範態度。然而，既有資料亦存在限制，例如：問卷無法完全依照研究者目的量身設計題項，因而影響研究結果的深度與細緻性。建議未來研究可自行設計問卷，以驗證與擴及本研究發現。

第五，本研究僅選擇 ChatGPT 做為代表以測量台灣民眾的生成式 AI 使用經驗。然而，現今的生成式 AI 科技的發展已如雨後春筍蓬勃發展，針對不同工作任務亦有不同更符合需求的 AI 工具 (例如：生成式圖像 Midjourney、即時簡報生成設計工具

Gamma等等），此外，隨著生成式 AI 的發展，目前出現新興程式設計模式：「vibe coding」。透過 AI 工具協助，開發者可透過向 AI 描述需求，由 AI 自動生成程式碼，強調完全順應感覺（vibe），不再關心程式碼的細節，而更關注於設計和想法表達（A. Yang, 2025）。因此，AI 科技在不同平台與工具的分界日趨模糊，因此未來在測量 AI 科技的使用經驗宜更謹慎去定義。最後，本研究僅以橫斷面資料進行調查，未來可嘗試以縱貫資料，以了解長期來看 AI 科技對台灣民眾評估與影響的趨勢變化。

參考文獻

- 中央社。(2025 年 4 月 25 日)。<〈韓國發布人工智慧基本法，促進人工智慧產業發展〉。取自<https://www.cna.com.tw/postwrite/chi/400092>
- 余弦妙、歐芯萌。(2025 年 6 月 11 日)。<〈AI 基本法 14 版本大亂鬥 仍具有三大分歧〉。《聯合新聞網》。取自 <https://tw.news.yahoo.com/ai%E5%9F%BA%E6%9C%AC%E6%B3%9514%E7%89%88%E6%9C%AC%E5%A4%A7%E4%BA%82%E9%AC%A5-%E4%BB%8D%E5%85%B7%E6%9C%89%E4%B8%89%E5%A4%A7%E5%88%86%E6%AD%A7-224044664.html>
- 林敬殷。(2025 年 6 月 11 日)。<〈立院審人工智慧基本法 決議要求數發部 7/15 提版本〉。《中央社》。取自：<https://www.cna.com.tw/news/aip/202506110172.aspx>
- 吳泰毅、鄧玉玲。(2023)。<〈初探台灣民眾對人工智慧產品與服務之採用經驗與信任感〉，《資訊社會研究》，44：97-128。[https://doi.org/10.29843/JCCIS.202301_\(44\).0004](https://doi.org/10.29843/JCCIS.202301_(44).0004)
- 徐子苓。(2024 年 8 月 19 日)。<〈AI 基本法草案出爐 黃彥男：年底推 AI 風險分級框架〉。《自由時報》。取自 <https://ec.ltn.com.tw/article/breakingnews/4773743>
- 財團法人台灣網路資訊中心 (2024)。<《2024 年台灣網路報告》。取自 <https://report.twNIC.tw/2024/>
- 許馥嘉、吳泰毅 (2025)。<〈AI 不釋手：初探台灣民眾對會話型 AI 的期望確認與持續使用意圖〉，《資訊社會研究》，48：23-58。[https://doi.org/10.29843/JCCIS.202501_\(48\).0002](https://doi.org/10.29843/JCCIS.202501_(48).0002)
- 經濟部台日產業合作推動辦公室。(2024年8月13日)。<〈日本政府考慮制定法規，以國家戰略平衡 AI 的使用和風險〉。取自：<https://www.tjpo.org.tw/NewsDetail.aspx?id=1172>
- 路森、羅文輝、魏然。(2023)。<〈新冠疫情虛假資訊的接觸頻率、預設影響與香港市民對虛假資訊的態度與行為〉，《傳播與社會學刊》，65：191-217。[https://doi.org/10.30180/CS.202307_\(65\).0008](https://doi.org/10.30180/CS.202307_(65).0008)

- 賴又豪。(2024年12月25日)。〈告別 DEI，創新競爭掛帥：從拜登到川普 2.0 的美國 AI 政策重構〉。《轉角國際》。取自 https://global.udn.com/global_vision/story/8663/8446105
- 關鍵評論網。(2023 年 9 月 25 日) 〈AI 時代的日本：探索民眾對人工智慧的理解與展望〉。取自 <https://www.thenewslens.com/article/192128>
- Ajzen, I. (1985). From intentions to actions: A theory of planned behavior (pp. 11–39). In J. Kuhl & J. Beckmann (Eds.), *Action control: From cognition to behavior*. Springer. https://doi.org/10.1007/978-3-642-69746-3_2
- Alkaissi, H., & McFarlane, S. I. (2023). Artificial hallucinations in ChatGPT: implications in scientific writing. *Cureus*, 15(2), e35179. <https://doi.org/10.7759/cureus.35179>
- Armstrong, K. (May 27, 2023). *ChatGPT: US lawyer admits using AI for case research*. BBC. <https://www.bbc.com/news/world-us-canada-65735769>
- BBC. (2017, August 27). *Meet the millennials: Who are Generation Y?*. <https://www.bbc.com/news/uk-scotland-41036361>
- Baek, Y. M., Kang, H., & Kim, S. (2019). Fake news should be regulated because it influences both “others” and “me”: How and why the influence of presumed influence model should be extended. *Mass Communication and Society*, 22(3), 301–323. <https://doi.org/10.1080/15205436.2018.1562076>
- Bensinger, G. (2024, August 27). *Big Tech wants AI to be regulated. Why do they oppose a California AI bill?*. Reuters. <https://www.reuters.com/technology/artificial-intelligence/big-tech-wants-ai-be-regulated-why-do-they-oppose-california-ai-bill-2024-08-21/>
- Boyle, M. P., McLeod, D. M., & Rojas, H. (2008). The role of ego enhancement and perceived message exposure in third-person judgments concerning violent video games. *The American Behavioral Scientist (Beverly Hills)*, 52(2), 165–185. <https://doi.org/10.1177/0002764208321349>
- Broussard, M., Diakopoulos, N., Guzman, A. L., Abebe, R., Dupagne, M., & Chuan, C.-H. (2019). Artificial intelligence and journalism. *Journalism & Mass Communication Quarterly*, 96(3), 673–695. <https://doi.org/10.1177/1077699019859901>

- Carroll, S. (2025, May 13). *Sam Altman's Gen Z brag: "They don't really make life decisions without asking ChatGPT"* Quartz. <https://qz.com/chatgpt-open-ai-sam-altman-genz-users-students-1851780458>
- Celik, I. (2023a). Exploring the determinants of artificial intelligence (Ai) literacy: Digital divide, computational thinking, cognitive absorption. *Telematics and Informatics*, 83, 102026. <https://doi.org/10.1016/j.tele.2023.102026>
- Celik, I. (2023b). Towards Intelligent-TPACK: An empirical study on teachers' professional knowledge to ethically integrate artificial intelligence (AI)-based tools into education. *Computers in human behavior*, 138, 107468. <https://doi.org/10.1016/j.chb.2022.107468>
- Chen, H., & Atkin, D. (2021). Understanding third-person perception about Internet privacy risks. *New Media & Society*, 23(3), 419-437. <https://doi.org/10.1177/1461444820902103>
- Chang, Y., Shahzeidi, M., Kim, H. & Park, M. C. (2012). Gender digital divide and online participation: A cross-national analysis. *The 19th ITS Biennial International Conference*, Bangkok, Thailand, November 18-21, 2012, <http://www.econstor.eu/dspace/handle/10419/72506>
- Chang, Y., Wong, S. F., & Park, M. C. (2016). A three-tier ICT access model for intention to participate online: a comparison of developed and developing countries. *Information Development*, 32(3), 226-242. <https://doi.org/10.1177/0266666914529294>
- Chung, M., & Kim, N. (2021). When I learn the news is false: How fact-checking information stems the spread of fake news via third-person perception. *Human Communication Research*, 47(1), 1-24. <https://doi.org/10.1093/hcr/hqaa010>
- Chung, M., & Wihbey, J. (2024). Social media regulation, third-person effect, and public views: A comparative study of the United States, the United Kingdom, South Korea, and Mexico. *New Media & Society*, 26(8), 4534-4553. <https://doi.org/10.1177/14614448221122996>
- Chung, M., Kim, N., Jones-Jang, S. M., Choi, J., & Lee, S. (2025). I see a double-edged sword: How self-other perceptual gaps predict public attitudes toward ChatGPT

- regulations and literacy interventions. *New Media & Society*, 14614448241313180. <https://doi.org/10.1177/14614448241313180>
- Cho, H., Shen, L., & Peng, L. (2021). Examining and extending the influence of presumed influence hypothesis in social media. *Media Psychology*, 24(3), 413-435. <https://doi.org/10.1080/15213269.2020.1729812>
- Common Sense Media. (2025, January 29). *Research Brief: Teens, Trust, and Technology in the Age of AI*. <https://www.common sense media.org/research/research-brief-teens-trust-and-technology-in-the-age-of-ai>
- Cook, R. D., & Weisberg, S. (1982). *Residuals and influence in regression*. Chapman & Hall.
- Courea, E. & Stacey, K. (2025, June 7). *UK ministers delay AI regulation amid plans for more “comprehensive” bill*. The Guardian. <https://www.theguardian.com/technology/2025/jun/07/uk-ministers-delay-ai-regulation-amid-plans-for-more-comprehensive-bill>
- Davison, W. P. (1983). The third-person effect in communication. *Public Opinion Quarterly*, 47(1), 1-15. <https://doi.org/10.1086/268763>
- Endert, J. (2024, March 26). *Generative AI is the ultimate disinformation amplifier*. DW Akademie. <https://akademie.dw.com/en/generative-ai-is-the-ultimate-disinformation-amplifier/a-68593890>
- European Parliament. (2023, August 06). *EU AI Act: first regulation on artificial intelligence*. <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence#what-parliament-wanted-in-ai-legislation-2>
- Faverio, M. & Tyson, A. (2023). *What the data says about Americans' views of artificial intelligence*. Pew Research Center. <https://www.pewresearch.org/short-reads/2023/11/21/what-the-data-says-about-americans-views-of-artificial-intelligence/>
- Field, A.P. (2018). *Discovering Statistics Using IBM SPSS Statistics*. 5th Edition, Sage, Newbury Park.

- Gupta, M., Akiri, C., Aryal, K., Parker, E., & Praharaj, L. (2023). From chatgpt to threatgpt: Impact of generative ai in cybersecurity and privacy. *IEEE Access*, 11, 80218-80245. <https://doi.org/10.1109/ACCESS.2023.3300381>.
- Ho, S. S., Goh, T. J., & Leung, Y. W. (2022). Let's nab fake science news: Predicting scientists' support for interventions using the influence of presumed media influence model. *Journalism*, 23(4), 910–928. <https://doi.org/10.1177/1464884920937488>
- Hoffner, C., Plotkin, R. S., Buchanan, M., Anderson, J. D., Kamigaki, S. K., Hubbs, L. A., ... & Pastorek, A. (2001). The third-person effect in perceptions of the influence of television violence. *Journal of Communication*, 51(2), 283-299. <https://doi.org/10.1111/j.1460-2466.2001.tb02881.x>
- Hong, Y. (2023). Extending the influence of presumed influence hypothesis: Information seeking and prosocial behaviors for HIV prevention. *Health Communication*, 38(4), 765-778. <https://doi.org/10.1080/10410236.2021.1975902>
- Huang, H. Y. (2023). Third-and first-person effects of COVID news in HBCU students' risk perception and behavioral intention: social desirability, social distance, and social identity. *Health Communication*, 38(13), 2956-2970. <https://doi.org/10.1080/10410236.2022.2129243>
- Ipsos. (2023, July). *Global Views on A.I. 2023: How people across the world feel about artificial intelligence and expect it will impact their life*. <https://www.ipsos.com/sites/default/files/ct/news/documents/2023-07/Ipsos%20Global%20AI%202023%20Report-WEB.pdf>
- Jang, S. M., & Kim, J. K. (2018). Third person effects of fake news: Fake news regulation and media literacy interventions. *Computers in Human Behavior*, 80, 295-302. <https://doi.org/10.1016/j.chb.2017.11.034>
- Karampelas, A. (2023, August 25). *The emergence of AI-natives*. Medium. <https://medium.com/@antonioskarampelas/the-emergence-of-ai-natives-6d67b2543561>
- King, J., & Meinhardt, C. (2024). *Rethinking privacy in the AI era: Policy provocations for a data-centric world*. Stanford Institute for Human-Centered Artificial Intelligence.

<https://hai.stanford.edu/white-paper-rethinking-privacy-ai-era-policy-provocations-data-centric-world>

- Lake, S. (2025, May 13). *OpenAI CEO Sam Altman says Gen Z and millennials are using ChatGPT like a “life advisor”—but college students might be one step ahead*. Fortune. <https://fortune.com/2025/05/13/openai-ceo-sam-altman-says-gen-z-millennials-use-chatgpt-like-life-adviser/>
- Lee, H., & Park, S. A. (2016). Third-person effect and pandemic flu: The role of severity, self-efficacy method mentions, and message source. *Journal of health communication*, 21(12), 1244-1250. <https://doi.org/10.1080/10810730.2016.1245801>
- Lee, S., Oh, J., & Moon, W. K. (2023). Adopting voice assistants in online shopping: Examining the role of social presence, performance risk, and machine heuristic. *International Journal of Human–Computer Interaction*, 39(14), 2978-2992. <https://doi.org/10.1080/10447318.2022.2089813>
- Lev-On, A. (2017). The third-person effect on Facebook: The significance of perceived proficiency. *Telematics and Informatics*, 34(4), 252–260. <https://doi.org/10.1016/j.tele.2016.07.002>
- Li, Z., Shi, J., Zhao, Y., Zhang, B., & Zhong, B. (2024). Indirect media effects on the adoption of artificial intelligence: The roles of perceived and actual knowledge in the influence of presumed media influence model. *Journal of Broadcasting & Electronic Media*, 68(4), 581-600. <https://doi.org/10.1080/08838151.2024.2377244>
- Lin, S. J. (2014). Media use and political participation reconsidered: The actual and perceived influence of political campaign messages. *Chinese Journal of Communication*, 7(2), 135-154. <https://doi.org/10.1080/17544750.2014.905867>
- Lo, V. H., & Wei, R. (2002). Third-person effect, gender and pornography on the Internet. *Journal of Broadcasting & Electronic Media*, 46, 13–33. https://doi.org/10.1207/s15506878jobem4601_2
- Long, D., & Magerko, B. (2020). *What is AI literacy? Competencies and design considerations*. *Proceedings of the 2020 CHI Conference on Human Factors in*

- Computing Systems*, 598–598. <https://doi.org/10.1145/3313831.3376727>
- Lyons, B. A. (2022). Why we should rethink the third-person effect: Disentangling bias and earned confidence using behavioral data. *Journal of Communication*, 72(5), 565-577. <https://doi.org/10.1093/joc/jqac021>
- Melnick, K. (2024, February 18). *Air Canada chatbot promised a discount. Now the airline has to pay it*. The Washington Post. <https://www.washingtonpost.com/travel/2024/02/18/air-canada-airline-chatbot-ruling/>
- Molina, M. D., & Sundar, S. S. (2024). Does distrust in humans predict greater trust in AI? Role of individual differences in user responses to content moderation. *New Media & Society*, 26(6), 3638-3656. <https://doi.org/10.1177/14614448221103534>
- Ng, D. T. K., Leung, J. K. L., Chu, K. W. S., & 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/pr2.487>
- NL Times. (2025, February 15). *Nearly half of Dutch distrust Chinese AI service DeepSeek, survey shows*. <https://nltimes.nl/2025/02/15/nearly-half-dutch-distrust-chinese-ai-service-deepseek-survey-shows>
- Perloff, R. M. (1999). The third-person effect: A critical review and synthesis. *Media Psychology*, 1, 353–378. https://doi.org/10.1207/s1532785xmep0104_4
- Perloff, R. M. (2009). Mass media, social perception, and the third-person effect. In J. Bryant, & M. B. Oliver (Eds.), *Media effects: Advances in theory and research* (3rd ed., pp. 252-268). Routledge.
- Pew Research Center. (2019). *Defining generations: Where Millennials end and Generation Z begins*. <https://www.pewresearch.org/short-reads/2019/01/17/where-millennials-end-and-generation-z-begins/>
- Pew Research Center. (2023). A majority of Americans have heard of ChatGPT, but few have tried it themselves. <https://www.pewresearch.org/short-reads/2023/05/24/a-majority-of-americans-have-heard-of-chatgpt-but-few-have-tried-it-themselves/>
- Prensky, M. (2001). Digital natives, digital immigrants, part I. *On the Horizon*, 9(5), 1-6.

<https://doi.org/10.1108/10748120110424816>

- Rosenthal, S., Detenber, B. H., & Rojas, H. (2018). Efficacy beliefs in third-person effects. *Communication Research*, 45(4), 554-576. <https://doi.org/10.1177/0093650215570657>
- Salwen, M. B., & Dupagne, M. (2001). Third-person perception of television violence: The role of self-perceived knowledge. *Media Psychology*, 3(3), 211-236. https://doi.org/10.1207/s1532785xmep0303_01
- Shen, L., Sun, Y., & Pan, Z. (2018). Not all perceptual gaps were created equal: Explicating the third-person perception (TPP) as a cognitive fallacy. *Mass Communication and Society*, 21(4), 399-424. <https://doi.org/10.1080/15205436.2017.1420194>
- Schmierbach, M., Andsager, J., Banning, S., Chung, M., Lyons, B., McLeod, D. M., ... & Sun, Y. (2023). Another point of view: Scholarly responses to the state of third-person research. *Mass Communication and Society*, 26(3), 359-383. <https://doi.org/10.1080/15205436.2023.2193512>
- Sun, Y., Shen, L., & Pan, Z. (2008). On the behavioral component of the third person effect. *Communication Research*, 35(2), 257-278. <https://doi.org/10.1177/0093650207313167>
- Sundar, S. S., & Kim, J. (2019, May). Machine heuristic: When we trust computers more than humans with our personal information. *Proceedings of the 2019 CHI Conference 57on human factors in computing systems*, 1-9. <https://doi.org/10.1145/3290605.3300768>
- Supantha, M. (2025, July 3). *Explainer: Will the EU delay enforcing its AI Act?* Reuters. <https://www.reuters.com/business/media-telecom/will-eu-delay-enforcing-its-ai-act-2025-07-03/>
- Toreini, E., Aitken, M., Coopamootoo, K., Elliott, K., Zelaya, C. G., & Van Moorsel, A. (2020, January). The relationship between trust in AI and trustworthy machine learning technologies. In *Proceedings of the 2020 conference on fairness, accountability, and transparency* (pp. 272-283).
- van Deursen, A. J., & van Dijk, J. A. (2019). The first-level digital divide shifts from inequalities in physical access to inequalities in material access. *New media & society*,

- 21(2), 354-375. <https://doi.org/10.1177/1461444818797082>
- van Dijk, J. A. (2002). A framework for digital divide research. *Electronic journal of communication*, 12(1). <https://www.cios.org/EJCPUBLIC/012/1/01211.html>
- van Dijk J.A. (2006). Digital divide research, achievements and shortcomings. *Poetics* 34(4), 221–235. <https://doi.org/10.1016/j.poetic.2006.05.004>
- van Dijk J.A. & Hacker, K. (2003). The digital divide as a complex and dynamic phenomenon. *The Information Society* 19(4): 315–326. <https://doi.org/10.1080/01972240309487>
- Vartak, M. (2023, June 29). *Six risks of generative AI*. Forbes. <https://www.forbes.com/sites/forbestechcouncil/2023/06/29/six-risks-of-generative-ai/?sh=3e94e3d73206>
- Wang, B., Rau, P. L. P., & Yuan, T. (2022). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324-1337. <https://doi.org/10.1080/0144929X.2022.2072768>
- Wang, C., Boerman, S. C., Kroon, A. C., Möller, J., & H de Vreese, C. (2024). The artificial intelligence divide: Who is the most vulnerable?. *New Media & Society*, 27(7) 3867–3889. <https://doi.org/10.1177/14614448241232345>
- Wang, S., Bolling, K., Mao, W., Reichstadt, J., Jeste, D., Kim, H.-C., & Nebeker, C. (2019). Technology to support aging in place: Older adults’ perspectives. *Healthcare*, 7, 60. <https://doi.org/10.3390/healthcare7020060>
- Wartgow, G. (2024, December 11). *Understanding Generational Differences in the Age of AI*. AEM. <https://www.aem.org/news/understanding-generational-differences-in-the-age-of-ai>
- Yang, A. (2025, May 13). *Noncoders are using AI to prompt their ideas into reality. They call it “vibe coding.”* NBC. <https://www.nbcnews.com/tech/tech-news/noncoders-ai-prompt-ideas-vibe-coding-rcna205661>
- Yang, J., & Tian, Y. (2021). “Others are more vulnerable to fake news than I Am”: Third-person effect of COVID-19 fake news on social media users. *Computers in Human*

Behavior, 125, 106950. <https://doi.org/10.1016/j.chb.2021.106950>

Zhao, X., & Cai, X. (2008). From self-enhancement to supporting censorship: The third-person effect process in the case of Internet pornography. *Mass Communication & Society*, 11(4), 437–462. <https://doi.org/10.1080/15205430802071258>

Zhou, S., & Zhang, Z. (2023). Impact of internet pornography on Chinese teens: the third-person effect and attitudes toward censorship. *Youth & Society*, 55(1), 83-102. <https://doi.org/10.1177/0044118X211040095>