

臺灣人工智慧發展的人權治理意見書

2025 年 12 月

目錄

壹、	前言	1
貳、	AI 發展涉及之國際人權公約及宣言	2
一、	世界人權宣言(UDHR).....	2
二、	公民與政治權利國際公約(ICCPR).....	2
三、	經濟社會文化權利國際公約(ICESCR)	2
四、	消除一切形式種族歧視國際公約(ICERD)	3
五、	消除對婦女一切形式歧視公約(CEDAW)	3
六、	兒童權利公約(CRC).....	3
七、	身心障礙者權利公約(CRPD).....	4
八、	人工智慧與人權、民主及法治框架公約.....	4
參、	聯合國及相關國際人權機構針對 AI 衝擊人權之因應措施	6
一、	AI 的定義	6
二、	人權風險識別.....	6
三、	AI 的人權治理	6
肆、	AI 發展的人權治理.....	9
一、	消除歧視.....	9
二、	保護隱私.....	12
三、	降低錯誤虛假訊息.....	14
四、	防制深偽色情與兒少性剝削.....	16
五、	有效的人工智慧治理.....	18
六、	平等參與數位轉型及永續發展.....	22
伍、	參考資料.....	25

臺灣人工智慧發展的人權治理意見書

壹、前言

人工智慧（AI）技術的快速發展，尤其是 2022 年生成式 AI 的崛起，加速改變人類社會的運作模式與生活型態。這股浪潮帶來前所未有的創新契機，卻同時伴隨著難以忽視的人權風險。

國際間從《世界人權宣言》到《公民與政治權利國際公約》(ICCPR)、《經濟社會文化權利國際公約》(ICESCR) 等核心人權文書，皆明確規範每個人免受歧視、享有隱私與言論自由、獲得工作保障以及平等參與文化與教育的基本權利。隨著 AI 技術快速迭代，這些保障正面臨前所未有的挑戰，包括可能加劇歧視與不平等、侵犯隱私、危害思想與言論自由，甚至衝擊勞動就業、教育機會與文化多樣性。對於國家 AI 相關法律制度、監管能力及政策方向提出嚴峻考驗。

聯合國及國際人權機構亦對科技發展中的風險提出警示，從聯合國人權事務高級專員辦事處（OHCHR）提出生成式 AI 人權風險分類，及聯合國兒童基金會（UNICEF）關注 AI 對兒童權利的衝擊，再到韓國、澳洲等國家人權委員會提出 AI 人權指導方針或建議，都凸顯了「以人權為核心治理 AI」的重要性。

臺灣在 AI 時代的浪潮中，必需同步落實人權保障與永續發展，確保科技進步不以犧牲平等與尊嚴為代價。因此，本文件以相關國際人權公約為基礎，盤點 AI 發展對人權的影響，並結合國際人權趨勢與在地觀察，分別就消除歧視、保護隱私、降低錯誤虛假訊息、防制深偽色情與兒少性剝削、有效的人工智慧治理、平等參與數位轉型及永續發展六大面向提出建議，協助我國在推動 AI 創新與數位轉型的同時，降低人權風險，確保公平與永續發展。

本文件為臺灣國家人權委員會（NHRC）對政府提出的政策建議，雖不具強制力，惟可作為政府制定或修正 AI 相關政策、法規、治理架構的重要參考文件。

貳、AI 發展涉及之國際人權公約及宣言

AI 發展可能造成歧視與平等權、隱私權、思想及言論自由、工作權、教育及文化權之衝擊，爰以這些權利為核心，盤點相關國際人權公約及宣言：

一、世界人權宣言(UDHR)

- (一) 人人皆得享有宣言所載之一切權利和自由，並有權享受法律的平等保護，不受任何歧視。(第 2 條及第 7 條)
- (二) 人人有權享有法律保護其私生活、家庭、住所及通訊不受任意干涉。(第 12 條)
- (三) 人人有思想、良心及宗教自由、表達意見不受干涉的自由、透過任何媒介尋求、接受及傳播資訊與想法的自由。(第 18 條及第 19 條)
- (四) 人人享有工作、自由選擇職業、享受公平、有利工作條件的權利，及有權享有免於失業之保障。(第 23 條)
- (五) 人人都有受教育的權利。(第 26 條)
- (六) 人人有權自由地參與社會的文化生活，共享科學所帶來的進步與益處。(第 27 條)

二、公民與政治權利國際公約(ICCPR)

- (一) 締約國承允無分種族、膚色、性別、語言、宗教、政見等一律享有本公約確認的權利。人人在法律上一律平等而不受歧視；兒童不因種族、膚色、性別、語言、宗教等而受歧視。(第 2 條、第 3 條、24 條及第 26 條)
- (二) 任何人之生活、家庭、住宅及通信不得無理或非法侵擾。(第 17 條)
- (三) 人人有思想、信念及宗教自由、保持意見不受干預之權利、發表自由之權利。(第 18 條及第 19 條)
- (四) 少數團體之人，與團體中其他分子共同享受其固有文化、使用其固有語言之權利，不得剝奪之。(第 27 條)

三、經濟社會文化權利國際公約(ICESCR)

- (一) 締約國承允不因種族、膚色、性別、語言、宗教、政見等而受歧視，確保本公約所載權利之享受。(第 2 條及第 3 條)

(二) 人人有工作之權利，包含有機會憑本人自由選擇或接受之工作謀生之權利；人人享有公平與良好之工作條件。(第 6 條及第 7 條)

(三) 人人有受教育及參加文化生活、享受科學進步及其應用之惠之權利。(第 13 條及第 15 條)

四、消除一切形式種族歧視國際公約(ICERD)

(一) 締約國承諾禁止並消除一切形式種族歧視，保證人人有不分種族膚色或原屬國或民族本源在法律上一律平等之權，尤得享受政治、經濟社會及文化等權利。(第 2 條及第 5 條)

(二) 人人享有思想、良心與宗教自由、主張及表達自由之權利。(第 5 條第卯款第 7 目及第 8 目)

(三) 人人享有工作、自由選擇職業、享受公平優裕之工作條件，免於失業之保障。(第 5 條第辰款第 1 目)

(四) 人人享受教育與訓練、平等參加文化活動之權利。(第 5 條第辰款第 5 目及第 6 目)

(五) 締約國承諾在講授，教育、文化及新聞方面以打擊導致種族歧視之偏見。(第 7 條)

五、消除對婦女一切形式歧視公約(CEDAW)

(一) 締約各國譴責對婦女一切形式的歧視，協議立即用一切適當辦法，推行消除對婦女歧視的政策。(第 2 條)

(二) 締約各國應採取一切適當措施，消除在就業方面對婦女的歧視，以保證她們在男女平等的基礎上享有相同權利，特別是人人有不可剝奪的工作權利等。(第 11 條)

(三) 締約各國應採取一切適當措施以消除對婦女的歧視，以保證婦女在教育方面享有與男子平等的權利。(第 10 條)

六、兒童權利公約(CRC)

(一) 締約國應確保其管轄範圍內之每一兒童均享本公約權利，不因兒童、父母或法定監護人之種族、膚色、性別、語言、宗教、政治或其他主張、國籍、身心障礙等之不同而有所歧視。(第 2 條)

(二) 兒童之隱私、家庭、住家或通訊不得遭受恣意或非法干預。(第 16 條)

(三) 締約國應尊重兒童思想、自我意識與宗教自由之權利，確保兒童可自國內與國際各種不同來源獲得資訊及資料。(第 14 條及第 17 條)

(四) 締約國確認兒童有接受教育之權利；少數人民或原住民之兒童應有與其群體的其他成員共同享有自己的文化、使用自己的語言之權利，此等權利不得遭受否定。(第 28 條及第 30 條)

七、身心障礙者權利公約(CRPD)

(一) 締約國應禁止所有基於身心障礙之歧視，保障身心障礙者獲得平等與有效之法律保護；採取適當措施消除任何個人、組織或私營企業基於身心障礙之歧視；應確保身心障礙者在與其他人平等基礎上，無障礙地利用資訊及通信，包括資訊與通信技術及系統，以及享有於都市與鄉村地區向公眾開放或提供之其他設施及服務。(第 4 條、第 5 條及第 9 條)

(二) 締約國應在與其他人平等之基礎上保障身心障礙者之個人、健康與復健資訊之隱私；締約國蒐集、保存與使用統計資料，應遵法定防護措施，及遵行國際公認之規範。(第 22 條及第 31 條)

(三) 公約之原則為尊重固有尊嚴、個人自主，包括自由進行個人選擇及個人自立；締約國應確保與行使法律行為能力有關之措施，尊重本人之權利、意願及選擇。(第 3 條及第 12 條)

(四) 締約國應確保身心障礙者能夠行使自由表達及意見自由之權利；促進身心障礙者有效與完整地參與公共事務之處理。(第 21 條及第 29 條)

(五) 身心障礙者享有與其他人平等之工作權利，包括於一個開放、融合與無障礙之勞動市場及工作環境中，身心障礙者有自由選擇與接受謀生工作機會之權利。(第 27 條)

(六) 身心障礙者享有受教育、與其他人平等基礎上參與文化生活之權利。(第 24 條及第 30 條)

八、人工智慧與人權、民主及法治框架公約

歐洲理事會 2024 年 9 月 5 日開放簽署《框架公約》，截至 2025 年 12 月已有歐盟及其他如美國、加拿大、日本等 17 個國家簽署。《框架公約》旨在規範 AI 系統的生命週期，以確保其發展和應用與人權、民

主和法治一致，與權利相關的概念如下：

- (一) 締約國應採納或維持措施，以確保人工智慧系統生命週期中各活動尊重平等，包括性別平等，以及適用於國際及國內法規定的禁止歧視原則。(第 7 條-人類尊嚴與個人自主權)
- (二) 在 AI 系統生命週期各項活動中，締約國應採納或維持措施以尊重人類尊嚴及個人自主權。(第 10 條第 1 項-平等與不歧視)
- (三) 締約國應對於 AI 系統生命週期中各活動採取或維持措施，以確保個人隱私權及個人資料受保護，及為個人提供有效的保障及保護措施。(第 11 條-隱私與個資保護)
- (四) 締約國應依其國內法及適用的國際義務，充分考慮身心障礙者與兒童的特定需求及脆弱性，並尊重其權利。(第 18 條-身心障礙者與兒童的權利)
- (五) 締約國應鼓勵並促進所有人口群體適當數位素養與技能培養，包括為負責識別、評估、防範及減緩人工智慧系統風險者提供專業技能。(第 20 條-數位素養與技能)

參、聯合國及相關國際人權機構針對 AI 衝擊人權之因應措施

一、AI 的定義

查歐盟《人工智慧法》及歐洲理事會《框架公約》對於 AI 系統之定義，皆參考「經濟合作暨發展組織」(Organisation for Economic Co-operation and Development, OECD) 之定義，即基於機器為基礎之系統，針對明確或隱含的目標，根據接收的輸入推斷如何產生輸出，例如預測、內容、建議或決策，從而影響實體或虛擬環境。不同的 AI 系統部署後的自主性和適應性程度有所不同。

二、人權風險識別

- (一) 聯合國人權事務高級專員辦事處(OHCHR)發布《生成式 AI 相關的人權風險分類》，探討可能受到生成式 AI 影響的人權，包含免受身心傷害的權利、法律之前人人平等並受保護免遭歧視的權利、隱私權、財產權、思想、宗教、信念與意見自由、表達與取得資訊自由、工作權及謀生權、兒童權利、文化、藝術及科學權利。
- (二) 韓國國家人權委員會(NHRCK)於 2022 年 5 月間發布《人工智慧開發與應用之人權指導方針》，旨在因應 AI 技術對社會與人權的廣泛影響，並確保 AI 發展與應用過程中能保障人類尊嚴與基本權利。指導方針於保障人類尊嚴、透明度與說明義務、保障自決權與隱私、禁止歧視、實施 AI 人權影響評估、風險分級與法規制度建立等原則下，給予指引供公私部門參考。
- (三) NHRCK 於 2024 年 7 月發布《人工智慧的人權影響評估工具》，包含 72 個問題，分為計畫與準備、分析與評估、改善與補救、公開與審查四階段協助 AI 開發與應用單位全面檢視技術風險與人權影響，並提供詳細說明以減輕評估負擔。該工具建議所有公部門及私部門使用之高風險 AI 均應進行自願性人權影響評估。

三、AI 的人權治理

- (一) 聯合國大會 2024 年 9 月通過《全球數位契約》(Global Digital Compact, GDC)，其 5 項目標中和 AI 與人權最為相關者，包

括目標 3「營造尊重、保障和促進人權的融合、開放、安全和可靠的數位空間」，與目標 5「加強人工智慧國際治理，以造福人類」：

1. 目標 3 之具體措施包括：促使相關公司與開發人員針對 AI 所生成內容中的仇恨言論和歧視，制定解決方案並公布所採取的行動。並將保障措施納入 AI 模型訓練過程、AI 生成材料識別、內容與來源真實性認證技術。
2. 目標 5 之具體措施，包括承諾在充分尊重包括國際人權法在內的國際法，並考量聯合國教科文組織(UNESCO)《人工智慧倫理問題建議書》等其他相關框架的情況下，推動採取公平和融合的辦法利用 AI 的正面效益並降低風險。

(二) OHCHR 於《全球數位契約意見書》(Submission to the Global Digital Compact)提出了 AI 治理的核心原則，包括自動化系統的設計與使用應受到透明度、人類監督、問責制與風險管理的約束。開發和部署人工智慧技術的公司有責任尊重人權，並識別、處理與減緩其商業活動所產生或相關的不利影響，依聯合國《商業與人權指導原則》(UNGPs)的人權盡職調查(HRDD)處理生成式 AI 造成的人權風險。對於可能對人權造成無法充分減緩風險的 AI 系統，應予以禁止。

(三) 聯合國教育、科學及文化組織(UNESCO)發布《人工智慧倫理問題建議書》中，提出 10 項原則：比例原則與不傷害、安全與保全、平等不歧視、永續、隱私權與資料保護、人類監督與決策、透明度與可解釋性、責任與問責制、意識與素養、多方利益相關者的適應性治理與協作，促使會員國制定 AI 技術倫理影響評估方法。

(四) 聯合國兒童基金會(UNICEF)與相關國際組織共同制定《AI 用於兒童之政策指引》，探討 AI 系統對兒童的影響，以促進政府及私部門 AI 政策及兒童權利保障。為了實現以兒童為中心的 AI，該指引提出 9 項具體要求，包含支持兒童的發展與福祉、確保兒童的包容性與參與、優先考量兒童之公平與反歧視、保護兒童的資料與隱私、提供透明度、可解釋性與問責制等。

(五) 澳洲人權委員會(AHRC)於 2024 年 5 月向政府提交「人工智慧在澳洲的發展與應用」意見書，內容係以 2021 年發布之「人權與科技最終報告」(Human Rights and Technology Final Report)為基礎，就 AI 可能造成偏見與演算法歧視、自動化偏差、錯誤與虛假訊息、外來干預、環境影響，闡述國內外趨勢及政策走向，並提出下列 6 項建議：

- 1.加強現有立法，必要時引入針對人工智慧的專門立法。
- 2.以人權為中心的方法來開發與部署人工智慧。
- 3.設立國家人工智慧專員作為獨立法定辦公室。
- 4.研究識別、預防和阻止人工智慧生成之虛假訊息及外來干預活動。
- 5.透過支持獨立研究來應對社群媒體干擾活動的能力。
- 6.加強與私部門合作，以更了解並減輕人工智慧對氣候的影響。

肆、AI 發展的人權治理

國際人權公約及宣言明確保障免受歧視、隱私、言論自由、工作、接受教育與參與文化等基本權利，爰綜合 AI 在全球加速應用之際，聯合國及相關國際人權機構關注重點，並結合本會辦理之論壇、工作坊及焦點座談內容，歸納出消除歧視、保護隱私、降低錯誤虛假訊息、防制深偽色情與兒少性剝削、有效的人工智慧治理、平等參與數位轉型及永續發展六大面向，提出觀察與建議：

一、消除歧視

(一) AI 系統生命週期分為 4 階段，包含規劃與設計、資料蒐集與輸入、模型訓練與驗證、系統部署與監控，各階段皆有可能引入偏見導致產出內容或結果產生歧視，包含 AI 系統開發中的人為偏見(Human Bias in AI Development)¹、資料偏見(Data Bias)²、演算法偏見(Algorithmic Bias)³、模型訓練與部署中的偏見(Bias in Model Training & Deployment)⁴及自動化或確認偏誤(Confirmation Bias)⁵。

(二) 聯合國種族歧視特別報告員 Ashwini K.P.指出，AI 之蓬勃發展應用引發對種族歧視的擔憂，資料、演算法偏見及演算法「黑箱」問題，助長了種族歧視，特別在執法、教育、醫療保健等領域，如預測性警務利用地點、事件及歷史犯罪數據間之關聯來預測可能發生犯罪行為之時間及地點，強化對少數種族及族裔社區的過度執法，形成警力過剩與偏見回饋循環，導致演算法預測更具偏見，持續加劇不公平現象；執法機構利用人臉識別技術，將圖像與資料庫比對以辨識身份，惟機器學習訓練多以白人男性圖像為主，缺乏種族、性別及文化多樣性，導致深膚色及處境不利群體易被錯誤匹配，進而加劇種族歧視，例如

¹ 規劃設計團隊缺乏多元背景，導致設計決策忽略少數群體需求；設計者主觀偏見或未考慮處境不利群體，導致系統設計存在歧視風險。

² 資料蒐集與準備階段，資料本身即隱含偏見，或缺乏目標族群的代表性時，進入模型訓練而產出的結果存在歧視風險。

³ AI 系統的設計方式或演算規則帶有偏見，如只考慮某些群體或參數、權重設定導致產出結果偏頗。

⁴ AI 模型如以帶有偏見的資料進行訓練，產出結果有歧視風險，AI 系統如沒有定期審核、更新，歧視可能會持續存在，甚至隨著時間的推移而加劇。

⁵ 過度依賴與相信 AI 產生的結果，而忽略自身的判斷，進而影響決策。例如使用 google maps 等全球定位系統時將汽車駛入大海。

巴西執法人員使用有偏見的人臉辨識系統進行逮捕，據 2019 年的研究顯示，巴西各城市逮捕人員 90% 是非洲裔。

(三) 非營利組織 Data-Pop Alliance 資料及技術總監 Zinnya del Villar 於 2025 年接受聯合國婦女署專訪時表示，AI 系統從充滿刻板印象之資料中學習，或依賴有偏見的演算法時，將加劇 AI 之性別歧視，尤其在決策、招聘、貸款審核及法律判決等領域，如 Amazon 曾發現 AI 招聘系統偏好男性履歷而停止使用。為了減少 AI 中的性別偏見，用於訓練 AI 系統的資料必須多樣化，能代表所有性別、種族和社群；AI 系統應由來自不同性別、種族和文化背景的人員組成多元化開發團隊進行規劃設計，藉以引入不同視角，減少導致系統偏差的盲點。

(四) 聯合國障礙者權利特別報告員 Gerard Quinn 於 2021 年提交給人權理事會的報告指出，AI 的發展可望為障礙者在就業、教育、自立生活等領域所需要的支持協助帶來突破與革新，然而 AI 已經存在某些歧視性運用案例，也正侵害障礙者的基本人權。例如：錯誤的 AI 風險評估，可能導致障礙者投保商業醫療保險被拒或保費增加；使用 AI 臉部辨識技術進行的面試，無法正確辨識唐氏症者、唇顎裂者的臉部特徵，使用 AI 情緒處理算法，也可能曲解自閉症者、帕金森氏症者的臉部表情，導致障礙者失去就業機會；過度依賴 AI 替代專業照顧人力，長期可能導致使用服務的障礙者被隔離及孤立，對其心理健康帶來嚴重風險。

(五) NHRC 建議：

1. 發展 AI 時，應避免因種族、階級、膚色、性別、性傾向、語言、宗教、政治或不同主張、國籍、族裔、原住民或社會背景、財產、出生地、年齡、職業、身心障礙等而予以歧視。
2. AI 系統生命週期各階段應建置反歧視機制，相關人員應接受反歧視教育訓練，以理解偏見如何導致模型訓練偏差進而產出歧視內容。規劃設計階段應讓多元群體參與其中，資料蒐集輸入階段進行審視資料來源、分析偏見與公平性，模型訓練驗證階段納入多元群體測試，部署監控階段建立持續監測與申訴救濟機制，

避免 AI 系統產出結果及決策對於特定群體產生歧視。

3. 政府應訂定指引，使資料蒐集處理、演算法設計及模型訓練之揭露方式有所依循，增加技術及流程透明度，減少資料偏誤，使外界了解演算法的運作方式，消除歧視同時兼顧透明度與商業機密。另應使利害關係人(包含受 AI 服務影響的使用者)理解 AI 產出內容或決策原因、哪些因素會影響其決策和輸出，提升可解釋性。
4. AI 系統產出的決策對人民權利有重大影響時，需由人為最終判斷的主體，並可建立第三方監測機制，對演算法獨立驗證，及對訓練資料蒐集及運用進行審查，持續定期監控與調整，尤其是在招聘、醫療保健、執法、社會福利等與民眾權益高度相關領域。

二、保護隱私

- (一) 公民團體曾於 NHRC 焦點座談指出，目前缺乏公開資訊，無法掌握政府開發或應用 AI 技術的情形。AI 相關計畫往往在接近完成或已經部署後才對外公布，使得公民團體難以即時監督或參與討論，例如衛福部於 2024 年 6 月與 Google 合作「AI 醫療照護研究計畫」，無法得知其數據來源是否涉及健保資料，或是否取得當事人知情同意，顯示政府在政策規劃及執行過程缺乏透明度及監督機制。
- (二) NHRC 於 2025 年舉辦數位人權國際工作坊指出，當前企業透過 AI 大量擷取網路資料進行分析及運算，經常未經個人同意即蒐集及處理個人資訊。此外，跨國企業擷取過多個人資訊並儲存，應用於 AI 訓練及精準廣告，此類行為不僅侵害隱私權，更違反資料最小化原則⁶。
- (三) 歐洲隱私權倡議組織 noyb 於 2025 年指出，Meta 未經用戶明確同意，使用 Facebook 及 Instagram 用戶之個人資料進行 AI 訓練，違反《一般資料保護規則》(GDPR) 之規定。儘管 Meta 主張其行為基於「合法利益」(legitimate interest)，僅提供用戶事後選擇退出 (opt-out)，卻未採行事前明確同意 (opt-in) 機制，noyb 批評此作法違背資料保護及資料自決權。
- (四) 義大利於 2025 年通過《人工智慧法》，建立明確且分級的同意機制保障兒少隱私。依該法規定，未滿 14 歲之兒少需經父母或監護人同意，方得使用 AI 系統或提供個人資料進行處理；年齡介於 14 至 17 歲之青少年則可自行表示同意，但 AI 服務提供者必須確保相關資訊以清楚且可理解之方式呈現，協助兒少作出知情決定，此制度設計展現義大利以兒少權益為核心的 AI 監管取向。
- (五) NHRC 建議：

1. 政府在蒐集、使用個人資料，尤其用於 AI 訓練時，應落實當事人知情同意原則。此應包括清楚揭示資料蒐集目的、處理方式及

⁶ 歐盟「一般資料保護規則」(GDPR) 第 5 條第 1 項 c 款指出，個人資料應適當、相關且限於處理目的所必要者。資料最小化原則為企業在處理個人資料時，應僅蒐集及處理營運確實需要的用戶資料。

可能用途，當用途超出原始目的範圍時，更應於各階段提供當事人可理解的資訊。此外，政府機關使用 AI 進行決策，造成個人權利受到重大影響時，當事人應有權要求人工審查，而非完全依賴 AI 判定。

2. 政府可參考歐盟近期爭議及監管實務，規範企業建立明確的 AI 訓練資料同意機制，以事前明確同意為原則，並賦予使用者可於事後選擇退出及刪除個資之權利，防止企業以模糊的合法利益為理由進行大規模資料蒐集及訓練，保障人民擁有充分的資料自決權及隱私權。
3. 為保障兒少最佳利益，建議政府針對特定年齡以下之未成年人，明確規範其在使用 AI 系統或提供個人資料時，需經父母或監護人同意。

三、降低錯誤虛假訊息

- (一) OHCHR 曾示警：「生成式 AI，包括其快速製造看似由人生產且具有權威性，實際上卻是虛假內容的能力，可能在各方面對言論自由構成風險。」根據法務部調查局函復 NHRC 之資料顯示，在 2024 年中華民國總統選舉期間，運用深偽技術(Deepfake)進行操控選舉之情形已有具體案例，正式立案調查件數為假訊息 9 件、境外勢力介選 1 件；該局另指出，多起案例為規避執法單位查緝，實施多元手法隱匿身分及來源，並建構多層次轉傳機制及人頭帳號進行擴散⁷，顯示虛假訊息正影響我國民主機制。
- (二) 根據瑞典哥德堡大學「多元民主中心」(V-Dem)資料庫，我國是全球最受境外虛假訊息影響的國家。台灣人工智慧實驗室也發現，在 2024 年總統選舉期間有大量協同帳號進行議題操作，手法包含利用 AI 批量製作音訊、視訊及文字內容，相關論述不僅與中國官方媒體說法呼應，真假混雜的誤導訊息更使第三方查核機構陷入「事實查核地獄」(Debunking Hell)。此外，AI 也被用於製造假帳號頭像，這些頭像通常具有相似的特徵，例如眼睛及嘴巴的位置一致、照片背景模糊等。而在選舉結束後，約有半數協同帳號隨即消失。
- (三) 虛假訊息在社群媒體及通訊軟體等平台上迅速傳播，但缺乏有效的監管機制。根據公民團體於 NHRC 焦點座談指出，企業推動媒體素養教育，然而部分行動意在降低政府監管壓力，而非真正承擔企業責任。此外，儘管部分跨國社群平台與第三方查核機構建立合作機制，不過合作關係隨時可能終止。
- (四) NHRC 建議：
1. 鑑於 AI 生成的虛假訊息正侵蝕公共討論品質，使基於事實為基礎的對話環境惡化，建議政府落實既有法規，明確要求社群平台為其傳播的資訊負責，特別是在其仰賴流量獲取廣告收益的商業模式之下，不實內容往往因其煽動性，較易在演算法推播

⁷ 法務部調查局(2023)，「境外敵對勢力介入我總統大選 國人宜謹慎識別網路假訊息」，取自：
<https://www.mjib.gov.tw/news/Details/1/953>

下快速擴散，造成公共討論失衡。

2. 政府應強化民眾及媒體的資訊素養，特別是兒少群體，以教育及宣導培養判斷資源來源、辨識虛假訊息能力，並且加強民眾參與公共政策制定及保障知情權，提供公平參與公共事務的機會，使政策資訊公開透明。

四、防制深偽色情與兒少性剝削

(一) NHRC 注意到，運用 AI 製作的深偽色情已由私領域擴及至公領域，成為威脅女性公眾人物之工具。針對女性之性暴力影像流通急劇上升，進而削弱女性在公共領域的參與空間及社會信任，加深性別不平等。

(二) 2024 年，美國國家失蹤及受剝削兒童中心（NCMEC）指出，該機構於過去兩年來收到 7,000 多份與 AI 生成相關之兒少性剝削案例，加害者不僅透過 AI 應用程式製作兒少裸照，造成受害者心理創傷，甚至以此進行脅迫，要求受害者提供更多私密內容或金錢。NCMEC 預測，隨著 AI 日漸普及，相關通報案件將持續增加。

(三) 根據衛福部保護服務司統計資料，近 6 年來兒少性剝削案件增加 3 倍，並於 2023 年首次超過 4,000 件次，案件類型以性影像最為普遍⁸，惟目前統計尚未區分與 AI 生成相關之性剝削類型。另根據 NHRC 焦點座談，公民團體從校園現場觀察到學生開始利用 AI 影像合成，例如製造同學的深偽裸照，構成數位性暴力風險，有些 AI 應用程式甚至可以「一鍵脫衣」（Nudify），輕易生成虛假的裸露影像並散播。

(四) NHRC 建議：

1. 社群平台的擴散效應加劇深偽色情的威脅，尤其針對兒少及女性公眾人物的惡意影片日益嚴重，建議政府強化防範措施，完善下架機制，明確要求社群平台承擔在數位性別暴力防制中的社會責任。
2. 儘管兒少群體是數位科技的主要使用者之一，不過其不利處境往往未被納入設計考量。立法院 2025 年 12 月 23 日三讀通過《人工智慧基本法》已規範應以兒少最佳利益為原則，建議政府後續開發或應用 AI 技術時，參考歐洲理事會《框架公約》及基本法之精神，適當考慮與兒童權利有關之具體需求及脆弱性。

⁸ 根據衛福部保護服務司統計，2024 年兒少性剝削通報案件共 4,486 件次，其中「拍攝、製造、散布、播送、交付、公然陳列或販賣兒童或少年性影像、與性相關客觀上足以引起性慾或羞恥之圖畫、語音或其他行為之物品」占 3,104 件。取自：<https://dep.mohw.gov.tw/dops/lp-1303-105-xCat-cat05.html>

3. 政府應於兒少性剝削通報案件中，增設與生成式 AI 相關之案件類型統計項目，以利掌握數位性剝削之新興趨勢，作為後續預防及修法之基礎。

五、有效的人工智慧治理

基於國際趨勢與臺灣 AI 發展情形，將 AI 治理分為風險管理(含監理沙盒、人權影響評估)、分工治理架構(含個資保護與資料治理)、申訴救濟、問責補償、監管與罰則等面向，提出觀察與建議：

(一) 現行國際間對於 AI 規範相關法令為：

1. 歐盟：2024 年 8 月實施《人工智慧法》，特色在於：

- (1) 將 AI 運用分為不可接受風險、高風險、有限風險及最低風險，進行分級管制，並要求使用高風險 AI 系統前應進行基本權利影響評估⁹。
- (2) 各成員國由國家資料保護機關擔任市場監督機關，可受理申訴，各成員國並設有國家層級 AI 監理沙盒規範，建立開發及上市前受控的實驗與測試環境，促進創新 AI 系統的開發、訓練、測試和驗證同時，識別風險、提出緩解措施並評估其有效性¹⁰。
- (3) 採取分階段實施，例如不可接受風險的 AI 系統於法案生效起 6 個月(2025 年 2 月 2 日)適用，最長於法案生效起 36 個月適用完畢¹¹。

2. 韓國：2024 年底通過《人工智慧發展與建立信任基本法》(Basic Act on the Development of Artificial Intelligence and the Establishment of Trust)，預計 2026 年施行，聚焦建立可信賴的 AI 應用基礎，特色在於：

- (1) 明確該法主管機關及審議風險規範、制定 AI 基本計畫、AI 政策及研究開發策略、AI 技術標準化及風險定義等之主政機關¹²，治理分工模式明確。
- (2) 規範高影響 AI(即高風險 AI)及生成式 AI 定義，要求企業履行相關措施，承擔透明度及安全義務，並實施基本

⁹ 歐盟《人工智慧法》第 5 條、第 6 條及第 27 條

¹⁰ 歐盟《人工智慧法》第 57 條、第 74 條、第 85 條，歐洲資料保護委員會(EDPB)於 2024 年 7 月 16 日發布聲明(Statement 3/2024 on data protection authorities' role in the Artificial Intelligence Act framework)，國家資料保護機關(data Protection Authorities, DPAs)應被指定為歐盟《人工智慧法》之市場監督機關

¹¹ 歐盟《人工智慧法》第 113 條

¹² 韓國《人工智慧發展與建立信任基本法》第 7 條、第 11 條及第 12 條

權利影響評估，及訂有違反透明度義務之罰則¹³。

3. 日本：日本亦於 2025 年 6 月施行《人工智慧技術研究開發及利用促進法》，特色在於：

(1) 設置 AI 戰略本部研擬 AI 基本計畫草案¹⁴，規範中央及地方政府在推進 AI 之相關權責，透過 AI 提升國家的競爭力和效率，採政府指導及業者自願合作方式，未設罰則。

(2) 法案附帶決議中要求政府執行政策時，應將風險降至最低，並儘早成立由 AI 倫理、法律和社會議題等領域專家組成的智庫機構。

(二) 臺灣《人工智慧基本法》

1. 臺灣《人工智慧基本法》草案於 2025 年 2 月底由國家科學及技術委員會(下稱國科會)移交予數位發展部(下稱數發部)主政，行政院前於 2025 年 8 月 28 日通過《人工智慧基本法》草案，立足於「鼓勵創新、兼顧人權」的核心理念，規劃由數發部推動風險分類框架，各目的事業主管機關視 AI 應用風險管理需要，以風險程度為基礎訂定層級管理規範。

2. 前述草案於立法院審議期間歷經多次協商及修正，經立法院 2025 年 12 月 23 日三讀通過，明定基本法的中央主管機關為國科會，地方主管機關為直轄市、縣(市)政府，並將由行政院成立國家 AI 戰略特別委員會，協調、推動及督導全國 AI 事務，已有初步的整體治理框架。惟仍屬綱要立法性質¹⁵：

(1) 未見對於政府所制定之風險管理與評估、高風險 AI 應用之問責、救濟、補償等規範之相關監管機制，以確認是否落實或有無相關問題。

(2) 各目的事業主管機關得建立或完備 AI 研發及應用服務之創新實驗環境，惟欠缺國家層級的監理沙盒制度。

3. 行政院前以數位政策法制協調專案會議作為協調各部會訂定相關作用法或指引等配套措式之用。倘行政院依《人工智

¹³ 韓國《人工智慧發展與建立信任基本法》第 2 條、第 34 條、第 35 條及第 43 條

¹⁴ 日本《人工智慧技術研究開發及利用促進法》第 18 條及第 19 條

¹⁵ 綱要立法性質即授權行政機關制訂相關法令

憲基本法》設立 AI 戰略特別委員會，由上而下協調、推動及督導全國 AI 事務，則因 AI 而受有權利侵害之問責、救濟、補償、整體監管等問題，需審慎評估處理。

(三) AI 發展與個資隱私、資料共享與再利用息息相關，原定於今年 8 月成立之個人資料保護委員會(下稱個資會)迄今尚未成立，且日前三讀通過之《個人資料保護法》部分條文修正草案僅賦予個資會相關執法權限，其他如數位時代下之隱私保護議題，尚待個資會成立後，本於主管機關立場賡續研擬修正；資料開放與共享相關機制則由數發部主政之《促進資料創新利用發展條例》草案進行規範，惟草案明定資料涉及個資時應依個資法辦理，並應避免不必要之個資蒐集、處理或利用，因兩法案主管機關不同，未來於制定過程中應確保資訊互相流通及同步並釐清分工權責。

(四) 曾有多位學者於 NHRC 焦點座談表示，AI 不同層次的應用會涉及不同法規，例如從使用者層次需考慮平台管制法規，從資料利用層次需考慮《個人資料保護法》，從應用層次亦需考慮平權與反歧視法規，另還有資安、數位服務或數位中介等等，各法律間的競合關係及 AI 在這些法律的定位及策略走向亦需一併探討釐清，因此建議需有整體的數位法規體系，才能有效應對人權衝擊。

(五) NHRC 建議：

1. 政府應整體思考數位法規框架，盤點與 AI 有關及受其影響層面的法規，包括個資保護、資訊安全、資料治理、數位中介服務、反歧視等領域，並釐清各法規的主管機關及其競合關係。
2. 政府應就下列事項建構完整且分工明確、具跨部會協調之治理機制：

(1) AI 風險管理與評估之監管。

(2) 以法令明確規範 AI 侵害人民權利(例如侵犯隱私、造成歧視性傷害)之問責、申訴、救濟(司法及非司法救濟途徑)與補償等之處理。

(3) AI 侵害權利案例與樣態之蒐集分析及定期發布，並據以動態

調整前述監管及處理機制。

3. 透過國際合作參與討論，借鏡國外經驗建立法律框架，例如歐盟的《人工智慧法》(AI Act) 具有高度指標性，指定 AI 之市場監督機關，並設有國家層級的監理沙盒機制；韓國《人工智慧發展與建立信任基本法》及日本《人工智慧技術研究開發及應用促進法》皆明定分工治理權責與主管機關，韓國更細緻化各單位治理措施、要求業者承擔義務及違反時之罰則規範。臺灣政府得藉此完善人工智慧相關法令，制定符合國情之 AI 政策規範。
4. 政府應參酌國際間 AI 風險分級規範，制定符合臺灣國情之風險分類方式，依分類區別 AI 系統管控強度，並建立動態調整機制，以適應科技快速發展。
5. 政府開發、部署或應用高風險 AI 系統時，應實施人權影響評估，納入技術、法律、人權、政策等多元領域專家，以事先辨識、評估及預測可能產生的人權風險及影響程度，採取相應的預防、緩解、補救措施，並持續監測該等措施的成效，依系統特性動態調整，必要時重新實施人權影響評估。過程中，需與受影響的利害關係人及公民團體合作，整合其意見評估納入影響評估程序。

六、平等參與數位轉型及永續發展

- (一) 國際勞工組織 (ILO) 研究指出，AI 最主要影響的可能為增強工作，而非完全取代該職位，對於勞動市場的潛在影響集中於文書人員、技術人員等知識工作，而女性在這些職業類別中的就業比例明顯高於男性。ILO 估計在全球範圍內，女性就業受到 AI 自動化影響之比例為男性的 2 倍以上¹⁶，此一性別差異可能削弱過去數十年來提升女性勞動參與的成果，並加深性別不平等。
- (二) 根據史丹佛大學《2025 年人工智慧指數報告》，AI 的商業應用正在加速普及，78% 的企業在 2024 年應用了 AI 技術，較前一年的 55% 有所提升；世界經濟論壇《2025 年未來就業報告》(Future of Jobs Report 2025) 也指出，全球 41% 的雇主計劃在未來 5 年內因 AI 自動化而裁員，預計到 2030 年，全球將有 1.7 億個工作被創造，同時有 9,200 萬個現有工作被取代。無法適應新職能的勞工面臨淘汰風險，尤其是無法使用 AI 者。
- (三) 雖然 AI 並未導致我國大規模失業，不過根據 NHRC 焦點座談，產業代表憂心部分勞工若無法即時適應新技術，將面臨失業風險，呼籲教育體系需加速職能訓練，以適應快速變化的就業趨勢。公民團體也指出，AI 可能擴大數位落差，例如都會區學校可支付 AI 專業版的費用，而偏鄉地區學校卻無此經費，導致城鄉學生學習機會不均，形成「AI 落差」；另老年、低技術人口也可能因為學習 AI 技術並與時俱進有困難，成為被邊緣化的對象，或無法進入勞動市場，進而加劇社會中之不平等與歧視。
- (四) 法國國家資訊自由委員會(CNIL)2024 年間因亞馬遜法國物流公司利用手持貨物掃描機追蹤員工工作績效及行為，對中斷、加速掃描行為精準計算，甚至要求員工對每次休息或中斷作出解釋，過度侵入監控員工行為及工作表現，造成員工身心壓力，違反歐盟《一般資料保護規則》(GDPR) 資料最小化及合法性¹⁷

¹⁶ ILO 指出，在高收入國家，具有高度自動化潛力的工作占女性就業的 8.5%，男性僅為 3.9%。在全球範圍內，3.7% 的女性就業可能受到 AI 自動化影響，男性僅為 1.4%

¹⁷ GDPR 第 6 條第 1 項 f 款規定，合法之資料處理，係為追求正當利益之目的所必須者。亞馬遜法國物

等原則，對其處以 3,200 萬歐元的罰鍰。顯示企業雖為提升營運效率，以 AI 自動化指標及演算法監控員工及評估績效，卻嚴重侵犯員工隱私及造成其身心壓力，不可不慎。

- (五) 非營利組織 AccessNow 執行長 Alejandro Mayoral Baños 與本會交流時指出，生成式 AI 係建構於模型訓練的基礎下，而模型訓練需仰賴資料與文字，惟原住民族語言很多都是口述傳承，沒有書寫文字，導致模型訓練需要更多資料建構；目前生成式 AI 對於原住民族的知識文化產出內容較不精確，容易造成誤解與刻板印象，也無法顯現各部落間的獨特性。
- (六) 聯合國原住民權利專家機制指出，數位化資訊可能使文化習俗遭到削弱，因為資料成為產品，而這些產品的呈現方式可能無法真實反映事實，也無法滿足原住民族的需求。原住民族對自身資訊缺乏控制，導致各種有害偏見的蔓延與文化挪用¹⁸。
- (七) 加拿大無障礙法案（Accessible Canada Act）授權訂定的無障礙標準，已於 2025 年 3 月發布〈CAN-ASC-6.2:無障礙與公平 AI 系統〉草案徵詢各界意見。該項標準強調 AI 開發從設計、測試到使用等各個環節，都必須有障礙者的參與；另外，AI 系統必須與科技輔具相互配合，避免為障礙者帶來新的問題。
- (八) AI 的快速發展亦帶來顯著的環境負擔，根據國際能源總署（IEA）預估，AI 資料中心的電力消耗將在 2026 年達到 1000 太瓦時（terawatt hours）¹⁹，相當於日本的用電量。我國經濟部則預估，AI 的用電需求在 2027 年將增加 200 萬瓩²⁰，比 2023 年成長約 8 倍（從約 24 萬瓩成長到約 228 萬瓩）。另有研究指出，AI 相關基礎設施所耗費的水資源將在 2027 年達到丹麥總人口的 6 倍。然而，即使目前國際間競相制定國家層次的人工智慧戰略，卻很少將永續發展考量在內。
- (九) 有專家學者於 NHRC 焦點座談表示，德國 2019 年 AI 燈塔計畫（Lighthouses of AI for Environment, Climate, Nature and

流公司所蒐集資料不符合正當利益，因其導致對員工的過度監控。

¹⁸文化挪用係指較強勢的文化群體在未經同意、未充分理解或尊重的情況下，採用或借用較弱勢文化群體的文化元素，如服裝、語言、音樂、傳統習俗等，可能導致對弱勢文化的誤解、歧視或嘲弄。

¹⁹「太瓦時」等於 1 兆（ 10^{12} ）瓦時（Wh），常用來衡量國家一年之總用電量。

²⁰「瓩」代表 1,000 瓦（W），1 度電就是 1,000 瓦耗電的用電器具，連續使用 1 小時所消耗的電量。

Resources) 支持資源高效利用的 AI 技術，策略性地將 AI 應用於環境和氣候保護(如生物多樣性保護、水資源管理等領域)，以環境友好跟 AI 交互應用為導向，展現政策前瞻性。

(十) NHRC 建議：

1. 政府應持續監測 AI 對於就業市場產生的影響，並與教育及勞動部門協作，建立完善的職能再培訓及就業安置體系，減輕 AI 自動化對勞工及不利處境群體造成的衝擊，並應避免 AI 所導致的勞動環境過度監控。
2. 政府加強數位基礎設施，確保偏鄉及弱勢族群享有平等的資訊近用權，並應將 AI 素養納入公民教育，涵蓋資安風險、資料自決權、演算法偏見、人機互動及數位倫理等議題，培養對 AI 的反思能力。
3. 政府應成立資料治理機制，鼓勵公私部門貢獻資料，建置原住民族語言等多語的 AI 語料庫，建置過程中確保原住民族共同參與治理，及行使對於文化、語言和決策等方面的集體權利，確保我國文化保存及傳承。
4. 政府應依據 CRPD 的通用設計原則，制定 AI 開發的無障礙標準，並且落實運用於政府採購與公共服務。
5. 政府應建立監管制度，提升大型或高耗能 AI 模型之耗能透明度，並研議使用 AI 工具協助解決環境問題，例如優化能源調度、監測空氣及水質污染、輔助資源循環管理、提升交通效率減少碳排放等，並結合 AI 提出環境永續策略，確保真正的永續發展。

伍、參考資料

一、國際文獻

- (一) AHRC, ‘Adopting AI in Australia’(Submission, May 2024)P.6
- (二) 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(Basic Act on the Development of Artificial Intelligence and the Establishment of Trust)
- (三) EU,AI Act
- (四) Government Digital Service ,‘Artificial Intelligence Playbook for the UK Government’(Playbook, February 2025)P.49
- (五) NHRCK, ‘Human Rights Guidelines on the Development and Use of Artificial Intelligence’(May, 2022)
- (六) NHRCK, ‘Human Rights Impact Assessment Tool for AI’(July,2024)
- (七) OECD,’Explanatory Memorandum On The Updated OECD Definition Of An AI System’(Paper, March 2024) P.4
- (八) OHCHR, ’Taxonomy of Human Rights Risks Connected to Generative AI’(Paper, November 2023)
- (九) OHCHR,’ Office of the United Nations High Commissioner for Human Rights Submission to the Global Digital Compact’(Submission, April 2023)
- (十) OHCHR, ‘Advancing Responsible Development and Deployment of Generative AI’(Headlines and Recommendations from UN B-Tech Foundational Paper, November 2023)
- (十一) The Council of Europe,’Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law’ (Explanatory Report, September 2024)
- (十二) UN, ‘Global Digital Compact’ (September 2024)

- (十三) UN, Ashwini K.P., 'Report of the Special Rapporteur on contemporary forms of racism, racial discrimination, xenophobia and related intolerance'(June 2024)
- (十四) UN, Gerard Quinn, 'Report of the Special Rapporteur on the rights of persons with disabilities on Artificial Intelligence and the rights of persons with disabilities'(December 2021)
- (十五) UN, The Expert Mechanism on the Rights of Indigenous Peoples,' Right of Indigenous Peoples to data, including with regard to data collection and disaggregation'(August 2025)
- (十六) UNESCO, 'Recommendation on the Ethics of Artificial Intelligence' (November 2021)
- (十七) UNICEF, 'Policy guidance on AI for children 2.0'(November,2021)
- (十八) World Economic Forum,' Future of Jobs Report 2025'(Report, January 2025)
- (十九) 人工知能関連技術の研究開発及び活用の推進に関する法律

二、網路資訊

- (一) Accessibility Standards Canada, 'CAN-ASC-6.2: Accessible and Equitable Artificial Intelligence Systems' Draft standard(March 2025), 取自：<https://accessible.canada.ca/creating-accessibility-standards/overview-asc-62-accessible-equitable-artificial-intelligence-systems>
- (二) CNIL, ' Employee monitoring: CNIL fined AMAZON FRANCE LOGISTIQUE €32 million'(January 2024), 取自：<https://www.cnil.fr/en/employee-monitoring-cnil-fined-amazon-france-logistique-eu32-million>

- (三) DataCentre MAGAZINE, Maya Derrick, 'What is the Truth About Future AI & Data Centre Emissions?'(July 2025), 取自：
<https://datacentremagazine.com/news/what-is-the-truth-about-future-ai-data-centre-emissions>
- (四) The Future of Privacy Forum (FPF), Dominic Paulger, 'Understanding Japan's AI Promotion Act: An "Innovation-First" Blueprint for AI Regulation'(July 2025), 取自：
<https://fpf.org/blog/understanding-japans-ai-promotion-act-an-innovation-first-blueprint-for-ai-regulation/>
- (五) Ha Dao Thu, 'Addressing AI Bias and Fairness: Challenges, Implications, and Strategies for Ethical AI'(April 2025), 取自：
<https://smartdev.com/addressing-ai-bias-and-fairness-challenges-implications-and-strategies-for-ethical-ai/>
- (六) International Labour Organization(2023), 'Generative AI and Jobs: A global analysis of potential effects on job quantity and quality', 取自：
<https://www.ilo.org/publications/generative-ai-and-jobs-global-analysis-potential-effects-job-quantity-and>
- (七) Li, P., Yang, J., Islam, M. A., & Ren, S. (2025). Making AI less "thirsty": Uncovering and addressing the secret water footprint of AI models, 取自：
<https://arxiv.org/pdf/2304.03271>
- (八) Megan Cerullo, 'AI is leading to thousands of job losses, report finds', CBS NEWS(August 2025), 取自：
<https://www.cbsnews.com/news/ai-jobs-layoffs-us-2025/>
- (九) Nature, 'Ethics and discrimination in artificial intelligence-enabled recruitment practices' (Review Article, September 2023), 取自：
<https://www.nature.com/articles/s41599-023-02079-x>
- (十) National Center for Missing & Exploited Children (2024), 'The Growing Concerns of Generative AI and Child Sexual

- Exploitation , 取 自 : <https://www.missingkids.org/blog/2024/the-growing-concerns-of-generative-ai-and-child-sexual-exploitation>
- (十一) Nestor Maslej, Loredana Fattorini, Raymond Perrault, Yolanda Gil, Vanessa Parli, Njenga Kariuki, Emily Capstick, Anka Reuel, ErikBrynjolfsson, John Etchemendy, Katrina Ligett, Terah Lyons, James Manyika, Juan Carlos Niebles, Yoav Shoham, Russell Wald, Tobi Walsh, Armin Hamrah, Lapo Santarlaschi, Julia Betts Lotufo, Alexandra Rome, Andrew Shi, Sukrut Oak. “The AI Index 2025Annual Report,” AI Index Steering Committee, Institute for Human-Centered AI, Stanford University, Stanford, CA,(April 2025)
- (十二) noyb,’noyb sends Meta ‘cease and desist’ letter over AI training. European class action as potential next step’ , 取 自 : <https://noyb.eu/en/noyb-sends-meta-cease-and-desist-letter-over-ai-training-european-class-action-potential-next-step>
- (十三) OECD.AI,GPAI(2025) , Lighthouses of AI for Environment, Climate, Nature, and Resources , 取 自 : <https://oecd.ai/en/dashboards/policy-initiatives/lighthouses-of-ai-for-environment-climate-nature-and-resources-1047>
- (十四) The Economist (2024) , The very real constraints on artificial intelligence in 2025 , 取 自 : <https://www.economist.com/the-world-ahead/2024/11/20/the-very-real-constraints-on-artificial-intelligence-in-2025>
- (十五) UN Environment Programme (2024) , AI has an environmental problem. Here’s what the world can do about that , 取 自 : <https://www.unep.org/news-and-stories/story/ai-has-environmental-problem-heres-what-world-can-do-about>

- (十六) UN Women, 'How AI reinforces gender bias—and what we can do about it' (February 2025), 取自：
<https://www.unwomen.org/en/news-stories/interview/2025/02/how-ai-reinforces-gender-bias-and-what-we-can-do-about-it>
- (十七) 台灣人工智慧實驗室 (2024), 2024 Taiwan Presidential Election Information Manipulation AI Observation Report, 取自：
<https://ailabs.tw/uncategorized/2024-taiwan-presidential-election-information-manipulation-ai-observation-report/>
- (十八) 資策會科技法律研究所, '韓國通過全球第二部「人工智慧基本法」, 提高國家人工智慧競爭力', 取自：
<https://stli.iii.org.tw/news2019-detail.aspx?d=664&no=57>
- (十九) 資策會科技法律研究所, '日本《人工智慧技術研究開發及應用促進法》簡介', 取自：
<https://stli.iii.org.tw/article-detail.aspx?tp=1&d=9379&no=64>

三、論壇、座談及活動

- (一) 2024.7.12 國家人權委員會「AI x 人權趨勢論壇」。
- (二) 2024.11.13 國家人權委員會與 Access Now 執行長交流活動。
- (三) 2025.1.20 國家人權委員會「探討及研究人工智慧對人權之影響」分眾座談第 1 場次：專家學者。
- (四) 2025.2.26 國家人權委員會 2025RightsCon 場邊活動「數位人權國際工作坊」。
- (五) 2025.3.31 國家人權委員會「探討及研究人工智慧對人權之影響」分眾座談第 2 場次：NGO。
- (六) 2025.4.14 國家人權委員會「探討及研究人工智慧對人權之影響」分眾座談第 3 場次：產業代表。
- (七) 2025.4.17 國家人權委員會「探討及研究人工智慧對人權之影響」分眾座談第 4 場次：政府機關。