

編號：3143

## 議題研析

### 一、題目：從比較法論 AI 發展下之兒少保護法制

### 二、議題所涉法規

人工智慧基本法、兒童及少年福利與權益保障法、個人資料保護法

### 三、背景說明

聯合國於 115<sup>1</sup> 年 7 月在日內瓦舉行首屆「全球人工智慧治理對話」(First Global Dialogue on AI Governance)，秘書長古特瑞斯提出<sup>2</sup>AI 正以各國制度與監理機關難以追趕速度發展，並倡議「人工智慧兒童安全承諾」<sup>3</sup>(AI Child Safety Pledge)，呼籲企業對兒少可接觸之 AI 系統進行兒少專屬安全測試及獨立監督，要求不得生成涉及兒童之性影像，且兒少出現危機訊號時停止一般互動及轉介真人支持之建議。此前，聯合國人工智慧獨立國際科學小組之初步報告<sup>4</sup>亦指出，AI 生成兒少性虐待素材、深偽技術促成之性暴力、迎合式回應(sycophancy)及情感依賴等風險已快速增加，且快於科學衡量與治理工具，引發外界擔憂。

### 四、問題爭點

查我國《人工智慧基本法》已於 115 年 1 月 14 日公布施行，國家科學及技術委員會為中央主管機關，並由數位發展部依同法第 16 條

<sup>1</sup> 本文有關年分之使用，原則以民國紀年表述，惟涉及外國法制或立法例部分，改採西元紀年表述。

<sup>2</sup> United Nations, Secretary-General's Remarks to the Opening of the First Global Dialogue on Artificial Intelligence Governance, 115 年 7 月 6 日，網址：<https://www.un.org/sg/en/content/sg/statements/2026-07-06/secretary-generals-remarks-the-opening-of-the-first-global-dialogue-artificial-intelligence-governance-delivered>，最後瀏覽日期：115 年 7 月 24 日。

<sup>3</sup> 鉅亨網，聯合國警示 AI「災難性危害」 日內瓦峰會呼籲建立全球治理機制，115 年 7 月 8 日，網址：<https://news.cnyes.com/news/id/6526943>，最後瀏覽日期：115 年 7 月 24 日。

<sup>4</sup> United Nations, The Preliminary Report of the Independent International Scientific Panel on AI，網址：<https://www.un.org/independent-international-scientific-panel-ai/en/preliminary-report>，最後瀏覽日期：115 年 7 月 24 日。

規定推動 AI 風險分類框架<sup>5</sup>；另依同法第 18 條規定，政府應於本法施行後 2 年內完成相關法規之制（訂）定、修正或廢止及行政措施之改進。惟現處法規調適期程內，目前規範多屬原則及政策指示，兒少可接觸 AI 之上市前測試、設計義務、持續監測、事故通報、年齡確認、危機介入與責任歸屬等事項尚未周全。是以，簡介各國相關法制規定及經驗，以供我國檢討修正參考。

## 五、探討研析

### （一）AI 發展下之兒少風險與我國現行法制現況

近年兒少使用生成式 AI（Generative AI）已非偶發現象，美國一項使用特定家長監護應用程式之家庭為樣本、並涵蓋 6,488 名 4 歲至 17 歲兒少之研究<sup>6</sup>顯示，約有三分之一兒少曾使用生成式 AI，且 15 歲至 17 歲者約達半數；另英國網路觀察基金會（Internet Watch Foundation, IWF）於 2025 年評估 AI 生成影像或影片呈現寫實之兒童性虐待內容<sup>7</sup>已達 8,029 件，數量與內容嚴重程度均呈快速升高趨勢。除色情、暴力及詐欺內容外，陪伴型聊天機器人亦可能透過記憶、擬人化、迎合式回應及長時間互動，形成依賴、錯誤信任或危機回應失靈；且部分以提升互動或停留時間為目標之推薦系統，亦可能反覆推送自傷、飲食失調或性化等內容。是以，目前僅以內容下架或事後刑罰處理，恐已不足涵蓋模型設計、個人化推送、資料蒐集及持續互動產生之風險。

查我國《人工智慧基本法》第 4 條揭示永續發展與福祉、人類自主、隱私保護與資料治理、資安與安全、透明與可解釋、公平與不歧視及問責等 7 項基本原則；第 5 條第 2 項以兒少最佳利益為原則；第 16 條至第 18 條並要求建立風險分類、責任救濟及於 2 年內完成法規

---

<sup>5</sup> 行政院，人工智慧基本法施行規劃，115 年 5 月 21 日，網址：<https://www.ey.gov.tw/Page/448DE008087A1971/f784469b-fe42-438b-8788-18846a53cec4>，最後瀏覽日期：115 年 7 月 24 日。

<sup>6</sup> Anne J. Maheux、Samir Akre-Bhide、Debra Boeldt，Generative Artificial Intelligence Applications Use Among US Youth，JAMA Network Open，115 年 2 月 2 日，網址：<https://jamanetwork.com/journals/jamanetworkopen/fullarticle/2844596>，最後瀏覽日期：115 年 7 月 24 日。

<sup>7</sup> IWF，AI-Generated Child Sexual Abuse: 2026 Report on Trends, Data & Human Impact，網址：<https://www.iwf.org.uk/about-us/why-we-exist/our-research/how-ai-is-being-abused-to-create-child-sexual-abuse-imagery/>，最後瀏覽日期：115 年 7 月 24 日。

調整。另數位發展部於 115 年公布之「AI 應用發展兒少影響評估報告」<sup>8</sup>，亦將深偽冒充身分詐騙或合成性影像勒索、演算法推送不適齡內容、AI 生成或散布有害身心健康之內容、情感陪伴型 AI 之操縱或欺騙，及蒐集兒少資料所生隱私風險等，列為優先關注之風險情境；再者，上開報告明載其屬階段性觀察，且未充分納入兒少及照顧者之第一手意見；現行法制對相關行為，於其適用範圍內仍有規範空間。惟查現行機制仍分散於網路內容申訴與自律、兒少性剝削防制及個案保護等面向，且《個人資料保護法》等尚未就兒少個人資料之處理，另定專門之行為追蹤限制及較嚴格之同意機制，允宜由相關主管機關持續研議，以完備兒少保護法制。

## （二）外國立法例與規範經驗

### 1. 歐盟：

查歐盟《人工智慧法》（Artificial Intelligence Act）<sup>9</sup>禁止利用自然人或特定群體因年齡等因素所生弱點，意圖或實際實質扭曲其行為，並造成或合理可能造成重大損害之 AI 應用行為；並將用於入學、學習成果評估等特定教育用途之 AI 系統列為高風險；另《數位服務法》<sup>10</sup>（Digital Services Act, DSA）規定，兒少可使用之平臺應採取適當且合比例之措施，以確保未成年人享有高度之隱私、安全及保障；平臺合理確定服務接受者為未成年人時，不得利用該服務接受者之個人資料進行剖析（profiling），並據以向其投放廣告。歐盟執委會於 2025 年發布兒少保護指引（Guidelines on the Protection of Minors）<sup>11</sup>，進一步處理誘拐、有害內容、成癮設計及網路霸凌，惟該指引僅作為執委會評估平臺遵循《數位服務法》情形之參考基準；另於同年公布年齡驗證方案（Age Verification

<sup>8</sup> 數位發展部，AI 應用發展兒少影響評估報告，115 年 3 月 26 日，網址：<https://www-api.moda.gov.tw/File/Get/moda/zh-tw/8bgQFPbXjgAjFSQ>，最後瀏覽日期：115 年 7 月 24 日。

<sup>9</sup> European Union，REGULATION（EU）2024/1689 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL，113 年 6 月 13 日，網址：<https://eur-lex.europa.eu/eli/reg/2024/1689/oj>，最後瀏覽日期：115 年 7 月 24 日。

<sup>10</sup> European Union，The Digital Services Act，網址：<https://digital-strategy.ec.europa.eu/en/policies/digital-services-act>，最後瀏覽日期：115 年 7 月 24 日。

<sup>11</sup> European Commission，Commission publishes guidelines on the protection of minors，114 年 7 月 14 日，網址：<https://digital-strategy.ec.europa.eu/en/library/commission-publishes-guidelines-protection-minors>，最後瀏覽日期：115 年 7 月 24 日。

Solution)<sup>12</sup>之藍圖、於2026年達到功能就緒並發布會員國建置建議，將年齡確認建構在與歐盟數位身分錢包（EU Digital Identity Wallets）相同技術規格之上，以證明使用者達特定年齡，並可不傳送其他身分資料，以符合資料最小化。

## 2. 英國：

英國《線上安全法》<sup>13</sup>（Online Safety Act 2023）要求受規範之服務提供者進行兒少風險評估，並採取適當之安全措施；其中，對色情等特定內容，要求採取高度有效之年齡確信措施（highly effective age assurance），以防止兒少接觸色情、自傷、自殺及飲食失調等內容；另英國教育部於2025年發布適用於教育場域之生成式AI產品安全標準（Generative AI: product safety standards）<sup>14</sup>，並於2026年更新，除揭示應建立活動紀錄及危險訊號警示機制、指定兒少保護聯絡窗口外，並持續監測使用者在認知卸載（Cognitive Offloading）、情感投入及使用時間等面向情形。惟查單純提供封閉式對話之陪伴型AI未必均屬該法規範圍，顯示以既有平臺法規處理AI仍可能出現適用之漏洞。

## 3. 美國：

美國聯邦貿易委員會（The Federal Trade Commission）於2025年修正《兒童線上隱私權保護規則》（Children's Online Privacy Protection Rule, COPPA rule）<sup>15</sup>規定，要求未滿13歲兒童之資料用於定向廣告或提供第三人時，需另行取得家長明確同意，並強化線上服務業者（Operator）之義務，例如資料保存限制及最小化要求。另2025年10月13日加州州長簽署《參議院第243號法案》（下稱SB

---

<sup>12</sup> European Commission, The EU Approach to Age Verification, 網址：<https://digital-strategy.ec.europa.eu/en/policies/eu-age-verification>, 最後瀏覽日期：115年7月24日。

<sup>13</sup> legislation.gov.uk, Online Safety Act 2023, 112年10月26日, 網址：<https://www.legislation.gov.uk/ukpga/2023/50>, 最後瀏覽日期：115年7月24日。

<sup>14</sup> GOV.UK, Generative AI: Product Safety Standards, 115年1月19日, 網址：<https://www.gov.uk/government/publications/generative-ai-product-safety-standards/generative-ai-product-safety-standards>, 最後瀏覽日期：115年7月24日。

<sup>15</sup> Federal Trade Commission, FTC Finalizes Changes to Children's Privacy Rule Limiting Companies' Ability to Monetize Kids' Data, 114年1月16日, 網址：<https://www.ftc.gov/news-events/news/press-releases/2025/01/ftc-finalizes-changes-childrens-privacy-rule-limiting-companies-ability-monetize-kids-data>, 最後瀏覽日期：115年7月24日。

243)<sup>16</sup>，制定陪伴型聊天機器人 (Companion chatbots) 安全規範，同日，加州州長並否決了限制範圍較廣之 AB 1064 法案<sup>17</sup>，以避免該法案可能無意間造成未成年人之全面禁用<sup>18</sup>，進而妨礙青少年學習安全使用及辨識 AI 風險。另查，SB 243 法案除要求揭露使用者與 AI 互動，對未成年人定期提醒休息，建立自傷危機防制與轉介程序外，並採取合理措施以防止生成描繪露骨性行為之視覺素材，或直接要求未成年人從事露骨性行為，且加州民間團體刻正推動兒少 AI 安全公民提案<sup>19</sup>，其後續發展仍待觀察。

### (三) 結論

綜上，AI 對兒少之危害具有跨內容、資料、產品設計與心理互動之特性，僅靠業者自律恐欠缺外部問責，而概括且無例外之全面禁止，亦可能衍生比例原則及兒童發展、教育與參與權之疑義。爰建議各主管機關允宜以《兒童權利公約》(Convention on the Rights of the Child, CRC) 所揭示之「兒童最佳利益」(the best interests of the child) 為優先考量<sup>20</sup>，兼顧比例原則、資料最小化、技術中立、正當程序及有效救濟等原則，並參酌前開各國法制及政策經驗，研議如何在兒少保護、隱私、發展權、教育利益及產業創新間取得衡平，透過建立事前預防、事中監測及事後救濟等治理機制，避免將保護責任轉由兒少或家長承擔，俾使 AI 創新得在可驗證、可問責及符合兒少權利之條件下發展。

撰稿人：黃珮瑜

---

<sup>16</sup> California Legislative Information, Senate Bill No. 243, 114 年 10 月 13 日，網址：[https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill\\_id=202520260SB243](https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202520260SB243)，最後瀏覽日期：115 年 7 月 24 日。

<sup>17</sup> Governor Gavin Newsom, OFFICE OF THE GOVERNOR, 114 年 10 月 13 日，網址：<https://www.gov.ca.gov/wp-content/uploads/2025/10/AB-1064-Veto.pdf>，最後瀏覽日期：115 年 7 月 24 日。

<sup>18</sup> Associated Press, California governor vetoes bill to restrict kids' access to AI chatbots, 114 年 10 月 14 日，網址：<https://apnews.com/article/california-chatbots-children-safety-ai-newsom-33be4d57d0e2d14553e02a94d9529976>，最後瀏覽日期：115 年 7 月 24 日。

<sup>19</sup> Parents and Kids Safe AI Coalition, PARENTS AND KIDS SAFE AI COALITION LAUNCHES CAMPAIGN TO PASS NATION'S STRONGEST CHILD AI SAFETY LAW, 115 年 3 月 17 日，網址：<https://parentsandkidssafeai.com/2026/03/17/parents-and-kids-safe-ai-coalition-launches-campaign-to-pass-nations-strongest-child-ai-safety-law/>，最後瀏覽日期：115 年 7 月 24 日。

<sup>20</sup> 聯合國兒童權利公約 CRC 資訊網，兒童權利公約，113 年 11 月 20 日，網址：<https://crc.sfaa.gov.tw/PublishCRC/CommonDetail?documentId=45733006-3802-4A3F-BBA1-91D23E1195AE>，最後瀏覽日期：115 年 7 月 24 日。