

AI時代的個資與隱私權新課題

黃允曄*

壹、AI浪潮下的法律新課題

一、從數位化到智慧化：個資與隱私權的演變

20世紀末的資訊化浪潮，使社會邁入數位化階段。數位化的特徵，在於運用資訊技術將現實世界中的各類資料轉換為可儲存、傳輸與運算的數位格式。當時的個人資料保護焦點，多半集中在「電腦安全」或是「資料庫安全」，社會關注個資外洩問題。而1990年代金融業及電信業的電腦化，使得消費者個人資訊集中儲存於中央資料庫，若遭入侵或內部人員濫用，將導致嚴重的隱私風險。因此，早期的個資保護措施，多以技術防護（防火牆、存取控制）與基本法律規範為主。

然而，進入21世紀後，人工智慧技術的突飛猛進，將人類社會從「數位化」推向「智慧化」時代。智慧化的特徵，在於資料不僅被儲存與處理，更能透過演算法進行模式辨識、行為預測，甚至自動決策。例如，電子商務能透過客戶對產品的點擊瀏覽軌跡，預測其個人偏好而精準推薦相關商品；智慧金融能依據交易模式，判斷是否涉及詐欺或洗錢；智慧醫療可透過病歷與影像資料，預測

個體的疾病風險。這些應用雖大幅提升效率與精準度，卻也引發了全新的隱私與個資風險：資料的蒐集規模遠超過過去，資料利用的範圍也更難預測。

在這樣的轉型下，個人資料與隱私權不再只是「不被外界干擾」的消極權利，而是涉及「如何在龐大的資料流通環境下，仍保有對個人資訊的自主控制」的積極權利。這種演變，也使傳統的法律框架逐漸顯得不足，需要新的制度性回應。

二、AI技術的雙面刃：效率提升與權利侵害的衝突

AI技術被譽為「雙面刃」。一方面，它能顯著提升社會運作效率，甚至挽救生命；另一方面，它也可能造成前所未有的基本權利侵害。

在正面效益方面，AI在醫療領域的應用，已展現出超越人類專家的診斷能力。例如，Google DeepMind所研發的AI眼科診斷系統，可在短時間內辨識糖尿病視網膜病變，其準確率與資深醫師不相上下。又如金融機構透過AI反詐騙系統，能即時攔截異常交易，保護消費者財產安全。這些應用都展現了AI對公共利益的重大貢獻。

然而，AI的負面效應也逐漸浮現。首先，

* 本文作者係中華民國內部稽核協會理事長

大規模資料蒐集與分析，可能使個體失去對自身資訊的掌控。例如，社群平台蒐集用戶的瀏覽習慣與貼文內容，進而推送特定廣告或政治訊息，這不僅涉及個資濫用，更可能影響民主程序。其次，演算法偏見問題亦日益嚴重。若AI模型以不平衡的數據進行訓練，結果可能出現歧視性決策，例如在招聘員工過程中可能不當排除少數族群。最後，AI系統的決策過程往往是黑箱，難以解釋，導致被影響的個人無法有效行使救濟權利。

因此，AI的價值不僅取決於技術本身，更取決於其治理方式。如果缺乏完善的法律與倫理框架，AI所帶來的風險可能遠超過其利益。

三、本文重點

在此發展背景下，本文將系統性探討AI時代的個資與隱私權議題，並特別關注在此脈絡下我們當前的挑戰與回應。重點疇涵蓋以下三個面向：

（一）歷史與理論基礎

回顧個資與隱私權保障的發展脈絡，說明其在憲法與人權上的定位，並解析從OECD原則到GDPR的演變。

（二）AI帶來的新挑戰

分析生成式AI在資料蒐集、再利用、偏見與外洩上的風險，並以實際案例說明其對個資與隱私權的衝擊。

（三）國際法規與台灣觀點

比較歐盟、美國、日本及對岸等主要國家的立法動態，並檢視我國現行法制的優劣，並提出修法與實務建議。

並試著回答以下核心問題：

- 1.在AI時代，個資與隱私權保障的核心挑戰為何？
- 2.現有法律制度是否足以因應這些挑戰？
- 3.在國際AI治理的脈絡中我們應如何定位？

貳、歷史回顧：個資與隱私權法規的發展軌跡

一、個資保護的緣起與發展：從資訊自決權到資料主體權

個人資料保護作為一項法律議題，其發展脈絡可追溯至20世紀中葉。當時，隨著電腦與資料庫技術的普及，政府與企業開始大規模蒐集與處理個人資料，引發社會對「資訊濫用」的憂慮。早期的討論多集中於「檔案隱私」問題，例如政府戶政資料或銀行信用紀錄，是否可能被不當利用，進而對公民權利造成侵害。

在此背景下，德國於1970年率先制定黑森邦資料保護法，被視為全球第一部個資保護法。該法以「保障公民免受國家機關過度監控」為核心，展現了個資保護與民主自由之間的緊密關聯。隨後，其他歐洲國家如瑞典（1973）、美國（1974年隱私法）、法國（1978年資訊與自由法）也相繼制定相關立法，逐漸形成跨國性的制度性潮流。

然而，這些早期法規多著重於行政管理，對於個人權利的明確建構尚屬不足。直到1983年，德國聯邦憲法法院在著名的「人口普查案」中，正式提出「資訊自決權」概

念，才真正奠定了個資保護的人權基礎。

二、資訊自決權的憲法基礎

1983年的「人口普查案」是個資保護史上的里程碑。當時，西德政府計畫進行全國性人口普查，要求公民提供包含家庭、職業、健康等多面向的資訊。民眾質疑此舉可能使國家對公民生活進行全面監控，於是提起憲法訴訟。

德國聯邦憲法法院最終判決認定，公民有權自主決定其個人資料是否被蒐集、如何使用，這種「對個人資訊命運自主決定的權利」，即為「資訊自決權」。法院指出，若公民無法控制自身資料，將導致「人格自由發展」受到威脅，進而危及民主制度的基礎。

此一判決的意義，在於將個資保護上升為憲法層級的基本權利，並賦予其與人格尊嚴、自由權等人權價值的連結。資訊自決權不僅影響德國後續立法，更啟發了歐洲人權法院及歐盟立法的方向，使個資保護逐漸成為國際人權體系的一環。

三、國際法規的萌芽：OECD隱私權保護八大原則

隨著跨境資料流通日益頻繁，僅依賴國內法規已難以解決問題。1980年，經濟合作暨發展組織（OECD）制定「隱私權與個人資料跨境流通指導原則」，提出八項核心原則，被視為國際個資保護的共同基礎：

（一）**蒐集限制原則（Collection Limitation Principle）**：資料應以合法且公平的方式蒐集，並經過當事人知情或同意。

（二）**資料品質原則（Data Quality Principle）**：蒐集之資料應正確、完整且隨時更新。

（三）**目的明確原則（Purpose Specification Principle）**：蒐集資料時應明確告知使用目的。

（四）**使用限制原則（Use Limitation Principle）**：資料不得用於蒐集目的以外之用途，除非獲得當事人同意或有法律依據。

（五）**安全保障原則（Security Safeguards Principle）**：應採取合理安全措施，避免資料外洩或不當存取。

（六）**公開原則（Openness Principle）**：資料蒐集與使用過程應具透明性。

（七）**個人參與原則（Individual Participation Principle）**：資料主體有權查詢、更正或刪除其個資。

（八）**責任原則（Accountability Principle）**：資料控制者應對遵循上述原則負責。

這八大原則具有高度前瞻性，後續的歐盟指令、GDPR，乃至於APEC跨境隱私規則（CBPR）皆受其影響。值得注意的是，OECD原則並非強制性法律，而是作為「軟法」促進各國政策協調。但其價值在於建立了跨國共通語言，使個資保護逐步走向全球化。

四、歐盟GDPR的誕生：全球個資保護的里程碑

1995年，歐盟頒布資料保護指令（Directive 95/46/EC），要求成員國制定國內法保障個資，並規範跨境資料傳輸。此舉雖具有重大意義，但由於指令必須透過各國轉化實施，導致各國標準不一，執行效果不彰。

2012年，歐盟委員會提出GDPR草案，經過多年辯論，最終於2016年通過，2018年正式生效。GDPR被視為全球最嚴格的個資保護法，其創新之處包括：

- (一) **域外效力**：不論企業是否設立於歐盟，只要處理歐盟居民的個資，即需遵循GDPR。
- (二) **高額罰款**：違反規範可處全球營業額4%或2000萬歐元的罰鍰。
- (三) **資料主體權利**：明確保障知情權、更正權、刪除權（被遺忘權）、資料可攜帶權。
- (四) **資料保護影響評估（DPIA）**：對高風險處理行為須進行事前評估。
- (五) **資料保護官（DPO）**：特定機構需設立專責人員負責。

GDPR的出現，不僅改變歐洲數位市場規範，更迫使全球企業提升個資保護標準。例如，美國科技巨頭Google、Facebook都因GDPR被歐盟監管機構鉅額裁罰，成為國際媒體關注焦點。GDPR因此被譽為「個資保護的世界標竿」。

五、隱私權的內涵演變：從「被遺忘權」到「知情權」

傳統隱私權通常被定義為「不受干擾的生活權」（The Right to be Let Alone），源於1890年美國學者Warren與Brandeis的經典論文。這種概念強調的是個人免於外界打擾的空間。然而，在數位時代，隱私權逐漸轉化為「資訊控制權」。

最具代表性的例子是「被遺忘權」（Right to be Forgotten）。2014年，歐洲法院在Google Spain SL v. AEPD判決中，裁定個人可

要求搜尋引擎刪除過時或不相關的資訊，避免其長期影響名譽。此案標誌著隱私權不再僅是「消極不干擾」，而是「積極控制資訊流通」。

此外，「知情權」亦成為新興權利。GDPR規定，資料主體有權知悉其個資如何被蒐集、處理與利用，並要求資料控制者提供清楚透明的告知。這不僅賦予公民更多資訊自主性，也迫使企業強化透明度。

六、隱私權與個資保護的關係辨析：異同與交集

在學理上，隱私權與個資保護雖重疊卻不完全相同。隱私權是一種憲法保障的基本權利，重點在於維護個人尊嚴與自由；個資保護則是具體化的制度性規範，重點在於規範資料流通與控制責任。

舉例而言，若某人遭受媒體偷拍，這涉及隱私權侵害；若某公司未經同意蒐集並販售客戶資料，則屬於個資保護的違反。兩者交集極大，尤其在AI時代，個資外洩往往同時侵害隱私權。這也是為何現代立法趨勢傾向將二者整合，例如GDPR既強調隱私權的尊重，也建構個資保護的制度化規範。

參、生成式AI帶來的個資與隱私權新挑戰

人工智慧的發展並非首次衝擊個資與隱私權，但「生成式AI」（Generative AI）的出現，確實將此議題推向前所未有的高度。與傳統的機器學習不同，生成式AI能透過龐大資料訓練，生成具創造性的文字、影像、語

音甚至影片。這種技術突破雖然帶來創新應用，但同時也讓資料蒐集、利用與再製的爭議加劇。以下將分層次探討生成式AI帶來的個資與隱私權新挑戰。

一、訓練資料庫的黑箱：合法性與透明度的爭議

生成式AI仰賴龐大的訓練資料庫。然而，這些資料的來源常常不透明。以OpenAI的ChatGPT為例，其訓練資料包括公開網頁、書籍、論文與論壇對話，但具體來源並未完全公開，導致外界難以判斷是否涉及未經授權的個人資料使用。

這種「黑箱式訓練」引發兩大爭議：

（一）合法性問題

若資料包含個人資訊（如姓名、聯絡方式、社群貼文），是否違反資料蒐集的合法依據？在GDPR的框架下，資料處理必須有明確法律基礎，例如同意、契約、法定義務或合法利益。但在AI模型訓練情境下，是否能以「合法利益」作為依據，仍具高度爭議。

（二）透明度問題

資料主體往往不知其資料是否被用於AI訓練，也無從得知資料如何影響模型的輸出。這不僅違反「知情同意」原則，也削弱了個人對自身資訊的控制力。

在美國，已有作者與藝術家對生成式AI公司提起訴訟，指控其未經同意使用創作資料進行訓練，侵犯著作權與隱私權，法院也已判決。此類案件顯示了資料合法性與透明度不足的風險。

二、大規模資料蒐集的合法性基礎：同意、公眾利益或合法權益？

AI系統經常需要蒐集規模龐大的個人資料，例如影像辨識系統需要大量人臉影像，語音辨識系統需要各種口音的語音檔案。問題在於：這些資料蒐集是否符合法律上的正當性？

在傳統個資保護法中，最常見的基礎是「同意」。然而，在生成式AI時代，要求逐一取得同意幾乎不可能。例如，若AI開發者從網路公開平台（如Twitter、YouTube）蒐集數以億計的貼文或影片，要逐一徵求使用者同意顯然不切實際。

於是，有些業者主張可以依據「合法權益」或「公眾利益」作為合法基礎。然而，這種主張在不同法域下有不同解釋：

- （一）**歐盟**：GDPR要求資料控制者必須平衡自身利益與資料主體基本權利，並非所有商業利益皆可作為正當基礎。
- （二）**美國**：多數情況下傾向容忍商業利用，除非涉及敏感個資或明顯侵權。
- （三）**中國**：個人資訊保護法則強調敏感個資須取得明示同意，國家安全與社會公共利益例外。

因此，生成式AI資料蒐集合法性往往陷於灰色地帶，也加劇了國際間規範落差。

三、數據再利用與目的限制原則的衝突

GDPR與台灣個資法均強調「目的限制原則」，即資料僅能用於蒐集時明確告知的目的。然而，生成式AI的特徵之一，就是將既有資料用於新的用途，例如：

- （一）原本為醫療研究蒐集的數據，被再利用於AI健康管理應用；
- （二）原本為社群交流的貼文，被再利用於情感分析模型。

這種「再利用」往往與原始目的不符，導致法律上的爭議。例如，歐洲監管機構曾質疑Clearview AI的人臉辨識系統，因其將社群平台的公開照片蒐集後再利用於執法，已明顯違反目的限制原則。

此處挑戰在於：若完全禁止再利用，AI創新將受到嚴重阻礙；但若完全開放，個人對資訊的控制權將蕩然無存。如何在「資料流通」與「資料主體控制」之間取得平衡，是當代法規的核心難題。

四、數據偏見（Data Bias）與歧視性決策的風險

AI系統的輸出結果取決於輸入數據。若訓練數據存在偏見，AI模型將複製甚至放大這些偏見，導致不公平或歧視性決策。

知名案例包括：

（一）亞馬遜招聘AI系統

該系統在訓練過程中，因歷史數據中男性求職者比例較高，導致AI模型自動對包含「女性」字眼的履歷扣分，被批評為性別歧視。

（二）美國刑事司法系統的「COMPAS」演算法

該系統用於預測再犯風險，研究發現其對非裔被告的誤判率顯著偏高，涉及種族歧視問題。

在生成式AI中，數據偏見可能以更隱晦的方式出現。例如，若模型常將「護理師」與「女性」關聯，將強化職業性別刻板印象；若模型將特定地區的名字與「犯罪」關聯，則可能導致族群歧視。

這些問題凸顯出，AI偏見不僅是技術問題，更是法律與倫理問題。若AI系統做出歧

視性決策，責任應由誰承擔？是開發者、使用者，還是資料提供者？這是法律必需釐清的課題。

五、AI模型的決策過程與個資保護：可解釋性與當責性

GDPR第22條保障「不受僅基於自動化決策影響的重要權利」，同時強調資料主體有權要求「有關邏輯的解釋」。這反映了歐盟對AI「可解釋性」（Explainability）的重視。

然而，在生成式AI中，決策過程往往難以完全解釋。深度學習模型具有高度複雜性，即使是開發者也無法清楚指出某一結果的具體原因。這造成兩大困境：

（一）透明度不足：資料主體難以理解AI為何做出某一輸出，無法判斷是否涉及偏見或錯誤。

（二）責任歸屬模糊：若AI產生錯誤結果（例如錯誤拒絕貸款申請），責任究竟在於模型開發者、資料蒐集者，還是金融機構？

這正是「演算法黑箱」問題的核心。部分學者建議透過模型稽核或可解釋AI技術（XAI）來提升透明度，但目前仍處於發展階段。

六、個資外洩與二次利用的風險

生成式AI帶來新的個資外洩型態，超越傳統的資料庫洩漏。過去，外洩通常意味著駭客入侵伺服器，竊取大量資料；如今，外洩可能直接發生在「模型」本身：

（一）模型反推攻擊（Model Inversion Attack）：攻擊者可藉由模型輸出，反推出原始訓練數據中的個人資訊。例如，若醫療影像資料用於訓練，攻

擊者可能重建出患者的臉部特徵。

(二) **資料提取攻擊 (Data Extraction Attack)**：研究顯示，透過特定查詢，生成式AI可能輸出原始訓練資料中的內容，例如電子郵件地址或身分證號。

這些新型外洩模式，使得「資料庫」與「模型」之間的界線模糊，傳統個資保護措施可能難以因應。

七、深度偽造 (Deepfake) 與個資再製：身分盜用與名譽損害

深度偽造技術 (Deepfake) 是生成式AI最具爭議的應用之一。它能以少量影像或語音樣本，生成高度擬真的偽造內容。這帶來嚴重的個資與隱私侵害：

- (一) **身分盜用**：攻擊者可利用Deepfake製作假冒聲音或臉部影像，進行詐騙。例如，2024年香港曾發生一起詐騙案件，犯罪集團利用Deepfake視訊冒充公司高層，成功騙取高達二千五百萬美元。
- (二) **名譽損害**：Deepfake常被用於製作不實影像，如假冒名人或一般人的色情影片，嚴重侵害人格尊嚴與名譽權。
- (三) **民主風險**：若Deepfake影片被用於政治領域，可能影響選舉結果，威脅民主制度。

目前，多數國家仍缺乏針對Deepfake的專門法規。雖然可援引誹謗、詐欺或著作權法進行處理，但在跨境網路環境下，實務執行困難重重。

八、小結

生成式AI帶來的個資與隱私新挑戰主要包

括：

- (一) **資料合法性與透明度不足**：訓練資料來源不清，缺乏有效告知與同意。
- (二) **目的限制的衝突**：資料再利用與既有法規矛盾。
- (三) **偏見與歧視風險**：AI模型可能強化社會不平等。
- (四) **演算法黑箱問題**：決策過程難以解釋，責任歸屬模糊。
- (五) **新型外洩樣態**：模型反推與數據提取攻擊。
- (六) **Deepfake帶來的名譽與安全風險**。

這些挑戰不僅是技術問題，更涉及法規設計、倫理規範與社會治理。AI時代的個資與隱私保障，必須從傳統的「防止外洩」轉向「全流程治理」，涵蓋資料蒐集、訓練、使用、刪除等各階段。

肆、國際法規的最新發展與回應

人工智慧技術的快速發展，使得各國法律制度面臨重大挑戰。個資與隱私權保護已非單一國家內部的議題，而是跨境流動下的全球治理難題。尤其生成式AI技術所引發的資料蒐集、再利用與責任歸屬問題，更迫使各國加速調整法規。以下將依序檢視歐盟、美國、中國、日本與新加坡的立法與政策動態，並比較不同模式的優劣。

一、歐盟GDPR的實務挑戰與因應

(一) GDPR在AI時代的限制

2018年正式生效的通用資料保護規範 (GDPR)，被譽為「全球最嚴格的個資保護

法」。它不僅規範個人資料的蒐集與利用，也賦予個人包括知情權、更正權、刪除權（被遺忘權）、資料可攜帶權等「資料主體權利」。然而，生成式AI的特性使GDPR的部分制度面臨挑戰：

- 1.被遺忘權的執行困難：若某人的資料被用於AI模型訓練，要完全「刪除」該資料並確保不再影響模型行為，技術上幾乎不可行。
- 2.資料可攜帶權的困境：AI模型往往透過大量非結構化資料訓練，如何以「可攜帶格式」提供資料主體其個資，難度極高。
- 3.自動化決策的規範不足：GDPR第22條禁止「僅基於自動化決策」對個人產生重大影響，但在現實中，大量AI應用（如金融信用評分、保險理賠、招聘篩選）仍屬「部分自動化」，因此落在灰色地帶。

（二）GDPR執法實務案例

歐盟各國監管機構已針對AI相關個資問題作出多起裁決：

- 1.Clearview AI案（2022）：義大利、法國與希臘監管機構均裁定Clearview AI違反GDPR，因其未經同意蒐集社群平台的人臉照片，並用於人臉辨識服務，要求停止處理並課以高額罰款。
- 2.ChatGPT案（2023）：義大利個資保護機構（Garante）曾暫時禁止ChatGPT在當地提供服務，理由是其未能充分保障資料主體權利，特別是兒童與青少年隱私。此案迫使OpenAI改進告知機制並增設年齡驗證功能。

這些案例說明，GDPR雖然設計完備，但在AI情境下的落實需要更多技術與政策配套。

二、歐盟人工智慧法案（AI Act）：全球首部AI專法

（一）立法背景與目標

2021年，歐盟委員會提出人工智慧法案（AI Act）草案，2024年底最終達成政治協議，並於2025年初正式通過。這是全球第一部專門針對AI制定的法律，其目的在於：

- 1.確保AI發展符合基本權利與安全標準；
- 2.建立風險導向的監管架構；
- 3.統一歐盟市場的AI規範，促進產業競爭力。

（二）風險分級監管模式

AI Act採用「風險分級」模式，依不同風險程度設定管制義務：

- 1.禁止性AI系統（Unacceptable Risk）：例如社會信用評分、大規模社會監控，全面禁止。
- 2.高風險AI系統（High Risk）：涉及醫療、金融、教育、就業、基礎設施等領域，必須符合資料治理、透明度、可追溯性、人工監督等嚴格要求。
- 3.有限風險系統（Limited Risk）：如聊天機器人，需遵循透明度義務，告知使用者其與AI互動。
- 4.最低風險系統（Minimal Risk）：例如遊戲或垃圾郵件過濾器，則無特別限制。

（三）高風險AI系統的個資保護義務

對於高風險AI，AI Act特別要求：

- 1.資料品質：訓練數據須具代表性、正確性與無偏性。
- 2.可追溯性：須建立完整文件，記錄模型訓練與運作過程。
- 3.透明度：必須提供清楚說明，讓使用者了解系統運作原理。
- 4.人類監督：確保重要決策最終由人類負責，而非完全交由自動化系統。

這些義務與GDPR相互補充，形成「雙重治理」：GDPR強調資料主體權利，AI Act則針對系統開發與使用階段設定義務。

（四）對全球的影響

由於歐盟市場規模龐大，AI Act與GDPR一樣具「域外效力」。全球AI企業若欲進入歐盟市場，必須遵守相關規範。換言之，AI Act將可能成為未來全球AI法規的標準。

三、其他國家立法動態：美國與中國的對比

（一）美國：州級法案與聯邦政策並行

美國在隱私與AI立法方面，傾向採取「分散治理」模式。

1.州級法案：

例如加州消費者隱私法（CCPA）與加州隱私權法（CPRA），賦予居民要求刪除與拒絕資料販售的權利，部分條文接近GDPR。其他如科羅拉多州也制定類似法律。

2.聯邦政策：

拜登政府於2022年提出AI權利法案藍圖（Blueprint for an AI Bill of Rights），強調五大原則：安全有效的

系統、演算法歧視保護、資料隱私、通知與解釋、人工替代選項。不過，美國在2025年7月的「AI行動計畫」是川普2.0政府在AI監管與發展方面所推動的一項政策。這項計畫與拜登政府的「AI權利法案藍圖」有著截然不同的方向。

「AI行動計畫」的主要特點在於推動去監管化，以加速AI發展並鞏固美國的全球主導地位。其核心理念包括：

- (1)反意識形態偏見：該計畫特別關注AI系統中的「意識形態偏見」，並試圖透過政策來消除此類偏見。一項名為「防止聯邦政府中的覺醒AI」的行政命令，要求與政府合作的AI模型必須「客觀並排除任何頂層意識形態偏見」。
- (2)鼓勵創新與去監管：計畫的核心目標是透過減少「不必要的繁文縟節」來加速AI的發展，並將過度監管視為阻礙創新的障礙。相關行政命令指示聯邦機構，應識別並廢除可能阻礙AI發展的法規。
- (3)基礎設施投資：計畫也旨在加速AI基礎設施的建設，包括簡化資料中心和半導體製造設施的許可程序，並推動電網現代化以應對AI運算對能源的巨大需求。
- (4)國際領導力：川普政府的策略是透過出口full-stack美國AI技術解決方案來鞏固其全球領導地位，同時也收緊對抗性技術的出口管制。

川普政府此計畫與拜登政府強調公民權利保護和演算法問責制的

「AI權利法案藍圖」相比，更著重於市場驅動、去監管化以及國家安全與經濟競爭力。同時，它對隱私和消費者保護的關注相對較少。

3.技術標準：

美國國家標準暨技術研究院（NIST）於2023年發布「AI風險管理框架」（AI RMF），推動業界自律與標準化。

美國模式的特徵是「市場驅動」與「產業自律」，法律強度相對較弱，尤其在川普政府執政下更在意保持創新彈性。

（二）中國：個人信息保護法（PIPL）與數據安全治理

中國則走向「國家主導」模式，三大基礎法律為：網絡安全法（2017）、數據安全法（2021）、個人信息保護法（2021）。其中，PIPL與GDPR類似，保障資料主體知情權、更正權、刪除權等，但更強調：

- 1.數據本地化：重要數據與大規模個人信息須儲存於境內，跨境傳輸需經安全評估。
- 2.國家安全優先：國家機關可基於公共安全、社會治理需要，廣泛蒐集與利用個人信息。
- 3.高額罰則：違法處理個人信息最高可罰5000萬人民幣或年度營收5%。

中國的模式兼具嚴格控制與政策靈活，強調數據主權與國家安全，與歐美模式並不相同。

四、亞洲其他國家經驗：日本與新加坡

（一）日本

日本於2003年制定個人情報保護法，經多

次修正，逐漸趨近GDPR模式。其特色在於設立「個人情報保護委員會」作為獨立監管機構，並加強跨境傳輸規範。此外，日本強調「產業應用」與「國際協調」，與歐盟簽訂互認協議，使日本被認定具「充分保護水準」。

（二）新加坡

新加坡在2012年制定個人資料保護法（PDPA），並於2022年推出「AI Verify」計畫，為全球首個AI系統測試框架。新加坡模式強調政府與產業合作，透過沙盒監理與標準制定，建立國際信任。其靈活性高，尤其適合中小企業導入AI。

五、我國法規借鏡之處

對我國而言，國際經驗具有以下啟示：

歐盟模式：強調資料主體權利與高風險AI的監管適合借鏡，但需兼顧產業發展。

美國模式：產業自律與標準化，可供在金融與科技產業推動「指引性規範」。

中國模式：雖強調國家安全，但對跨境資料流動的嚴格規範，提醒在兩岸資料傳輸上必須特別謹慎。

日本與新加坡模式：展現「中庸之道」，兼顧法遵與產業發展，特別適合作為我們推動「區域性數據合作」的參考。

六、小結

國際法規的最新發展，展現了不同模式：

- 歐盟採「權利本位」與「風險導向」模式；
- 美國偏向「市場自律」與「標準引導」；
- 中國則是「國家控制」與「安全優先」；
- 日本與新加坡則走向「合作治理」與「國際協調」。

在AI時代，我們若要保持國際競爭力並確保權利保障，必須積極借鏡歐盟與亞洲鄰近國家的經驗，並建立具有本土特色的法制架構。

伍、在地法律挑戰與因應

生成式AI快速進入日常生活及政府與企業應用，對個資保護制度提出新的要求。我國近期修法動作大幅加速，已開始回應司法改革與AI時代隱私治理的迫切需求。

一、2025年個資法修正與獨立監管機制

（一）個人資料保護委員會（PDPC）設立

2025年3月27日，行政院提出個人資料保護法部分條文修正草案與個人資料保護委員會組織法草案，在於建立獨立監督機制，賦予PDPC對公私部門執法權，統籌事故通報與政策制定等職能。此係因憲法法庭於健保資料庫案要求於2025年8月前設立獨立監督機構，惟立法尚未完成。

（二）通報與事故應變制度加強

修正草案明定公私部門在發生重大個資事故時，須採取應變措施、保存事故紀錄，並依一定條件通報主管機關。此外，企業不得以檢視調查等理由延遲通知當事人。

PDPC將協調建立子法，具體規範通報與通知程序、時間與方式。

（三）設置個人資料保護長（PDPO）與稽核人員

公務機關須設置PDPO，由首長指派，統籌推動與督導個資保護事務。預告版本曾要求

設置稽核人員，但最新調整僅保留PDPO規定。

（四）過渡期安排與高風險行業檢查

非公務機關仍由現有主管機關監管，但PDPC將在六年內逐步接管監督責任。

修正案亦授權PDPC評估高風險行業，優先進行行政檢查。

二、制度變革下的實務衝擊與瓶頸

（一）制度集中但缺AI專章

雖然建立獨立機構是重要進展，但目前修法內容仍集中於監管架構與基本通報制度，未包含針對生成式AI特有風險（如演算法透明度、訓練資料合法性、可解釋性等）的專門規範。

（二）跨境資料傳輸仍有落差

現行跨境規範仍局限於主管機關限制命令，缺乏GDPR充分性認定及標準契約等制度，使AI雲端應用在跨境資料流動上仍具法遵不確定性。

（三）PDPO職責與能量不足風險

公務機關雖設置PDPO，但若缺乏足夠資源與獨立性，僅是形式機制，難以因應AI特殊挑戰。

（四）通報義務與民間界定模糊

法規要求「重大危害風險」才須通報，但界定仍曖昧，有待細化界線與標準，避免法遵不一致。

三、在地案例與相關風險仍存在

（一）政治人物手機信令議題

2024年立法院相關討論中，有政黨使用手機信令分析民眾年齡層引發監控爭議，此類

資料處理雖稱民間分析，但凸顯缺乏透明與合法查驗機制作為漏洞。

（二）罷免案個資涉嫌違法公開

2025年3月，某政黨發言人公開罷免團體成員資訊與群組對話，引發違法調查與個資法規疑慮，凸顯社會仍需穩固制度保障。

這些案例都在提醒制度設計必與強化法律意識與落實執法並進。

四、小結與更新後之修法建議

本次修法雖使我國朝向建立獨立個資監管機構的制度性躍進，為AI時代的資料治理奠定基礎。但若欲真正回應生成式AI治理挑戰，仍必須進一步：

- （一）在後續修法階段加入AI專章，明定AI模型透明度、訓練資料合法性與目的限制、可解釋性、偏見防範等規範；
- （二）強化PDPO職能與獨立性，賦予責任與適合的稽核權或其他機制；
- （三）制定跨境資料傳輸準則，包括「充分性」與標準契約；
- （四）清楚界定「重大危害危險」標準並發布指引；
- （五）透過社會教育與政務透明提升公民數位隱私素養。

陸、在創新與權利保護間尋求平衡

一、AI時代下個資與隱私權的核心挑戰

生成式AI與機器學習的廣泛應用，使個人資料與隱私權保護面臨前所未有的挑戰。與傳統資訊科技不同，AI的特徵在於大規模資料蒐集、持續學習、跨境傳輸、黑箱決策，

導致既有法律原則（例如目的限制、同意機制）逐漸失效：

（一）同意機制的弱化

AI系統往往透過巨量資料進行訓練，當事人難以事前知悉用途，更難提供有意義的知情同意。

（二）資料再利用的模糊性

資料常被重複使用於與原始目的不同的AI模型，挑戰「目的限制原則」。

（三）黑箱與歧視風險

演算法缺乏透明度，難以檢視偏見或差別待遇，導致權益受侵害卻難以舉證。

（四）跨境資料治理不足

AI雲端與跨境傳輸頻繁，但各國規範不一，產生管轄與法遵的落差。

（五）資料外洩新態樣

除了傳統資料庫外洩外，AI模型本身亦可能洩漏個資（例如透過反向推理訓練資料）。

這些特徵顯示著個資與隱私權不再僅是單純的資訊管理議題，而已成為一個橫跨法律、技術與倫理的「治理議題」。

二、法律、技術與倫理的三重協同

AI治理無法單靠法律規範解決，必須結合技術措施與倫理文化，才能形成完整的防護網：

（一）法律層面

國際趨勢：歐盟GDPR與AI Act已建立完整的權利基礎與風險分級框架；美國雖採取分散立法，但透過州法與白宮行政命令逐步推動AI原則；中國則以數據安全與國家治理視角切入，強調國家控制。

我國進展：2025年修法成立個人資料保護

委員會（PDPC），奠定獨立監管的基礎；而AI基本法雖已在八月由行政院院會通過送立法院，但就個資保護僅於該基本法第十三條原則性地提及，尚未能充分回應生成式AI的特殊風險。

（二）技術層面

- 1.隱私增強技術（PETs）：如差分隱私、同態加密、聯邦學習，可降低個資外洩風險。
- 2.模型治理工具：模型可解釋性（Explainable AI, XAI）、公平性檢測工具，可提升透明度與可稽核性。

（三）倫理層面

- 1.企業責任：AI發展者與應用者需遵守「隱私內建」（Privacy by Design）與「倫理審查」原則。
- 2.公民社會：需要提升公民對AI資料風險的認知，並強化監督角色。

三、我國未來角色與挑戰

我國雖已積極進行個資法修正，並即將成立PDPC，但若要在AI治理上成為亞洲區域標竿，仍需進一步努力：

（一）建立AI專章或專法

在個資法或另立AI法中，明文規範AI系統透明度、訓練資料合法性、可解釋性、風險分級監管機制。

（二）強化跨境資料傳輸規範

參照GDPR「充分性認定」、標準契約（SCCs）制度，確保台灣企業在國際供應鏈與雲端運算中具法遵合規基礎。

（三）推動產業與政府協力治理

由主管機關、產業公協會、研究機構共同制定AI實務準則，建立「沙盒機制」讓AI技術測試能兼顧創新與風險控制。

（四）強化司法體系實務AI能力

對於法官、檢察官與行政官員進行強化訓練，以及律師公會協助律師使其能理解AI技術特性，避免司法實務在AI案件出現技術落差。

（五）提升公民參與與社會對話

鼓勵NGO、媒體與教育機構參與AI與隱私議題討論，建立多元監督機制。

四、結語：平衡的治理願景

AI技術帶來的效率與創新價值無庸置疑，但若忽視個資與隱私權的保護，將侵蝕社會信任基礎，最終反噬產業發展。

因此，未來的關鍵在於「平衡治理」，需要在創新與保障權利之間找到制度化的中間點；在國際規範與在地特色之間，建構在地版本的AI資料治理模式；在法律、技術與倫理三者之間，建立多層次協同框架。

唯有如此，我國才能在AI治理的全球舞台上，不僅追隨國際，更能提出務實的制度創新，成為亞洲AI與隱私權治理的重要推手。