



經濟部

Ministry of Economic Affairs

經濟部產業技術司科技專案研發成果



stli 資策會科技法律研究所叢書

AI 法制的3H：

人工智慧這樣管就對了



主 編 王自雄

執行編輯 韓佳盈

作者群 周晨蕙、張腕純、柯亦儒
陳 箴、許嘉芳、王德瀛

購書請上：<https://www.angle.com.tw/Book.asp?BKID=15616>

AI 法制的3H： 人工智慧這樣管就對了

主 編 王自雄

執行編輯 韓佳盈

作 者 群 周晨蕙、張腕純、柯亦儒

陳 箴、許嘉芳、王德瀛



元照出版提供 請勿公開散布。

財團法人資訊工業策進會科技法律研究所 出版

» 推薦序

FOREWORD »

在亞瑟·李·塞繆爾（Arthur Lee Samuel）於1950年代以機器學習方式讓IBM 701電腦執行跳棋程式後，人工智慧開始展露頭角；其後歷經1997年IBM電腦「深藍」（Deep Blue）擊敗世界西洋棋冠軍加里·基莫維奇·卡斯帕羅夫（Garry Kimovich Kasparov）、2016年AlphaGo擊敗韓國職業棋士李世石，乃至2022年11月Open AI研發的ChatGPT問世，AI技術已不斷突破，相關應用逐漸深入日常生活各個面向，成為社會不可或缺的一部分。

有關AI法制相關議題，早期圍繞對AI倫理及道德議題的探討，如歐盟於2019年公布之「可信賴之AI倫理指引」、聯合國教科文組織於2021年制定之「人工智慧倫理建議書」等。伴隨AI技術發展，目前常見AI法制議題包括：機器學習蒐集大量資料之行為，可能侵害他人著作權、隱私或個資；AI生成物是否享有智慧財產權保護；AI對於勞動權益之衝擊等等，不一而足。

有鑑於AI對社會各領域帶來之衝擊與影響，各國均紛紛提出不同的監管策略，期能以順應國情的監管方式，

購書請上：<https://www.angle.com.tw/Book.asp?BKID=15616>

帶動技術及整體產業的發展。在AI浪潮席捲而來的現代社會，如何對其進行規範才能平衡創新與風險，誠為重大挑戰。財團法人資訊工業策進會作為「數位轉型化育者」，將持續關注AI法制議題，以協助政府進行立法或政策推動，提升AI及相關產業的發展。

財團法人資訊工業策進會副執行長

蕭博仁



元照出版提供

請勿公開散布

謹誌

中華民國112年11月

» 推薦序

FOREWORD »

科技日新月異，人工智慧為目前各國重點發展項目之一，各國亦就AI定有不同的監管方案。如歐盟刻正針對《人工智慧法》草案進行研商，希望最快可於2023年底前通過；美國總統拜登在2023年10月30日發布「總統行政命令」（Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence），以確保美國可在「掌握AI發展前景」及「AI風險控管」上，處於世界領先地位；我國亦研擬《人工智慧基本法》草案，並已於2023年8月31日公布「行政院及所屬機關（構）使用生成式AI參考指引」，顯示我國對此議題之重視。

近年來，AI因其可於短時間內蒐整及條列資訊的特性，常被用於改善工作流程及進行決定支援，已對現代人日常生活領域帶來一定的「數位創新」。本書亦介紹與資料流通、資料治理有關的規範，如歐盟《資料治理法》、《資料法》草案及〈健康資料空間規則〉草案，期能對我國資料整體經濟價值的促進，有所助益。

購書請上：<https://www.angle.com.tw/Book.asp?BKID=15616>

財團法人資訊工業策進會科技法律研究所身為國內最大法制智庫單位，長期研析並觀測國際新興科技法律議題，期透過本書整理AI法制三大議題，從AI監管方向、可信賴之AI及資料活用出發，協助產業了解其進行AI化的創新轉型階段所可能面臨的挑戰，並將隨著AI的發展，持續關注AI法制議題。

再請各位先進不吝撥冗指教，以共同創建可信賴、負責任且風險可控之AI生態系。

財團法人資訊工業策進會科技法律研究所所長
元照出版提供 請勿公開散布

萬宏宜

謹誌

中華民國112年11月

» CONTENTS

推薦序

蕭博仁

推薦序

蕭宏宜

第一單元 // How to make AI under the law : 讓管制成爲助力而非阻力

❖ 全球第一部AI法：歐盟AI法案

壹、前 言	3
貳、AIA立法背景及規範重點	4
一、立法背景與沿革	5
二、人工智慧之定義	7
三、人工智慧應用依風險高低分級	9
參、後續發展與協商重點	14
肆、結 論	17

❖ 先政策後法律：英國AI國家戰略

壹、前 言	19
貳、英國國家AI戰略（National AI Strategy）	20
一、投資AI生態系長期需求	20
二、確保所有產業領域與區域之AI利益	20
三、有效治理AI	21

參、支持創新之AI監理方法	
(Pro-Innovation Approach to Regulating AI)	22
一、當前法制環境與挑戰	23
二、基礎思維	26
三、跨領域原則 (cross-sectoral principles)	27
肆、結 論	30
❖ 生成式AI：在管制過度與管制麻木間求取平衡	
壹、前 言	33
貳、歐盟透過AIA規範生成式AI	35
參、美國尊重自由市場運作，業者成立自律組織	38
肆、日本以現行法和指引進行監管	40
伍、廣島AI進程	42
陸、結 論	44



元照出版提供 請勿公開散布

第二單元 // How to make AI trustworthy： 被信賴、負責任、控風險

❖ 落實AI倫理：OECD倡議

壹、「值得信賴AI」之負責任管理原則	51
一、包容性成長、永續發展與福祉	51
二、以人為本價值與公平	52
三、透明與可解釋性	52
四、穩健、資安與安全	52
五、可歸責性	53

貳、「值得信賴AI」之實踐工具	53
一、技術類工具	53
二、程序類工具	56
三、教育類工具	57
參、於AI生命週期中落實「值得信賴AI」	58
一、界定（Defining）	59
二、評估（Assessing）	61
三、處理（Treating）	62
四、治理（Governance）	63
肆、結 論	65
❖釐清AI責任：歐盟AI責任相關規範草案之影響	
——以自駕車侵權責任為例	
壹、歐盟人工智慧相關責任指令草案	70
一、人工智慧責任指令草案	70
二、產品責任指令修正草案	76
三、兩部責任指令草案之異同	84
貳、自駕車AI應用之歸責建立	86
一、AI及自駕車技術	86
二、自駕車之侵權責任——以我國為例	90
參、代結論：歐盟人工智慧責任相關指令可能 產生之法制變遷	95
❖掌握AI風險：美國AI風險管理框架	
壹、前 言	99
貳、人工智慧風險管理框架	100
一、AI風險管理可能遇到之挑戰	100
二、從AI生命週期檢視各階段參與者應關注事項	103

參、從治理、映射、量測、管理四個面向降低AI 技術應用潛在風險.....	106
一、治 理.....	106
二、映 射.....	108
三、量 測.....	110
四、管 理.....	112
肆、結 論	113

第三單元 // How to make AI useful : 讓資料成為餵養AI的養分

❖ 讓資料動起來：資料中介服務

壹、前 言	117
貳、資料中介的型態與案例	120
一、個人資料管理系統（Personal Information Management System, PIMS）	120
二、資料保管機構（Data Custodians）	124
三、資料交換（Data Exchange）服務.....	129
四、產業資料平台（Industrial Data Platforms）	132
五、資料協作（Data Collaboratives）	133
六、可信任之第三方（Trusted Third Parties）	134
七、資料合作社（Data Cooperatives）	136
八、資料信託（Data Trusts）	138

參、政策與研究建議.....	140
一、研析適合我國需求之資料中介服務類型， 選擇適當機制重點推動.....	140
二、持續鼓勵研發應用隱私強化技術PETs，提高 民眾與企業資料共享意願.....	142
肆、結 論	145
❖從資料共享到價值創造：歐盟資料治理法 及資料法案	
壹、前 言	147
貳、全球首部針對資料中介服務的立法 ——歐盟資料治理法.....	150
一、立法背景與歷程	151
二、資料中介服務之規範	154
三、資料利他主義機構之規範.....	161
四、公部門受保護資料之再利用與資料中介之 關聯性.....	166
參、輔助資料中介服務運作的立法 ——歐盟資料法草案.....	167
一、立法背景與歷程	168
二、B2C與B2B資料共享之規範與權利義務.....	170
三、資料處理服務之切換	178
四、互操作性標準與智能合約規範	180

肆、法制政策建議.....	181
一、強化資料主體之資料主控權，打造適合資料 中介服務之環境	182
二、推動制定資料互操作性標準，降低資料共享 之技術障礙與成本	183
三、參考DGA認證機制，把關資料中介服務的 安全性以強化大眾信任，促進資料分享	184
伍、總 結	186
❖ 建立資料生態系：歐盟健康資料空間規則提案	
壹、前 言	189
貳、歐洲健康資料空間規則草案簡介	193
一、健康資料之原始利用	193
二、健康資料之二次利用	196
參、代結論：發展中的健康資料生態管理規範	202

第一單元

How to make AI under the law： 讓管制成爲助力而非阻力

➤ 全球第一部AI法：歐盟AI法案

／周晨蕙

➤ 先政策後法律：英國AI國家戰略

／張腕純

➤ 生成式AI：在管制過度與管制麻木間
求取平衡／周晨蕙

“

全球第一部AI法： 歐盟AI法案

”

資策會科技法律研究所／組長 周晨蕙

壹 ▶ 前言

有關人工智慧監管議題，早期世界各國聚焦於對AI倫理及道德議題的探討，在經過一段時間討論及發展後，已逐步演進到針對特定領域利用AI之行爲加以規範。澳洲於2023年6月1日公布「支持負責任之AI：意見徵詢」文件（Supporting responsible AI: discussion paper），便指出目前世界各國主要透過一般性法規（如消費者保護法）或特定領域之現行法規（如醫療、金融服務等），以及自願性倡議或協議等作法來規範AI¹。

在2022年底ChatGPT橫空出世後，其強大的功能和表現，引發民眾強烈的不安。近期好萊塢爆發之罷工潮中，編劇和演員便十分擔心資方會利用AI生成劇本或角色，主張應限制片廠運用AI取代勞工²。在上述背景之下，各國進一步展

¹ Supporting responsible AI: discussion paper, Department of Industry, Science and Resources, <https://consult.industry.gov.au/supporting-responsible-ai/submission> (last visited Aug. 25, 2023).

² 〈好萊塢罷工僵局若持續：AI會完全取代真人的編劇和演員嗎？〉，端傳媒，2023年7月28日，<https://theinitium.com/article/202>

4 ■ AI法制的3H：人工智慧這樣管就對了

開是否全面性規範AI產品及服務之討論。由此可知，人工智慧監管議題不再只停留於紙上談兵階段，而是伴隨著技術發展，正式成為迫切需要解決的課題之一。

為落實AI監管，歐盟執委會於2021年4月21日公布「人工智慧共同規則建議草案（Proposal for a Regulation laying down harmonised rules on artificial intelligence）」，亦即《人工智慧法》（The Artificial Intelligence Act, AI Act）草案（以下簡稱AIA），以風險為基礎，將AI應用分為不同等級，並給予不同的監管密度，嘗試建立AI監管框架。

由於AIA為全球首部針對AI之立法，世界各國均十分關注AIA進展，而我國目前亦刻正研擬《人工智慧基本法》草案³，並已於2023年8月31日公布「行政院及所屬機關（構）使用生成式AI參考指引」⁴，故本文擬以歐盟AIA為對象，簡介草案重點及立法最新進度，以供我國後續立法參考。

貳 ▶ AIA立法背景及規範重點

歐盟執委會於2021年4月21日提出AIA，旨在確保歐盟市場AI系統之安全性及符合歐盟價值、基本權和現行法規，並

30728-culture-ai-and-hollywood（最後瀏覽日：2023/10/18）。

³ 草案目前因考量AI技術發展快、應用面太廣而延後時程，詳細可參考：〈鄭文燦：考量技術發展快應用廣 AI基本法延後提出〉，中央社，2023年8月1日，<https://www.cna.com.tw/news/afe/202308010228.aspx>（最後瀏覽日：2023/10/18）。

⁴ 〈行政院通過AI參考指引草案 十大指引方向一次看〉，經濟日報，2023年8月31日，https://money.udn.com/money/story/7307/7406794?list_ch2_index（最後瀏覽日：2023/10/18）。

促進適法、安全且可信賴之AI應用發展⁵。以下本文將簡要整理AIA立法背景及沿革，以及草案中有關AI之定義、不同風險等級及對應規範，以及AIA最新進展和討論。

一、立法背景與沿革

近年AI技術快速發展，常被用於改善工作流程、支援決策、提供個性化服務及資源分配等事項，雖有助於提高各產業競爭力，卻也同時為個人和社會帶來新的風險和負面影響。有鑑於此，為確保技術發展符合歐盟價值觀、基本權和現行法規，並維持歐盟在技術上之領導地位，歐盟於2020年2月19日發表「AI白皮書：邁向卓越與信賴之歐洲路徑」（White Paper on Artificial Intelligence: A European approach to excellence and trust），指出未來將兼顧「監管」與「投資」，促進人工智慧之應用並同時解決該項技術帶來之風險⁶。

上開白皮書提到未來將制定明確之歐洲監管框架，而歐洲議會及歐洲理事會亦表明應立法規範安全、可信賴之AI開發，並兼顧AI風險與利益衡平。在白皮書發表後，歐洲議會亦陸續於同年10月達成多項與AI有關之決議，包括倫理、法

⁵ Proposal for a Regulation laying down harmonised rules on artificial intelligence, European Commission, <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence> (last visited Oct. 22, 2023).

⁶ European Commission, *White Paper On Artificial Intelligence: A European approach to excellence and trust* (Feb. 19, 2020), https://ec.europa.eu/info/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust_en (last visited Aug. 4, 2023).

6 ■ AI法制的3H：人工智慧這樣管就對了

律責任及著作權；並於2021年作出AI刑事相關之決議，以及教育、文化和視聽覺相關之決議；與此同時，歐洲委員會亦針對立法提案，提出應注重AI及機器人技術開發、發展和利用之倫理原則之勸告⁷。

在上述背景下，歐盟執委會於2021年4月21日提出AIA（下稱執委會版），希望能達成確保歐盟市場使用之AI系統安全性，以及AI系統符合歐盟價值觀、基本權和現行法規；促進AI投資和創新之法安定性；強化有關AI治理、基本權及適用AI系統之安全性要件；促進合法、安全且可信賴之AI應用於單一市場之發展等目的⁸。

在執委會版草案提出後，歐盟理事會於2022年12月達成共通立場（common position/general approach），提出理事會版本之AIA（下稱理事會版）；2023年6月14日，歐洲議會達成協商立場（negotiating position）並公布議會版本之AIA（下稱議會版）。由於理事會版、議會版及委員會版草案內容略有差異，故目前歐盟理事會、歐洲議會及執委會刻正針對草案內容進行三方協商，草案預計最快於2023年底公布⁹。

⁷ 轉引自：總務省，《人工知能に関する調和の取れたルールを定める規則の提案》，https://www.soumu.go.jp/main_content/000826706.pdf，頁2-3（2021）。

⁸ 同前註。

⁹ MEPs ready to negotiate first-ever rules for safe and transparent AI, European Parliament, <https://www.europarl.europa.eu/news/en/press-room/20230609IPR96212/meps-ready-to-negotiate-first-ever-rules-for-safe-and-transparent-ai> (last visited Aug. 12, 2023).

二、人工智慧之定義

執委會版AIA第一章規範AIA適用對象、範圍及AI定義，其對AI採取正面表列之定義方式，並考量到AI技術及市場快速發展特性，需要保留一定彈性，將AI定義為「使用附件1所列之特定技術或方法¹⁰，根據人類定義之目標，產出內容、預測、建議或影響互動環境之決策。」¹¹

在執委會提出草案後，歐洲經濟和社會委員會、區域委員會、歐洲中央銀行等陸續提出對AIA之意見，其中由於執委會版採用之定義過於廣泛，使AI定義成為討論重點。為回應各界意見，執委會隨後提出妥協版本，除微調AI定義外，並指出AI系統應得以附件所列技術和方法進行學習、推理或模擬，決定如何達成人類定義之目標，以避免將通常不被認為是AI之傳統軟體納入監管範圍¹²。

有鑑於執委會版對於AI定義過於廣泛而引發之爭議，理事會版將AI定義限縮為「透過機器學習方法，以及基於邏輯和知識方法開發的系統」¹³。2023年3月3日，歐洲議會就AI

¹⁰ 包括機器學習（machine learning）、邏輯式智慧式方法論（logic-and knowledge-based approaches）、統計方法（statistical approaches）、貝氏推論（Bayesian estimation）、搜尋與優化方法等。

¹¹ Commission Proposal for a Regulation laying down harmonised rules on artificial intelligence, at 39, COM (2021) 206 final (Apr. 21, 2021).

¹² Commission Proposal for a Regulation laying down harmonized rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts, at 3, 2021/0106 (COD) (Nov. 29, 2021).

¹³ Artificial Intelligence Act: Council calls for promoting safe AI that respects fundamental rights, European Council, <https://www.consilium.europa.eu/en/press/press-releases/2022/12/06/artificial-intelligence-act->

8 ■ AI法制的3H：人工智慧這樣管就對了

定義達成政治協議，同意採用經濟合作與發展組織（OECD）對於AI之定義，以結束有關定義之爭論¹⁴。OECD將AI定義為「一種基於機械的系統，旨在針對特定目標，輸出影響環境的預測、建議、或決策。其透過機器或以人為主之資料和輸入，用以感測真實或虛擬環境、以自主方式或手動方式將感測資料化為模型，並使用模型推演出結果選項。」確保類似ChatGPT之生成式AI不會成為監管上的漏洞¹⁵。

因應生成式AI發展，議會版更進一步將生成式AI定義為「一種具有不同程度自主性，專門用於生成複雜文本、圖像、音訊或影片等內容之基礎模型（foundation model）」並要求生成式AI供應商應揭露AI生成的內容、設計模型避免AI產出違法內容，以及公開訓練AI所用受著作權保護資料的摘要等，以避免生成式AI侵害他人智慧財產權，以及利用生成式AI產出操縱性內容¹⁶。

綜觀上述定義，可以發現議會版和理事會版對於AI之定義較為相近，均將範圍予以限縮。此種定義方式反映出輿論對AI定義過於廣泛之擔憂，惟定義太過明確亦可能失去彈

council-calls-for-promoting-safe-ai-that-respects-fundamental-rights/ (last visited Aug. 12, 2023).

¹⁴ EU lawmakers set to settle on OECD definition for Artificial Intelligence, Euractiv, <https://www.euractiv.com/section/artificial-intelligence/news/ai-act-meps-extend-ban-on-social-scoring-reduce-ai-office-role/> (last visited Mar. 17, 2023).

¹⁵ OECD AI Principles overview, OECD.AI, <https://oecd.ai/en/ai-principles> (last visited Oct. 17, 2023).

¹⁶ European Parliament, *supra* note 9.

性，使法律難以跟上技術變革之腳步。有鑑於此，最終歐盟會採用何種版本之定義，仍需持續關注。

三、人工智慧應用依風險高低分級

AIA依照風險高低，將AI分為禁止使用之「不可接受風險」(Unacceptable risk)、上市前需符合嚴格規範之「高風險」(High risk)、負有一定程度透明性義務之「有限風險」(Limited risk)，以及不受規範之「小或低風險」(Low or minimal risk) AI等4種，針對AI使用容許性、使用規範等面向，給予不同程度之規範¹⁷。本文將先以執委會版為主，簡述各類型AI之規範重點，再補充理事會版及議會版調整之處。

(一)不可接受風險之AI

所謂「不可接受風險之AI」，亦即被禁止使用之AI。此種類型之AI應用包括：利用年齡或缺陷操縱特定人或弱勢群體的認知或行為、利用AI進行社會信用評分，以及在公眾場所執法為目的之遠端生物辨識系統等，其目的在於避免AI被用於操控人類及監控社會，以保障基本權及弱勢族群權益¹⁸。

¹⁷ Proposal for a Regulation laying down harmonised rules on artificial intelligence, European Commission, <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence> (last visited Feb. 22, 2023).

¹⁸ EU AI Act: first regulation on artificial intelligence, European Parliament, <https://www.europarl.europa.eu/news/en/headlines/society/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence> (last visited Aug. 12, 2023).

購書請上：<https://www.angle.com.tw/Book.asp?BKID=15616>

國家圖書館出版品預行編目(CIP)資料

AI 法制的 3H：人工智慧這樣管就對了/王德瀛, 周晨蕙, 柯亦儒,
許嘉芳, 張腕純, 陳箴作 ; 王自雄主編.-- 初版.--
臺北市：經濟部, 財團法人資訊工業策進會科技法律研究所, 2023.11
面 ; 公分
ISBN 978-986-533-393-5 (平裝)
1.CST: 人工智慧 2.CST: 資訊法規 3.CST: 產業政策
312.83 112018737

AI法制的 3H：人工智慧這樣管就對了

5L053PA

主 編	王自雄
執行編輯	韓佳盈
作 者 群	王德瀛、周晨蕙、柯亦儒、許嘉芳 張腕純、陳箴（依姓名筆畫順序排列）
編 輯 所	元照出版有限公司
出版機關	經濟部 地址：臺北市福州街15號 網址： https://www.moea.gov.tw 電話：(02)2321-2200
出 版 者	財團法人資訊工業策進會科技法律研究所 地址：臺北市大安區敦化南路二段216號22樓 網址： https://stli.iii.org.tw 電話：(02)6631-1000／傳真：(02)6631-1001
經 銷 商	元照出版有限公司 地址：臺北市館前路28號7樓 網址： https://www.angle.com.tw 電話：(02)2375-6688／傳真：(02)2331-8496 郵政劃撥：19246890元照出版有限公司

政府出版品展售門市：

國家書店松江門市
臺北市中山區松江路209號1樓／電話：(02)2518-0207
五南文化廣場臺中總店
臺中市西區臺灣大道二段85號／電話：(04) 2226-0330
著作權管理資訊：經濟部產業技術司保有所有權利。
欲利用本書全部或部分內容者，須徵求經濟部產業技術司
同意或書面授權。

其他類型版本說明：本書未同時發行其他版本

著作權所有，非經經濟部書面同意，不得翻印、轉載或以任何方式重製。

■2023年11月初版第1刷

定價：350元

ISBN：9789865333935

GPN：1011201573

AI法制的3H： 人工智慧這樣管就對了



AI進入2.0的生成式AI時代，已經成為驅動產業數位轉型的核心，本書用心整理國際最新AI法制動態，可以作為我國產業落地時的重要參考依據。

—吳漢章・AI大聯盟會長、華碩雲端暨台智雲總經理

AI應用之前提，在於建立可信任、負責任之AI環境，同時掌握超越AI的思考架構。本書正是在介紹國際間重要法制，為打造可信賴AI生態系奠定厚實基礎。

—詹婷怡・前國家通訊傳播委員會主任委員
・人工智慧科技基金會常務董事

資策會科法所長期致力於臺灣新科技管制研究。這次，本書則為台灣AI治理途徑的探索，提供了清晰的參考文本。

—侯宜秀・台灣人工智慧學校基金會秘書長

AI是否需要管制，如何管制，近來在國際間引起激烈論辯。資策會科法所作為我國法制智庫，透過本書提供了深入淺出的觀察與分析。

—李宏毅・國立臺灣大學電機工程學系教授

 元照出版公司

地址：臺北市館前路28號7樓
電話：(02)2375-6688
網址：www.angle.com.tw



9 789865 333935

5L053PA

GPN：1011201573

定價：350元