

歐盟發佈首份通用人工智慧實踐準則

以因應歐盟人工智慧法對通用人工智慧模型之監管

黃可激 編譯

摘要

歐盟執委會於今 (2025) 年 7 月發布通用人工智慧實踐準則，聚焦於透明性、著作權保護與安全性等三大面向，協助通用人工智慧模型供應商遵循歐盟「人工智慧法」對於透明性、著作權保護與安全性之規範。該準則雖無法律拘束力，然歐盟表示其將作為通用人工智慧模型供應商對於歐盟人工智慧法的法遵實務參考。雖然歐盟希望業界採用此準則以遵循人工智慧法之規定，但目前業界對於是否採用通用人工智慧實踐準則意見分歧。未來該準則的接受度可能影響其是否轉化為業界標準，或歐盟主管機關將採取更強硬之手段使通用人工智慧模型供應商履行人工智慧法之法定義務。

(取材資料：Jan-Jaap Koningsveld, Minke Reijneveld & Auke-Frank Tadema, *EU's GPAI Code of Practice: The World's First Guidance for General Purpose AI Model Compliance*, STIBBE (July 14, 2025), <https://www.stibbe.com/publications-and-insights/eus-gpai-code-of-practice-the-worlds-first-guidance-for-general-purpose>; Madelaine Harrington, Stacy Young & Alberto Vogel, *AI Office Publishes Final Version of the Code of Practice for General-Purpose AI Models*, COVINGTON & BURLING LLP. (July 30, 2025), <https://www.globalpolicywatch.com/2025/07/ai-office-publishes-final-version-of-the-code-of-practice-for-general-purpose-ai-models/>.)

歐盟執委會 (European Commission，下稱執委會) 於今 (2025) 年 7 月 10 日發布「通用人工智慧實踐準則 (General-Purpose AI Code of Practice，下稱 GPAI 準則)」¹。其屬於不具拘束力的軟法，旨在協助人工智慧開發者遵守歐盟「人工智慧法 (EU AI Act，下稱 AI 法)」²中新訂有關透明性、安全性與智慧財產權之規定。這份自願性準則旨在引導例如大型語言模型與其他基礎模型等「通用人工

¹ European Commission Press Release IP/25/1787, The Commission, General-Purpose AI Code of Practice Now Available (July 10, 2025).

² Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 Laying Down Harmonised Rules on Artificial Intelligence and Amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), 2024 O.J. (L Seris) 144 [hereinafter EU AI Act].

智慧模型 (General-Purpose AI Models, 下稱 GPAI 模型)」³ 供應商, 如何履行 AI 法第 53 條與第 55 條下之法定義務⁴。即使 AI 法中有關 GPAI 模型供應商的條款已於今年 8 月 2 日實施, GPAI 準則仍係 GPAI 模型供應商重要的法遵輔助標準⁵。GPAI 模型供應商可透過自願簽署與遵守 GPAI 準則來證明他們遵守 AI 法, 以減輕法遵成本並增加 AI 法之法律確定性⁶。

本文將介紹 GPAI 準則之內容與現況, 了解其對 GPAI 模型與其供應商有何具體規範要求, 並於結尾做一結論。

壹、GPAI 準則之三大重點

GPAI 準則總共有三個章節, 分別涉及透明性、著作權保護與安全性等三大內容, 以協助 GPAI 模型供應商遵循 AI 法, 以下分別對三個章節進行說明。

一、通用人工智慧模型之透明性要求

關於 GPAI 模型供應商為遵守 AI 法第 53 條第 1 項 (a) 與 (b) 款所應保存的模型文件⁷, 準則提出以下三項具體措施:

(一) 編製並即時更新模型文件以供查閱

GPAI 模型供應商須製作並即時更新模型相關技術文件, 且在歐盟 AI 辦公

³ 「通用人工智慧模型 (general-purpose AI models)」在歐盟 AI 法中係指經由大規模自我監督學習並以大量資料訓練之模型, 該模型展現出顯著的通用性, 能夠執行多種不同類型的任務, 且不受其進入市場方式影響, 亦可整合至多種下游系統或應用程式中, 常見的通用人工智慧模型包含 OpenAI 所開發的 GPT-4 與 Google 所開發的 Gemini 等。EU AI Act, art. 3(63) (“[G]eneral-purpose AI model’ means an AI model, including where such an AI model is trained with a large amount of data using self-supervision at scale, that displays significant generality and is capable of competently performing a wide range of distinct tasks regardless of the way the model is placed on the market and that can be integrated into a variety of downstream systems or applications, except AI models that are used for research, development or prototyping activities before they are placed on the market.”).

⁴ *The General-Purpose AI Code of Practice*, EUR. COMM’N, <https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai> (last updated Sep. 9, 2025).

⁵ European Commission Press Release IP/25/1787, *supra* note 1.

⁶ *Id.*

⁷ NURIA OLIVER & RISHI BOMMASANI, CODE OF PRACTICE FOR GENERAL-PURPOSE AI MODELS: TRANSPARENCY CHAPTER 2 (2025); EU AI Act, art. 53(1)(a)(b) (“1. Providers of general-purpose AI models shall: (a) draw up and keep up-to-date the technical documentation of the model, including its training and testing process and the results of its evaluation, which shall contain, at a minimum, the information set out in Annex XI for the purpose of providing it, upon request, to the AI Office and the national competent authorities; (b) draw up, keep up-to-date and make available information and documentation to providers of AI systems who intend to integrate the general-purpose AI model into their AI systems. Without prejudice to the need to observe and protect intellectual property rights and confidential business information or trade secrets in accordance with Union and national law, the information and documentation shall: (i) enable providers of AI systems to have a good understanding of the capabilities and limitations of the general-purpose AI model and to comply with their obligations pursuant to this Regulation; and (ii) contain, at a minimum, the elements set out in Annex XII . . .”).

室 (AI Office) 要求時提供該些文件，亦應可供下游供應商 (即在該 GPAI 模型基礎上建構 AI 系統者) 查閱。簽署方可選擇使用本準則所附的「文件模範表單 (Model Documentation Form)」來提供這些資訊，但準則並無強制 GPAI 模型供應商採用此份表單⁸。

(二) 提供相關資訊予下游供應商與 AI 辦公室

當下游供應商或 AI 辦公室提出要求時，GPAI 模型供應商應提供必要資訊。為實現此承諾，簽署方須公開其聯絡資訊，以便 AI 辦公室與下游供應商能請求存取相關文件內容⁹。

(三) 確保資訊的品質、完整性與安全性

簽署方必須確保其控管所記載資訊文件之品質、完整性與安全性，並加以留存以作為簽署方已履行 AI 法義務之證據，且應避免未經授權之變更¹⁰。

二、著作權規範與資料使用政策

有關著作權保護之部分，亦即 AI 法第 53 條第 1 項 (c) 款要求 GPAI 模型供應商必須制定著作權使用政策以確保其遵守歐盟著作權法，尤其是尊重著作權人得以限制他人使用其作品之權利¹¹。然而，遵循本準則並不能保證該模型一定符合歐盟著作權法的規定¹²。準則規定簽署方須承諾採取以下五項措施：

(一) 制定、即時更新並實施著作權使用政策

GPAI 模型供應商必須制定並持續更新著作權使用政策，以確保所有在歐盟市場上架的 GPAI 模型均符合歐盟有關著作權與著作鄰接權之法律規範¹³。

(二) 在進行網路爬蟲時僅複製與擷取合法可存取之受著作權保護內容

模型在進行網路爬蟲¹⁴時，不得規避「資訊社會中著作權及著作鄰接權調和

⁸ *Id.* at 5.

⁹ *Id.*

¹⁰ *Id.*

¹¹ ALEXANDER PEUKERT & CÉLINE CASTETS-RENARD, CODE OF PRACTICE FOR GENERAL-PURPOSE AI MODELS: COPYRIGHT CHAPTER 3 (2025); EU AI Act, art. 53(1)(c) (“[P]ut in place a policy to comply with Union law on copyright and related rights, and in particular to identify and comply with, including through state-of-the-art technologies, a reservation of rights expressed pursuant to Article 4(3) of Directive (EU) 2019/790.”).

¹² ALEXANDER & CÉLINE, *supra* note 11.

¹³ *Id.* at 4.

¹⁴ 網路爬蟲是一種電腦程式，能以系統化、自動化的方式瀏覽全球資訊網 (World Wide Web) 並抓取資訊。Mohammad Abu Kausar, Vijaypal Singh Dhaka & Sanjeev Kumar Singh, *Web Crawler: A Review*, 63(2) INT’L J. COMPUT. APPLICATIONS 31, 31 (2013).

指令」第 6 條第 3 項所定義的有效科技措施 (Technological Measures)¹⁵。這些措施旨在防止或限制他人利用未經授權的著作權資料，特別是應尊重訂閱模式或收費機制所設置的存取限制¹⁶。GPAI 模型不得透過網路爬蟲程式抓取未經授權之受著作權保護資料。即使模型未將該等資料直接呈現於輸出結果，若其在抓取資料的過程中規避前述的有效科技措施，仍可能構成違反歐盟著作權法之行為。簽署方亦須承諾，GPAI 模型不得將經歐盟或歐洲經濟區之法院或主管機關認定為持續、大規模且達到商業規模侵權之網站，納入模型之爬蟲範圍¹⁷。

(三) 在進行網路爬蟲時識別並遵守權利保留聲明

使用網路爬蟲的簽署方僅使用能夠讀取並遵循「機器人排除協議 (Robots Exclusion Protocol)」¹⁸中所表達之指令的爬蟲程式，以及該協議的任何後續版本——前提是網際網路工程任務組 (Internet Engineering Task Force)¹⁹已證明該版本在技術上對 AI 模型供應商與內容提供者 (含著作權人) 而言是可行且可實施的²⁰。此外，簽署方應承諾其 AI 模型能識別並遵守依據「數位單一市場著作權指令」第 4 條第 3 項所制定的其他適當可機讀協議²¹，以確保著作權人保留權利之意旨得以落實²²。此承諾不影響著作權人以任何適當方式，依「數位單一市場

¹⁵ 有效科技措施係指權利人用以防止或限制他人對著作或其他受著作權或著作鄰接權保護之標的物進行未授權之利用，常見方式為加密、存取限制與防盜拷措施等。Directive 2001/29/EC of the European Parliament and of the Council of 22 May 2001 on the Harmonisation of Certain Aspects of Copyright and Related Rights in the Information Society, art. 6(3), 2001 O.J. (L 167) 10, 17 (“...the expression ‘technological measures’ means any technology, device or component that, in the normal course of its operation, is designed to prevent or restrict acts, in respect of works or other subject-matter, which are not authorised by the rightholder of any copyright or any right related to copyright as provided for by law...”).

¹⁶ ALEXANDER & CÉLINE, *supra* note 11, at 4.

¹⁷ *Id.* at 4-5.

¹⁸ 機器人排除協議 (Robots Exclusion Protocol, REP, 亦簡稱 robots.txt) 係 1994 年由荷蘭工程師馬泰恩·科斯特 (Martijn Koster) 提出，使網站擁有者得以對爬蟲程式提出「告示」，明確指示爬蟲程式不得造訪該網站，亦不得抓取該網站之資訊。相關介紹請參閱：許家銘，AI 橫行，30 年前寫給「君子」的 robots.txt 擋得住今日的爬蟲巨獸嗎？，地球圖輯隊，2025 年 7 月 8 日，<https://dq.yam.com/post/16602> (最後瀏覽日：2025 年 10 月 13 日)。

¹⁹ 網際網路工程任務組 (Internet Engineering Task Force, IETF) 係一開放的網路標準制定機構，透過開放流程開發網路技術標準。其未訂立正式的會員資格，任何人皆可參與制定網路標準。相關介紹請參考：IETF 組織，財團法人台灣網路資訊中心，<https://ietf.tw/ietfGroup> (最後瀏覽日：2025 年 10 月 13 日)。

²⁰ ALEXANDER & CÉLINE, *supra* note 11, at 5.

²¹ Directive (EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on Copyright and Related Rights in the Digital Single Market and Amending Directives 96/9/EC and 2001/29/EC, art. 4(3), 2019 O.J. (L 130) 92, 114 (“The exception or limitation provided for in paragraph 1 shall apply on condition that the use of works and other subject matter referred to in that paragraph has not been expressly reserved by their rightholders in an appropriate manner, such as machine-readable means in the case of content made publicly available online.”)；可機讀協議為一種技術性手段，使著作權人以電腦程式可讀取與可解釋之格式，表達其反對網路爬蟲程式辨識其受著作權保護資料之意圖。Charles Brecque, *An Introduction to Machine-Readable Documents*, TEXTMINE (May 24, 2024), <https://textmine.com/post/an-introduction-to-machine-readable-documents>.

²² ALEXANDER & CÉLINE, *supra* note 11, at 5.

著作權指令」第 4 條第 3 項之規定，明確表示不得將其作品或其他受著作權保護內容用於文字與資料探勘。簽署方並承諾向受影響之著作權人提供所使用之網路爬蟲程式，以及為尊重權利保留而採取措施的相關資訊。

(四) 降低模型產生侵害著作權內容之風險

簽署方必須實施適當且相稱的技術防護措施，防止其模型在輸出結果時，以侵權方式重製受歐盟著作權及相關權利保護之訓練資料²³。同時，簽署方亦須在其模型之合理使用政策 (Acceptable Use Policy)²⁴、使用條款或其他具同等法律效力之文件中，明確禁止以侵害著作權之方式使用該模型²⁵。

(五) 指定聯絡窗口並建立申訴機制

簽署方必須指定單一聯絡窗口，以便受影響的著作權人以電子方式聯繫，並應提供易於取得的單一聯絡窗口資訊²⁶。此外，簽署方必須建立一套申訴處理機制，使得受影響的著作權人可以透過電子方式提交具體且經充分佐證的申訴，說明簽署方未履行所承諾之義務。簽署方亦應對外提供易於取得的相關資訊，以協助著作權人有效使用該申訴機制²⁷。

三、高風險人工智慧模型之安全保障措施

GPAI 準則中有關安全性之內容係適用於可能構成系統性風險的 GPAI 模型，亦即對社會具有廣泛或重大影響風險的模型²⁸。此類模型的供應商必須建立一套全面性的風險管理架構，以識別並控制這些風險。

這項義務包括持續評估模型可能造成的重大損害，例如誤用風險²⁹、不可預

²³ *Id.* at 6.

²⁴ 合理使用政策 (Acceptable Use Policy) 係指在數位服務中常見的規範文件，用以界定使用者在哪些情形下可以或禁止使用該服務或產品。在 AI 模型中，合理使用政策將規定 AI 使用者在使用該 AI 系統時可以或不可以做什麼事，旨在防止將 AI 模型用於有害或非法活動。Barry Scannell & Leo Moore, *The Role of Acceptable Use Policies in AI*, WILLIAM FRY (May 23, 2024), <https://www.williamfry.com/knowledge/the-role-of-acceptable-use-policies-in-ai/>.

²⁵ ALEXANDER & CÉLINE, *supra* note 11, at 6.

²⁶ *Id.*

²⁷ *Id.*

²⁸ 在歐盟 AI 法中，「系統性風險」係指該風險影響範圍廣泛，可合理預期此種風險可能會對歐盟之公共衛生、安全、公共秩序、基本權利，或整體社會造成重大負面結果。EU AI Act, art. 3(65) (“[S]ystemic risk’ means a risk that is specific to the high-impact capabilities of general-purpose AI models, having a significant impact on the Union market due to their reach, or due to actual or reasonably foreseeable negative effects on public health, safety, public security, fundamental rights, or the society as a whole . . .”).

²⁹ 此處「誤用」指的是於歐盟 AI 法中所稱的「合理可預見的誤用」，係指模型使用者以違反系統原始設計目的之方式操作 AI 系統，但該誤用情形係模型開發者在設計階段可以合理預見會發生的情況。EU AI Act, art. 3(13) (“[R]easonably foreseeable misuse’ means the use of an AI system in a way that is not in accordance with its intended purpose, but which may result from

測行為，或其他大規模負面結果（如資料外洩），並確保這些風險維持在可接受範圍內³⁰。GPAI 準則期待開發者進行結構化的系統性風險評估，並在模型演進或新風險出現時，持續更新相關風險分析。

本準則中的一項重要概念是風險分級（tiers）或門檻（thresholds）的設計。GPAI 模型供應商應為系統性風險建立明確之分級標準，並事先決定當模型能力逐步進階時應實施之對應安全措施³¹。對於每一項識別出的風險情境，供應商應設定可接受風險等級，並記錄當模型達到某特定風險等級時，應啟用哪些額外防護機制³²。

對於具有系統性風險的 GPAI 模型，其供應商必須部署最先進的安全措施與資安控管機制，以有效降低風險。其中包括在模型生命週期中持續實施各項技術性防護手段³³。GPAI 準則同時要求對這類高衝擊模型的安全機制進行外部或獨立審查。引入第三方或外部專家進行稽核，有助於補強對模型風險控管與成效測試之監督，以驗證模型風險管理是否有效且符合產業最佳實務標準（Best Practice）³⁴。

GPAI 準則要求在模型發布後持續進行上市後監測，即持續觀察模型在現實世界的使用情況與表現，是否有出現新風險或其他非預期後果³⁵。簡言之，高風險 AI 模型在其生命週期中應伴隨持續的監管，模型供應商應與主管機關開誠佈公地溝通，並具備快速應對能力以面對異常情況。若發生嚴重事件或重大故障，GPAI 模型供應商有義務即時且在無不當延遲的情況下，向歐盟 AI 辦公室與相關國家主管機關通報事件情況與所採取的修正措施³⁶。

為促進監管成效，準則特別強調 GPAI 供應商應與新設立的歐盟 AI 辦公室保持合作關係。選擇遵循 GPAI 準則，即表示模型開發者有意願與監管機關攜手合作，共同監督並規範高階 AI 發展³⁷。這種協作導向的方式，旨在建立社會對高衝擊 AI 系統開發與後續監督之信任。

貳、GPAI 實踐準則之展望

遵循 GPAI 準則可作為 GPAI 模型供應商展現其符合 AI 法的方式之一，但

reasonably foreseeable human behaviour or interaction with other systems, including other AI systems.”).

³⁰ MATTHIAS SAMWALD ET. AL, CODE OF PRACTICE FOR GENERAL-PURPOSE AI MODELS: SAFETY AND SECURITY CHAPTER 6 (2025).

³¹ *Id.* at 15.

³² *Id.* at 16.

³³ *Id.* at 3.

³⁴ *Id.* at 14.

³⁵ *Id.* at 13.

³⁶ *Id.* at 25.

³⁷ *Id.* at 4.

GPAI 模型供應商仍可透過其他方式證明其符合 AI 法之規定³⁸。此外，執委會於其問答集中指出，在明 (2026) 年 8 月 2 日以前，對於在簽署 GPAI 準則後未能即時且全面落實所有承諾之簽署方，歐盟 AI 辦公室不會因此認定其違反承諾，也不會據此指責其違反 AI 法。相反地，在這類情況下，AI 辦公室將認定該簽署方係出於誠信行事，並願意與其合作，尋找達成全面合規的方法³⁹。

參、結論

至關重要的是，GPAI 準則的成功與否將取決於產業界的採納程度。執委會正積極鼓勵所有主要的 AI 模型開發者簽署並遵守本準則。無論是大型科技公司，還是開源模型實驗室，這些生成式 AI 領域的重要參與者都必須決定其是否願意承諾遵守準則的各項要求。截至目前，GPAI 供應商如 Google、OpenAI 與 Mistral 皆已表態願意簽署 GPAI 準則，反觀 Meta 則尚未同意簽署⁴⁰。他們的決定可能會受到諸多因素影響，包含主管機關對準則是否持正面態度，以及簽署方是否能因為遵循準則而取得競爭優勢。若有眾多開發者加入，GPAI 準則有可能實質上成為產業標準；反之，若採納率偏低，主管機關可能將採取更為嚴格的立場，直接依據 AI 法中的規定進行執法。

³⁸ European Commission Press Release IP/25/1787, *supra* note 1.

³⁹ *Questions and Answers on the Code of Practice for General-Purpose AI*, EUR. COMM'N, <https://digital-strategy.ec.europa.eu/en/faqs/questions-and-answers-code-practice-general-purpose-ai> (last updated July 11, 2025).

⁴⁰ Alice Hancock, *Google Set to Sign EU Code of Practice on Development of AI*, FIN. TIMES (July 30, 2025), <https://www.ft.com/content/97ee867c-64c2-44ee-a519-36ca670ed565>.