



國會應用 AI 指引

Guidelines for AI in Parliaments

編輯

Fotios Fitsilis
Jörn von Lucke
Franklin De Vrieze

2024年7月



WFD



版權與免責聲明

姓名標示－非商業性－相同方式分享（CC BY-NC-SA）

本著作的所有權利，包括著作權，均由西敏寺民主基金會（Westminster Foundation for Democracy Limited, WFD）所有，並受英國及國際法律保護。本著作依據「姓名標示－非商業性－相同方式分享」創用 CC 授權條款授權。此授權允許您在非商業的前提下重製、改作並建立衍生作品，條件是您需標示出處為 WFD，並以相同條款授權您的新作品。若有超出本授權範圍之使用，須向 WFD 另行申請授權。

本出版品中所載資訊與觀點均屬作者個人意見，未必反映 WFD、其資助者或英國政府的官方立場。WFD 或其代表亦不對本出版品內容的使用所產生的任何後果承擔責任。

本指引之現行版本完成於 2024 年 4 月 29 日。

主筆團隊已盡最大努力與專業確保截至報告日期內容準確無誤。

然而，人工智慧在國會中的潛在應用範圍多元且難以做長期的完全預測，亦無法保證人工智慧是否會促進民主與有效治理，或反其道而行。在此背景下，主筆團隊對任何依據本出版品內容所造成的損失不承擔責任，亦不對國會導入與使用 AI 工具與服務所產生的影響與後果負責。

ISBN (electronic version, CN)

978-1-0685174-8-8

致謝

本編輯團隊誠摯感謝全球國會社群在撰寫本指引 1.0 版過程中，就引言及國會工作中的人工智慧應用，給予熱烈響應及寶貴指教，該社群的集體洞見充實豐富了本出版品內容，使其成為全球各國會機關的參考資源，面面俱到。

本團隊也感謝波隆那大學（University of Bologna）的 Monica Palmirani 教授及其團隊於本文件撰寫過程中，全程提供專業知識與用心投入。他們對法律與科技議題的深刻理解，對於本指引內容的成成功不可沒。

我們亦向西敏寺民主基金會致上誠摯謝意，感謝其高瞻遠矚，全心擁抱本計畫，並給予充足協助進行推廣。該基金會致力於強化民主實務，有助於本指引在全球的推行，培力世界各地國會有效應對 AI 整合的各種挑戰。

內容索引

1. 引言

10

引言介紹了人工智慧（AI）與生成式人工智慧（generative AI），並說明為何需要指引、在國會情境中使用 AI 所面臨的挑戰，以及 AI 在國會中的可能應用方式。

2. 指引

22

本文件第二部分為指引。繼要點說明之後，詳細內容分為六大節，涵蓋多項關鍵議題：

- 倫理原則
- 通用人工智慧（AGI）
- 隱私權
- 治理
- 系統設計
- 能力建構

全篇共 40 項指引，以結構化的方式回應以下三大提問：

- 本指引為何重要？
- 是否已有實施案例？
- 如何實施本指引？

每項指引皆包含簡要延伸考量及建議。

3. 前行之路

80

第三部分扼要說明了各國國會未來製定人工智慧指引的著手之道。

4. 延伸閱讀

82

第四部分收錄縮寫一覽表、術語詞彙表以及參考書目。

主筆團隊

本出版品由國會學者及專業人士組成之國際團隊共同撰寫：

- Fotios Fitsilis 博士，希臘國會
- Jörn von Lucke 教授，德國齊柏林大學（Zeppelin University）
- Franklin De Vrieze，西敏寺民主基金會
- George Mikros 教授，卡達哈馬德·本·哈里發大學（Hamad Bin Khalifa University）
- Monica Palmirani 教授，義大利波隆那大學
- Alex Read，首席技術專家，聯合國開發總署（UNDP）
- Günther Schefbeck 博士，奧地利國會
- Alicia Pastor y Camarasa 博士，瑞士洛桑大學（University of Lausanne）
- Stéphane Gagnon 教授，加拿大魁北克大學渥太華河區分校（Université du Québec en Outaouais）
- João Alberto de Oliveira Lima，巴西聯邦參議院
- Antonino Nielfi 博士，澳洲國會
- Georgios Theodorakopoulos，希臘國家法律顧問委員會（Hellenic State Legal Council）
- Marina Cueto Aparicio，西班牙參議院
- Juan de Dios Cincunegui 教授，阿根廷奧斯特拉爾大學（Universidad Austral）
- Ari Hershowitz，Govable.ai
- Ahto Saks，愛沙尼亞國會
- Jonas Cekuolis，國會發展專家
- Jonathan Ruckert，NovaWorks Australia
- Elhanan Schwartz，以色列司法部
- Zsolt Szabó 教授，匈牙利正教會卡洛里大學（Károli Gáspár University of the Reformed Church）、塞切尼·伊什特萬大學（Szechenyi Istvan University）
- Nicola Lupo 教授，羅馬路易斯大學（LUISS University）
- Marci Harris，POPVOX 基金會

前言

我們已經目睹了人工智慧（AI）在國會工作場域所帶來的影響。在不遠的將來，我們可能會看到 AI 系統與以 AI 為基礎的服務順暢輔助國會議員，協助其處理國會事務及選區職責。

想像一下，有一套以 AI 服務輔助的可靠支援決策系統，能促成資訊充足的判斷；再設想一下，有一套智慧檢視機制，確保法律提案與既有法規不產生衝突，再加上 AI 主導的社群媒體平台政治論述監測工具。

以上皆不是科幻小說情節。即使在今日的技術能力下，此類數位解決方案不僅能開發出來，也能導入國會資訊系統之中，對機構運作與代表民意的功能產生深遠影響。

20 多位國會學者與實務工作者所組成的工作團隊共同努力之下有了本出版品。本指引涵蓋了倫理原則、通用人工智慧與人類自主、隱私與安全、治理與監督、系統設計與操作、能力建構與教育等面向。

本指引的出版不僅促進我們對 AI 的理解，更為國會以「負責任且促進共融」方式導入 AI 奠定了基礎。

西敏寺民主基金會倡導 AI 的民主化應用以及 AI 與國會機關的整合，並以此自豪。新科技應為民主所用，而非破壞民主。全球首套 AI 國會指引的問世，正彰顯本基金會引領國會創新方面的堅定決心。我們遍佈世界各地的專家團隊，將持續與研究人員與有志一同的各國國會攜手合作，發展並治理新技術，以促進全球民主發展。

Anthony Smith

執行長

西敏寺民主基金會

2024 年 7 月

編者序

擁抱數位工具及服務的國會機關日漸增多，人工智慧（AI）的崛起有望更近一步加速此趨勢，並發揮巨大作用，推動國會從傳統紙本運作組織轉型為資料驅動機構。

本份指引望能協助民意代表機關為國會工作場域導入與應用 AI 做好準備。本出版品由來自世界各地的國會學者與專業人士共同撰寫，立足於過去此領域的研究之上，自 2023 年 9 月至 2024 年 4 月完成，歷時八個月。

有鑒於技術與制度環境不斷變化，本指引仍需要持續完善，同時依然有潛力促進知情監理，在政策制定、公共參與、能力建構等方面，培力國會，確保 AI 的負責任整合，兼顧政治與行政流程中的透明度與倫理問題，進而加強公眾信任，並維護公共利益。此外，這套指引可協助 AI 工具及服務與民主原則及社會需求同步，也能大幅促進成功經驗與符合倫理之行為的分享，最終促成國會社群間的知識成長與跨國合作。

在多層級治理環境下，本指引相當切合地方、區域、國家乃至超國家議會之需求。本指引以全方位角度，處理倫理、隱私、安全、監督、系統設計與教育等議題，探討 AI 在國會內應用的明確面向，包括成功實施的範疇、案例及要素。因此本指引一方面不僅有助於處理當前議題，另一方面，對於評估更偏理論的問題，如通用人工智慧（artificial general intelligence，簡稱 AGI）對立法機構的影響，也相當切要。

科技日新月異。因此，本指引在設計上為技術中立，換句話說，並不處理特定的 AI 技術。然而，依然勾勒出主要科技趨勢的可能發展，例如生成式與混合式 AI。

我們希望本出版品能廣為流傳，觸及每一個國會、議員、行政人員以及所有關心如何在立法機關中，充分發揮 AI 正面效益同時又將潛在風險減至最小的熱心人士。正因如此，本編者團隊與執筆者懷抱堅定決心，深入國會，並與社會利害關係人合作，推動本指引的繼續發展。指引、合作及因地制宜的溝通有助於確保有效的落實，並針對多元的機構環境而應變調整。

我們誠摯歡迎有意者與我們合作，翻譯指引、開發訓練教材、協助指引之落實或分享實務案例，加速 AI 與國會工作場域有效整合。無論是單邊、雙邊，或多邊的概念驗證（PoC）與示範計畫，皆有助於在多元環境中進行實際測試並精進本指引內容，我們也期待從中汲取教訓。

Fotios Fitsilis

希臘國會

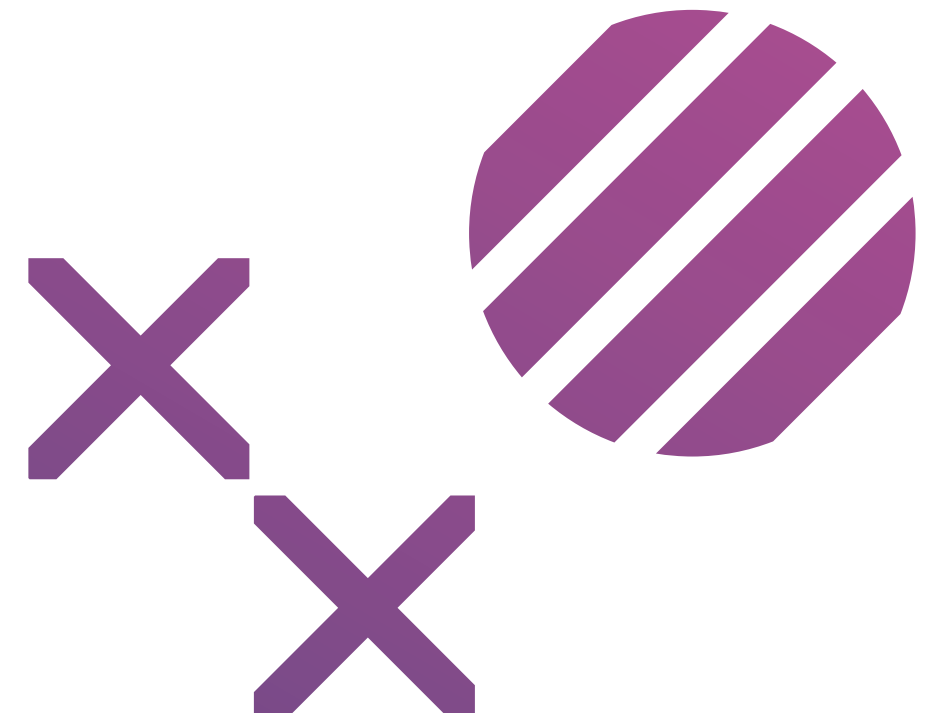
Jörn von Lucke

德國齊柏林大學

Franklin De Vrieze

英國西敏寺民主基金會

2024 年 7 月



重點摘要

背景說明

人工智慧（AI）為國會運作流程帶來轉型契機。AI 用途越來越廣，例如辯論逐字稿與翻譯、文件摘要、協助起草法律文件，以及與公民溝通等。已有數個高瞻遠矚的國會開始實驗 AI 或應用 AI 程式，潛在效益充分，涵蓋國會職權各種不同面向。

儘管 AI 對起草法案的影響仍在研究中，目前已有助於分析大量法律文件、找出規律並提出改進建議。此外，AI 演算法也可摘要冗長報告、法案與委員會調查結果，使立法者與公民更易理解與使用國會文件。

這些應用有助於提升透明度，並促進知情決策。此外，AI 驅動的聊天機器人亦可提供即時國會活動資訊與公民互動，促進公民參與。各種 AI 模型亦能趨勢預測、政策潛在影響及民意動向，給予預測型洞察分析。因此，藉由這類前瞻分析，立法者可主動回應新興議題，提升國會工作效率。

自 2022 年底以來，我們見證人們快速採用生成式預訓練轉換器（generative pre-training transformers，簡稱 GPT），這種 AI 技術具備前所未見的潛力，能提升國會職能。

雖有幾間機構反應迅速，但大多數機構依然沒有明確策略，不確定如何開發、執行與善用 AI 工具。本指引目的在於刺激數位創新及負責任的科技應用，另一方面也要預防 AI 對於民主及人類在眼下及未來可能帶來的威脅。

本出版品由來自 16 國、共 22 位國會專家學者與專業人士組成的技術工作團隊，歷時八個月（自 2023 年 9 月至 2024 年 4 月）共同撰寫完成。內容探討多項與國會相關的 AI 技術及用途、其導入所面臨的挑戰與障礙，以及 AI 監理的沿革。

指引內容

以下共 40 項指引，依主題劃分為六大領域，提供未來國會量身設計監管架構之通用指南。每項指引皆回應一組核心問題：此指引為何重要？是否已有已知實例？該如何落實？每則指引結尾提供應用建議，說明利害關係人如何在國會 AI 專案中採用及調整此技術。

各領域指引數目

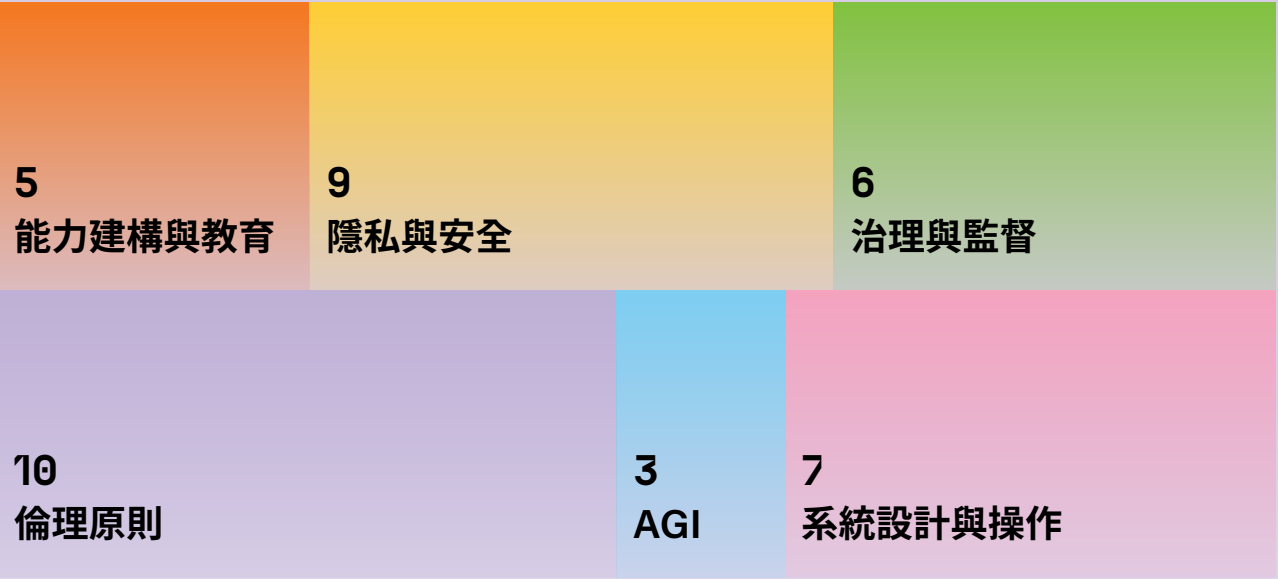


圖 1 顯示各領域指引的分布情形，說明專家們特別著重於倫理架構（共 10 項指引），同時亦未忽略通用人工智慧（AGI）議題，儘管此種情況發生的可能不太高。

本套指引重視的倫理原則，涵蓋確實問責、透明度與公平。同時強調尊重人類尊嚴、隱私權與文化多樣性的重要，同時因應並處理資料與演算法中的偏見問題。指引亦強調應促進人類自主與決策能力，並正視通用人工智慧（AGI）可能帶來的潛在影響。隱私與安全亦是重要考量，須採取穩健措施保護個人資料、預防網路攻擊。

有效治理與監督，是確保 AI 應用符合民主價值、維持透明度的關鍵，指引即為此提出說明。系統設計與操作應優先考量可互通、

透明度、可靠度與安全，同時也要對 AI 系統實施監管與監測。本指引亦重視能力建構與教育，協助國會議員及其工作人員掌握負責任應用 AI 所需的技能與知識。

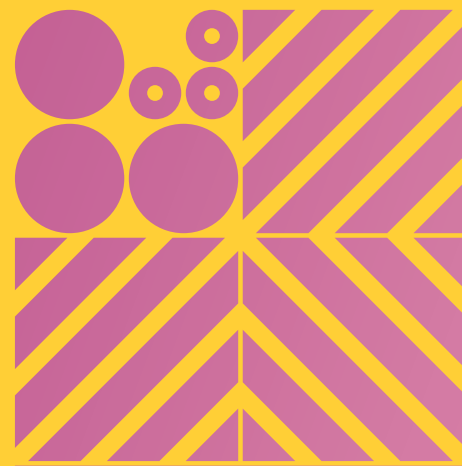
本指引另鼓勵與利害關係人合作、推動公共教育，以促進各方對 AI 應用於國會流程的理解與接納。本文件亦主張若要分享經驗及資源以加速 AI 技術實施，則跨國會及與議會組織間的合作不可或缺。

Part 1.



引言

Introduction



隨著人工智慧（AI）工具與服務迅速演變並被廣泛採用（包括國會流程在內），我們必須建立一套兼顧倫理與實作的指引，以確保確實問責、資訊透明與人類自主，同時促進永續發展目標，並保護隱私、安全與多元性。

為達成此目標，首版指引（v1.0）於 2023 年完成。¹本出版品即是立足於第一版的基礎之上，為全球各國國會發展出一套全面且實用的應用架構，協助國會探討這些科技及應用，並發展出自有的監理架構。

本出版品不僅探討 AI 在國會工作場域中的應用，更採取了宏觀角度，闡述 AI 對國會議員、國會行政與國會機關本身的潛在影響。

本文件內容著重於限制與預防措施。此舉並非為了阻擋 AI，反而是鼓勵國會在特定條件下，深入認識 AI，並在擁抱及善用 AI 的同時，對 AI 的潛在風險與機會保持警覺。眼下極需開始關於 AI 與民主的公共與政治討論，我們必須時時銘記 AI 在形塑國會的未來上展現了可觀潛力。

應用 AI 是國會數位轉型中的最新發展方向。數位科技對國會²與整體民主制度各種不同層面的影響，已有充分紀錄文獻。³然而一旦涉及到 AI，學界在研究其對機構發展的影響時，多採取較保守，而非顛覆性的觀點。然而，這種審慎態度往往忽略了 AI 與生成式 AI（AI 最顯著的子類別）已成為國會工作場域的潛在「轉捩點」，有望使國會在效率、效益與透明運作上皆更上一層樓。本指引除單純作為輔助文件外，亦概述介紹這些新興科技所帶來的正向潛能與相關挑戰。

AI 與生成式 AI 介紹

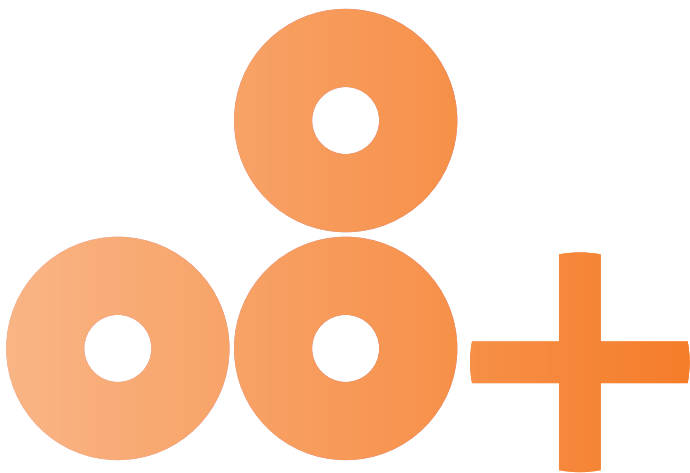
AI 是一個複雜且日新月異的領域，過去已有許多人嘗試說明。⁴與其提供簡潔定義，本出版品選擇採用較為概括式的說明，勾勒出 AI 技術、概念、風險與國會工作場域內導入 AI 所帶來的效益。⁵

「人工智慧」一詞泛指綜合不同科技、學習方法、系統架構、演算法與運用操作，運用電腦運算能力來模擬人類智慧，以獨立或依指令執行特定任務，包括：自主系統、機器學習、深度學習、神經網絡、模式辨識、自然語言處理、即時翻譯、聊天機器人與機器人等。

AI 的功能是為了輔助或自動化人類行為與流程。模式與文字辨識、語音與說話者辨識、影像與空間辨識、臉部與手勢辨識等，開啟了各式各樣的可能應用。以 AI 為本的文字、聲音、語音、影像、空間與影片生成系統以及寫程式功能，也同樣擴大了應用範圍。這些技術將帶來更多新系統、應用與流程，促成以 AI 為本的即時感知、通知、建議、預測、防範、決策與狀態意識（situational awareness）。

生成式人工智慧（generative artificial intelligence），有時簡寫為 GenAI，則以既有訓練素材為基礎產生全新內容，生成式 AI 並非純粹隨機運作，而是靠著已識別出來與學習到的模式產生合成資料。如 ChatGPT 這樣的大型語言模型（large language model，簡稱為 LLM），就支援生成文字與程式碼，而 AI 翻譯服務則可將文本轉換為其他語言。其他應用領域包括簡報製作、IT 系統程式設計與工作流程規劃等。文字也能用來生成不同音調的語音與聲音序列。生成圖片與影片的能力亦日益重要，許多人特別擔憂的是利用圖片資料與錄音檔生成對嘴影片的風險（deepfakes 深偽技術）。

現階段已有多種開源與封閉的大型語言模型可供使用，從方法論的角度來看，如何選擇最適合特定用途的模型，本身即是一項關鍵課題。有些模型規模龐大，有些則可本機安裝。然而，國會於應用時，還需考量是否適用，一方面也要保障基礎設施主權、避免外部行為者干預、保障資料所有權、確保過程可追溯性，以及維護整體程序的合法性。⁶



國會應用 AI 為何需要指引？

指引可提供明確的架構、條理一致的作法與方向，有助於實務交流，提高解決方案與方法在其他機構內複製的可能性，並確保遵循倫理行為守則，促進知識成長與研究者間的合作。在資安⁷與個人資料保護領域早已出現此類文件。⁸

國會應用 AI 指引可確保國會以負責任方式整合 AI 應用，處理機構行政與決策流程所涉及的透明度與倫理議題，同時促進公眾信任。此外，指引亦可確保 AI 工具及服務與民主原則及社會需求同調。從法律角度來看，制定這類指引對法律理論的發展或許也會有實質貢獻。⁹

下頁的表格羅列了適用於國會情境，最為重要的若干原則，以及在國會工作場域中的可能應用。2024 年 2 月，義大利眾議院亦發表了一套類似的原則。¹⁰可以從表格清楚看到，AI 具備潛力，能為國會生態系帶來許多正面改變。

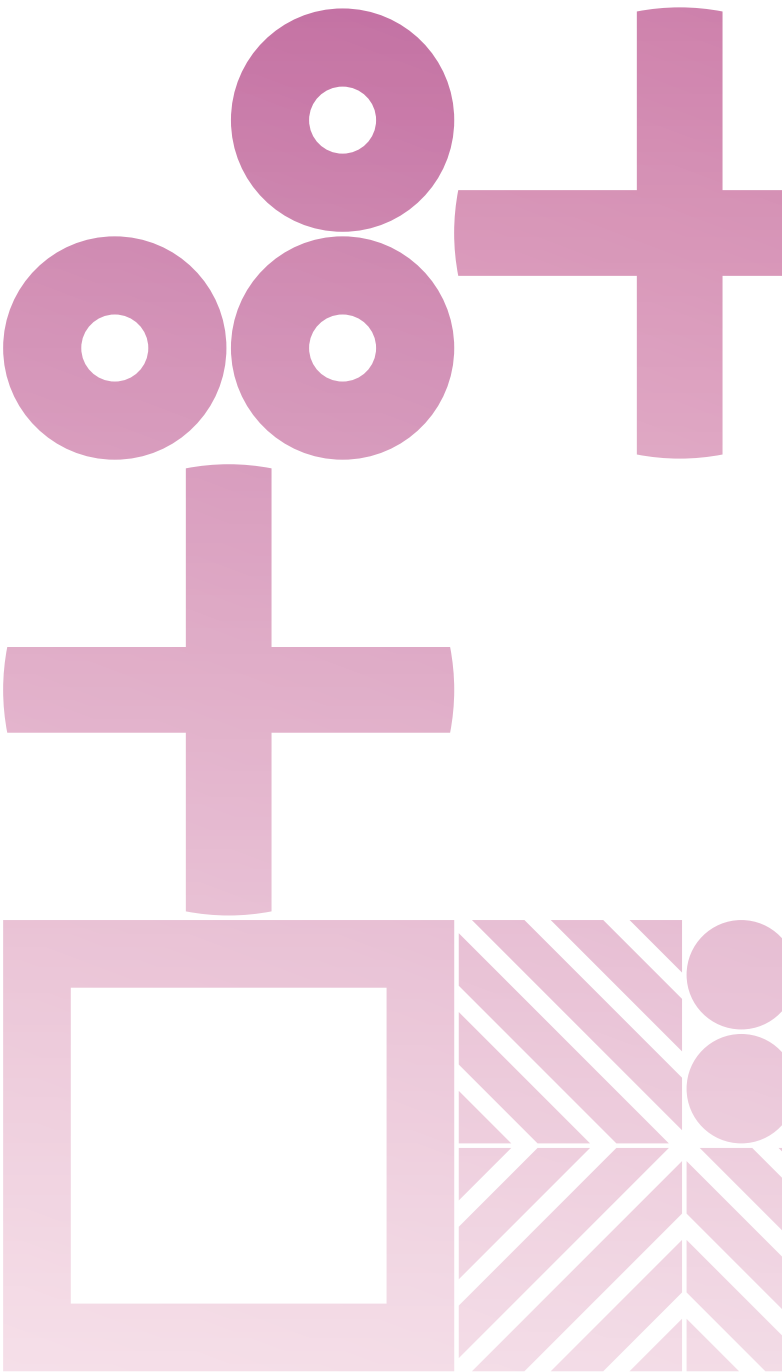
本指引第二部分已羅列這些原則，其中亦包含與資安、資料隱私相關的指引。

國會應用 AI 相關原則	在國會工作場域中的應用
問責與透明	確保 AI 決策及應用，可理解、追溯、於理有據
決策者自主權	維持決策者自主權，防止操弄
符合倫理與負責任的 AI 應用	維持倫理標準，避免 AI 應用時發生誤用或偏見
人類監督與相關解釋	維持人類掌控 AI 系統的能力，同時也有能力向不同受眾（例如法律從業者、公民）解釋說明
緩解風險與基本權利影響評估（fundamental rights impact assessment，簡稱 FRIA）	辨識並處理與 AI 導入相關並由基本權利影響評估（FRIA）偵測到的風險
公眾信任	建立並維持大眾對國會機關使用 AI 工具與服務的信心
共融與多元	促進國會行政與決策過程中的公正與平等
適應科技進展	協助國會運用 AI 進步技術提升運作效率與成效
跨國會合作	促進各國國會調和全球 AI 政策與法規，使其步調一致
公眾參與	鼓勵公民與社會利害關係人參與國會導入 AI 與其應用之討論與決策
符合法規	確保國會使用 AI 符合相關法律與規範

國會 AI 系統與解決方案

AI 具備為國會生態系帶來正向變革的潛力，也相當切合不同類型的國會服務需求。執筆團隊選擇類型學分類法，全面羅列出各種以 AI 為基礎的應用領域，並強調在加強國會程序，確保效率、透明度與反應能力方面，應用 AI 的各種方法。¹¹ 下頁表格列舉了國會應用 AI 方法實例。

最上方的國會 AI 應用依其相關屬性劃分為數個群組，此分類方式係根據專家建議與數國國會機關的實證資料整理而成，這些國會分別為：希臘國會、阿根廷國家眾議院，以及加拿大國會。¹²



國會 AI 應用類型

類型	AI 應用
國會議員	• 國會議員發言即時字幕 • 於全體會議與委員會中使用的可靠投票系統 • 口頭與書面質詢內容生成 • 支援資訊檢索
立法	• 檢視法案草案與其他法規的交互作用 • 依據已知漏洞、問題及相關法律提供立法建議 • 草擬初步法律提案 • 強化監管以及數位化友善的政策推動
國會監督與國會外交	• 分析與國會活動有關之媒體報導 • 與國會活動相關之社群媒體資料分析 • 偵測資訊環境受操縱之狀況 • 排除 AI 產出提案所包含的偏見與歧視內容
公民教育與國家文化	• 國會網站前端智慧搜尋功能 • 透過（連結式）開放資料促進透明度 • 辯論與論點的視覺化呈現 • 促進公眾參與國會程序
國會行政、國會建築、 駕駛服務與警力	• 服務身心障礙者的虛擬助理 • 資安軟體 • 生成會議記錄與翻譯服務
國會各局處以及選舉相關	• 偵測由 AI 生成且試圖操弄民主程序之偽造內容 • 流程自動化 • 專案管理
研究／科學服務	• 智慧文件檢索 • 進階知識管理 • 事實查核

這類廣泛的 AI 應用範疇，突顯了此種科技輔助、精簡、甚至加強國會各項程序的方式相當多元。

世界各地的國會已在使用類似系統，成熟度不一。¹³ 大多數系統運用了自然語言處理（natural language processing，簡稱 NLP）演算法，其中最常派上用場的功能包括語音轉文字、文字分類與模式辨識，而模式辨識亦涵蓋語音、影像、物件與臉部辨識等技術。

這類系統應用有兩點值得關注。¹⁴ 首先，許多國會似乎優先精簡立法程序相關的流程，包括審議、院會與委員會會議。其次則著重於給予公民的數位服務，例如提升公民資訊近用，以及分析公眾諮詢工具所收到的回饋意見。

還有一項新趨勢，同時結合多種 AI 技術，以降低單一方法的風險；並採取象徵式（symbolic）、次象徵式（sub-symbolic）與神經象徵式（neuro-symbolic）AI 的混合式策略。¹⁵

國會應用 AI 的挑戰與阻礙

國會導入 AI，既帶來前所未有的機會，也伴隨著不可小覷的挑戰。

目前，國會使用 AI 的相關法律與規範尚付之闕如。此種監理真空所產生的未知數，可能導致大眾無法信任 AI 服務及其提供者。此外，AI 工具可能存在的資安漏洞，也引發外界對國會系統安全性與健全的疑慮。另一方面，對 AI 的認識仍顯有限，即便是在工程技術領域亦然，而國會相關人員尚未受過充分訓練。知識不足不僅有礙有效整合與操作，也使國會利害關係人更容易受外界影響。

本文主張 AI 是一股轉型之力，希望提供具體指引探索新領域，協助國會駕馭其益處，同時也能防範潛在陷阱。

隨著 AI 進入國會領域，建立相應的防範措施與規範監管制度已刻不容緩。¹⁶ 要建構有效的監理架構，需處理以下多項關鍵考量：

- 資料隱私與資訊安全，以及資料存取權限與資料所有權問題。
- AI 系統的不同部署選擇，例如選擇地端安裝與雲端服務的利弊風險。¹⁷
- 服務與資料的可攜性（portability）。
- 確保 AI 服務提供者的可信度與明確的所有權結構。

- 倫理疑慮與偏見問題，特別是訓練資料的品質。
- 透明度、可解釋與問責機制，這些均為建立公眾對於國會 AI 系統信任的關鍵支柱。
- 決策者自主性，這是讓眾人接納法律實務工作者使用 AI 為輔助工具的根本前提。
- 若要共融與有效實施 AI，多語言處理能力必不可少。
- 公共參與，可落實民主價值並隨時廣納外界觀點。

此外，國會日常業務中導入 AI 技術，也需有標準與架構可循。例如，需訂定明確規範，內容涵蓋資料儲存與刪除的範圍、倫理監督、持續監測，確保國會 AI 系統符合最高標準。這也再度點出了建立國會 AI 系統品質基準之必要。

然而，由於甚少有國會具備足夠專業能力與資源來處理上述課題，因此，本文亦主張應加強機構間與國會間的合作。

整體而言，這份指引的目標在一方面駕馭 AI 的轉型潛力，一方面也維護國會制度之完整，並在兩者間取得平衡。

AI 監理沿革

規範國會 AI 的議題目前尚未獲得各國國會的重視。AI 在各國國會的應用情況，從全面整合至拒之門外皆有。阻礙及規範都有可能限制 AI 開放所帶來的優勢與機會。這種分歧突顯了此技術尚在演變當中的過程，因此制定指引推動國會以負責任的方式擁抱 AI，更加有必要。

相對於不具法律約束力的文件或「軟法」（soft law），例如決議、行為準則與指南等，具法律拘束力的文件或「硬法」（hard law）包含了法規、命令與法律。

有兩項具法律約束力的文件值得一提：首先，歐洲議會在通過多項相關決議後，最終於 2024 年 3 月通過《AI 法案》（AI Act）。該法案採用以風險為基礎的方式，要求 AI 系統的開發者與部署者履行諸多義務，包括針對高風險應用進行基本權利影響評估（fundamental rights impact assessment，簡稱 FRIA¹⁹）。此法亦將國會所使用的部分 AI 系統列為高風險技術，針對此類應用羅列特定義務。

其次，歐洲理事會已完成《人工智慧、人權、民主及法治架構公約》（Framework Convention on AI, Human Rights, Democracy, and the Rule of Law），即將開放使用與批准。²⁰ 本公約為全球首部涵蓋 AI、人權與法治，且同時具有法律拘束力的法律文件。不過，目前此公約尚未針對國會機關使用 AI 技術制定額外義務。2020 年，歐洲理事會議員大會（Parliamentary Assembly of the Council of Europe，簡稱 PACE）

通過多項決議與建議，自此開始探討 AI 對人權、民主與法治的可能影響，為此公約奠定基礎。²¹

同時，於 2024 年 3 月，聯合國大會也踏出重要的一步，為推動 AI 應用促成全球共好，而通過一項重要決議。決議的目的在於促進「安全、可信且保障資訊安全」的 AI 系統發展，以加速實現《2030 永續發展議程》。²² 儘管如同《世界人權宣言》，本決議並不具法律效力，但各地區與國家法規文件經常視其為用來達成總體目標的「道德指南」。

儘管目前各國已展現監理 AI 的可觀努力，截至 2024 年初，尚無針對國會內部應用 AI 的具體指引或原則，而國會正是民主體制的最高機構之一。²³ 根據 2022 年底一項調查（在 OpenAI 開放免費 ChatGPT 基本服務之前），全球已有 10 國國會機關正在使用共 39 項 AI 解決方案。²⁴ 隨著 ChatGPT 問世，眾人對於生成式 AI 對立法的間接或直接影響也產生高昂興趣。²⁵ 尤其是 2023 年美國國會採購了 40 組 ChatGPT Plus 授權，讓內部人員探索生成式 AI，國會分發使用授權給國會議員辦公室，供立法者與工作人員於內部使用此項轉型技術進行實驗。²⁶ 2024 年 4 月，美國國會院務管理委員會（Committee on House Administration，簡稱 CHA）公布了一套國會內部使用 AI 或其他科技時的總方針。²⁷

2017 年，英國國會正式登記成立「人工智慧跨黨派議會小組」（All-Party Parliamentary Group，簡稱 APPG），率全球之先討

制定與改進指引的方法

論 AI 技術之應用與影響。2023 年 3 月，英國政府發表了一份白皮書，提出以創新為導向的 AI 監理架構。該架構注重「比例適當、不過時且支持創新」，目的為兼顧技術發展與風險管理。²⁸ 隨後於 2023 年 11 月，上議院提出了一項 AI 監理的法案，目前正處於委員會審查階段。除本法案外，全球各地皆在討論類似法案，顯示對於此一強大技術，確實需要設下合理限制。

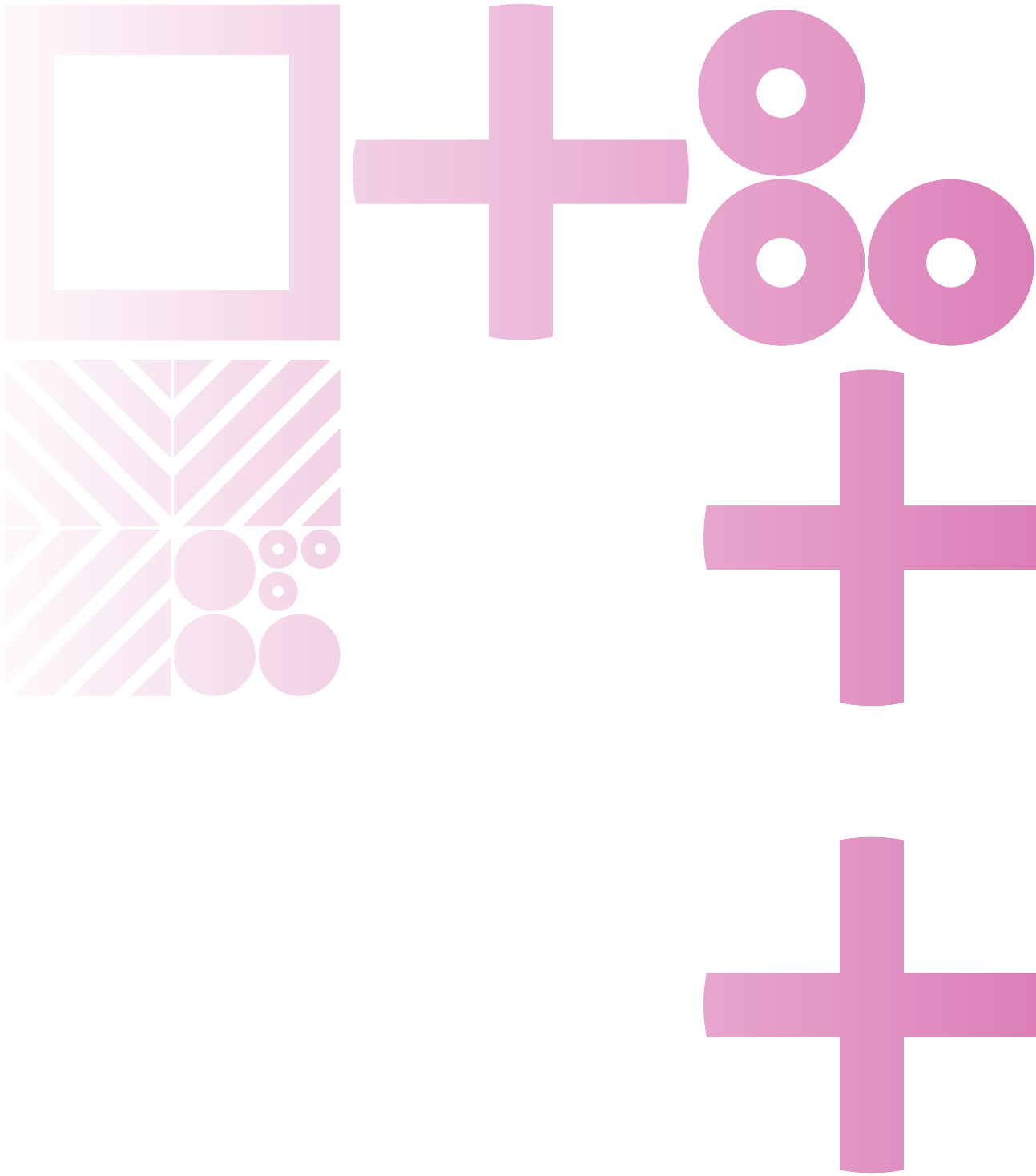
因預期未來 AI 工具與服務會進一步導入國會工作場域，已有人著手研擬指引與法規。2023 年 4 月，即有臨時工作小組成立，針對國會工作場域導入 AI 與應用完成初版指引。本文件即承襲初版的基礎，是為 2.0 版。

技術工作團隊制定指引時所遵循的工作方法來自於現有知識、文獻分析，同時結合國會事務專家的見解。2023 年 9 月開始，至 2024 年 4 月完成，同時也透過互動工作坊來輔助更新指引的修訂過程。

由於共有超過 20 位專家參與，指引撰寫從一開始要達成全體共識，就遇上艱鉅挑戰，處處折衷皆為必要，本出版品即是努力平衡之下的成果。但最終，各國國會依然要扛起責任，根據這套指引自行界定適用範疇、擬定策略、辨別輕重緩急。

本研究方法的基礎，結合人類智慧與先進 AI 能力，涵蓋協作式編輯工具 text pads 與大型語言模型（LLMs）。專家們在多元腦力激盪會議中，除傳統人類腦力激盪外，還輔以 LLM 腦力激盪的創意。使用如 OpenAI 的 ChatGPT 模型（GPT-3.5、GPT-4）生成指引，進行比較分析，使得團隊對於解決方案的瞭解更豐富與評估更詳細。

由此制定出的 40 項指引，經過詳細分析，納入設計思維（design thinking），強化以使用者為本的原則，共歸納成六大領域。



Part 2.

國會應用 AI 指引

Guidelines for AI in Parliaments

指引摘要

1 倫理原則	24	2 通用人工智慧與人類自主性	36
1.1. 問責與透明	26	2.1. 促進人類自主	38
1.2. 尊重人類尊嚴、權利與隱私	27	2.2. 對設計者與開發者的倫理要求	39
1.3. 公平、平等與不歧視	28	2.3. 正視 AGI 可能成真	40
1.4. 處理資料與演算法內之偏見	29		
1.5. 維護智慧財產權	30	3 AI 隱私與安全	42
1.6. 守護人類共同價值及文化多元性	31	3.1. 納入安全性與強健的資安功能	44
1.7. 評估與減輕非預期後果之影響	32	3.2. 納入隱私保護設計 (privacy-by-design) 的概念	45
1.8. 公眾參與及互動	33	3.3. 對個資進行安全處理	46
1.9. 遵循法治與民主價值	34	3.4. 外包考量	47
1.10. 推動政策目標	35	3.5. 資料主權議題考量	48
		3.6. 確保來源資料的完整度	49
		3.7. 過度依賴 AI 之風險	50
		3.8. 確保訓練與測試資料之安全	51
		3.9. 確保安全決策由人為進行	52

4 AI 治理與監督	54	5 AI 系統設計與運作	62
4.1. 納入整體國會數位化策略	56	5.1. 推動標準化資料格式與流程	64
4.2. 有效資料治理與管理規範	57	5.2. 強調 AI 演算法之可解釋性	65
4.3. 建立國會倫理監督單位	58	5.3. 建立穩健可靠的 AI 系統	66
4.4. 評估國會 AI 之影響	59	5.4. 監理 AI 系統的使用與部署	67
4.5. 確保資料存取與控管之安全	60	5.5. 評估風險	69
4.6. 與利害關係人合作	61	5.6. 監測與評估 AI 系統	69
		5.7. 達成對於最低準確度之共識	70
		6 AI 培力與教育	72
		6.1. 建立專家團隊	74
		6.2. 規劃訓練課程	75
		6.3. 支持知識交流與合作	76
		6.4. 紀錄 AI 相關活動	77
		6.5. 針對 AI 在國會的應用與限制 進行公眾教育	78

1.

倫理原則 Ethical principles

民主制度建立在問責制度與透明運作之上，這是全球各地國會機關的金科玉律。倫理原則給予國會 AI 系統的開發與部署一套可供依循的架構，確保系統值得信賴、透明，並符合人類共同的價值。如此可充分發揮 AI 效益，同時又降低其潛在危害。

在開發、實施、與使用 AI 技術過程中，人類尊嚴與隱私至上，同時也反映公平、平等與不歧視等價值。國會必須正視資料與演算法中的偏見，並協助維護人類價值與文化多樣性，例如透過謹慎的模型訓練與部署方式。為此，將需制定新的評估標準，以因應並減輕人工智慧可能帶來的非預期後果。為確保人工智慧的應用能因應各國國會的情境並取得共識，公眾的參與與投入將愈發關鍵。最終，這關係到對法治與民主價值的維護。



1.1. 落實問責機制與透明作業

► 為何重要？

在國會導入與部署 AI 系統時，落實問責機制與透明作業，為必要步驟，才能保障民主程序完整，並保護公民的權利與利益。

► 是否有已知實例？

2020 年，歐洲理事會議員大會（Parliamentary Assembly of the Council of Europe，簡稱 PACE）通過多項決議與建議，探討 AI 對民主、人權與法治的影響，²⁹ 同時 PACE 也支持人工智慧應用的開發和實施時，必須堅守一套基本倫理原則，包括透明作業、人類為演算法決策負責等等。

► 如何落實？

為了促進國會 AI 系統的問責機制、可稽核度和透明度，國會需要實施明確的應用政策，優先考慮倫理原則，並建立獨立的稽核機關進行監督。此外，議會需要建立透明的資料實務處理法以及演算法問責制，並定期發表系統性能和演算法報告。在這方面，使用可解釋的人工智慧很重要，值得鼓勵，但仍有技術侷限。³⁰ 同時也應兼顧與利害關係人、專家的互動，並處理偏見問題。最終，立法機關積極持續監督系統需要成為常態。

► 補充建議與考量

- 鼓勵研究和學術機構對國會程序所使用的 AI 系統進行獨立評估。
- 在國會環境中培養問責和透明的文化，鼓勵議員和工作人員接納這些原則。

1.2. 尊重人性尊嚴、人權與基本權利及資料保護規範

► 為何重要？

國會機關應確保 AI 技術的運用合乎倫理、具備責任感，AI 技術的開發和實施，在各方面皆應尊重人的尊嚴和隱私，這點十分重要，才能保障所有涉及國會流程或是會受流程影響民眾的權利。

► 是否有已知實例？

歐洲理事會議員大會在 2020 年提出的原則涵蓋了正義、公平、隱私等，在開發和部署 AI 應用時皆應遵守此原則。³¹

► 如何落實？

國會可透過採取嚴格的資料保護規範與政策、落實倫理人工智慧指引，以及定期進行隱私影響評估，來維護 AI 應用中的人性尊嚴與隱私權。此外，透明的 AI 系統可以確保個人資訊獲得謹慎處理，並尊重個人的權利和尊嚴。

► 補充建議與考量

- 在國會制度下設立資料保護官（data protection officer，簡稱 DPO）或隱私權倡導專員有助於監督是否遵循隱私規定，也能提供指引。在歐盟，DPO 職位法源為《一般資料保護規則》（General Data Protection Regulation，簡稱 GDPR）³²
- 此外，可以考慮制定一套國會使用人工智慧的倫理守則，涵蓋與隱私和人類尊嚴相關等原則。

1.3. 落實公平、公正與不歧視的原則

► 為何重要？

在使用和部署國會 AI 系統時，如果要防範這些技術導致政治或機關程序中出現偏見或不平等情況，那麼落實公平、公正和不歧視的原則就非常重要。

► 是否有已知實例？

2020 年，歐洲理事會議員大會提倡一套 AI 應用的發展與落實的基本倫理原則，³³涵蓋了正義與公平等內容，同時通過一項決議，防範因使用 AI 而導致歧視。³⁴

► 如何落實？

國會可藉由維持 AI 開發團隊的多樣性、進行偏見稽核，制定明確的指導方針來降低決策過程中的偏見，以落實這些 AI 原則。而定期檢查 AI 系統是否存在落差，並即時處置，可進一步加強落實原則。

► 補充建議與考量

偏見是政治程序中固有存在，實際上也是無可避免的組成成分。AI 提供了寶貴的工具，有助於辨識各種偏見，建立結構分明的政治論證。此指引的重點在於訓練資料不足或不平衡，會導致的不當歧視型偏見。若要解決這些問題，需要與弱勢族群和倡議團體合作，收集 AI 系統影響相關回饋意見並做出相應的改進。此外，可在國會內部培養符合倫理應用 AI 的文化，以公平、公正和不歧視為核心價值。此外，國會可與研究機構及公民社會組織合作，研究 AI 對國會程序是否公平和公正的影響。

1.4. 瞭解並解決來源資料和演算法中的潛在偏見

► 為何重要？

理解並處理訓練資料中可能存在的偏見³⁵，是確保人工智慧系統在國會程序中遵循公平、公正與不歧視原則的關鍵一步。國會機構可主動針對訓練資料與演算法中可能出現的偏誤進行處理，以確保人工智慧系統在支持政治或制度性決策過程時，更有可能產生公平且無偏的結果。

► 是否有已知實例？

AI 的應用存在各種偏見風險，這也與基礎模型的訓練和開發有關，美國第 14110 號行政命令就處理了偏見等相關問題。³⁶

► 如何落實？

國會可以透過徹底檢視和稽核訓練資料來源來落實此項原則，以偵測和降低偏見。此外也可採用透明的資料蒐集方法，確保資料集的多樣性和代表性，並定期評估 AI 系統輸出的資料，找出和修正所使用的資料和演算法中的潛藏偏見。然而，某些針對資料偏見的因應措施，可能在倫理上有爭議，或需要開發新方法和技術。

► 補充建議與考量

- 在國會工作場域推廣資料倫理³⁷文化，強調處理偏見的重要性。
- 促進與學術機構及研究組織的合作，以掌握偏誤偵測與消除的最新、最佳作法。
- 考慮發表透明度報告，詳細說明採取了哪些措施，處理 AI 系統偏見及對公平和公正的影響。

在此脈絡下，重要之處在於我們要正視使用歷史資料時，不論如何，固有偏見幾乎無可避免，通常需要介入處理，不僅在訓練資料層面，演算法層面也需要介入。然而，如果「介入」意味著因為政治立場而排除某些資料，那麼就會陷入困難且危險的境地。

1.5. 避免使用侵害智慧財產權的訓練資料

► 為何重要？

避免使用侵犯智慧財產權（IP）的訓練資料，不僅是出於倫理之必要，也是法律的要求。使用第三方非國會產出的資料，且其使用方式偏離原始公告目的時，可能引發問題。在開發用於國會的AI系統時，必須遵守相關的智慧財產法與規範。國會機關可確保其AI開發過程遵守智慧財產權，並遵守倫理與法律標準，以降低侵權風險。

► 是否有已知實例？

已有相關報導提及基礎模型訓練侵犯智慧財產權的指控，儘管尚未波及國會。值得注意的是，2023年，OpenAI 因其聊天機器人 ChatGPT 疑似未經作者同意使用書籍進行訓練，而在舊金山聯邦法院面臨集體版權訴訟。同年，《紐約時報》控告 OpenAI 和微軟，指控其使用報社的專有內容訓練聊天機器人，並與報社直接競爭。³⁹

► 如何落實？

國會可先取得正式許可、使用開源或授權資料，並進行盡職調查來避免侵犯智慧財產權，確認資料來源符合版權與授權協議。對於國會或地方層級的議會AI系統，與政府機關、出版商、媒體或大數據所有權人簽署通用協議或許不失為可行之道。

然而，對於使用全球資料訓練的系統而言，此做法難度較高，甚至可能不切實際。這個思想實驗引發了人們對國會應用的一般範圍AI模型的質疑。

國會文件本身也可應用於訓練，因其原則來說不受智慧財產權保護。

► 補充建議與考量

對於智慧財產權的尊重應深植於機構文化中。除了倫理層面外，還涉及法律層面。因此，需與專精於智慧財產與科技法的法律專家合作，確保各方面皆符合智慧財產權法規。這些內部或外部專家需持續掌握不斷演變的智慧財產法規與AI，特別是大型語言模型（LLM）發展的最佳實務作法，以調整相關政策與操作流程。

1.6. 守護人類共同價值與文化多元性

► 為何重要？

在國會AI的設計與執行過程中，守護人類價值與文化多樣性，是確保AI技術符合其所服務社會的倫理與文化規範的關鍵，可促進國會共融環境與文化敏感度。

► 是否有已知實例？

2020年歐洲理事會議員大會的決議和建議專門探討了AI為人權帶來的機會和風險。⁴⁰ 人權與價值觀密切相關，但與人權不同的是，價值觀未必具普世性或法律效力，並且在不同文化與社會中可能差異甚大。

► 如何落實？

國會AI設計若要維護人類價值與文化多樣性，就要招募具備共融意識的開發團隊，確保廣納多元觀點，並展現文化敏感度。開發團隊可接受文化敏感度訓練，以理解各種文化細節及相關倫理架構。

然而，不同社會對價值觀可能存在深刻分歧，且價值觀並未完全成為法條，因此難以辨識或描述。因此，在評估各國國會實施指引的情境時，應優先參照憲法規範與已有成文法規範的普世人權，而非模糊的「價值觀」概念。

► 補充建議與考量

- 與文化組織、專家及學術機構合作，深入理解AI設計與部署的文化面向。
- 針對AI應用於國會程序所帶來文化的影響，鼓勵進行研究與學術探討。
- 與多元文化群體溝通管道保持暢通，確保能持續獲得意見回饋並對回應疑慮。

1.7. 評估並減少非預期後果或附帶損害

► 為何重要？

國會機關可主動評估與減輕 AI 系統使用時產生的非預期後果或附帶損害，確保國會程序以負責任及可問責的方式部署 AI 技術。

► 是否有已知實例？

尚無已知實例。

► 如何落實？

處理因為國會 AI 系統造成的非預期後果，包含數個複雜步驟。首先，建立一套完整的評估架構十分重要，架構應包括定期的影響評估，並輔以第三方稽核取得中立意見。再來為持續監測 AI 系統，以及時介入處理。此外，導入使用者意見回饋機制，使用者若能直接提供意見，有助於調整系統，降低負面影響並提升整體效能與問責能力。

► 補充建議與考量

- 參考既有系統評估所得之建議與結論。
- 持續掌握新興 AI 研究、最佳案例與倫理準則，以因應不斷變化的挑戰並降低潛在後果。
- 鼓勵國會工作人員與議員參與 AI 系統及使用結果的相關訓練。
- 推動國會內部負責任使用 AI 的文化，鼓勵人員回報疑慮並提出改進建議。
- 推動積極公民參與與民主參與的文化，鼓勵個人積極參與國會 AI 政策與實務的形

1.8. 鼓勵公眾參與國會 AI 系統的發展、推動與監督工作

► 為何重要？

於初期鼓勵大眾回應及參與國會 AI 系統的開發、執行與監督，有助於從一開始就確保共融意識、透明度與民意代表性，如此一來可反映國會所服務之公眾的價值觀、需求與觀點，促進更具共融意識與更具民意代表性的民主程序。

► 是否有已知實例？

尚無已知實例。

► 如何落實？

國會可設立專門平台以蒐集大眾意見，舉辦 AI 政策的公眾諮詢或聽證會，甚至成立有公民成員加入的顧問委員會。專家與公民應可取得資料集、模型與流程的資訊，使其能主動參與及互動。國會亦可公開 AI 相關的文件供公眾審閱與回饋意見，確保在發展、執行與監督 AI 系統的過程海納百川。

► 補充建議與考量

- 推動積極公民參與與民主參與的文化，鼓勵個人積極參與國會 AI 政策與實務的形塑。
- 運用科技促進線上參與，使不同地區的公民皆能參與討論與協商。
- 肯定並表揚積極參與國會 AI 負責任應用規劃的公民與組織的貢獻。
- 投資 AI 素養教育也是一種方式，可讓公民參與共創過程。

1.9. 尊重法治與民主價值

► 為何重要？

國會AI的開發與應用尊重法治與民主價值，是維護民主程序完整度與國內外法律原則的關鍵，促進國會內培養符合法律規範的民主環境。

► 是否有已知實例？

2020年，歐洲理事會議員大會（PACE）通過一系列決議與建議，分析AI對民主與法治所帶來的潛在益處與危險。⁴¹

► 如何落實？

國會應確保AI系統遵循現有法律與憲政架構以及相關的AI指引，包括倫理指引。作為代議機關，國會亦可建立透明的問責機制、定期稽核AI流程，並納入立法監督，確保AI系統符合民主價值、法律與憲政標準，並保障公民權。

國會在使用AI系統的同時，若要保障公民權有一些方法，例如可考慮調整像是《公民與政治權利國際公約》⁴²這種現有工具，或採用目前正在制定中、專門處理此類議題的法律文件。⁴³

► 補充建議與考量

- 在國會工作場域培養法律與倫理意識文化，強調維護民主價值與法治的重要。
- 與法律專家、學術機構及關注AI治理與民主價值的公民社會組織合作。
- 持續掌握AI治理相關法律發展與全球最佳實務，以調整政策與實務操作。

1.10. 運用AI去推動並監督全球性、國家級或區域型政策目標

► 為何重要？

當我們需要推動與監督像是SDG這種國家、區域、或全球性的目標時，應用國會AI系統可以發揮關鍵作用，因應各種挑戰。國會有監督與節制之則，監督這類目標即為其核心職能之一。此種做法本質上與倫理原則密切相關，因其有助於推動邁向更永續與更公平的未來。

► 是否有已知實例？

在數位化政策制定的架構下，希臘國會的科學服務部有個工作小組，正在探討如何運用AI工具來監測國家級的永續發展目標（SDGs）。⁴⁴

► 如何落實？

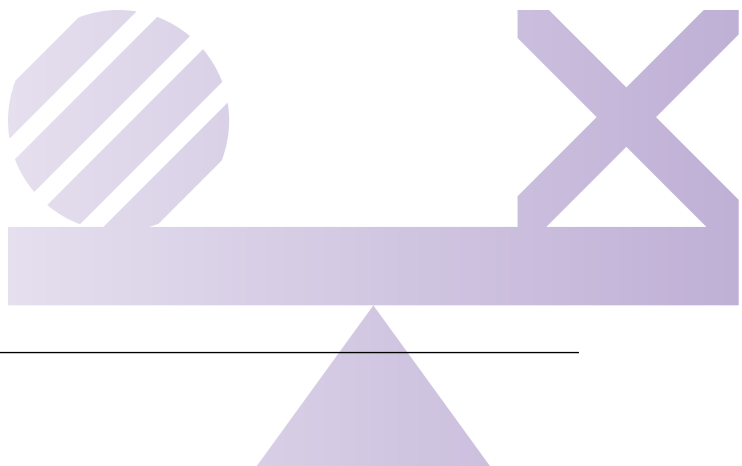
國會AI工具也有助於《京都議定書》與《巴黎協定》等國際協議的落實工作，AI可提供資料為本的洞察分析，蒐集證據，並提供國會議員與政策制定者各項協議相關資訊。如此，國會可運用AI系統分析並強化政策制定、監測進展，以及處理相關公共政策議題，以推動國際協議的落實。

► 補充建議與考量

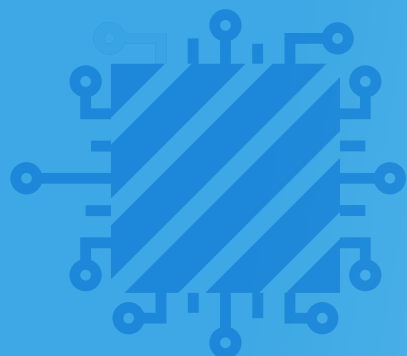
- 鼓勵並與AI開發者與研究人員合作，主力開發能直接因應國際協議與條約實施挑戰的AI解決方案。
- 在國會工作人員與議員間推廣AI素養與能力建構，以促進AI主導的有效行動，因應與上述目標有關的議題。因此AI相關專案可能需要另尋資金機會與夥伴支持。

國會應用AI的倫理課題

隨著國會應用AI指引持續發展，納入健全倫理思維勢在必行。為達成此目標，國會機關應找出內部的AI倫理提倡者，設立AI倫理研究經費，以促成行動方案。此外，亦須探討如何調整現有已獲認可的倫理架構，如聯合國教科文組織的架構⁴⁵，以應用於國會情境中。採納上述步驟可強化以符合倫理原則方式應用AI的決心，並在數位時代維護民主原則。



2.



通用人工智慧 與人類自主性

(能動性與真實性)

Artificial general
intelligence and
human autonomy

(agency and authenticity)

通用人工智慧 (AGI)，係指「整體智慧超越人類」的技術。⁴⁶ 若 AGI 得以成真，將具備前所未見的潛力，能輔助甚至取代人類認知能力。雖然國會運作極為複雜，但政治辯論與政策議題的細膩度已逐漸被 AI 學習，而 AGI 可能開啟輔助民主發展的新局。

國會應立即正視 AGI 與人類自主性之間複雜的關係，兼顧能動性 (agency) 與真實性 (authenticity)。⁴⁷ 現有科技已經威脅到了人類自主性。取決於 AGI 系統的設計方式及涵蓋領域，使用 AGI 系統可能會增強或削弱人類的自主性。

因此，在啟動任何 AGI 技術計畫之前，必須嚴格評估設計者與開發者的倫理責任。同時，最好先視 AGI 可能會成真，藉此克服恐懼、從錯誤中學習，並奠定成功經驗，開始發展。

2.1. 促進人類自主性，確保國會 AI 用途在於輔助高等認知能力而非取而代之

► 為何重要？

國會使用 AI 時，應促進人類自主性，AI 應為輔助工具而非替代方案，這對於維護民主原則與守護人類在治理中的判斷價值十分重要。國會機關可以身作則，在駕馭 AI 效益與保有人類決策與民主治理核心角色之間取得平衡。

► 是否有已知實例？

尚無已知實例。

► 如何落實？

國會可以 AI 為輔助與加強決策的工具，促進人類自主性。前提是建立明確的 AI 定位、訓練立法者，並制定以人類監督與倫理為本的 AI 應用指引。此外，優先採用以人為本的 AI 設計、確保健全的人類監督機制並推動透明決策機制，皆可加強落實民主原則之決心與對個人自主的保障。

► 補充建議與考量

- 某些國會部門，例如法律文本撰寫或資訊檢索領域，AI 很可能取代部分（例行性）人工作業。不用視這一轉變為對人類自主性的威脅，這反而可促成組織內人力資源重新分配。目前，複雜的細緻認知任務仍需仰賴人類自主性，而重複型工作則可由 AI 工具與服務有效處理。
- 有鑑於此，國會工作場域應建立一種負責任使用 AI 的文化，決策過程仍以人類判斷為核心。並須與 AI 倫理專家、學術機構與公民社會組織合作，確保應用方式與最佳實務保持一致。國會利害關係人亦須持續掌握 AI 技術發展，因其可能影響國會流程中的人類自主性與決策機制。

2.2. 針對國會 AI 設計與開發人員訂定特殊規範

► 為何重要？

國會 AI 系統的設計者與開發者在面對通用人工智慧（AGI）時，肩負特別的倫理責任，須防止濫用，並限縮強人工智慧（strong AI）或奇點（singularity）對機構、社會與公民的衝擊。這類責任與機構在執行徵才流程或委外廠商時所進行的標準審查一樣，需確保其價值觀與倫理觀與機構一致，方能維持協調的合作關係。因此，對參與國會 AGI 系統的設計者與開發者要進行倫理責任評估與安全背景調查，這一步驟能確保 AGI 系統以最高水準的誠信、問責與安全標準開發與維護。國會機關得以藉此確保 AGI 技術的開發符合其價值與倫理原則，同時保護國會程序的安全性。

► 是否有已知實例？

目前尚無針對 AGI 的明確規範。有鑑於專門的 AGI 指引付之闕如，立法機關可比照標準採購程序或勞務程序處理。

► 如何落實？

AGI 系統開發流程可添加一環倫理評估程序，確保其符合產業最佳實務與機構的倫理標準。國會可推動倫理指引、安全背景調查，設定嚴格的資格條件。設計者與開發者應展現倫理責任決心，包括：評估對社會的潛在影響、確保透明度、以及遵循最佳實務，守護機構與社會，以防範 AI 帶來的風險。

► 補充建議與考量

負責制定指引的技術工作團隊最終討論階段時發現，本條指引明顯不僅在落實方面可能極度困難，同時在各倫理方面也可能引發重大爭議。然而，此處依然收錄 AGI 指引，以支持未來發展 AGI 時，應優先以倫理責任為重的原則。

2.3. 促使大眾意識到 AI 技術仍在進步，以及 AGI 確實可能成真

► 為何重要？

AGI 有一天可能成真。不論哪一個國家，國會都是國家決策中樞，應未雨綢繆。任何國家都需要充分知情的決策能力，並為高等 AI 可能帶來的社會衝擊做好準備，才能建立負責任且有準備的治理體系，迎接高等 AI 帶來的潛在挑戰與機會。

► 是否有已知實例？

尚無已知實例。

► 如何落實？

國會可以推動持續教育、提升大眾意識，並與專家討論 AI 技術尚在進步的本質，來促進對 AI 的理解。可考慮設立專責任務小組落實前述工作。強調 AGI 未來可能會成真，才能促成積極規劃與納入倫理考量，以面對最後的發展。

補充建議與考量

- 持續掌握 AI 與 AGI 的進展，定期向國會工作人員與議員報告最新發展及其可能影響。
- 與智庫、研究機構與關注 AGI 的國際組織合作，以善用集體專業能力。

AGI 的考量

有鑑於目前針對 AGI 可行性與成真時間依然爭論不休，國會需主動把握先機，採取措施。儘管種種未知讓人無法確定 AGI 是否會成真，本指引採取未雨綢繆立場，必須評估其對民主制度的潛在影響並有所因應，才能為未來「奇點型」科技進展預作準備。

3.

AI 隱私與安全

AI privacy and security



監理國會 AI 應用時，由於 AI 系統經常處理敏感資料，例如個人資訊或國家安全資料，隱私與安全相當重要，若無適當的隱私與安全措施，可能導致資料外洩、身分盜用及其他有害後果，而破壞公眾對國會程序的信任，並不利民主制度。

資安與隱私（包括個人資料保護）應從設計階段即納入考量（by-design），因此，在國會採用 AI 之前，其模型訓練、微調與部署過程必須提供足夠保障。安全地處理個人可識別資訊（personally identifiable information，簡稱 PII）極為重要，由於國會程序具跨司法轄區與國際性質，亦須納入資料主權（data sovereignty）的考量（即資料受其蒐集或儲存地所在國法律規範之概念）。

過度依賴 AI 也是一種風險，僅能透過嚴謹的 AI 策略與應用治理架構加以因應。整體而言，在安全決策中，人類監督絕對是首要之務。

3.1. 國會 AI 系統嵌入安全性與強健的資安功能

► 為何重要？

國會 AI 系統嵌入安全性與穩健資安功能十分重要，可保護個人、內部網路與機構本身免受潛在危害與網路安全威脅。全面性的「設計納入資安」(security-by-design) 原則可提升立法機構 AI 系統的資安與安全。

► 是否有已知實例？

2020 年，歐洲理事會議員大會 (PACE) 通過一系列 AI 相關決議與建議，其中包括隱私與資安原則。⁴⁸

► 如何落實？

國會可要求嚴格測試、加密處理，以及遵守網路資安標準，確保 AI 系統的安全與資安。應建立持續監控機制、漏洞評估流程與應變協議，預防對個人的危害、保護內部網路，並防範潛在威脅與資料外洩。

► 補充建議與考量

- 在國會工作場域中培養資安意識文化，使人員保持警惕，主動辨識與通報安全疑慮。
- 設立專責資安團隊或單位，持續監控並加強 AI 系統的安全性。
- 與政府資安機構與專家合作，獲得有效保護國會 AI 系統的指導。
- 請注意，許多國會議員希望使用文字生成式 AI，以此類服務來說，應使用內部聊天機器人並設立存取權限限制，以避免機密資料意外洩露給未授權第三方。

3.2. 國會 AI 系統的開發設計納入隱私保護

► 為何重要？

國會 AI 系統的設計與部署納入隱私保護，可保護敏感資訊，並確保負責任使用 AI。其設計與部署應尊重個人隱私權，同時遵守資料保護法規，以促進負責任與符合倫理的 AI 應用。

► 是否有已知實例？

2020 年，歐洲理事會議員大會 (PACE) 通過一系列決議與建議，其中包括隱私與資料保護。⁴⁹

► 如何落實？

國會可採用強大資料加密、存取權限控管及定期資安稽核等措施來整合隱私保護措施。AI 設計應遵守「設計納入隱私」(privacy-by-design) 原則，遵守現行資料保護法規，以確保國會 AI 系統具備最高水準的隱私保障。

► 補充建議與考量

- 與隱私專家、法律專業人士及資料保護主管機關合作，確保符合隱私相關法規。
- 定期為國會工作人員與議員舉辦 AI 相關的隱私與資料保護訓練與宣導活動。
- 隨時掌握隱私威脅的發展趨勢，並因應調整 AI 系統與操作實務。

3.3. 確保 AI 系統處理的個人可識別資訊（PII）具備適當的保護措施

► 為何重要？

涉及 AI 系統時，保護個人可識別資訊（PII）相當重要，以保障個人隱私並遵守資料保護法規。本指引為前一條的延伸，具體探討個人資料的保護。

► 是否有已知實例？

2020 年，歐洲理事會議員大會（PACE）通過了一系列決議與建議，包括隱私與資料保護。⁵⁰

► 如何落實？

國會須知道 AI 系統會處理個人可識別資訊，因此必須建立嚴格的資料保護機制，包括使用強大加密措施、存取權限控管與稽核措施。此外，國會應建立內部與外部監督機制，以確保符合資料保護法規與倫理標準，例如導入自動化 PII 假名處理（pseudonymisation），以保護 AI 系統所處理的敏感資訊。

► 補充建議與考量

- 處理 PII 時，應遵守相關資料保護法律與規範，如《一般資料保護規則》（GDPR）⁵¹、《健康保險可攜與責任法案》（Health Insurance Portability and Accountability Act，簡稱 HIPAA）⁵²，或其他適用的區域性法規。
- 定期向工作人員與議員報告最新有關 PII 安全最佳實務與資料保護政策。
- 與隱私與資安專家合作，確保 AI 系統對 PII 的處理符合產業標準與最佳實務。

透過遵循上述步驟與考量，國會機關可為 AI 系統的運作建立健全的保護機制，守護 AI 系統使用的個人可識別資訊（PII），一方面確保資料安全與個人隱私，另一方面也符合資料保護法規的要求。

3.4. 瞭解外包 AI 系統中所儲存、處理與擷取的資料內容

► 為何重要？

國會將 AI 工具外包時，首要前提是全面了解該 AI 系統所儲存、處理與擷取的資料內容，更要特別注重與隱私、資料保護與機密有關的部分。

► 是否有已知實例？

原則方面，請參閱 PACE 2020 AI 關於隱私和資料保護內容。⁵³ 2023 年 6 月，美國眾議院在實驗使用 GPT-4.0 兩個月後，眾議院總行政官（Chief Administrative Officer，簡稱 CAO）呼籲國會辦公室限制商業大型語言模型（LLM）的使用，改回使用 ChatGPT，同時提供保護敏感資料的操作指引。⁵⁴

► 如何落實？

國會應要求廠商提供透明的資料處理操作實務、建立詳細的資料目錄，並執行嚴格的隱私風險評估。合約中應明確規定資料使用方式與保護機制，並要求第三方廠商遵守嚴格隱私與資安標準，以保護國會敏感資訊並確保符合法規。廠商亦須比照高維安需求產業對服務提供的嚴格標準執行。

► 補充建議與考量

- 聘請法律與隱私專家審閱與外包廠商之間的合約與協議，確保合約充分回應隱私與機密考量。
- 持續掌握資料保護法規的發展，並因應調整外包安排。

3.5. 處理資料時，認識並接受資料與基礎設施主權相關問題

► 為何重要？

處理資料時，特別是外包 AI 服務的情境下，認識並接受與資料及基礎設施主權有關的問題，非常重要，以確保符合資料保護法規並因應潛在法律與地緣政治風險。

► 是否有已知實例？

使用由某國公司開發並在該國運作的商業 AI 系統，可能對其他國家的國會構成風險。主要例子為美國 OpenAI Inc. 推出的 ChatGPT⁵⁵。雖然美國國會已在使用，但出於國安考量，其他國家的國會可能要考慮選擇開源或特定國家的模型，並在安全無虞且資安可靠的基礎設施環境中執行系統。

► 如何落實？

國會應進行影響評估、釐清資料所有權並建立明確的司法管轄規則，找出資料主權相關問題。協議書與政策必須明確說明資料如何處理，確保符合本地與國際法規，以建立國會應用 AI 時資料主權共識。最後，各國國會應研究是否可利用本國高效能運算（high-performance computing，簡稱 HPC）基礎設施來運作 AI 系統。

► 補充建議與考量

- 持續掌握資料主權法規以及可能影響資料處理安排的地緣政治發展。
 - 跨境傳輸資料時，應考慮使用加密與安全通訊協定，以降低資料被攔截或未授權存取的風險。
- 國會機關透過上述步驟與考量，能有效探討資料主權議題，確保資料處理符合法律要求，同時因應 AI 系統外包所帶來的跨境資料傳輸挑戰。在此脈絡中，也可另外進一步探討訓練與測試資料的主權概念。

3.6. 確保 AI 不得以合成資料取代原始資料來源

► 為何重要？

有鑒於國會 AI 不應以生成（亦即合成）內容取代原始資料，而是以有意義的內容補充國會語料庫，這是一項負責任且有效應用 AI 的基本準則。如此一來，可避免立法、程序與行政文件不會隨著時間流逝遭到竄改，可保護歷史與當代資料的正確度與完整度。AI 勒索軟體攻擊會加密並覆寫國會資料，就是一種應極力防範的典型高風險情境。

► 是否有已知實例？

並無已知實例。

► 如何落實？

國會必須確實知曉 AI 目的應在於補充原始資料，而非替代。國會應制定明確的指引與工作流程，以人類監督與決策主導，定位 AI 為資料分析與增強的輔助工具，以確保其作用為補充，而非取代國會工作成果。

補充建議與考量

- 在國會工作場域中培養負責任使用 AI 的文化，強調人類判斷與原始資料的重要性。
- 與 AI 倫理專家及組織合作，制定可加強 AI 作為輔助工具功能的指引及操作實務。

3.7. 正視過度依賴 AI 可能帶來的風險

► 為何重要？

國會正視過度依賴 AI 系統所帶來的風險非常重要，避免產生虛假安全感，並確保由人類判斷主導。國會機關應在善用 AI 效益與保持必要質疑態度之間取得平衡，以防止過度依賴及虛假安全感。

► 是否有已知實例？

目前已有多國國會機關工作使用大型語言模型（LLMs），研究已指出，國會運作過度依賴 AI 會讓人誤以為安全無虞，此為潛在風險，因此建議審慎使用，不應無條件信任 LLMs 及其產出結果，同時應正視 AI 可能產生幻覺與錯誤。⁵⁶

► 如何落實？

國會應保持警覺心，不過度依賴 AI，正視其可能造成自滿與虛假安全感，因此，必須持續重視人類參與及決策，積極管理 AI 系統，避免過分依賴 AI，而損害國會的誠信與效能。

補充建議與考量

- 在國會工作場域中培養批判思考的文化，人類智慧與 AI 系統並行，積極參與。
- 定期針對國會工作人員與議員進行問卷調查與意見回饋訪談，評估其對 AI 在國會工作流程所扮演的角色有何看法與感受。

3.8. 保障國會 AI 系統訓練與測試所使用之資料安全，防範遭受旨在操控系統互動方式的資安攻擊

► 為何重要？

保護國會 AI 系統的訓練資料十分重要，能防範敵人試圖操縱或重新訓練這些系統以達到惡意目的。此舉可維護 AI 在國會環境中所生成洞察與建議的完整性與可信度。

► 是否有已知實例？

已有多起針對國會系統的攻擊的詳細報導，但截至目前，尚無關於國會 AI 系統遭攻擊的紀錄或公開文件，國會也沒有公開 AI 的資料管理機制。

► 如何落實？

國會應採取強大的資安措施，包括加密技術與存取權限控管，以保護訓練資料不落入惡意人士之手。定期進行資安稽核、部署入侵偵測系統與嚴格的資料存取權限規範，有助於防止未經授權的重新訓練操作，確保國會 AI 系統互動的健全，並防範惡意竄改。

► 補充建議與考量

- 與資安專家合作，持續評估並加強訓練資料與 AI 系統的資安防護。
- 建立專門應變計畫，因應訓練資料外洩或資料安全事件。

3.9. 確保安全決策由人為進行

► 為何重要？

負責任應用 AI 的關鍵環節強調人類監督的重要性，並確保所有安全相關決策須上報給人類操作員，國會更是如此。

► 是否有已知實例？

歐洲理事會議員大會（PACE）於 2020 年所通過的決議與建議中，明確指出 AI 原則包括由人類負責決策。⁵⁷

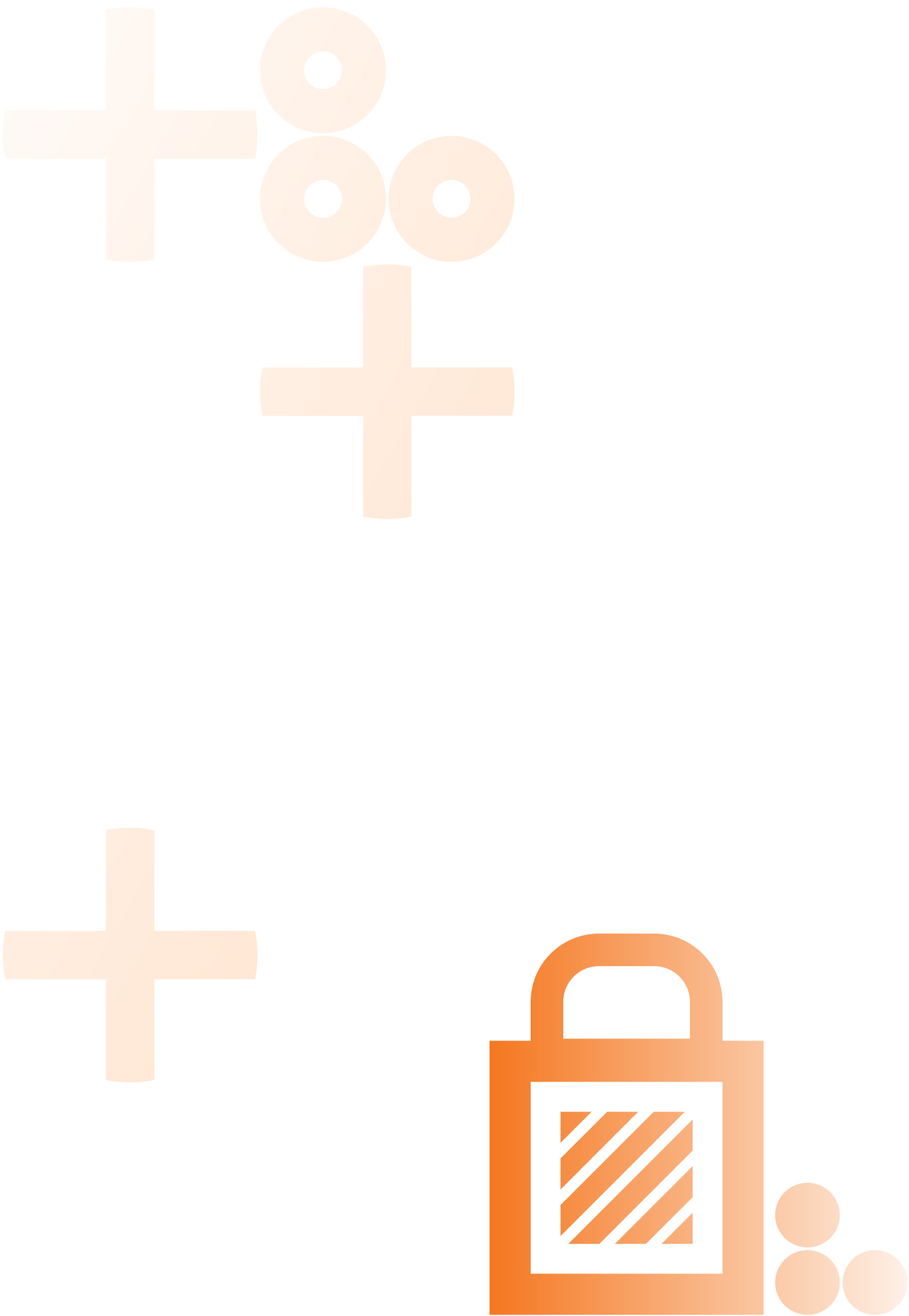
► 如何落實？

國會應制定程序與規則，要求安全相關決策上呈由人類操作員處理。AI 系統可協助進行威脅偵測，但應由人類作出關鍵資安判斷，才能確實問責，兼顧倫理考量，具備能力有效因應複雜且不斷演變的威脅。

► 補充建議與考量

- 在國會工作環境中建立警覺心與責任感的文化，鼓勵人類操作員積極投入 AI 系統，有必要時，質疑其輸出結果。
- 與資安專家與專業人員合作，加強人類監督角色，並提升 AI 資安防護措施。

遵循上述步驟與考量，國會機關可維持由人類監督主導 AI 資安機制，有效回應安全事件，並確保在國會應用 AI 時確實問責及保障安全。



4.

AI 治理與監督

AI governance and oversight



AI 系統的開發與部署應符合民主價值與程序。國會的監督角色可以賦予 AI 生成結果正當性，而有效的 AI 治理機制則有助於推動創新並促進公共利益。

AI 計畫迅速湧現，成熟水準不一，因此需要謹慎將其納入格局更大的國會數位策略之下。在眾多科技挑戰中，有效率的資料治理與管理規範需因應近期 AI 的普及應用而與時俱進。建立國會內部的 AI 倫理監督機制，也有助於確保策略與操作實務符合一致。治理團隊可以負責評估國會 AI 應用對各種操作實務的影響。此外，與 AI 利害關係人合作制定政策，也有助於確保國會成為變革的推動者與採用技術的領導者，充分發揮 AI 全面應用於社會的潛力。

4.1. 將 AI 系統的設計與導入納入更廣泛的數位國會發展策略之中

► 為何重要？

將 AI 系統的設計與實施納入格局更大的國會數位策略中，可確保 AI 有效輔助國會的政策目標與目的，同時與機關整體的數位轉型方向一致，進而提升效率、透明度與確實問責。

► 是否有已知實例？

2024 年 2 月，義大利眾議院（Italian Chamber of Deputies）紀錄活動監督委員會（Supervisory Committee on Documentation Activities）公布了一套運用 AI 支援國會業務的原則。⁵⁸ 巴西眾議院也已將國會 AI 系統納入其 2021–2024 年數位策略中。⁵⁹

► 如何落實？

國會可將 AI 納入數位策略的大框架之下，做法包括使應用 AI 的目的與整體國會目標方向一致，強調跨職能部門合作、確保系統可擴展規模，並調整 AI 成為現有數位計畫的輔助工具，如此一來 AI 就能成為國會數位生態系的重要一環。

► 補充建議與考量

- 邀請 AI 專家、數位策略規劃師與科技領袖提供整合流程建議與專業意見。
- 定期檢視並更新國會數位策略，確保其與 AI 領域發展趨勢保持一致。

4.2. 採用高效的資料治理與管理規範，確保 AI 系統使用資料的正確、完整與安全

► 為何重要？

國會程序必須伴隨一套有效率的資料治理與管理規範，以確保 AI 系統所使用資料的正確、完整與安全，進而促進國會工作流程應用 AI 時，過程透明、確實問責且卓有成效。

► 是否有已知實例？

資料治理本身已是一項成熟概念，⁶⁰ 但目前尚無制定出國會專用的具體資料治理架構，此外，國會 AI 系統管理資料的完整規範目前也付之闕如。

► 如何落實？

國會可建立嚴謹的資料治理與管理規範，以維護 AI 系統資料的準確、完整與安全。過程包括資料品質檢查、加密處理、存取權限控管、定期稽核，以及遵守資料保護法規，進而確保 AI 應用時資料的可靠與完整。

► 補充建議與考量

- 與資料治理專家與專業人士合作，設計並導入有效的資料治理規範。
- 與國會議員與工作人員交流，徵詢其對資料治理與管理實務的意見與回饋。在此情況下，可考慮採用「可查找（Findable）、可取得（Accessible）、可互通（Interoperable）與再利用（Reusable）」（簡稱 FAIR）資料（管理）原則與方法。⁶¹

4.3. 設立具權限的國會倫理監督單位，或將職責併入既有之審查國會 AI 系統與應用之監督委員會

► 為何重要？

設立具權限的國會 AI 倫理監督單位，或將此職責併入現有監督委員會，這種積極作為能確保國會程序負責任且合乎倫理地應用 AI。

► 是否有已知實例？

目前關於國會是否設有 AI 倫理監督機制的資訊依然不完善，可能是因為許多國家的國會對於該領域的專業認識有限，因此十分謹慎。

► 如何落實？

國會可設立專責倫理監督單位，或培力現有委員會來審查 AI 系統。此機構應由專家、立法者與相關利害關係人組成，確保 AI 應用能獲得透明審查，定期評估、遵循倫理指引與落實公共問責，這有助於國會場域以負責任、公平中立方式使用 AI。

► 補充建議與考量

- 鼓勵監督單位與專精 AI 倫理的國際組織或機構合作，掌握全球最佳實務。
- 公開監督單位的工作與成效，以建立國會在利害關係人與大眾間的信任感與公信力。

4.4. 監測 AI 對各種關鍵議題的影響

► 為何重要？

持續監測 AI 對各種議題的影響，例如智慧財產權、法律責任與問責、就業與勞動、社會經濟議題、隱私與資料保護、偏見與歧視、國家安全與國防、倫理治理與監督，以及環境問題，有助於理解其潛在影響，就是否推行做出合理判斷。

► 是否有已知實例？

多國國會系統經常進行「立法影響評估 (impact assessment)」，但 AI 影響的評估尚未落實。

► 如何落實？

國會可透過持續的研究、諮詢與影響評估，來檢視 AI 對上述議題的影響。與專家合作、邀請利害關係人參與、並定期檢討 AI 應用情形，能確保評估的全面性，使立法者得以調整政策與法規，因應各領域變化不斷的挑戰。

► 補充建議與考量

- 考慮成立國會專責委員會或任務小組，負責監督與統籌 AI 對各議題影響之評估工作。
- 制定一套全面的影響評估架構，包括標準化的方法學與回報機制。針對關鍵議題進行全面評估，國會機關就能全盤掌握 AI 的實質影響，據此做出充分知情的決策，以駕馭 AI 之益處，同時也降低潛在風險與挑戰。

4.5. 確保對國會 AI 系統所用資料的安全存取與權限管控

► 為何重要？

要確實問責、保護資料與資訊安全，保障 AI 系統使用資料的安全存取與權限管控，就十分重要，如此一來，國會也能監督 AI 系統決策過程。

► 是否有已知實例？

並無已知實例。

► 如何落實？

國會可實施嚴格資料管理規範、存取權限控管與加密措施，確保資料存取安全。同時，為達同樣目的，國會也可制定明確的資料共享政策、依必要性授權存取權限，並定期稽核資料使用，以在資料安全存取與使用控管之間取得平衡。

► 補充建議與考量

機構可聘用資料隱私與資安專家，或與其合作，設計並實施健全的資料存取權限與控管機制。同時應持續掌握資料保護法規的發展，以確保符合不斷變動的法律要求。

4.6. 與多元利害關係人合作，制定兼顧創新與保障人權的政策與法規

► 為何重要？

為發展具備韌性的政策與法規，就要與各界利害關係人合作打下基礎，對象包括其他國家國會、學界、公民社會及產業界，如此一來，在應用國會 AI 系統時，便能在促成創新及保障人權兩者間達成平衡。

► 是否有已知實例？

2017 年成立的希臘光學字元辨識（Optical Character Recognition，簡稱 OCR 團隊）⁶² 是一項全球的科學群眾外包合作計畫，促進全球代議機關、國會學者與專業人士間的合作。^{63 64}

► 如何落實？

國會可透過公開對話、協同工作小組及知識交流，與多元利害關係人合作，同時廣納來自學界、公民社會、產業界及國際國會網絡的意見，也有助於制定充分知情的政策，在開發及監理 AI 時，一方面可鼓勵創新，一方面又能維護人權與倫理原則。

► 補充建議與考量

- 定期檢討並更新 AI 政策與法規，以因應不斷變化的技術與社會需求。
- 公開草案、提案與影響評估供審閱與指教，促進政策制定過程的透明度。

AI 治理與監督的重點建議

AI 治理與監督的重要建議包括指派國會專責人員負責 AI 治理與法令遵循事務。此外，要建立 AI 透明資訊平台以加強確實問責與大眾信任，還要提供國會議員 AI 訓練，才能保障立法程序可做出知情決策，並有效應用 AI 科技。

5.

AI 系統設計與運作

AI system design and operation



設計與操作指引為在國會工作環境中導入 AI 提供了基本框架。這些指引強調規範 AI 系統使用、風險評估與影響監測的重要性，也指出系統的準確性及倫理考量的不可或缺，並納入所有相關利害關係人參與決策過程。

AI 計畫引出多項技術議題，對國會而言創新契機也伴隨著風險。為因應國會資訊的政治特質，導入標準化的資料架構與處理流程十分重要，AI 演算法也必須具備可解釋性 (explainability)，才能確保民選官員所做出的決策有一套透明可見的標準及依據。要打造健全可靠的 AI 系統，也需更加重視決策的可重現性 (reproducibility)，並能從成功案例中持續學習。國會亦可透過擔任先行使用者的角色，協助規範 AI 在其機構內部及整個社會中的使用與部署。對 AI 系統進行監測與評估時，也需建立開放式架構，以在取得同意的情況下，讓監督團隊能更有效接觸最終使用者。最後，國會利害關係人需就最低準確度、決策品質與機構績效等標準達成共識。

5.1. 建立標準化的資料架構與流程，以確保不同 AI 平台與程式彼此互通相容

► 為何重要？

要確保國會 AI 系統中不同平台與應用程式彼此互通與相容，必要措施就是導入標準化的資料架構與流程，而採用國際標準組織（International Organization for Standardization，簡稱 ISO）的標準最為理想。

► 是否有已知實例？

發展結構分明、經過驗證且為開放格式的資料集十分重要，並應優先採用標準化格式。像資訊交換標準格式「Akoma Ntoso」（AKN）這類法律標準，能促進法律資訊的調和與系統互通，具長期效益。⁶⁵ 目前，歐洲議會、義大利參議院、巴西參議院、烏拉圭國會、阿根廷眾議院、智利眾議院、英國相關機構與美國眾議院皆已常態使用 AKN。⁶⁶

► 如何落實？

國會可設立中央監理機構，負責制定並執行標準化的資料架構與流程，此監理機構應與技術專家合作，訂立明確指引，並要求所有 AI 平台與應用遵循指引，以促進彼此互通與相容，同時符合資料安全與隱私標準。

► 補充建議與考量

- 定期檢討並更新標準化資料架構與流程，以因應不斷變化的資料需求及技術發展。
- 廣納資料管理與互通性領域專家及利害關係人意見，持續精進標準化作業。

5.2. 強調演算法的可解釋性

► 為何重要？

重視國會應用 AI 時，其演算法的可解釋性，能確保所有根據 AI 而來的決策與建議，背後思維對利害關係人而言是清楚、可理解的。這對於建立信任、促進理解與透明度十分重要，亦有助應用國會 AI 系統做出知情決策。

► 是否有已知實例？

並無已知實例。

► 如何落實？

國會可要求開發者採用可解釋的演算法，強制使用透明的 AI 系統，包括使用可解讀的模型、提供可理解的文件，並建立監督機制以確實問責。

補充建議與考量

- 清楚溝通說明 AI 演算法的侷限，以管理期待，避免產生誤解。
- 制定 AI 演算法說明的標準格式或指引，確保內容一致與清晰。

5.3. 建置穩健可靠的國會 AI 系統，並具備偵測與修正錯誤的能力

► 為何重要？

建立具備錯誤偵測與修正能力的健全可靠國會 AI 系統，有助於維護系統的完整性與效能。

► 是否有已知實例？

並無已知實例。

► 如何落實？

國會可強制執行嚴格測試、持續監測與建立故障防護機制，以確保 AI 系統的健全。定期稽核、回饋機制與專責監督單位能即時發現與修正錯誤，提升系統可靠度並維持國會 AI 系統的健全。

► 補充建議與考量

- 定期進行系統稽核與執行後檢討，以找出系統穩定性與錯誤處理機制上的改善空間。
- 與軟體工程與穩定性工程領域的專家合作，確保遵循最佳操作範例。

5.4. 規範國會 AI 系統的使用與部署，包括風險評估、授權要求與安全標準

► 為何重要？

透過具法律約束力與非約束力的工具，規範國會應用與部署 AI 系統非常重要，可確保國會程序以責任且合乎倫理的方式採用 AI 技術。

► 是否有已知實例？

截至 2024 年中，尚無發現具體的規範措施。本指引可作為國會制定相關措施時的參考基礎。

► 如何落實？

國會可建立完整架構、進行風險評估及明訂授權要求來規範 AI 系統，也可執行安全標準、定期稽核與專家合作，確保 AI 系統以負責任方式部署，為國會 AI 系統建立安全且確實問責的使用環境。

► 補充建議與考量

- 與 AI 倫理、法律及科技監理領域的專家合作，確保監理架構完整並與時俱進。
- 在制定與修正 AI 系統規範時，廣泛徵詢公民社會、學術界與產業界等相關利害關係人的看法與意見回饋。

5.5. 對含有 AI 功能的雲端服務或軟體即服務（SaaS）產品進行風險評估

► 為何重要？

使用包含 AI 功能的 SaaS 產品或雲端執行 AI 時，要進行風險評估，以確保符合倫理考量與其他保障措施。

► 是否有已知實例？

目前尚無國會架構案例，特別針對 AI 國會系統進行風險評估。廣義來說，歐盟《AI 法案》已針對特定功能強大且具影響力的系統訂出相應條文。

► 如何落實？

國會可針對 SaaS 或雲端 AI 服務進行完整風險評估，內容包括資料隱私、安全措施、倫理考量與廠商透明度。建立監管指引、認證制度與持續監測機制，可在部署此類技術時確保符合倫理與維持完整保障。

► 補充建議與考量

- 與 AI 倫理與負責任 AI 領域的專家合作，進行全面風險評估，確保已納入符合倫理的保障措施。
- 鼓勵與廠商開誠布公對話，針對已知風險或疑慮進行討論，並爭取其承諾負責任使用 AI。

5.6. 監測與評估國會 AI 系統的運作與輸出內容

► 為何重要？

定期且按部就班地監測與評估國會 AI 系統為必要之舉，才能準確評估其對國會流程與工作結果的影響。持續監測內部 AI 系統的輸出，才能做出知情決策，有能力調整法規，讓國會能以負責任的方式部屬 AI。此舉也能加強各方對該工具的信任，促使國會議員與行政人員多加使用。

► 是否有已知實例？

並無已知實例。

► 如何落實？

國會可設立監督委員會，或與外部專家合作，定期進行公正評估。此外，亦可分配資源與人力執行此類評估工作。

► 補充建議與考量

國會機關應持續監測並評估其 AI 系統的運作與輸出內容。此種積極作法有助於持續改善、負責地使用 AI，並確保其與更廣泛的社會目標保持一致。

5.7. 在採用 AI 系統前，與所有利害關係人達成對於最低準確度的共識

► 為何重要？

各別國會 AI 系統所需的準確程度，取決於具體應用情境與用途而有所不同。要確保 AI 系統達成預期目標並有效運用，關鍵在於相關利害關係人就執行 AI 系統的最低準確度達成共識。

► 是否有已知實例？

並無已知實例。

► 如何落實？

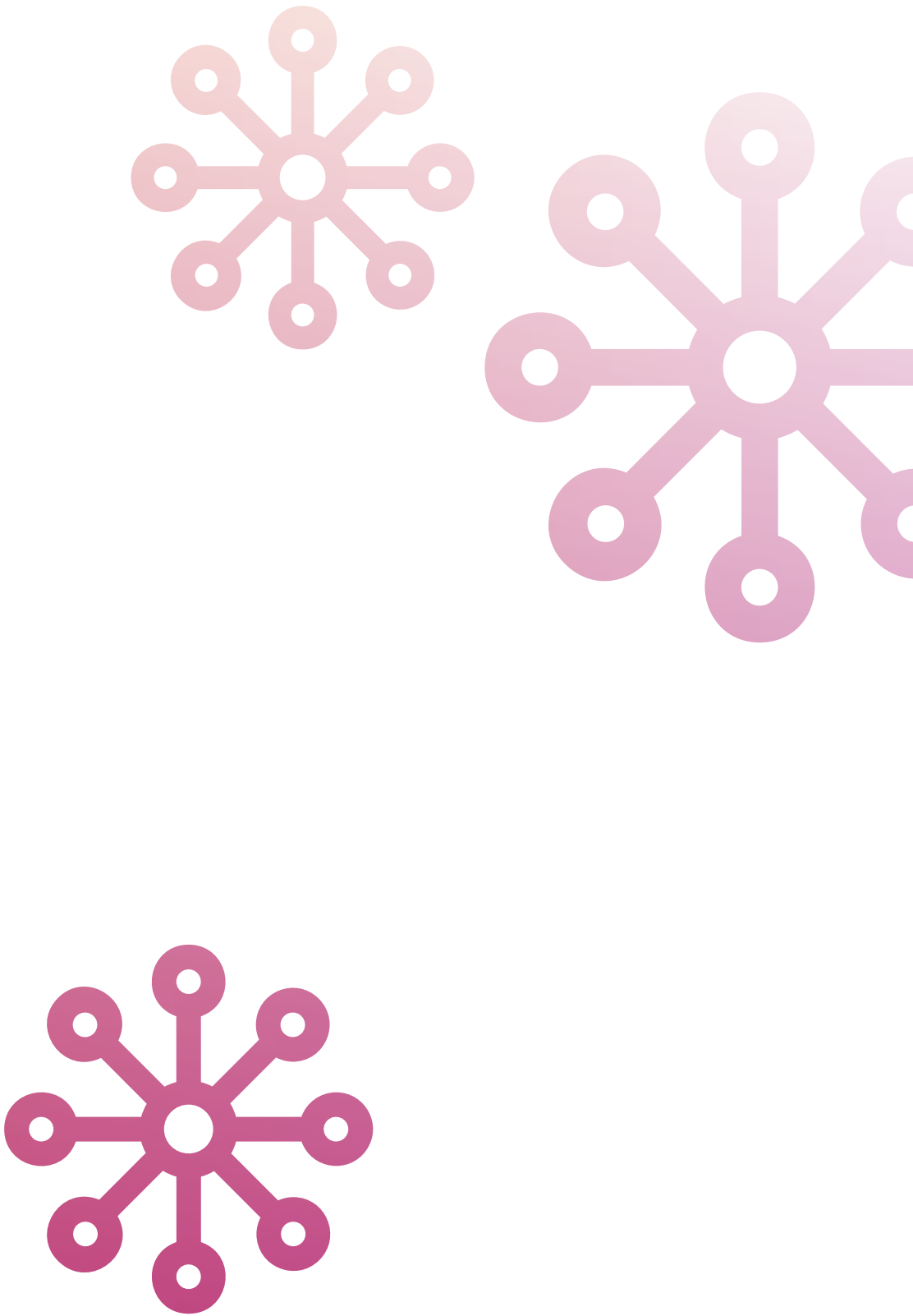
國會可設定績效基準、進行獨立評估，並徵詢多元利害關係人意見，達成最低準確度共識。執行嚴格測試、保持驗證過程透明，並主動蒐集意見回饋，促成知情決策，以執行 AI 系統，可於利害關係人間建立可靠信任。

► 補充建議與考量

- 在設定準確度目標時，應考量誤判為真（false positives）與誤判為假（false negatives）可能導致的後果，因為在不同的使用情境中，它們所造成的影響可能大不相同。
- 全程鼓勵與利害關係人進行開放透明的溝通，以建立信任並確定準確度目標共識相同。

沙盒（sandbox）環境與創新實驗室

開發國會科技時，於 AI 沙盒環境中作業有助於在控制情境下進行實驗，既可以探索 AI 應用，又不會有中斷運作的風險。此外，設立創新實驗室可提供專門空間，促進協作解決問題與 AI 解決方案的更新開發，回應國會的個別需求。這類作法可加強系統的敏捷力與創新力，同時確保 AI 技術順利整合至國會系統的設計與運作之中。



6.

AI 培力與教育 AI capacity building and education



國會工作場域中導入 AI 時，進行培力與訓練非常重要，因為這能協助立法委員與國會工作人員建立必要的知識與技能，對有效且負責任地運用 AI 再重要不過。培力與教育訓練包含認識 AI 技術、其潛在應用及對社會的影響，同時也有倫理與法律議題需要考量。國會投入培力與教育訓練，就能獲得足夠能力，探索 AI 帶來的機會與挑戰。培力與教育訓練也包含提供國會議員與國會工作人員所需資源，使其能就 AI 技術及國會應用與大眾溝通互動。

在開放環境邀請各種利害關係人組成專家團隊，有助於學習交流與傳遞最佳實例。組織 AI 主題的訓練課程，對於推動國會內外部的培力與教育會越來越重要。AI 熱潮席捲全球，國會可善用此一趨勢，交流知識並與社會各界合作，教育公眾 AI 在國會的用途與其侷限也有助於維持觀感與期待。

6.1. 建立專責專家團隊，以掌握 AI 及其他領域的技術創新趨勢

► 為何重要？

建立並擴大專家團隊以掌握 AI 及其他技術領域的創新發展，有助於國會機關掌握最新資訊，做出知情決策，並有效運用 AI 所帶來的效益。

► 是否有已知實例？

此類任務可由前瞻型機構負責。例如芬蘭國會設有「未來委員會」，實質上相當於國會內部智庫。⁶⁷該委員會於 2021 年舉辦過一場創新的 AI 系統國會聽證會。⁶⁸

► 如何落實？

國會可投入長期的培訓、與外部專家合作，以及與教育機構及 AI 產業建立夥伴關係等方式，建立專家團隊。定期更新知識、進行跨領域人才招募，以及培養創新文化，能讓團隊持續掌握 AI 技術演變趨勢。

► 補充建議與考量

- 鼓勵團隊成員發表研究論文、報告或文章，促進 AI 領域的整體知識體系發展。
- 在團隊內部培養創新文化，鼓勵針對國會程序的 AI 應用進行實驗與創意探索。
- 研究如何與專業知識機構建立關係，例如，透過研究部門的專業人員以及與處理 AI 問題的委員會聯繫。

6.2. 為國會官員及行政人員定期舉辦 AI 訓練課程

► 為何重要？

經常為國會官員與行政人員舉辦 AI 訓練課程，有助於培養關鍵 AI 素養，並促進國會機關以負責任與合乎倫理方式應用 AI。

► 是否有已知實例？

相關訓練課程案例包括美國參議員的「洞察 AI 論壇 (AI Insight Forum)」⁶⁹，以及國際國會聯盟 (Inter-Parliamentary Union，簡稱 IPU)「國會創新中心」(Centre of Innovation in Parliament，簡稱 CIP) 所舉辦的線上研討會。⁷⁰

► 如何落實？

國會可與教育機構及產業合作、舉辦實體工作坊，並開發線上學習單元，來舉辦 AI 訓練課程。課程強調倫理考量、資料隱私以及批判思考力的養成，使得國會官員與行政人員能培養基本 AI 素養，藉此推動負責任及符合倫理的 AI 應用之道。

► 補充建議與考量

- 善用數位學習平台與資源，促進遠距或自主學習機會，並鼓勵各國國會間溝通與分享經驗。
- 在國會內部設立 AI 學院或卓越中心 (Center of Excellence)，以培養技術專業並促成合作。
- 可採用線上學習平台⁷¹與大型開放式線上課程 (Maasive Open Online Courses，簡稱 MOOCs) 來提供普及資源，以持續培養技能，確保國會議員及工作人員能及時即時掌握新興功能、風險與危害，採取「觀測站」式的長期策略。
- 鼓勵課程參與者在所屬團隊或部門內分享 AI 知識與洞察觀點，推動知識分享的文化。
- 可考慮提供國會議員工具與資源，使其能向公眾進行 AI 議題教育，提升國會應用 AI 流程的透明度，並促進大眾理解。

6.3. 支持與外部利害關係人進行知識交流，並參與雙邊及多邊合作機制

► 為何重要？

支持與外部利害關係人進行知識交流，並參與雙邊與多邊合作機制，是國會能持續掌握新興科技領域最新資訊、促進合作、並善用專業知識的關鍵策略。這些行動有助於深化對 AI 與其他新興科技的理解，並促進國會在推動負責任與符合倫理的 AI 治理過程中發揮積極作用。

► 是否有已知實例？

推動 AI 議題知識交流的實例如下：全球人工智慧夥伴關係（Global Partnership on AI，簡稱 GPAI）⁷²、聯合國「人工智慧造福人類」倡議（AI for Good）⁷³，以及狹義一點來說，國際國會聯盟 IPU 的國會創新中心（CIP）即透過跨國會方式促進人工智慧議題的知識交流。

► 如何落實？

國會可設立論壇、建立夥伴關係，並啟動與外部利害關係人的合作案，推動知識交流。

我們期望國會參與者彼此分享資訊，或諮詢專家瞭解適合的提示詞（prompt）設計與方法，積極參與雙邊與多邊合作機制，可促進資訊共享、技術進展與政策協調同步，推動全球掌握資訊，相互串聯來因應國會挑戰，包括 AI 相關議題。

► 補充建議與考量

廣納多元背景及地區利害關係人，促進共融，確保 AI 治理與倫理議題具備豐富觀點，並鼓勵積極參與國際 AI 治理倡議。

6.4. 紀錄 AI 相關活動的推動歷程與成果

► 為何重要？

紀錄 AI 相關工作的推動歷程與成果，可建立機構記憶庫，在國會機關內部傳遞知識。

► 是否有已知實例？

在 AI 相關活動方面，希臘國會的「科學文獻紀錄與科學服務監督部門」（Department of Scientific Documentation and Supervision of the Scientific Service）曾以機構立場公開表達其致力於加強機構記憶庫及與內部利害關係人共享知識的決心。巴西眾議院則開發了「Caggle」這套協作型數位平台⁷⁴，用於記錄、分析並分享各項資料主導的專案與實驗。該工具可促進議員間有效協作，詳細記錄 AI 相關活動的洞察分析與成果，並可立即作為組織持續學習與發展之用。

► 如何落實？

國會可建立詳實紀錄、制定標準化報告架構，並導入知識管理系統，來紀錄 AI 活動。定期向國會內部成員傳達最新進展，可為國會的 AI 工作累積機構記憶庫，加強透明度及知情決策。

► 補充建議與考量

- 可考慮導入文件管理軟體或知識管理平台，以提升 AI 相關文件紀錄的高效率儲存、檢索與分享。
- 鼓勵員工積極參與紀錄工作並表揚其對機構記憶的貢獻。

6.5. 以易於理解的方式向大眾說明國會 AI 系統的應用與侷限

► 為何重要？

以通俗易懂的方式向大眾說明 AI 在國會中的應用與侷限，對於提升機構透明度與大眾信賴十分重要。公民能深入了解立法與監督過程，進而對於機構以負責任及確實問責的方式部署 AI 的承諾產生信心。向大眾說明 AI 在國會中的應用與限制，也能提升 AI 相關事務的透明度、落實問責並促進公民參與。

► 是否有已知實例？

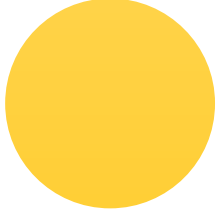
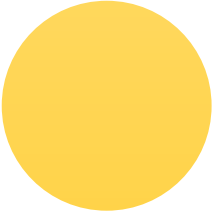
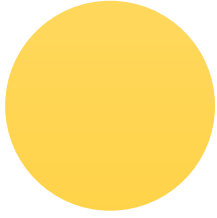
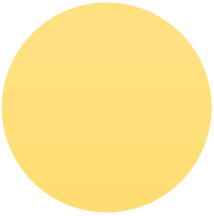
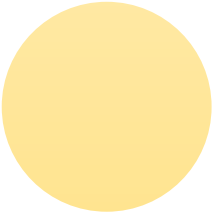
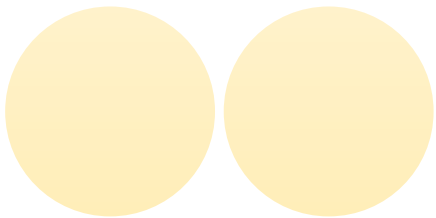
美國眾議院的「現代化委員會專門小組」(Modernization Subcommittee) 已開始定期發布「快報」，列舉國會支援單位已採用或打算採用的 AI 系統。⁷⁵

► 如何落實？

國會可以推動提升意識活動、舉行公開論壇，並建立使用者友善資源，向大眾說明。透明溝通、淺白語言說明以及各種不同媒體管道有助於以平易近人的方式傳遞國會 AI 系統的用途與侷限。

► 補充建議與考量

國會針對 AI 對社會、經濟、政治等影響的整體推廣及公共參與計畫，也可包含向大眾說明 AI 導入國會的用途與侷限。在此架構下，國會（例如由新聞處負責）可定期檢討與更新對外溝通資料，反映 AI 使用狀況的最新變化。也可強調國會以負責任與合乎倫理方式使用 AI 的決心，以建立社會大眾對國會應用 AI 的信任。



Part 3.

前行之路

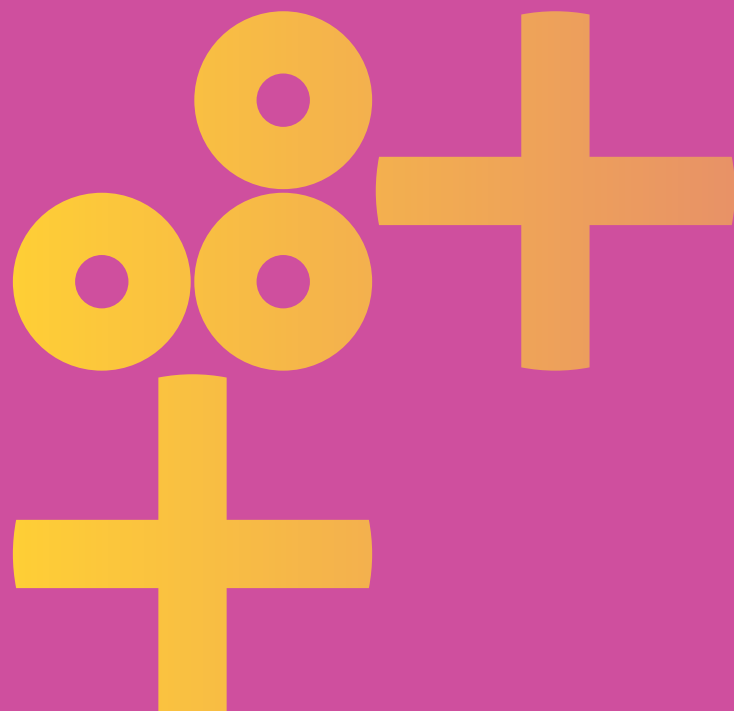
The way forward

不論哪一國國會，在實際推動這套指引時，需經歷數個獨特的關鍵步驟，其中必定包括開啟內部討論，或許還要公開討論規範的範圍、優先順序及初衷。接著，應思考並採取去制定策略、釐清優先事項與執行技術，同時兼顧治理面向。這些步驟因為涉及改變現行程序與流程以整合 AI 技術，因此可能也需要改變組織文化。

有鑒於國會系統中 AI 技術不斷演變的趨勢，在公布本指引後，批判式反思能為討論及未來版本鋪路。透過協作式的 SWOT 分析，將有助於深入瞭解本指引及 AI 技術的優勢、弱點、機會與風險。

或許將國會應用 AI 指引擴大成「持續調整的文件」(living document) 也很有價值，例如透過管理式線上平台進行持續更新。此舉的優勢在於可以納入真實世界經驗與 AI 發展，而持續演化與調整。

指引的最終目標，是能轉為被廣泛接受的標準與規範，並確立國會導入 AI 所需遵循的負責任基準。現有工作團隊將持續發展，一方面回應當前實際問題，另一方面為長期挑戰做好準備。團隊的核心承諾始終如一，即是打造不僅能因應當前 AI 發展情勢，亦能促進符合倫理、共融與透明的國會治理，這樣子的一套指引。



Part 4.

延伸閱讀 Useful reading

詞彙表

Artificial general intelligence，通用人工智慧 (AGI)：一種能以類似人類智慧的方式理解、學習並應用知識的人工智慧。不同於針對特定任務設計的專門 AI 系統，AGI 目標為掌握廣泛類型的認知能力，可執行多種任務並適應不同情境，而無需針對每項任務個別撰寫程式。簡言之，AGI 追求模仿人類心智的靈活變通與解決問題能力，最終可能發展出能在多個領域中像人類一樣思考、推理並解決問題的機器。從長遠來看，AGI 在各種認知任務上的表現可與人類相當，甚至更佳。

Artificial intelligence，人工智慧 (AI)：運用電腦運算能力，複製人類智慧，獨立或根據指令執行特定任務的技術、學習方法、系統架構、演算法與方法。例如：自主系統、機器學習、深度學習、類神經網絡、圖形識別、自然語言處理、即時翻譯、聊天機器人與機器人。AI 能力之目的在於支援或自動化人類的活動與流程。

AI System，AI 系統：人工智慧 (AI) 技術的電腦系統或軟體應用程式，執行通常需要人類智慧才能處理的任務。AI 系統的設計目的是模擬或複製人類的認知能力，例如學習、推理、解決問題、感知與語言理解，因此能分析資料、做出決策，並自主或在最少人為介入下採取行動。

Autonomous AI systems，自主 AI 系統：能夠感知環境、下決策並獨立行動的智慧代理系統，無需持續的人為監督或介入。這類系統仰賴先進演算法、機器學習技術與資料輸入，分析複雜情境、因應變化，並優化表現以達成預設目標。自主 AI 系統的案例包括自駕車、機器人流程自動化與智慧個人助理。發展自主 AI 系統的目的是打造能在現實環境中高效且有效運作的技術，有潛力藉由自動化任務並根據大量資料做出判斷，在不同產業掀起革新，並改善人類生活。

Bias detection，偏見偵測：偏見是指有意或無意的傾向，會使判斷或決策偏向某一方向。在人工智慧中，當演算法因資料不完整、假設錯誤、設計或訓練過程內建先入為主的觀念等種種因素，無意間偏袒、歧視特定群體或結果，就可能產生偏見。偵測並修正 AI 系統中的偏見，對於確保其輸出結果公平、平等、正確，避免強化社會不平等十分重要。

Explainable AI，可解釋的 AI (XAI)：AI 系統就其決策與行動，提供可讓人理解原因的能力。XAI 的目的是提升 AI 系統透明度，可以解釋說明，使人類能夠理解 AI 如何及為何做出某項特定決策。這一能力在國會這種情境下特別重要，因為 AI 決策可能對個人或整體社會造成重大影響。

Fairness，公平性：AI 公平性是一項關鍵原則，確保人工智慧系統平等對待所有個人與族群，避免基於種族、性別、年齡或社經地位等原因產生偏見與歧視。為達到 AI 公平性，須考量不歧視、機會平等、結果與代表性的公平性、透明度與落實問責。達成 AI 公平性是一項持續挑戰，AI 開發流程的每一步皆需審慎考量。在公平與其他目標之間取得恰當平衡，是在社會中建立信任並推動負責任使用 AI 的關鍵。

Fundamental rights impact assessment，基本權利影響評估 (FRIA)：FRIA 是一項工具，目的為處理高等 AI 系統可能帶來的潛在危險，不僅僅只是遵守《歐盟 AI 法案》這類法規的規定而已。相較於《AI 法案》強調技術要求並確定 AI 系統符合特定標準，FRIA 關注這些系統對人民基本權利的影響，以及對隱私、言論自由和平等權等權利的衝擊。

Generative AI，生成式 AI (GenAI)：生成式 AI 是一種能以既有學習內容為基礎，而產生新內容的技術。這項技術根據已識別與學習到的模式產出合成資料。大型語言模型 (large language models，簡稱 LLMs) 支援生成文本，而 AI 翻譯服務則能將文本轉換成其他語言。其他應用包含生成簡報、程式與流程的生成。文字也可用來生成具不同音調的語音或聲音片段。圖像與影片的生成亦日益重要，特別是根據圖像與語音錄音產出對嘴影片。

Human autonomy，人類自主性：人類自主性是指個人在無外力影響或脅迫情況下，能獨立做出選擇與決策的能力。它是人類尊嚴與自決能力的根本，使人得以依照自身價值觀與偏好掌握人生、追求目標與利益。自主性涵蓋決策、自由、自我治理與尊重權利等層面，是民主社會的基石，使個人的權利和自由得到保障及尊重。以人工智慧與自動化而言，維護人類自主性是關鍵考量，須確保技術系統設計與部署方式能培力個人、尊重其權利及選擇，並提升個人經營有意義的自主人生的能力。

Human-centred AI，以人為本的 AI：以人為本的 AI 目標為提升人類能力，回應人類社會需求，並以人類為靈感來源。其研究並為人類打造有效的夥伴和工具，例如給長者的機器人助手和陪伴者。應用以人為本的 AI 對國會也非常重要，才能確保 AI 系統的決策流程以人類福祉、民主價值、社會需求為重。

Hybrid AI，混合式 AI：混合式 AI 結合規則基礎系統與統計學習方法，以這種方式打造調適力及能力皆更為強大的 AI。

Intellectual property，智慧財產權 (IP)：智慧財產指的是心智創作產物，如發明、文學與藝術作品、設計、符號、名稱與圖像，這些均受法律保護。智慧財產權賦予創作者或所有權人在一定期間內對其創作的獨家使用與控制權。智慧財產權類型包括專利權、著作權、商標權、營業祕密與設計權。智慧財產權對於促進創新、創意與經濟成長十分重要，因為智慧財產權給予個人與組織投資研發的誘因。若 AI 服務提供者在未經許可的情況下，以受智慧財產權保護的資料訓練語言模型，即屬侵權。國會不應使用此類 AI 服務。

Natural language processing，自然語言處理 (NLP)：自然語言處理是人工智慧的一個分支，專攻電腦與人類語言之間的互動。此技術涉及開發能讓電腦理解、解析並生成人類語言（文字或語音形式）的演算法與模型。NLP 涵蓋各種任務，例如情緒分析、機器翻譯、專有名詞辨識 (named entity recognition)、文本摘要與問答等等。NLP 的目標為縮短人類溝通與電腦理解之間的差距，使電腦能處理與分析大量非結構化語言資料，並促進人機之間更自然且高效率互動。

Singularity，奇點：一個假想的未來時間點，在這個時間點，AI 的智慧超越人類智慧，因此引發快速科技成長及文明根本的變化。這樣的結果會孕育出獨立的超級智慧存在，使爆發式的進展無法逆轉。可能創造一個新的超人類時代，人類與高等 AI 實體的互動越來越深。這些機器將做出什麼行為，關鍵取決於其當初所設定的目標與價值觀。在這樣的未來情境下，國會將發揮關鍵重要性，因其負責處理此類先進科技所帶來的倫理與社會影響。

Training data，訓練資料：用來訓練演算法或機器學習模型的資料，是 AI 系統得以發展的基礎。訓練資料必須為人類產出，來自其工作或過去成果的資料。品質越好，AI 系統輸出的準確度就越高。公部門，包括國會在內，會需要統一的資料管理方式，以利 AI 系統的應用。同時要注意訓練資料可能內含偏見或是受到智慧財產權保護。

縮寫一覽表

AGI	Artificial general intelligence，通用人工智慧
AI	Artificial intelligence，人工智慧
AKN	Akoma Ntoso
CAI	Committee on Artificial Intelligence，人工智慧委員會
CAO	Chief Administrative Officer，總行政官
CHA	Committee on House Administration，美國國會院務管理委員會
CIP	Centre of Innovation in Parliament，國會創新中心
DPO	Data protection officer，資料保護官
FAIR data	Findable, accessible, interoperable, and reusable data，可查找、可取得、可互通與再利用資料
FRIA	Fundamental rights impact assessment，基本權利影響評估
GDPR	General Data Protection Regulation，《一般資料保護規則》
GenAI	Generative artificial intelligence，生成式人工智慧
GPAI	Global Partnership on AI，全球人工智慧夥伴關係
GPT	Generative pre-trained transformer，生成式預訓練轉換器
HIPAA	Health Insurance Portability and Accountability Act，《健康保險可攜與責任法案》
HPC	High-performance computing，高效能運算
International IDEA	International Institute for Democracy and Electoral Assistance，國際民主及選舉協助研究所
IP	Intellectual property，智慧財產
IPU	Inter-Parliamentary Union，國際國會聯盟
ISO	International Organization for Standardization，國際標準組織
LLM	Large language model，大型語言模型
MOOC	Massive Open Online Course，大規模開放式線上課程
MP	Member of parliament，國會議員
NLP	Natural language processing，自然語言處理
OCR	Optical character recognition，光學字元辨識
PACE	Parliamentary Assembly of the Council of Europe，歐洲理事會議員大會
PII	Personally identifiable information，個人可識別資訊
SDG	Sustainable Development Goal，永續發展目標
SWOT	Strengths, weaknesses, opportunities, and threats，優勢、弱點、機會與威脅
UNDP	United Nations Development Programme，聯合國開發總署
WFD	Westminster Foundation for Democracy，西敏寺民主基金會
XAI	Explainable artificial intelligence，可解釋的人工智慧

參考資料

AI for Good initiative, <https://aiforgood.itu.int>

Akoma Ntoso use cases, http://www.akomantoso.org/?page_id=275

Akoma Ntoso Version 1.0.Part 1: XML Vocabulary, <http://docs.oasis-open.org/legaldocml/akn-core/v1.0/akn-core-v1.0-part1-vocabulary.html>

Altman, S. (2023) Planning for AGI and beyond. <https://openai.com/index/planning-for-agi-and-beyond/>

Aul Tremblay, et al. v. OpenAI, Inc., et al.(Case No. 3:23-cv-03223-AMO).

Brazilian Congress, Chamber of Deputies, Digital Strategy 2021-2024, <https://www2.camara.leg.br/a-camara/estruturaadm/gestao-na-camara-dos-deputados/gestao-estrategica-na-camara-dos-deputados/gestao-estrategica-de-tic/estrategia-digital>

Bresciani, P.F., & Palmirani, M. (2024). Constitutional Opportunities and Risks of AI in the law-making process.FEDERALISMI.IT 2, 1-18. <https://hdl.handle.net/11585/953858>

Committee on House Administration (2024). Flash report - Artificial Intelligence Strategy ; Implementation. https://cha.house.gov/_cache/files/a/d/ad4d1279-c8f8-439b-9e3b-a95b01d61d03/56078B0226EDF1EAF76D863A2E7765A5.cha-q1-flash-report.pdf

Council of Europe: Council of Europe's Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law (AI Convention).

Dai, X., & Norton, P. (Eds.).(2013).The Internet and parliamentary democracy in Europe: A comparative study of the ethics of political communication in the digital age. Routledge.

Data Ethics Commission (2019).Opinion of the Data Ethics Commission.Berlin. <https://www.bmi.bund.de/SharedDocs/downloads/EN/themen/it-digital-policy/>

datenethikkommission-abschlussgutachten-lang.pdf?blob=publicationFile;v=5

European Data Protection Board (2023). Guidelines 01/2022 on data subject rights - Right of access. https://edpb.europa.eu/system/files/2023-04/edpb_guidelines_202201_data_subject_rights_access_v2_en.pdf

European Parliament (2024).Artificial Intelligence Act. https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_EN.html

European Parliament resolution of 19 May 2021 on artificial intelligence in education, culture and the audiovisual sector (2020/2017(INI)), https://www.europarl.europa.eu/doceo/document/TA-9-2021-0238_EN.html

European Parliament resolution of 20 May 2021 on shaping the digital future of Europe: removing barriers to the functioning of the digital single market and improving the use of AI for European consumers (2020/2216(INI)), https://www.europarl.europa.eu/doceo/document/TA-9-2021-0261_EN.html

European Union Agency for Cybersecurity, Milenkovic, G., & Dekker, M. (2020).Guideline on security measures under the EEECC, Publications Office. <https://data.europa.eu/doi/10.2824/44013>

Fitsilis, F. (2019).Imposing regulation on advanced algorithms.Cham: Springer.

Fitsilis, F. (2021). Artificial Intelligence (AI) in parliaments–preliminary analysis of the Eduskunta experiment. The Journal of Legislative Studies, 27(4), 621-633. <https://doi.org/10.1080/13572334.2021.1976947>

Fitsilis, F. (2024). The parliamentary perspective of better regulation in Greece. GRNET Tech Day in Digital Ready Policy Making, 1 February 2024. https://events.grnet.gr/event/138/attachments/393/669/Fotis%20Fitsilis_The%20parliamentary%20perspective%20of%20better%20regulation%20in%20Greece.pdf

Fitsilis, F., & Costa, O. (2023). Parliamentary administration facing the digital challenge. In The Routledge Handbook of Parliamentary Administrations (pp. 105-120). Routledge.

Fitsilis, F., & de Almeida, P.G.R. (2024). Artificial Intelligence and its Regulation in Representative Institutions. In: Research Handbook on Public Management and AI. Edward Elgar: Cheltenham. <https://www.e-elgar.com/shop/gbp/research-handbook-on-public-management-and-artificial-intelligence-9781802207330.html>

Fitsilis, F., & Mikros, G. (2021). Development and Validation of a Corpus of Written Parliamentary Questions in the Hellenic Parliament. Journal of Open Humanities Data, 7, p.18. <https://doi.org/10.5334/johd.45>

Fitsilis, F., Koryzis, D., & Scheffbeck, G. (2022). Legal informatics tools for evidence-based policy creation in parliaments. International Journal of Parliamentary Studies, 2(1), 5-29.

Fitsilis, F., von Lucke, J., Mikros, G., Ruckert, J., Alberto de Oliveira Lima, J., Hershowitz, A., Philip Todd, B., & Leventis, S. (2023). Guidelines on the Introduction and Use of Artificial Intelligence in the Parliamentary Workspace (Version 1). figshare. <https://doi.org/10.6084/m9.figshare.22687414.v1>

Global Partnership on Artificial Intelligence, <https://gpai.ai/>

GO FAIR: FAIR Principles, <https://www.go-fair.org/fair-principles>

Harris, M., and Wilson, A. (2024). Representative Bodies in the AI Era: Insights for Legislatures. Volume 1. POPVOX Foundation. <https://www.popvox.org/ai-vol1>

Health Insurance Portability and Accountability Act (HIPAA) (1996). Pub. L. No. 104-191

Heine, M., Dhungel, AK, Schrills, T., & Wessel D. (2023) Künstliche Intelligenz in öffentlichen Verwaltungen. Springer Gabler Wiesbaden. <https://link.springer.com/book/10.1007/978-3-658-40101-6>

Holdsworth, J. & IBM Inc. (2023). What is AI bias? <https://www.ibm.com/topics/ai-bias>.

IPU (2024). Using generative AI in parliaments. Geneva: IPU.

IPU CIP AI webinars: <https://www.ipu.org/innovation-tracker/story/2023-transforming-parliaments-webinar-series>

Italian Chamber of Deputies (2024). Using Artificial Intelligence to support parliamentary work, https://comunicazione.camera.it/sites/comunicazione/files/notiz_prima_pag/allegati/Rappporto_IA_ENG_WEB_V2.pdf

Kaggle: Level up with the largest AI & ML community. <https://www.kaggle.com>

Khatri, V., & Brown, C. V. (2010). Designing data governance. Communications of the ACM, 53(1), 148-152. <https://doi.org/10.1145/1629175.1629210>

KI Campus: <https://ki-campus.org/https://ki-campus.org/front?locale=en>

Koryzis, D., Dalas, A., Spiliotopoulos, D., & Fitsilis, F. (2021). ParlTech: Transformation Framework for the Digital Parliament. Big Data and Cognitive Computing 5(1):15. <https://doi.org/10.3390/bdcc5010015>

Koskimaa, V., & Raunio, T. (2020). Encouraging a longer time horizon: the Committee for the Future in the Finnish Eduskunta. The Journal of Legislative Studies, 26(2), 159-179. <https://doi.org/10.1080/13572334.2020.1738670>

Miller, G. (2023). US Senate AI ‘Insight Forum’ Tracker. <https://www.techpolicy.press/us-senate-ai-insight-forum-tracker/>

MOOC and Open Access Teaching Book Artificial Intelligence in Government on the German eGov-Campus: https://egov-campus.org/courses/kiverwaltung_uzl_2021-1

PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law., <https://pace.coe.int/en/pages/artificial-intelligence>

PACE Resolution 2343 (2020). Preventing discrimination caused by the use of artificial intelligence. <https://pace.coe.int/pdf/263ef53d02a37aaf8fbcdc7a706ec811d15244/res.%202343.pdf>

Palmirani, M., Vitali, F., Van Pyumbroeck, W., & Nubla Durango, F. (2022). Legal Drafting in the Era of Artificial Intelligence and Digitisation. Brussels: European Commission. <https://joinup.ec.europa.eu/sites/default/files/document/2022-06/Drafting%20legislation%20in%20the%20era%20of%20AI%20and%20digitisation%20%E2%80%93%20study.pdf>

Prunkl, C. (2022). Human autonomy in the age of artificial intelligence. Nature Machine Intelligence, 4(2), 99–101. <https://doi.org/10.1038/s42256-022-00449-9>

Read, A. (2023). A Democratic Approach to Global Artificial Intelligence (AI) Safety. London: WFD. <https://www.wfd.org/what-we-do/resources/democratic-approach-global-ai-safety>

Regulation (EU) 2016/679

Russell, S.J., & Norvig, P. (2021). Artificial Intelligence: A Modern Approach (4th ed.). Hoboken: Pearson.

Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Artificial intelligence hallucinations. Critical Care, 27(1), 180. <https://doi.org/10.1186/s13054-023-04473-y>

UNESCO (2021). Recommendation on the Ethics of Artificial Intelligence. <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>

US Executive Order 14110 of October 30, 2023 (Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence), <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>

US House of Representatives' Modernization Subcommittee: <https://cha.house.gov/modernization>

UK Government (2023). AI regulation: a pro-innovation approach. Policy paper. <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach>

United Nations (2024). Seizing the opportunities of safe, secure and trustworthy artificial intelligence systems for sustainable development, United Nations General Assembly, New York. <http://www.undocs.org/A/78/L.49>

United Nations (General Assembly). (1966). International Covenant on Civil and Political Rights. Treaty Series, 999, 171.

Vale, D., El-Sharif, A. & Ali, M. Explainable artificial intelligence (XAI) post-hoc explainability methods: risks and limitations in non-discrimination law. AI Ethics 2, 815–826 (2022). <https://doi.org/10.1007/s43681-022-00142-y>

von Lucke, J. (2024). Wie verändert künstliche Intelligenz die Verwaltungsarbeit? PDV-News 2024.

von Lucke, J., Fitsilis, F.; Etscheid, J. (2023). Research and Development Agenda for the Use of AI in Parliaments, in: David Duenas Cid et al (Eds.): DGO '23 - Proceedings of the 24th Annual International Conference on Digital Government Research, Association for Computing Machinery (ACM), 423-433. <https://doi.org/10.1145/3598469.3598517>

Wolff, J. (2021). How Is Technology Changing the World, and How Should the World Change Technology?. Global Perspectives, 2(1), 27353.

注釋

1

Fitsilis, F., von Lucke, J., Mikros, G., Ruckert, J., Alberto de Oliveira Lima, J., Hershowitz, A., Philip Todd, B., & Leventis, S. (2023). Guidelines on the Introduction and Use of Artificial Intelligence in the Parliamentary Workspace (Version 1.0). figshare. <https://doi.org/10.6084/m9.figshare.22687414.v1>

2

Fitsilis, F., & Costa, O. (2023). Parliamentary administration facing the digital challenge. In The Routledge Handbook of Parliamentary Administrations (pp. 105-120). Routledge.

3

Dai, X., & Norton, P. (Eds.). (2013). The Internet and parliamentary democracy in Europe: A comparative study of the ethics of political communication in the digital age. Routledge.

4

See, for instance, Russell, S.J., & Norvig, P. (2021). Artificial Intelligence: A Modern Approach (4th ed.). Hoboken: Pearson.

5

Translation and adaptation after von Lucke, J. (2024). Wie verändert künstliche Intelligenz die Verwaltungsarbeit? PDV-News 2024.

6

IPU (2024). Using generative AI in parliaments. Geneva: IPU.

7

European Union Agency for Cybersecurity, Milenkovic, G., & Dekker, M. (2020). Guideline on security measures under the EEC, Publications Office. <https://data.europa.eu/doi/10.2824/44013>

8

European Data Protection Board (2023). Guidelines 01/2022 on data subject rights - Right of access. https://edpb.europa.eu/system/files/2023-04/edpb_guidelines_202201_data_subject_rights_access_v2_en.pdf

9

Bresciani, P.F., & Palmirani, M. (2024). Constitutional Opportunities and Risks of AI in the law-making process. FEDERALISMI. IT 2, pp. 1 - 18. <https://hdl.handle.net/11585/953858>

10

Italian Chamber of Deputies (2024). Using Artificial Intelligence to support parliamentary work, https://comunicazione.camera.it/sites/comunicazione/files/notiz_prima_pag/allegati/Rappporto_IA_ENG_WEB_V2.pdf

11

Such comprehensive lists and research agendas have already been published: Fitsilis, F., Koryzis, D., & Schefbeck, G. (2022). Legal informatics tools for evidence-based policy creation in parliaments. International Journal of Parliamentary Studies, 2(1), 5-29; von Lucke, J., Fitsilis, F., & Etscheid, J. (2023). Research and Development Agenda for the Use of AI in Parliaments, in: David Duenas Cid et al (Eds.): DGO '23 - Proceedings of the 24th Annual International Conference on Digital Government Research, Association for Computing Machinery (ACM), 423-433.

12

von Lucke, J., Fitsilis, F., & Etscheid, J. (2023). Research and Development Agenda for the Use of AI in Parliaments, in: David Duenas Cid et al (Eds.): DGO '23 - Proceedings of the 24th Annual International Conference on Digital Government Research, Association for Computing Machinery (ACM), 423-433.

13

Fitsilis, F., & de Almeida, P. (2024). Artificial Intelligence and its Regulation in Representative Institutions. In Charalabidis, Y., Medaglia, R., & van Noordt, C. (Eds.), Research Handbook on Public Management and Artificial Intelligence (pp. 149-167). Edward Elgar Publishing.

14

Ibid.

15

Palmirani, M., Vitali, F., Van Pyumbroeck, W., & Nubla durango, F. (2022). Legal Drafting in the Era of Artificial Intelligence and Digitisation. Brussels: European Commission.

16

Fitsilis, F. (2019). Imposing regulation on advanced algorithms. Cham: Springer.

17 Wolff, J. (2021). How Is Technology Changing the World, and How Should the World Change Technology?. *Global Perspectives*, 2(1), 27353.

18 European Parliament resolution of 19 May 2021 on artificial intelligence in education, culture, and the audiovisual sector (2020/2017(INI))
https://www.europarl.europa.eu/doceo/document/TA-9-2021-0238_EN.html;
European Parliament resolution of 20 May 2021 on shaping the digital future of Europe: removing barriers to the functioning of the digital single market and improving the use of AI for European consumers (2020/2216(INI))
https://www.europarl.europa.eu/doceo/document/TA-9-2021-0261_EN.html

19 European Parliament (2024). Artificial Intelligence Act. https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_EN.html

20 Framework Convention on the design, development, and application of artificial intelligence systems based on the Council of Europe's standards on human rights, democracy and the rule of law, and conducive to innovation, in accordance with the relevant decisions of the Committee of Ministers (CAI).

21 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

22 United Nations (2024). Seizing the opportunities of safe, secure and trustworthy artificial intelligence systems for sustainable development, United Nations General Assembly, New York. <http://www.undocs.org/A/78/L.49>

23 Nonetheless, ambitious approaches do exist. Read, A. (2023). A Democratic Approach to Global Artificial Intelligence (AI) Safety. London: WFD. <https://www.wfd.org/what-we-do/resources/democratic-ap-proach-global-ai-safety>

24 Fitsilis, F., & de Almeida, P. (2024). Artificial Intelligence and its Regulation in Representative Institutions. In Charalabidis, Y., Medaglia, R., & van Noordt, C. (Eds.), *Research Handbook on Public Management and Artificial Intelligence* (pp. 149-167). Edward Elgar Publishing.

25 von Lucke, J., Fitsilis, F.; Etscheid, J. (2023). Research and Development Agenda for the Use of AI in Parliaments, in: David Duenas Cid et al (Eds.): DGO '23 - Proceedings of the 24th Annual International Conference on Digital Government Research, Association for Computing Machinery (ACM), 423-433. <https://doi.org/10.1145/3598469.3598517>

26 Harris, M., and Wilson, A. (2024). Representative Bodies in the AI Era: Insights for Legislatures. Volume 1. POPVOX Foundation. <https://www.popvox.org/ai-vol1>

27 Committee on House Administration (2024). Flash report - Artificial Intelligence Strategy & Implementation. https://cha.house.gov/_cache/files/a/d/ad4d1279-c8f8-439b-9e3b-a95b-01d61d03/56078B0226EDF1EAF-76D863A2E7765A5.cha-q1-flash-report.pdf

28 UK Government (2023). AI regulation: a pro-innovation approach. Policy paper. <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach>

29 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

30 Vale, D., El-Sharif, A. & Ali, M. Explainable artificial intelligence (XAI) post-hoc explainability methods: risks and limitations in non-discrimination law. *AI Ethics* 2, 815–826 (2022). <https://doi.org/10.1007/s43681-022-00142-y>

31 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

32 Regulation (EU) 2016/679

33 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

34 PACE Resolution 2343 (2020)

35 Holdsworth, J., & IBM (2023). What is AI bias? <https://www.ibm.com/topics/ai-bias>

36 U.S. Executive Order 14110 of October 30, 2023 (Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence), <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>

37 Data Ethics Commission (2019). Opinion of the Data Ethics Commission. Berlin. https://www.bmi.bund.de/SharedDocs/downloads/EN/themen/it-digital-policy/datenethik-kommission-abschlussgutachten-lang.pdf?__blob=publicationFile&v=5

38 Aul Tremblay, et al. v. OpenAI, Inc., et al. (Case No. 3:23-cv-03223-AMO).

39 The New York Times Company v. Microsoft Corporation, OpenAI, Inc., et al. (Case No. 1:23-cv-11195)

40 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

41 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

42 United Nations (General Assembly). (1966). International Covenant on Civil and Political Rights. Treaty Series, 999, 171.

43 See, e.g., ongoing Council of Europe's Convention on Artificial Intelligence, Human Rights, Democracy, and the Rule of Law (AI Convention).

44 Fitsilis, F. (2024). The parliamentary perspective of better regulation in Greece. GRNET Tech Day in Digital Ready Policy Making, 1 February 2024. https://events.grnet.gr/event/138/attachments/393/669/Fotis%20Fitsilis_The%20parliamentary%20perspective%20of%20better%20regulation%20in%20Greece.pdf

45 UNESCO (2021). Recommendation on the Ethics of Artificial Intelligence. <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>

46 Altman, S. (2023) Planning for AGI and beyond. <https://openai.com/index/planning-for-agi-and-beyond/>

47 Prunkl, C. (2022). Human autonomy in the age of artificial intelligence. *Nature Machine Intelligence*, 4(2), 99–101. <https://doi.org/10.1038/s42256-022-00449-9>

48 PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law, <https://pace.coe.int/en/pages/artificial-intelligence>

49

PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

50

PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

51

Regulation (EU) 2016/679

52

HIPAA (1996). Pub. L. No. 104-191

53

PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

54

Harris, M., and Wilson, A. (2024). Representative Bodies in the AI Era: Insights for Legislatures. Volume 1. POPVOX Foundation. <https://www.popvox.org/ai-vol1>

55

<https://chat.openai.com>

56

Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Artificial intelligence hallucinations. *Critical Care*, 27(1), 180.

57

PACE (2020). Artificial Intelligence: Ensuring respect for democracy, human rights and the rule of law. <https://pace.coe.int/en/pages/artificial-intelligence>

58

Italian Chamber of Deputies, Using Artificial Intelligence to support parliamentary work, https://comunicazione.camera.it/sites/comunicazione/files/notiz_prima_pag/allegati/Rappporto_IA_ENG_WEB_V2.pdf

59

Brazilian Congress, Chamber of Deputies, Digital Strategy 2021-2024, <https://www2.camara.leg.br/a-camara/estruturaadm/gestao-na-camara-dos-deputados/gestao-estrategica-na-camara-dos-deputados/gestao-estrategica-de-tic/estrategia-digital>

60

Khatiri, V., & Brown, C. V. (2010). Designing data governance. *Communications of the ACM*, 53(1), 148-152.

61

FAIR stands for Findable, Accessible, Interoperable, and Reusable data, see GO FAIR, <https://www.go-fair.org/fair-principles/>

62

<https://hellenicocrteam.gr>

63

Fitsilis, F., & Mikros, G. (2021). Development and Validation of a Corpus of Written Parliamentary Questions in the Hellenic Parliament. *Journal of Open Humanities Data*, 7, p.18. <https://doi.org/10.5334/johd.45>

64

Koryzis, D., Dalas, A., Spiliotopoulos, D., & Fitsilis, F. (2021). ParlTech: Transformation Framework for the Digital Parliament. *Big Data and Cognitive Computing* 5(1):15. <https://doi.org/10.3390/bdcc5010015>

65

Akoma Ntoso Version 1.0. Part 1: XML Vocabulary, <http://docs.oasis-open.org/legaldocml/akn-core/v1.0/akn-core-v1.0-part1-vocabulary.html>

66

Akoma Ntoso use cases: http://www.akomantoso.org/?page_id=275

67

Koskimaa, V., & Raunio, T. (2020). Encouraging a longer time horizon: the Committee for the Future in the Finnish Eduskunta. *The Journal of Legislative Studies*, 26(2), 159-179. <https://doi.org/10.1080/13572334.2020.1738670>

68

Fitsilis, F. (2021). Artificial Intelligence (AI) in parliaments–preliminary analysis of the Eduskunta experiment. *The Journal of Legislative Studies*, 27(4), 621-633. <https://doi.org/10.1080/13572334.2021.1976947>

69

Miller, G. (2023). US Senate AI ‘Insight Forum’ Tracker. <https://www.techpolicy.press/us-senate-ai-insight-forum-tracker/>

70

IPU CIP AI webinars: <https://www.ipu.org/innovation-tracker/story/2023-transforming-parliaments-webinar-series>

71

See, for instance, German KI Campus: <https://ki-campus.org> / <https://ki-campus.org/front?locale=en>.

72

GPAI: <https://gpai.ai/>

73

AI for Good initiative, <https://aiforgood.itu.int>

74

Caggle is similar to Kaggle, <https://www.kaggle.com/>

75

U.S. House of Representatives’ Modernization Subcommittee: <https://cha.house.gov/modernization>

西敏寺民主基金會 (Westminster Foundation for Democracy, 簡稱 WFD)
為英國的公共機構，專責推動全球民主發展。
WFD 活躍於國際，與各國國會、政黨、公民社會團體合作，
致力於促進政治制度的公平、共融與確實問責。



www.wfd.org



@WFD_Democracy



@WFD_Democracy



Westminster Foundation for Democracy (WFD)



掃描此處訂閱 WFD 新消息



西敏寺民主基金會是由英國外交、國協及
發展部 (Foreign, Commonwealth and
Development Office) 資助的非部會公共
組織 (Non-departmental Public Body)。



Foreign, Commonwealth
& Development Office