

負責任 AI 執行中心 FAQ

一、政府推動「負責任 AI 中心」的目的為何？是否有其必要性？

人工智慧 (AI) 已廣泛應用於臨床診斷、治療輔助及醫療決策，因此，世界衛生組織 (WHO)、世界醫學會 (WMA)，以及歐盟、美國等國家與國際組織，均要求醫療 AI 的使用必須建立在倫理治理 (AI Governance) 架構下，以保障病人安全、醫療品質及社會信任。我國推動「負責任 AI 中心」，目的就是建立醫療 AI 的治理能力，協助醫院在安全、透明、可監督的前提下導入 AI，與國際發展方向一致，並非我國獨有的制度。

二、目前已有 AI 倫理原則，為何還要成立「負責任 AI 中心」？

目前世界衛生組織 (WHO)、世界醫學會 (WMA) 及歐美各國均提出醫療 AI 應遵循的倫理原則，但多數仍屬於原則性的倡議，例如公平性、透明性、可解釋性及問責性等，醫院在實際導入 AI 時，往往缺乏具體可執行的制度、流程及管理工具。

因此，我國率先將國際 AI 倫理原則轉化為醫院可實際落實的治理制度，建立涵蓋 AI 導入前評估、風險管理、本地驗證 (Local Validation)、偏差評估 (Bias Assessment)、持續監測及 AI 登錄管理等完整治理架構，使 AI 倫理真正落實於醫療現場，而不只是停留在原則宣示。

這套治理制度已於世界衛生大會 (WHA) 期間舉辦的國際論壇向各國衛生官員及專家發表，同時也於世界醫學會 (World Medical Association, WMA) 國際會議進行專題報告，獲得來自多國政府代表、醫學會及醫療 AI 專家的高度肯定。與會專家普遍認為，台灣成功將抽象的 AI 倫理原則轉化為可操作、可落實、可持續改善的治理模式，不僅具有高度實務價值，更可作為各國推動醫療 AI 治理的重要參考，建立國際醫療 AI 治理的新典範。因此，推動「負責任 AI 中心」並非台灣自行創設新的管理要求，而是回應國際醫療 AI 治理趨勢，並進一步將國際倫理原則具體制度化，讓台灣在全球醫療 AI 治理領域走在前端，為國際社會提供可供借鏡的實踐經驗。

三、AI 已取得食品藥物管理機關許可，為何還需要醫院再做治理？

AI 與傳統醫療器材不同，即使取得主管機關許可，也不代表在所有醫院、所有病人族群及所有臨床情境下，都能維持相同的安全性與效能。醫療 AI 的安全與有效性，仍須透過醫院持續進行本地驗證 (Local Validation)、偏差評估 (Bias Assessment)、風險管理及持續監測。因此，美國及歐洲許多頂尖醫院均已成立醫療 AI 治理聯盟 (Trustworthy and Responsible AI Network, TRAIN, <https://jamanetwork.com/journals/jama/fullarticle/2830340>)，共同分享最佳實務，建立治理制度。我國推動「負責任 AI 中心」，正是與國際發展趨勢接軌。

四、政府推動此制度，是否會增加區域醫院、地區醫院及中小型醫院負擔？

負責任 AI 中心定位為「卓越中心 (Center of Excellence) 」，目的在建立示範模式，累積治理經驗及最佳實務，再提供全國醫院參考。本計畫並非強制參與，不是醫院評鑑項目，也未與任何法律責任或行政裁罰連結，而是一項鼓勵醫院自主提升治理能力的輔導計畫。政府也提供經費補助、輔導機制及全國 AI 登錄管理平台，協助醫院逐步建立 AI 治理能力，而非增加醫院負擔。

五、醫院是否需要增加許多行政程序？是否流於形式？

治理程序的目的不是增加行政作業，而是降低 AI 對病人可能造成的風險。

例如，本地驗證、偏差測試、風險評估及持續監測等程序，都是確保 AI 能安全應用於醫療的重要措施。若省略這些程序，一旦 AI 發生錯誤，可能直接影響病人安全，而相關風險往往無法事後完全補救。因此，這些程序是導入 AI 所必要的治理成本，也是保障病人安全不可或缺的一環，使用 AI 單位負起善良管人責任。

六、若 AI 發生醫療事故，責任仍由醫師承擔，政府如何保障醫師？

正如委員所提，AI 的醫療責任最終仍由醫療機構或醫師承擔。因此，醫師必須在 AI 具備透明性 (Transparency)、可解釋性 (Explainability) 及可追溯性等條件下，才能合理判斷是否採納 AI 建議。負責任 AI 中心建立的治理架構，正是提供醫院及醫師必要的制度支持與風險管理工具，使 AI 成為醫師可信賴的輔助工具，而非增加醫師執業風險。

七、中小型醫院資源有限，是否有能力推動 AI 治理？

考量不同層級醫院資源差異，若中小型醫院尚未建立完整治理能力，建議優先導入行政管理、醫院營運管理等低風險 AI 應用，再逐步建立治理制度，循序擴展至臨床應用。政府亦將持續提供輔導資源與最佳實務分享，協助各級醫院逐步建立治理能力。

八、目前的治理架構是否也適用於生成式 AI？

目前公布的「負責任 AI 中心」治理架構，主要適用於判別式 AI (Discriminative AI) 及傳統推論型 AI 系統。生成式 AI (Generative AI) 因具有幻覺 (Hallucination)、自主生成內容及持續演化等不同風險，目前尚未納入本管理架構。未來將透過 AI 沙盒 (AI Sandbox) 機制，在可控環境下協助醫院導入、驗證及管理生成式 AI，逐步建立符合臨床需求與國際趨勢的治理模式。

九、政府推動「負責任 AI 中心」最終目標為何？

推動「負責任 AI 中心」並非限制 AI 發展，而是建立安全、可信賴且可持續發展的醫療 AI 治理制度。透過卓越中心示範、政府輔導及全國最佳實務分享，協助各級醫院安全導入 AI，保障病人權益、保障醫師執業安全，同時建立我國在醫療 AI 倫理治理及 AI Governance 的國際領導地位，讓台灣成為全球推動負責任 AI 的重要典範。