



星球永續健康線上直播

AI 醫療治理(1)

可信任智慧醫療輔助

2026 年 7 月 1 日

人工智慧快速發展改變醫療照護模式，從疾病篩檢、影像判讀到臨床決策支持，AI 逐漸成為醫療輔助工具。隨著 AI 醫療應用日益普及，資料隱私、模型偏誤、資安風險與責任歸屬等議題也受到高度關注。本週我們將探討人工智慧治理，以及智慧醫療 LDCT 篩檢的治理應用實例。

健康科學新知

中東地緣破冰 海峽安全待穩：「危中求衡」

美國與伊朗上周在瑞士及區域斡旋框架下推進的臨時協議，正逐步從停火安排進入經濟、能源、核查與海上通行管理的實質談判階段。國際觀察顯示美伊衝突進入為期約 60 天的過渡期：美方暫時放寬部分對伊朗的能源與金融限制，伊方則被要求維持荷姆茲海峽通航，並就核查、制裁解除及區域安全議題繼續談判。這項安排一方面帶來油市與航運市場的短暫緩和，另一方面也暴露出雙方對協議內容的解讀仍高度分歧。然而臨時協議脆弱性十分明顯，美伊雙方近期仍因新加坡籍船舶遇襲、沿岸雷達與軍事設施遭打擊、無人機攻擊及海峽航道爭議相互指責對方違反停火。雖然油價因市場期待海峽維持開放而回落，部分肥料、石油與商船運輸也逐步恢復，但保險市場、船運業與能源交易商仍保持高度警戒。只要核查、資金用途、制裁解除與海峽管理權沒有明確制度化，任何局部軍事事務都可能迅速破壞停火氣氛。能源制裁方面，美國財政部授權伊朗在 60 天內恢復原油、石化產品及相關衍生品的生產、交付與銷售，並允許以美元進行付款。此舉等同暫時鬆動多年來壓制伊朗能源出口的核心制裁工具，也使伊朗有機會以較接近市場價格出售石油，而非過去透過折扣、轉運與規避制裁網絡銷售。美方將此描述為談判進展後的「臨時授權」，強調伊朗必須在通航、核查與後續談判上履行相對承諾；伊方則將其視為自身在戰後談判中取得的具體經濟成果。伊朗另稱美方同意釋放約 120 億



美元凍結資產，但雙方對資金用途仍有歧見。荷姆茲海峽船運雖逐步恢復，管理權、航道規範與是否收費仍未定案；國際原子能總署核查人員能否重返伊朗，也因雙方說法不一而充滿變數。與此同時，美伊仍就船舶遇襲、雷達及軍事設施遭攻擊互控違反停火，顯示局勢仍處「危中求衡」。卡達、巴基斯坦、阿曼與瑞士正協助建立海上溝通與技術工作小組，但任何局部衝突，都可能使停火再度破裂。

英國首相請辭 議員柏南接班呼聲高：「政局重整」

施凱爾上週宣布辭去首相與工黨黨魁後，英國正式進入權力交接期，政局焦點轉向熱門接班人、曼徹斯特市長安迪·柏南。柏南以親民形象與「曼徹斯特主義」受到工黨支持者期待，主張將中央權力與資源下放地方，並兼顧經濟發展、公共服務與勞工權益。政權更替也使原定英歐峰會延後，施凱爾推動的英歐關係重置、貿易與安全合作暫時停滯。德國總理梅爾茨、法國總統馬克宏及烏克蘭總統澤倫斯基等盟友，均關注新政府能否延續對烏克蘭、北約及歐洲安全的支持。柏南若接任，還須處理與川普政府的關係、國防支出、財政規則及生活成本壓力，在內政改革與國際承諾間取得平衡。

委內瑞拉遭逢百年強震：「震災映治」

委內瑞拉 6 月 24 日晚間接連發生規模 7.2 與 7.5 強震，兩震僅相隔約 39 秒，形成罕見「雙震」，重創北部地區並波及哥倫比亞。截至 6 月 27 日，至少 1,430 人死亡、3,238 人受傷，約 860 萬人暴露於中度至嚴重震動，直接損失估達 67 億美元，約占 GDP 6%。拉瓜伊拉州與首都卡拉卡斯災情最重，超過百棟建築倒塌或嚴重受損，許多居民流離失所，醫院、電力與基礎設施也承受巨大壓力。聯合國協調 44 支國際搜救隊、2,245 名專家與 140 隻搜救犬投入救援，但偏遠災區仍缺乏重型機具、人力與資訊。這場「震災映治」不僅考驗黃金救援，也暴露委內瑞拉長期醫療、治理與防災體系的脆弱。

氣候-能源-地緣政治多重危機：「複合風險」

近期歐洲與亞太同時面臨極端高溫、聖嬰與中東能源危機交織的「複合風險」。英國氣象局對英格蘭中南部及威爾斯發布罕見紅色高溫警報，部分地區可能超過 38°C，醫療、交通與弱勢族群風險升高。法國、西班牙、義大利也出現逾 40°C 熱浪，造成溺



水、鐵路受阻、景點關閉、核電廠降載與戶外勞工保護措施。氣候與能源壓力於亞太地區造成能源與糧食安全的壓力。中東危機與荷姆茲海峽的不確定性，使高度依賴中東能源的東南亞國家受到衝擊。報導指出東南亞約 60% 原油進口與三分之一天然氣進口來自中東。荷姆茲海峽航運穩定情況以及伊朗對通行船隻收取費用，進一步推高燃料、航運、保險與進口成本。能源壓力已使部分東南亞國家開始採取節能措施，能源供應不穩已開始影響政府行政、企業運作與民眾生活。能源危機與聖嬰現象疊加影響亞洲糧食生產輸出。聖嬰是赤道中東太平洋海水異常升溫所形成的氣候現象，會影響全球降雨、風場與氣溫分布，並可能加劇乾旱、洪水與熱浪。對亞太農業而言，聖嬰可能造成降雨不穩與作物減產。聯合國糧農組織警告，如果中東危機導致能源與肥料價格上升，再加上聖嬰造成的乾旱與高溫，亞太地區可能有 700 萬至 800 萬噸稻米產量面臨風險。其中泰國因乾旱跡象明顯，又高度暴露於能源與肥料價格上漲，為高複合風險國家，印尼、菲律賓與部分太平洋島國也可能受到影響。

身心健康與頂尖足球賽：「競技永續」

國際職業足球員協會（FIFPRO）醫學主任 Vincent Gouttebarger 指出，球員常被視為「超人」，卻面臨每週多達兩至三場比賽且缺乏恢復期的困境，極易引發焦慮、憂鬱及肌肉骨骼傷害。此外，本屆世足由三國協辦帶來的跨時區長途飛行，以及北美午後的高溫高濕環境，不僅大幅增加體能負荷，更會加劇心理倦怠。專家亦強調應重視腦震盪及反覆頭球造成的**「亞腦震盪」衝擊**，尤其是守門員長傳後高速球的影響，亟需規則改革與實證防護。他呼籲應打破心理健康污名，將專業心理人員納入醫療團隊，確保球員在追求競技巔峰的同時，獲得全方位的健康保障。

AI 診斷罕病邁進臨床：「智診新境」

罕見疾病患者全球估計約有 3 億人，但由於已知罕見疾病超過 7,000 種，且多數疾病發生率極低，患者往往必須經歷多年求診、反覆轉介、誤診與不必要治療，才可能獲得正確診斷，這段艱辛過程常被稱為「診斷奧德賽」。DeepRare 的重要性不僅在於能提出診斷答案，更在於它能整合臨床資料、基因資訊、醫學文獻與病例資料庫，產生有



排序的診斷假設，並提供可追溯的推理過程。其架構類似一個 AI 協調團隊，由中央推理引擎整合分析，不同模組則分別負責從病歷文字萃取標準化症狀、搜尋醫學文獻、比對相似病例資料庫，以及在具備 DNA 定序資料時分析基因變異。這種多模組分工方式類似罕見疾病中心的多專科團隊討論，但 AI 能以更快速度與更大規模處理資料。在臨床驗證中，DeepRare 表現優於可使用搜尋引擎的罕見疾病專科醫師。其排名第一的診斷預測正確率達 64%，若納入前五個診斷建議，正確率可達 78.5%，高於醫師的 65.6%。更重要的是，醫學專家對 DeepRare 推理鏈的認可程度達 95%，顯示其解釋過程具有臨床一致性與可驗證性。文章特別強調，DeepRare 的核心價值在於「可追蹤推理」：罕見疾病診斷並非只要猜中答案，醫師還需要理解某個診斷為何合理，以及背後有哪些證據支持。DeepRare 能引用具體文獻、病例報告與資料庫條目，建立可檢視的證據路徑，使其相較於黑箱式 AI 輸出更容易被醫療人員審查，也更有助於建立臨床信任。

DeepRare 高表現的關鍵在於兩項機制：第一，它能從網路與醫學資料來源擷取較新的知識；第二，它具備「自我反思迴圈」，可重新檢查初步結論並修正可能錯誤。這對罕見疾病尤其重要，因為新的基因疾病與臨床知識持續增加，診斷工具必須能跟上醫學知識的快速更新。然而，DeepRare 仍存在限制，例如少部分 AI 產生的參考文獻被發現屬於「幻覺」引用，也就是看似合理但實際不存在的來源，這反映大型語言模型在臨床應用上仍需嚴格驗證與風險控管。單靠大型語言模型並不足以完整掌握罕見疾病診斷所需的臨床特徵與基因關聯，DeepRare 的真正優勢在於將 LLM 的語言理解與推理能力，結合整理過的專業知識來源、病例資料庫與基因分析模組，形成一種更接近臨床多專科判斷流程的 AI 輔助診斷模式。

智慧語模彌補健康資訊缺口：「智療新途」

越來越多人使用一般型 AI 聊天機器人詢問健康與醫療問題，反映出醫療資訊取得、照護可近性與健康溝通上的缺口。Costa-Gomes 等人在 Nature Health 發表的研究分析 2026 年 1 月約 61.8 萬筆使用者與 Microsoft Copilot 的健康相關對話，觀察大眾如何透過 AI 處理健康資訊、症狀、情緒福祉與醫療體系導航等問題。健康相



關查詢中，最大宗為一般健康資訊與教育，約占五分之二，內容包括藥物作用、疾病成因與生物醫學概念等。另有部分查詢與學術研究、醫療文件及行政流程有關，顯示使用者不僅包含一般民眾，也可能包含醫療或研究專業人員。約四分之一的對話涉及個人健康問題，包括症狀、特定疾病照護與情緒福祉。這些個人化查詢更常透過手機提出，也更常出現在夜間，尤其情緒福祉相關問題在夜晚比例明顯增加，顯示當傳統醫療服務較難取得時，AI 聊天機器人可能成為即時、低門檻的替代資訊來源。許多查詢並非只關於使用者本人，而是替家人、照顧對象或其他親近者詢問健康狀況，反映照顧者支持與健康溝通仍有不足。另有部分使用者詢問保險、就醫流程與醫療體系導航，顯示醫療制度對一般人而言仍然複雜且不易理解。AI 聊天機器人的普及不只是科技使用現象，也揭示醫療服務在時間、可近性、溝通與制度透明度上的結構性缺口。這項研究並未評估聊天機器人的回答是否正確，也未證明其能帶來臨床效益。大型語言模型雖然在醫學考題或標準化診斷推理任務中表現良好，但面對一般民眾語境不完整、個人化且開放式的健康提問時，仍可能產生事實錯誤、不當建議，甚至在是否應就醫或急診等關鍵判斷上帶來風險。一般型 AI 聊天機器人也可能迎合使用者原有想法，而非像醫師一樣透過結構化問診、臨床判斷與實證依據進行評估。大眾使用 AI 詢問健康問題已成為現實，重點不再只是是否應該使用，而是何時使用、為何使用、替誰使用，以及使用後可能產生什麼風險。醫療專業受到倫理、臨床指引、監管制度與醫療器材標準約束；相較之下，一般型 AI 聊天機器人若未被明確定位為醫療產品，往往缺乏同等程度的有效性驗證、問責機制與可靠性保障。未來治理重點不僅在提升模型準確性，也在於釐清平台責任、建立風險提示、強化轉介醫療專業的機制，並補足現有醫療體系在資訊可近性與照護連續性上的不足。

AI 限制存取治理模式：「開放有界」

Anthropic 於 2026 年 4 月宣布，其開發的旗艦模型 Claude Mythos 因具備發現作業系統與瀏覽器漏洞等強大能力，存在高度資安與國家安全風險，決定不對大眾開放。該公司啟動「Project Glasswing」，僅限量提供給約 50 個經審查的可信任組織，



顯示先進 AI 的釋出模式正從「公開透明」轉向「受限存取」與「分級授權」。此趨勢已擴散至生命科學領域，如 OpenAI 的 GPT-Rosalind 及 Google 的 AI 共研系統均採申請制或受控存取，以防範 AI 被濫用於生物武器研發。然而，限制使用恐加深科研不平等，且面臨開源社群快速追趕的挑戰。專家呼籲，單靠企業自律已不足，政府應介入建立如 NIST 旗下的 AI 標準中心或國際協調機制，針對具備「雙重用途」的技術實施類似出口管制的監管，以平衡技術創新與公共安全。

人工智慧治理

在《黑色止血鉗》第二季影集中，東城大學附設醫院希望發展新的心血管醫療中心，因此積極尋找能夠帶領醫院邁向下一個世代的領導者。原有的院長因年事已高，加上健康因素影響，醫院希望邀請世界頂尖的心臟外科專家加入主持新設心血管醫療中心。因此邀請來自澳洲的天才外科醫師天城雪彥。天城醫師憑藉高超的醫療技術與精準的臨床判斷，發展出冠狀動脈疾病的精準治療方法，血管直接吻合術。傳統的冠狀動脈繞道手術，通常是取用其他部位的血管，繞過阻塞的位置重新建立血流繞道手術。然而天城醫師的技術更進一步能直接將健康血管與病灶血管進行吻合，如同重新建構原本的血流循環，因此被譽為世界上唯一能完成這項手術的醫師。憑藉這項獨步全球的技術，天城醫師受邀來到東城大學附設醫院，甚至成為候選院長。不過，雖然他成功治癒許多困難個案，但其特立獨行的行事風格與個人作風，也在院內引起不少爭議與討論。

另一方面，東城醫院的競爭對手維新大學附設醫院則選擇了另一條發展路線，也就是人工智慧醫療技術。他們開發了一套名為「艾卡諾 (Elcano)」的 AI 診療系統，並結合達文西機器手術手臂，打造精準自動化的智慧手術平台。艾卡諾其實就像現在的 AI 一樣，可以分成「大腦」與「雙手」兩個部分。AI 的大腦負責分析影像、病歷資料以及醫學文獻，透過電腦視覺與決策演算法，判斷手術過程中應採取哪些策略與處置方式，甚至能與醫師進行互動討論，共同評估哪一種治療方式對病患最有利。而 AI 的雙手，則是透過機器手術手臂執行精細操作，就像現實世界中的達文西手術系統一樣。劇中的開發者早川醫師曾展示艾卡諾的精準度，利用 AI 控制機械手臂，在米粒大小的範圍內



描繪出蒙娜麗莎圖像，展現極高的穩定性與精細度。不過，再強大的 AI 仍然需要真實世界資料與專家經驗進行訓練。由於天城醫師掌握著世界唯一的冠狀動脈直接吻合技術，因此維新大學希望能透過臨床合作，讓艾卡諾學習這項高難度的手術技巧。劇中有一段非常有趣的情節，天城醫師與艾卡諾並非互相競爭，而是共同討論病患最適合的治療策略。原本天城醫師打算採用直接動脈吻合手術，雖然他的技術十分高超，但心臟手術畢竟屬於高侵入性、高風險的治療方式。利用人機共同討論與評估後，艾卡諾提出另一種侵入性更低的治療策略，能夠在不進行開胸手術的情況下，達到相近的治療效果。不過兩家醫院合作的過程中出現了轉折，艾卡諾的開發者早川醫師突然在醫院昏倒，經過影像檢查後，發現罹患了冠狀動脈瘤，需要接受手術治療。艾卡諾讀取影像、病歷資料以及相關學術文獻，分析各種治療方案的風險與效益，並提出建議的治療策略。艾卡諾認為最佳的治療方式，是先利用橈動脈作為繞道血管建立冠狀動脈繞道，之後再將冠狀動脈瘤進行縫合隔離，避免動脈瘤持續擴大或破裂。原本大家希望由天城醫師親自主刀，畢竟他是世界上唯一能執行冠狀動脈直接吻合術的專家。不過，由於天城醫師手部受傷，無法獨自完成手術，因此最終決定由艾卡諾結合機器手術平台協助完成治療，而早川醫師也因此成為這項技術的重要臨床案例。手術開始後，艾卡諾先操作機器手臂協助完成橈動脈血管摘取，並引導醫療團隊進行手術流程。隨後，系統接手完成冠狀動脈繞道與動脈瘤縫合隔離手術，整個過程看似十分順利。然而，就在手術即將完成時，艾卡諾卻出現了異常判斷。系統誤將已經完成縫合的血管辨識為尚未處理的病灶，做出了錯誤操作，導致早川醫師在術中心臟與血管受損造成大量出血。此時天城醫師因為先前手部受傷，左手無法正常施力，無法單獨完成高難度的冠狀動脈直接吻合修補手術。於是劇中安排由天城醫師負責判斷與決策，而艾卡諾則成為他的「右手」，透過機器手臂執行精細的操作。最終，由天城醫師、助手以及艾卡諾共同協作，成功完成受損血管的修補，挽救了早川醫師的生命。呈現的是智慧醫療最重要的概念之一，也就是人機協作。AI 並沒有取代醫師，而是在醫師無法獨自完成任務時，成為醫師能力的延伸與擴增，讓人類的專業判斷與機器的精準操作彼此互補。



後續調查發現，艾卡諾之所以誤判，是因為它將已經完成縫合的動脈瘤辨識為尚未處理的病灶，因此試圖再次進行縫合，最終造成周圍心臟組織與血管受損。更重要的是，造成誤判的原因並非演算法本身失效，而是其訓練資料與資料庫遭到惡意竄改與污染。劇中揭露，過去參與開發的研究人員因個人因素，刻意將錯誤資料植入艾卡諾的學習資料庫，使系統在特定情境下產生錯誤判斷。這其實反映了現代人工智慧治理中的重要議題，資料安全與資料治理。即使演算法再先進，如果輸入的資料受到污染、遭到入侵或被惡意操弄，人工智慧仍然可能做出錯誤決策，甚至在高風險醫療場域中造成病人傷害。

在人類與 AI 的時代，我們應該如何重新定位彼此的角色，也就是所謂的「人機同鑑」。75 年前，圖靈提出問關於機器是否能夠展現與人類相當的認知能力？但 75 年後的今天，隨著大型語言模型的出現，這個問題已經不再只是理論，而是我們每天都在面對的現實。現在的大型語言模型，無論在知識的廣度或推理的深度上，都已經展現出接近甚至在某些領域超越人類的能力。

過去的醫療模式是以醫師為中心，醫療品質主要建立在醫師個人的專業知識與臨床經驗之上；但隨著 AI 的快速發展，我們正逐漸走向另一種模式，也就是所謂的「擴增智慧 (Augmented Intelligence)」。擴增智慧強調的並不是 AI 取代醫師，而是 Human 與 AI 的協作。AI 擅長大量資料分析、影像辨識與風險預測，而醫師則負責整合病人的價值觀、臨床情境與倫理判斷，兩者彼此互補，共同完成醫療決策。對此，世界衛生組織 (WHO) 的立場非常明確，也就是所謂的 Human-in-command 原則。AI 可以提供建議與分析結果，但最終的決策與責任，仍然必須由人類來承擔。這理念正符合台灣《人工智慧基本法》第一條所強調的精神。我們推動 AI 的目的，並不是取代醫療專業，而是在保障人民基本權利與提升社會福祉的前提下，促進人工智慧的創新應用與擴增智慧發展。

因此，在醫療場域中，除了醫師與 AI 的協作之外，未來更重要的是醫師、病人與 AI 三方共同參與決策，也就是我們現在所強調的共享決策 (Shared Decision Making)。為了因應人工智慧快速發展所帶來的新挑戰，台灣也開始推動人工智慧治理架構。《人



《智慧基本法》草案主要從五大面向來建構我國未來 AI 發展的基礎架構。第一個面向是總則，包含立法目的、主管機關、AI 定義以及法規適用範圍等基本架構。第二個面向則是治理原則與組織架構，其中最核心的內容就是人工智慧治理七大基本原則，以及未來國家 AI 戰略與政府部門使用 AI 的風險評估與內控制度，這也是後面我們會進一步介紹的重點。第三個面向是發展與創新促進，希望透過教育與數位素養提升、產官學合作、產業扶植以及國際合作，讓人工智慧真正成為促進人類創新與社會進步的重要工具。無論是前面提到的罕見疾病診斷，或是健康聊天機器人的應用，都屬於這個面向的重要發展方向。第四個面向則是風險管理與安全。尤其是人工智慧的風險分級管理、高風險應用的監督機制，以及當 AI 發生錯誤時的責任歸屬與救濟制度，都是目前國際間最受到關注的議題。第五個面向則是權益保障與法制調適，包括個人資料保護、勞動權益保障，以及現有的法規如何因應人工智慧發展而進行修訂與調整。因此，台灣人工智慧法所希望建立的，並不是單純追求技術發展，而是在發展、治理、安全與權益保障之間取得平衡，進一步建構一個以人為本、可信任且負責任的人工智慧環境。換句話說，未來的 AI 不只是 Trustworthy AI，也必須是 Responsible AI，不只是 Technology-centered 更必須是 Human-centered，最終仍然回到以人為中心、以病人利益為核心的智慧醫療發展方向。

人工智慧治理的核心建立於七大基本原則之上，包括永續發展與福祉、人類自主、隱私保護與資料治理、資安與安全、透明與可解釋、公平與不歧視，以及問責機制。其中，永續發展與福祉為所有治理原則的核心目標，人工智慧的發展應以提升人類福祉與社會公益為最終目的，而非僅追求技術能力的提升。隨著人工智慧能力快速提升，人類自主、隱私保護與資料治理的重要性也日益增加。同時，資料污染、惡意攻擊、生物安全風險，以及演算法偏誤等問題，均可能對高風險應用領域造成重大影響，因此透明、可解釋與可問責的治理機制成為人工智慧發展的重要基礎。此外，人工智慧發展的典範亦逐漸由資料驅動（Data-driven AI）轉向治理驅動（Governance-driven AI）。過去強調模型效能最大化的發展模式，正逐漸轉向兼顧人類自主、隱私保護、資料治理與智



慧財產權保障的新治理架構。因此，未來人工智慧發展除了追求技術創新外，更必須兼顧社會公益、數位平權、創新研發與國家競爭力，建立可信任、負責任且以人為本的人工智慧環境。

隨著人工智慧能力快速提升，資安治理的重要性也持續上升。未來 AI 競爭除模型能力與效能外演變為資安能力、監管能力以及國際影響力之間的競爭。以歐盟為例，近年透過《AI Act》以及 AI Office 等制度設計，希望建立全球人工智慧治理與監管標準。然而，即使建立了法規與監管機構，實際的模型存取權與技術能力，仍然高度掌握在少數大型科技公司手中。因此，人工智慧治理也逐漸從技術議題，延伸至資安、產業與外交層面的戰略競爭。近年來，網路安全漏洞與資安事件的通報數量持續攀升，顯示數位系統的攻擊面正隨著 AI 技術普及而快速擴大。先進的 AI 模型能夠協助發現系統弱點與提升防禦能力，但同樣也可能被攻擊者利用來加速漏洞探索、攻擊設計與自動化滲透，因此模型能力的提升也意味著資安風險同步升高。

因此，模型的存取權與使用規範，逐漸成為人工智慧治理的重要議題。例如 OpenAI 向歐盟開放網路安全模型能力進行測試與合作，也反映出未來 AI 發展將更加重視安全治理與監管合作。另一方面，人工智慧競賽也正快速從「模型競賽」走向「生態系競爭」。隨著 AI 從軟體模型逐步擴展至機器人、自動駕駛與實體世界應用，算力、晶片、資料中心與硬體供應鏈的重要性持續提升。龐大的算力需求也帶動了半導體與高效能運算產業的快速成長，AI 的競爭焦點逐漸從單純的模型效能，擴展到晶片設計、運算基礎設施、供應鏈韌性與能源管理等更廣泛的領域。同時，各國政府也開始要求大型模型在公開發布之前進行安全測試與風險評估，以降低模型誤用、濫用或造成重大社會風險的可能性。因此，未來人工智慧產業的競爭，不再只是模型能力的競爭，而是算力、供應鏈、安全治理與監管能力的全面競爭。即使擁有最先進的模型與最強大的算力，如果缺乏完善的安全治理與風險管理機制，人工智慧的快速發展仍可能帶來巨大的社會與安全風險。

隨著人工智慧逐漸進入醫療場域，健康數位資料治理的重要性也日益提升，而資料



共治 (Data Co-governance) 正逐漸成為未來發展的重要方向。目前健康資料治理面臨幾項重要瓶頸。首先是資料割裂問題，健康資料分散於醫院、研究機構以及相關產業之間，缺乏互通機制，導致資料共享與整合困難，也限制了人工智慧模型的訓練與應用。其次是共享機制失效。現行資料共享多建立於機構間的合作關係與自願參與，缺乏明確的法律規範與經濟誘因，因此難以形成長期且穩定的資料流通模式。此外，現有法規也面臨新的挑戰。例如傳統隱私法規在面對基因資料、穿戴裝置資料以及連續生理監測資料等新型態健康資料時，逐漸出現適用上的限制。因此，未來健康資料治理的核心概念，首先是將健康資料視為公共利益相關的基礎設施，而非單一機構的私有資產。其次，資料共享模式將逐漸由資料集中轉向聯邦式架構，例如聯邦學習 (Federated Learning) 等技術，透過模型移動而非資料移動的方式，在兼顧隱私保護、資料安全與研究需求之間取得平衡。第三，則是建立透明且合理的補償機制，透過定價與回饋制度，確保資料提供者與參與者獲得合理的利益回饋。因此，健康資料治理的核心，其實是在個人隱私與公共利益之間尋求動態平衡。同時透過標準化資料格式與互通機制，例如 FHIR API 等國際標準，提升跨機構資料交換能力與系統互通性。此外，健康資料治理也需要納入病人、研究者、醫療機構以及政府等多方利害關係人，共同參與治理與監督機制的建立。未來發展方向則包括完善隱私保護法制、降低資料壁壘、建立合理的財務誘因，以及培養兼具人工智慧技術與倫理素養的治理人才，以建立可信任的健康資料共治模式。

然而，即使建立了完善的資料治理制度，高風險人工智慧應用仍然面臨另一項重要挑戰，也就是責任歸屬問題。以醫療人工智慧為例，當 AI 輔助手術或臨床決策導致病人傷害時，責任究竟應由誰承擔？是醫師、醫院、軟體開發商、人工智慧公司、機械手臂製造商，還是資料提供者？隨著人工智慧系統逐漸涉及多方參與者，傳統醫療責任架構也變得愈來愈複雜。因此，未來法律討論的重點，已不再只是「人工智慧是否犯錯」，而是如何建立一條完整、透明且可追溯的責任鏈。因此，《人工智慧基本法》第十七條特別強調，高風險人工智慧應用必須建立明確的責任歸屬、補償制度以及保險機制。尤其在醫療領域，若缺乏完善的責任分配與風險保障制度，即使人工智慧技術已經成熟，



也可能因法律風險與責任不明，而影響醫療機構與產業的導入意願。因此，建立可問責 (Accountable)、可補償 (Compensable) 以及可追溯 (Traceable) 的治理架構，將成為高風險醫療人工智慧發展的重要基礎。

智慧醫療 LDCT 篩檢治理應用實例

智慧醫療輔助工具在數位醫療平台助力下快速發展進入醫療應用領域，對臨床照護以及疾病防治模式也造成影響。以智慧平台結合數位醫療資料應用為例，全球廣泛使用的 Epic 敗血症預測模型，雖宣稱可依電子病歷即時偵測敗血症風險，但外部驗證結果顯示，其敏感度僅 33%、AUC 僅 0.63，遠低於原先公布的表現，代表許多真正的敗血症病人可能無法被及時辨識。此案例顯示 AI 模型若缺乏充分外部驗證便大規模導入臨床，可能造成誤判，醫療決策仍須由醫師審慎評估。除臨床照護電子平台智慧應用，疾病預防與篩檢如肺癌與乳癌等影像篩檢應用在智慧工具快速發展也建立創新模式。以肺癌篩檢為例，假定 64 歲陳先生曾有 35 包年的吸菸史，雖已戒菸 3 年，仍因年齡、吸菸暴露風險史符合肺癌高風險族群篩檢資格。該民眾去年 LDCT 結果為 Lung-RADS 2，在定期追蹤檢查接受 LDCT 篩檢。AI 輔助系統可快速分析 LDCT 影像，產生標準化草稿報告，協助醫師提升判讀效率並縮短作業時間。然而 AI 模型入侵以及資料保護也使此智慧診斷輔助工具成為易受如思維鏈挾持等入侵目標。當 AI 判讀擬稿與其他病例差異不明顯時，臨床人員若過度依賴系統內容，可能直接簽署報告，使病人誤以為病灶為良性，錯失進一步檢查或及早治療的時機，凸顯 AI 仍需結合臨床專業共同判讀。

前述陳先生肺癌篩檢該例於症狀出現後 22 個月才因乾咳加劇與血痰就醫，肺部診斷 CT 發現右上葉 28x23 mm 腫塊，並有肺門及縱膈淋巴結侵犯。FDG-PET-CT 顯示病灶高代謝，病理確診為中度分化肺腺癌，EGFR L858R 突變陽性，臨床分期為 Stage IIIA。此案例顯示早期可處理的肺部病灶若 AI 輔助診斷工具表現偏移則可能延誤追蹤，病灶進展為局部晚期肺癌。原可於 Stage IA3 早期確診並接受 VATS 肺葉切除與淋巴廓清，五年存活率約 80%，但在 22 個月後進展為 Stage IIIA 晚期病灶則則需同步放化療與標靶治療，五年存活率降至約 36%至 42%。顯示 AI 醫療工具必須重視資安、驗證與醫師覆



核。AI 醫療應用須兼顧安全、警示、醫師自主判讀與法遵機制。智慧診斷輔助工具應納入完整生命週期治理，從識別、保護、偵測、回應到復原，建立多層防線。流程上需盤點 AI 資料與分級、進行輸入過濾與人力規劃，並透過 PACS 持續比對、矛盾告警與人工覆核降低風險。對高風險結果應強制獨立判讀，並保存稽核軌跡，以達成監測、複核、究責與救濟，確保 AI 輔助診斷安全可靠。

以上內容將在 2026 年 7 月 1 日(三) 10:00 am 以線上直播方式與媒體朋友、全球民眾及專業人士共享。歡迎各位舊雨新知透過[星球永續健康網站專頁](#)觀賞直播！

- 星球永續健康網站網頁連結：

<https://www.realscience.top/7>

- Youtube 影片連結：<https://reurl.cc/o7br93>
- 漢聲廣播電台連結：<https://reurl.cc/nojdev>
- 不只是科技：<https://reurl.cc/A6EXxZ>



講者：

陳秀熙教授/英國劍橋大學博士、許辰陽醫師、陳立昇教授、嚴明芳教授、林庭瑀博士

聯絡人：

林庭瑀博士 電話：(02)33668033 E-mail：happy82526@gmail.com

劉秋燕 電話：(02)33668033 E-mail：r11847030@ntu.edu.tw