

人工智慧與醫療

金旻辰律師

前言

急診室螢幕上累積著數十個檢驗數值。醫師在這些數值、患者的表情、家屬的話語以及檢驗室傳來的影像之間作出判斷。人工智慧進入這一場景，快速讀取資訊並揭示容易忽視的訊號。但它並不代為承擔診療責任。

醫療 AI 已滲透至影像解讀、紀錄撰寫、風險預測、治療設計、新藥探索、醫院營運、教育與研究等領域。本書不僅止於羅列技術名稱，更深入探討醫院內部發生了什麼變化，以及需要哪些安全機制。

書中收錄的資料係根據多國語言講座與最新文獻整理而成。各項數值與建議在實際應用於患者之前，必須經由醫療團隊的判斷及最新臨床指南再次確認。

金炅辰律師

目錄

前言

第一部 了解 AI 醫療

- 第一章 AI 醫療導論 | 核心概念、臨床現場與安全邊界
- 第二章 AI 醫療的核心技術與資料 | 學習模型與醫療資料

第二部 AI 醫療的實踐

- 第三章 診斷與風險預測 | 診斷、預測與早期警示
- 第四章 治療計畫與精準醫療 | 個人化照護、手術與新藥開發
- 第五章 醫院業務與患者體驗 | 醫院營運與患者體驗
- 第六章 醫療教育與研究 | 教育、研究與臨床試驗

第三部 AI 醫療的爭議與未來

- 第七章 信任、安全與責任 | 信任、安全、隱私與責任
- 第八章 導入之壁與下一步 | 基礎設施、制度與下一步

第一部

了解 AI 醫療

第一章 AI 醫療導論

1.1 AI、機器學習、深度學習、自然語言處理

人工智慧（AI）在醫療領域過去數年間持續取得進展。初期從基於規則的專家系統開始，如今已大幅擴展至統計與神經網絡模型，乃至生成式 AI，其範圍與能力顯著提升。此種發展主要歸功於電子健康紀錄（EHR）、物聯網（IoT）設備、基因體資料庫等所產生的龐大資料量，以及處理這些資料之運算能力的大幅提升。AI 在醫學診斷、治療計劃、醫療紀錄管理以及與患者溝通等多元醫療領域帶來創新變革，成為提高醫療服務效率與準確性的核心技術。

機器學習（ML）與深度學習（DL）的角色

機器學習是 AI 的子領域，為專注於透過資料集經驗自我學習與改進的特徵識別及分析技術。深度學習則是機器學習的子類別，利用神經網絡（neural networks）學習更複雜的特徵並生成預測。這些技術透過分析龐大的臨床資料集，實現對疾病進展、治療結果及患者風險因素的預測，最終支援個人化治療與診斷。在醫療領域中，ML 與 DL 在影像診斷、病理學、基因體學等各個領域展現出卓越性能，能識別並預測人眼難以察覺的細微特徵。

自然語言處理（NLP）與生成式 AI

自然語言處理（NLP）是 AI 系統理解、評估及生成人類語言的技術，在醫療領域中對於從非結構化臨床文件中擷取寶貴洞察至關重要。近年來，隨著生成式 AI（Generative AI）及大型語言模型（LLM）的發展，NLP 的應用範圍進一步擴大。LLM 經過數十億個詞彙與影像的訓練，能夠識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能不知疲倦地處理龐大資料。藉此，可從醫師筆記、出院摘要、放射科報告及科學出版物等資料中自動擷取資訊，並實現如臨床試驗配對與醫療論文摘要等高級應用。

在醫療現場診斷領域的應用

在醫療紀錄與文件化領域的應用

AI 能減輕醫療團隊過重的行政業務負擔，並大幅提升紀錄效率。特別是生成式 AI 能大幅縮短醫療紀錄撰寫時間。在醫師或護理師與患者交談時，AI 將語音轉換為文字，並自動記錄為符合醫療資訊需求格式的 AI 聽寫（AI Scribe）技術正導入臨床現場。這使得醫療團隊能減少用於文件作業的時間，更專注於患者診斷與互動。例如，出院摘要或轉診單等醫療文件可由 AI 在 20 秒內生成，並可用於將診察室內的對話轉為文字及生成包含圖表的說明。護理紀錄亦可透過 AI 予以減輕，基於智慧型手機的 AI 可將護理紀錄撰寫時間減少高達 70%。AI 還能提供 ICD（國際疾病分類）編碼建議，以簡化保險申報相關業務。LLM 還可用於臨床試驗摘要或 FDA 法規文件（dossier）撰寫，FDA 甚至自 2025 年 5 月起批准利用 LLM 模型撰寫最終文件。

在患者溝通與教育領域的應用

AI 聊天機器人與虛擬助理在改善患者溝通與教育方面扮演重要角色。患者可透過 AI 聊天機器人提出健康相關問題，並獲得個人化的健康資訊與建議。有趣的是，一項研究顯示，患者評估 AI 聊天機器人提供的回答比人類醫師更有同理心且具憐憫心。AI 可考量患者的年齡、背景、識字程度、社會經濟條件等，以個人最適化的方式提供資訊，從而提高理解度。特別是在夜間時段對個人健康問題的諮詢需求較多，AI 聊天機器人在填補健康資訊落差方面發揮著重要作用。

AI 護理師助理可為患者執行導覽、預防跌倒、情緒支持等重複性業務，協助實際護理師為患者提供更多人性化照護。AI 亦能超越醫療專家的解釋能力，回答患者關於檢驗結果或疾病的常見問題。然而，AI 所提供資訊的準確性與可靠性始終至關重要，且因 AI 的幻覺（hallucination）現象而存在提供錯誤資訊的風險。AI 可能犯下 1% 的致命錯誤，這使得醫療專家的批判性審查與監督成為必需。由於 AI 目前仍難以捕捉人類的非語言資訊與身體性，因此在理解複雜的人類情感與語境方面存在局限。

挑戰課題與倫理考量事項

醫療 AI 的發展雖為醫療領域帶來革命性變化，但仍存在資料隱私、安全、演算法偏見以及法律與倫理問題等待解決的課題。當 AI 做出誤診或給出錯誤建議時，最終法律責任歸屬於醫療專家。因此，AI 系統的透明度與可信度至關重要，需要努力解決難以理解 AI 系統如何做出決策的「黑盒子（black box）」問題。AI 並非取代醫療專家，而是作為「聰明且不知疲倦的同僚」來輔助醫療團隊的工具。當人類的專業知識與 AI 的能力相結合時，能創造出最佳結果，醫療團隊應學習 AI 的使用方法並培養批判性評估其產出的能力。

[參考資料]

1. 標題：Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者・機構：MédicosÉlite®

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'MédicosÉlite®' (來源中未提供影片直接 URL)

支持的主張：AI 經歷了傳統基於規則的 AI、機器學習、生成式 AI（generative AI）三個階段的發展，LLM（大型語言模型）經過數十億個詞彙與影像訓練，能夠識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能處理龐大資料。AI 並非取代醫師，而是作為聰明且不知疲倦的同僚支援醫療團隊的工作，最終決定與責任始終由醫師承擔。

2. 標題：生成AI事例☒討会part2 ～診療・研究・醫療DXへの実☒編～

作者・機構：Kamijo Kyosuke

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'Kamijo Kyosuke' (來源中未提供影片直接 URL)

支持的主張：AI 可在 20 秒內生成出院摘要（discharge summary）文件，並能將醫師的語音紀錄立即轉換為文件，連同圖表轉化為說明資料提供給患者。

3. 標題：首創AI醫療影像3D標準化 ☒ MR 應用成果發表大會

作者・機構：智捷醫學科技

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel '智捷醫學科技' (來源中未提供影片直接 URL)

支持的主張：智捷醫學科技（Zijie Smart Medical Technology）的 AI「Anatomy Cloud」可在 3 至 5 分鐘內處理 3D 醫療影像並辨識 200 個以上的器官，未來還包含自動識別病變或腫瘤的功能，實現從 2D 黑白影像向 3D 彩色影像的轉變。

4. 標題：Conferencias IA y Salud.- IA en la clínica. El fin del médico que conocíamos

作者・機構：Careexpand Media

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'Careexpand Media' (來源中未提供影片直接 URL)

支持的主張：聊天機器人能提供比人類醫師更有同理心的回答，在 16 歲至 19 歲的年輕人中，約有 90% 利用 AI 來解決健康問題。

1.2 醫療 AI 的歷史與現狀

人工智慧（AI）在醫療領域持續發展，其歷史從早期的基於規則系統到如今高階的統計與神經網絡模型，再到生成式 AI，充斥著變革與創新。患者資料的爆發性增長與運算性能的提升，是加速 AI 技術發展與應用的核心動力。特別是醫療 AI，在提升診斷準確度、最佳化治療計劃、減輕行政業務、改善患者溝通等多元領域，為提高醫療服務品質作出貢獻。

基於規則的專家系統之胎動

醫療 AI 的萌芽可追溯至 1970 年代開發的基於規則專家系統。代表性例子包括 1971 年的 INTERNIST-1 與 1970 年代的 MYCIN，旨在支援診斷決策。這些早期 AI 系統以特定規則為基礎（例如「若體溫在 38 度以上則為發燒」），具有固定的邏輯結構。此類系統宛如遵循預先編程的決策樹，雖在靈活性方面有所限制，卻展現了 AI 應用於醫療現場的可能性。

電子健康紀錄（EHR）與醫療影像資料的普及

隨著時間推移，AI 開始從基於知識的程式轉型為統計與神經模型。在此轉型的背後，電子健康紀錄（EHR）、物聯網（IoT）設備、基因體資料庫等產生的龐大資料量，以及處理這些資料之運算能力的大幅提升發揮了關鍵作用。EHR 系統將患者的醫療履歷數位化，在建立廣泛臨床資料庫方面起到了決定性作用。以往由護理師手寫紀錄的血糖與血壓資料，改由感測器設備自動傳送至 HISS 系統，使醫師能遠端確認患者資料，此案例展現了此種變化如何提高醫療資料管理的效率。在美國大學醫院中，AI 整合至 EHR 中以支援醫療團隊的診療活動已十分常見。

醫療影像資料亦成為 AI 發展 重要 的基礎。隨著 AI 分析 X 光、CT、MRI、病理切片等各種影像資料，學習複雜特徵與感知人眼容易忽略的細微異常徵兆的能力顯著提升。傳統上需要 10 至 30 分鐘的影像分析作業，透過 AI 縮短至 3 至 5 分鐘以內，極大化了診斷過程的效率。如智捷醫學科技（Zijie Smart Medical Technology）的「Anatomy Cloud」等 AI 解決方案可在 3 至 5 分鐘內處理 3D 醫療影像並辨識 200 個以上的器官，未來甚至計畫包含自動識別病變或腫瘤的功能。這使得從 2D 黑白影像向 3D 彩色影像的轉變成為可能。

機器學習（ML）與深度學習（DL）的臨床導入

機器學習（ML）是 AI 的子領域，為專注於透過資料集經驗改進機器的特徵識別及分析技術。深度學習（DL）是機器學習的子類別，利用神經網絡學習複雜特徵並生成預測。這些技術從龐大的臨床資料集中挖掘隱藏特徵，預測疾病進展、治療結果、患者風險因素等，實現個人化治療與診斷。例如，學習了數百萬張 X 光影像的 AI，能以高於人眼的準確度辨識乳癌。此種深度學習技術的演進自 2000 年代起尤為顯著，機器學習與深度學習的出現標誌著 AI 第三階段的開始。

生成式 AI 與大型語言模型（LLM）的臨床導入

近年來 AI 領域最具革新性的發展是生成式 AI（Generative AI）與大型語言模型（LLM）的出現。這些技術賦予 AI 自然理解、生成人類語言並掌握複雜醫療語境的能力。LLM 經過數十億個詞彙與影像訓練，能夠識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能不知疲倦地處理龐大資料。

在醫療現場的應用案例：

診斷：AI 可從 X 光、CT、MRI 等醫療影像中感應細微特徵，從而早期發現疾病。例如，梅奧診所（Mayo Clinic）的 AI 模型已被證實可用於比臨床診斷最多提前 3 年檢測出胰臟癌。在腦中風等急診狀況下，AI 可在 3 至 5 分鐘內處理複雜的影像資料並向醫療團隊發出警報，協助迅速做出治療決定。研究結果亦顯示，在特定類型皮膚癌診斷中，AI 展現出高於皮膚科專科醫師的準確度。AI 還能根據個人健康狀況生成個人化風險輪廓，並透過對疾病進展與治療結果的預測，支援早期干預與預防性治療。

紀錄與文件化：生成式 AI 能大幅減輕醫療團隊的行政負擔。在醫療團隊與患者交談時，AI 將語音轉換為文字，並自動記錄為符合醫療資訊需求格式的 AI 聽寫技術已導入臨床現場，協助節省文件作業時間並更專注於患者診療。出院摘要或轉診單等醫療文件可由 AI 在 20 秒內生成，並用於將診察室內的對話轉為文字及生成包含圖表的說明。這有助於提高效率與紀錄品質。此外，LLM 還可用於臨床試驗摘要或法規文件撰寫。

患者溝通與教育：AI 聊天機器人與虛擬助理用於改善患者溝通與教育。患者可透過 AI 聊天機器人提出健康相關問題，並獲得個人化的健康資訊與建議。AI 考量患者的年齡、背景、識字程度等提供最適化資訊，從而提升患者的理解度。例如，亦有研究結果顯示，接受抗癌治療患者的不安與憂鬱感可透過與 AI 聊天機器人的對話予以降低。AI 護理師助理為患者執行導覽、預防跌倒、情緒支持等重複性業務，協助實際護理師為患者提供更多人性化照護。AI 亦應用於支援自我管理的應用程式中，如向患者提供用藥提醒及監測治療順從性等。

挑戰課題與未來展望

醫療 AI 的發展雖帶來醫療服務的革新，但仍存在資料隱私、安全、演算法偏見以及法律與倫理問題等待解決的課題。特別是在生成式 AI 的情況下，因幻覺（hallucination）現象而存在提供錯誤資訊的風險，這使得醫療專家的批判性審查與監督成為必需。AI 並非取代醫療專家，而是應作為增強其能力並減輕業務負擔的聰明且不知疲倦的同僚發揮作用。當人類專業知識與 AI 能力和諧結合時，醫療 AI 將創造最佳結果並為實現以患者為中心的醫療作出貢獻。

[參考資料]

1. 標題：Artificial Intelligence in Clinical Trials | Drug Development, AI Tools & Careers

作者・機構：IICRS (Clinical Research and Studies)

可確認的發行日・上傳日：2025 年 5 月（內容中提及 FDA 批准）

URL: YouTube channel 'IICRS (Clinical Research and Studies)' (來源中未提供影片直接 URL)

支持的主張：LLM（大型語言模型）可支援臨床試驗摘要或 FDA 法規文件（dossier）撰寫。

2. 標題：Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者・機構：MédicosElite®

可確認的發行日・上傳日：來源摘錄中未標明

URL: YouTube channel 'MédicosElite®' (來源中未提供影片直接 URL)

支持的主張：AI 經歷了傳統基於規則的 AI、機器學習、生成式 AI（generative AI）三個階段的發展，生成式 AI 可在數秒內閱讀醫療紀錄，連結複雜症狀並提出新診斷與實證協定。LLM 經過數十億個詞彙與影像訓練，能夠識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能不知疲倦地處理龐大資料。AI 並非取代醫師，而是作為聰明且不知疲倦的同僚支援醫療團隊的工作。

3. 標題：人工智慧基本法與智慧醫療治理趨勢（上）李崇僖教授

作者・機構：李崇僖教授

可確認的發行日・上傳日：來源摘錄中未標明

URL: YouTube channel '理律文教基金會' (來源中未提供影片直接 URL)

支持的主張：生成式 AI 在醫療團隊與患者交談時將語音轉換為文字以自動生成醫療紀錄，從而減輕紀錄業務負擔。護理紀錄業務亦可透過 AI 予以減輕。AI 透過歷史資料預測未來的能力非常出色。

4. 標題：生成AI事例☒討会part2 ～診療・研究・醫療DXへの実☒編～

作者・機構：Kamijo Kyosuke

可確認的發行日・上傳日：來源摘錄中未標明

URL: YouTube channel 'Kamijo Kyosuke' (來源中未提供影片直接 URL)

支持的主張：AI 可在 20 秒內生成出院摘要（discharge summary）文件，並能將醫師的語音紀錄立即轉換為文件，連同圖表轉化為說明資料提供給患者。

1.3 醫療現場使用 AI 的原因

在醫療現場加速導入人工智慧（AI）的背景下，存在著多種複合因素的共同作用。急劇增加的醫療需求、有限的醫療資源以及醫療專家過重的業務負擔等，更加凸顯了 AI 技術的必要性。AI 在提高診斷與治療效率、實現以患者為中心的醫療，以及提升醫療系統整體品質方面具有貢獻的潛力。

高齡化與慢性疾病的增加

全球邁入高齡化社會，患有慢性疾病的患者人數急劇增加。這是導致醫療服務需求暴增的主要原因之一。隨著患者擁有更複雜且長期的醫療需求，僅憑傳統醫療系統已難以應付此類需求。AI 能為此類問題提供解答。AI 可透過分析龐大的臨床資料，生成符合個人健康狀況的個人化風險輪廓，並透過對疾病進展與治療結果的預測，支援早期干預與預防性治療。例如，基於 AI 的預測分析能成功識別慢性疾病發病風險高的患者，並引導其進行有效的健康計畫。此外，AI 聊天機器人可回答健康相關問題並提供個人化健康資訊與建議，協助患者自主管理健康。特別是在夜間時段對個人健康問題的諮詢需求較多，AI 聊天機器人在填補健康資訊落差方面發揮著重要作用。

醫療團隊不足與業務負擔減輕

醫療團隊短缺是全球性現象，導致留任醫療團隊的業務負擔日益加重。醫師與護理師除患者診療外，還飽受龐大行政業務困擾，這導致職業倦怠與工作滿意度下降。AI 能減輕醫療團隊的過重業務，協助其將更多時間分配給患者。例如，AI 能自動化重複性的診療紀錄撰寫、出院摘要、轉診單生成等工作，大幅減少醫療團隊的文件作業時間。藉此，醫療團隊能更專注於與患者溝通及直接診療。此外，AI 護理師助理執行患者導覽、預防跌倒、情緒支持等重複性護理業務，協助實際護理師為患者提供更多人性化照護。

紀錄業務的效率化

醫療現場龐大的紀錄業務消耗醫療團隊大量時間，這可能影響患者診療品質。AI 在自動化與簡化此類紀錄業務方面發揮核心作用。AI 聽寫技術在醫療團隊與患者交談時將語音轉換為文字，並自動記錄為醫療資訊系統（EHR/EMR）所需的格式。藉此，醫療團隊可減少投入文件作業的時間，目視患者進行溝通。例如，如出院摘要等文件可由 AI 在 20 秒內生成，這大幅減輕了行政負擔。護理紀錄亦可透過 AI 予以減輕。患者管理軟體搭載 AI 以分析患者的歷史紀錄，自動摘要下次診療所需的資訊，從而簡化文件化流程。

解決診斷延誤問題

消除區域醫療落差

區域間醫療資源的不均衡引發嚴重的醫療可及性問題。特別是在偏遠地區或醫療匱乏地區，由於缺乏專業醫療團隊與先進設備，患者難以即時獲得適當診療。AI 在消除此類區域醫療落差方面可發揮重要作用。將 AI 與可攜式診斷設備相結合，由 AI 解讀現場拍攝的醫療影像以輔助診斷，並將其傳送至合作醫院的專業醫療團隊，從而實現遠距診療與診斷。藉此，即使在醫療資源匱乏的地區，也能讓更多人獲得高品質的醫療服務。AI 透過提高診斷準確度並縮短疾病發現時間，協助以有限的醫療人力為更多患者提供必要的診療。

醫師的最終責任與 AI 的互補角色

AI 雖為醫療領域帶來革新，但醫療專家的最終判斷與責任依然至關重要。AI 並非取代醫師，而是作為「聰明且不知疲倦的同僚」輔助醫療團隊的工具。AI 生成的診斷或治療建議必須經過醫療團隊的批判性審查與監督。AI 具有資料隱私、安全、演算法偏見以及幻覺（Hallucination）現象等局限，

[參考資料]

1. 標題：Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者・機構：MédicosÉlite®

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'MédicosÉlite®' (來源中未提供影片直接 URL)

支持的主張：AI 並非取代醫師，而是作為聰明且不知疲倦的同僚支援醫療團隊，最終決定與責任始終由醫師承擔。大型語言模型（LLM）不知疲倦地全天候 24 小時處理龐大資料，識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能提出每 73 天更新一次的最新實證協定，提供人類難以跟上的最新資訊。

2. 標題：人工智慧基本法與智慧醫療治理趨勢（上）李崇僖教授

作者・機構：李崇僖教授

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel '理律文教基金會' (來源中未提供影片直接 URL)

支持的主張：AI 在醫療團隊與患者交談時將語音轉換為文字，自動生成符合醫療資訊需求格式的紀錄，從而大幅減輕醫師與護理師的紀錄業務負擔。護理紀錄亦可透過 AI 予以減輕。AI 卓越的預測能力可提高醫院內動線、病床分配等院務指揮中心的效率。隨著 AI 導入提升醫療水準，若醫療團隊疏忽了 AI 警告的風險而發生事故，醫療團隊的責任範圍可能會擴大。

3. 標題：Ponente Internacional: Seminario Internacional Tecnología médica e IA 2025

作者・機構：PhD. Antony Paul Espiritu Martinez

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'PhD. Antony Paul Espiritu Martinez' (來源中未提供影片直接 URL)

支持的主張：考慮到高患者需求與醫療系統的弱點，AI 有助於節省時間與資源，醫療正從反應式接近方式轉型為由 AI 所賦能的預測性與個人化接近方式。AI 根據基因分析預測未來疾病，最佳化包含病床管理及患者轉診/回診流程在內的醫院工作流程，並能透過超精密影像演算法實現早期癌症發現。

4. 標題：【即便如此最好還是不要向AI進行健康諮詢】精通AI應用的醫師・大塚篤司／品質與同理心AI皆超越醫師／但因AI指示發生中毒患者／無法將「1%致命錯誤」降為零【lon1 Health】

作者・機構：TBS CROSS DIG with Bloomberg

可確認的發行日・上傳日：（來源中未標明）

URL: YouTube channel 'TBS CROSS DIG with Bloomberg' (來源中未提供影片直接 URL)

支持的主張：AI 的診斷能力正大幅提升，在特定皮膚疾病診斷中甚至展現出高於人類皮膚科專科醫師的準確度。患者評估 AI 聊天機器人提供的回答比人類醫師更有同理心且具憐憫心，AI 諮詢有助於減輕不安感與憂鬱感。然而，AI 因「幻覺（hallucination）」現象可能提供錯誤或有害資訊（例如建議攝取毒性物質），對患者診療的最終責任依然由人類醫師承擔。

第二章 AI醫療的核心技術與資料

2.1 監督學習、非監督學習、深度學習、LLM

人工智慧（AI）在醫療領域歷經長期發展，初期從如 INTERNIST-1 (1971) 及 MYCIN (1970年代) 等基於規則的醫療專家系統開始，演變至今日的統計與神經網路模型。在此發展背景下，除了來自電子健康紀錄（EHR）、物聯網（IoT）裝置、基因體資料庫等產生的龐大資料量外，還有算力的革新性提升。AI 現今在醫療診斷、治療計畫、紀錄管理、病患溝通等多個醫療領域扮演關鍵角色，並為提高醫療服務的效率與準確性作出貢獻。

機器學習（Machine Learning, ML）的基本概念

機器學習是 AI 的子領域，專注於透過資料集的經驗自行學習與改善，並識別與分析模式的技術。在醫療領域中，ML 被廣泛應用於疾病分類、病患風險分層、治療結果預測等。ML 模型主要分為兩種學習方式。第一，監督學習（Supervised Learning）是使用手動標記的訓練資料來訓練演算法的方式。例如，利用在 MRT 影像中手動標註腫瘤的資料，訓練演算法精準識別腫瘤。第二，非監督學習（Unsupervised Learning）是在沒有標籤的資料中獨立識別模式或異常徵兆的方式。這對於發現新的疾病亞型或病患群體（cohort）非常有幫助。此外，雖然也存在強化學習（Reinforcement Learning）等其他學習範式，但在臨床環境中使用頻率較低。ML 透過審視龐大且複雜的醫療資料集，提升醫療決策的準確性與個人化。

深度學習（Deep Learning, DL）的出現與影響

深度學習是機器學習的子類別，是使用深層神經網路（deep neural networks）來學習複雜模式並生成預測的技術。DL 模型從原始資料中逐步擷取特徵，以學習複雜且階層式的表示。雖然這主要在帶有標籤資料的監督學習環境中使用，但近期也開發出了較少需要或完全不需要明確標籤的自監督及非監督方式。深度學習方法的開發為醫療影像分析領域帶來了革命，在 X 光、CT 掃描、MRI、病理切片等多種影像模態中檢測異常徵兆方面取得了令人矚目的成功。DL 使醫療影像的精準診斷成為可能，並有助於識別肉眼難以察覺的微妙模式。

自然語言處理（Natural Language Processing, NLP）與大語言模型（Large Language Models, LLM）

自然語言處理（NLP）是 AI 系統理解、評估及生成人類語言的技術，對於從非結構化臨床文件中擷取有價值的洞察至關重要。NLP 工具已透過醫療紀錄自動化、工作流程最佳化、結構化臨床決策支援等應用於護理領域。近期，利用基於 Transformer 的神經網路的大語言模型（LLM）的發展尤為顯著。如 ChatGPT 或 GPT-4 等 LLM 經過數十億個詞彙與影像的訓練，能夠識別醫療文獻、臨床案例、研究協定，理解複雜的醫療語境，並能不知疲倦地處理龐大資料。

生成式 AI（Generative AI）的臨床應用

生成式 AI 是指能夠基於所學習資料生成新內容的模型，並在各種臨床應用領域中崛起為具有前景的工具。生成式 AI 被應用於 AI 語音聽寫（AI scribe）技術中，能在醫師與病患對話時將語音轉換為文字，並自動記錄為符合醫療資訊需求的格式，大幅減少醫療紀錄撰寫時間。此外，它還能在 20 秒內生成出院摘要或轉診單等醫療文件，並用於將診察室對話轉錄為文字及生成包含圖表的說明。LLM 亦可用於支援臨床試驗摘要或 FDA 監管文件的撰寫。

LLM 在診斷及病患溝通中的角色

LLM 在診斷決策支援方面亦展現出強大潛力。學習了數百萬筆真實醫療資料的 AI 醫療助手，能針對病患的症狀（例如：腹瀉 3 次）提供合理的解答與治療方法（例如：糞便檢查、少食多餐、補充水分、制酸劑、抗病毒藥物）。此外，對於心臟疼痛或打嗝等各種症狀，也能提出治療方案或緩解方法。LLM 可用於病患提出健康相關問題並獲取個人化資訊，部分病患甚至評價 AI 聊天機器人比人類醫師更有同理心。LLM 能超越醫療人員的說明，回答病患關於檢查結果或疾病的問題，在夜間時段發揮彌補健康資訊落差的重要作用。

挑戰課題與倫理考量事項

儘管 LLM 功能強大，但仍存在幻覺（hallucination）現象，即生成不準確或捏造資訊的風險。此類錯誤使得醫療專業人員的批判性審視與監督成為必需。AI 系統可能犯下 1% 的致命錯誤，這暗示著醫療人員不應盲目遵循 AI 的建議，而應始終優先考慮自己的專業判斷。AI 並非取代醫師，而是作為「聰明且不知疲倦的同事」來輔助醫療人員的工具。在技術發展的同時，資料隱私、安全、演算法偏見以及法律與倫理問題等待解決的課題依然存在。

[根據資料]

1. 標題: 大模型醫療問答机器人問目（大模型/deepseek基問知問/人工智慧/AI/rag/大模型基問路問）LangChain//LangChain入門LangChain問用問发

作者・機構: 大模型入门

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel '大模型入门' (來源未提供影片的直接 URL)

支持的主張: AI 醫療助手能基於數百萬筆真實醫療資料，針對病患的症狀（例如：腹瀉 3 次、心臟疼痛、打嗝）提出合理的解答與治療方案。由於 LLM 可能引發幻覺（hallucination），因此不應盲目信任 AI 提供的解答，醫師或使用者必須進行批判性審視。LLM 可用於支援臨床試驗摘要或 FDA 監管文件的撰寫。

2. 標題: Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者・機構: MédicosElite®

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel 'MédicosElite®' (來源未提供影片的直接 URL)

支持的主張: AI 歷經傳統基於規則的 AI、機器學習、生成式 AI（generative AI）三個階段發展而來，LLM（大語言模型）經過數十億個詞彙與影像的訓練，能夠識別醫療文獻、臨床案例、研究及協定，理解複雜的醫療語境，並能不知疲倦地處理龐大資料。生成式 AI 能在數秒內閱讀醫療紀錄，連動複雜症狀並提出新診斷與基於實證的協定。

3. 標題: KI: Forschung, Innovation und klinische Praxis

作者・機構: thieme-connect.de

可確認的發行日・上傳日: (來源未明示)

URL: <http://www.thieme-connect.de/products/ejournals/pdf/10.1055/a-2741-9717.pdf>

支持的主張: AI 模型可分為監督學習（使用 annotated data）、非監督學習（自行學習模式）、深度學習（使用深層神經網路）。深度學習主要用於監督學習環境中，但亦開發出了自監督及非監督方式。LLM（ChatGPT、GPT-4 等）屬於基於 Transformer 的網路，用於處理複雜語言。

2.2 電子病歷、影像、基因體、穿戴式裝置資料

人工智慧（AI）透過龐大醫療資料的運用發揮其最大潛力，此類資料是 AI 模型學習與改善的核心動力。電子健康紀錄（EHR）、醫療影像、基因體資料以及從穿戴式裝置收集的資料，是 AI 推動醫療診斷、治療計畫、病患管理革新的必備要素。AI 整合並分析這些多元資料來源，為醫療專業人員提供寶貴洞察，從疾病早期發現到個人化精準醫療，改善整體醫療過程。然而，此類資料本質上十分複雜，預處理過程不可或缺，缺漏值與錯誤以及對臨床語境的理解，對於 AI 模型的成功應用具有決定性影響。

電子病歷（Electronic Health Records, EHR）資料

資料性質：電子病歷（EHR）是包含病患全面醫療紀錄、檢查結果、治療文件等的數位資料庫。其中包含醫師的診診筆記、出院摘要、放射科報告等非結構化文字資料，亦包含診斷代碼、藥物處方、生命徵象測量值等結構化資料。這些資料被整合並應用於醫院的臨床決策支援系統中。

預處理：為了運用 EHR 資料，精細的預處理過程不可或缺。特別是為了從非結構化文字資料中擷取有意義的資訊，使用了自然語言處理（Natural Language Processing, NLP）技術。NLP 能從醫師的診診紀錄中擷取症狀說明或藥物調整等臨床相關資訊，以支援實時臨床決策。此外，為了保護病患隱私，必須經過資料匿名化（anonymization）過程。匿名化是去除非公開個人識別資訊，使資料無法連結至特定個人的過程。

缺漏值與錯誤：EHR 資料常因異質（heterogeneous）基礎設施而分散，系統間缺乏可互操作性（interoperability）是一大問題。這使得資料的存取與整合變得困難。此外，EHR 資料還存在品質低下、偏見（bias）以及缺漏值（missing values）等問題。例如，特定人口群體（例如：女性、特定種族）在資料集中代表性不足，或因錯誤的診斷編碼或不完整紀錄而導致資料品質下降。此類資料偏見在 AI 演算法應用於真實臨床環境時，可能會降低診斷準確度。

臨床語境：AI 能分析 EHR 與實時資料，比傳統方式更快速、更精準地檢測疾病，並能監測藥物交互作用並向醫療人員發出實時潛在副作用警告。此外，基於 AI 的 EHR 能改善疾病監測、支援人口健康管理，並透過預測住院人數或工作人員工作量需求來最佳化資源分配。AI 能內建於病患管理軟體中，分析病患過去紀錄，自動摘要下次診診所需資訊，從而簡化文件化過程。藉此減輕醫療提供者的文件化負擔，大幅提升醫療工作流程。

醫療影像資料

資料性質：醫療影像是 AI 發揮最強大效能的領域之一。包含 X 光、CT 掃描、MRI、病理切片、超音波影像等多種影像資料。這些資料具有高維度且複雜，且常存在標註不足（poorly annotated）的情況。特別是 3D 醫療影像包含更多資訊，AI 能對其進行分析，突破 2D 黑白影像的局限。

預處理：醫療影像資料需經過旨在減少雜訊（noise）與偽影（artifacts）的正規化（normalization）、切割特定區域的分割（segmentation），以及提升影像品質的重構（reconstruction）等預處理過程。為了訓練 AI 模型，標註（annotation）工作（即手動標示病變位置或特徵）不可或缺，這需要大量時間與專業知識。此外，醫療影像資料同樣出於個人隱私保護而進行匿名化（anonymization）處理。

缺漏值與錯誤：醫療影像資料集可能存在資料不平衡（imbalance）或偏見（bias）問題。例如，若皮膚癌診斷 AI 模型的訓練資料集主要由淺色皮膚影像組成，則在診斷深色皮膚病患時可能會出現效能下降或犯錯。此類種族偏見（racial bias）會降低 AI 診斷結果的可信度。醫療影像資料本身的雜訊、低解析度、病患間的可變性等，亦對 AI 方法論提出了挑戰。

臨床語境：AI 能檢測醫療影像中的微妙模式，早期發現肉眼容易忽略的疾病。AI 能將傳統上需要 10 分鐘至 30 分鐘的醫療影像處理工作在 3 分鐘至 5 分鐘內完成，協助在腦中風等急診情況下快速作出治療決策，有助於將病患腦神經細胞損傷降至最低。AI 在肺結節分類與檢測、骨折檢測、胰臟癌早期偵測等方面已成功應用。Zijie Smart Medical Technology 的「Anatomy Cloud」可在 3 至 5 分鐘內處理 3D 醫療影像，識別 200 個以上的器官，並自動識別病變或腫瘤，實現從 2D 黑白影像向 3D 彩色影像的跨越。AI 能協助醫師進行術前演練、規劃精準的手術路徑，並掌握病變的精準位置。

基因體資料

資料性質：基因體資料包含個人的 DNA 序列資訊、基因變異（variants）等，在精準醫療（precision medicine）的發展中扮演關鍵角色。透過大規模基因體測序計畫，已發現數百萬個基因體變異，其中相當一部分可能與疾病相關。

預處理：AI 能分析基因體資料以識別致病變異體，並實現基於分子標記的病患分層。能提高基因分析中尋找變異體的速度，部分系統的準確度達到 93%。此外，亦能協助透過基因檢測判斷變異的嚴重程度。基因體資料屬於敏感個人資訊，因此確保資料安全的透明資料治理與匿名化不可或缺。

缺漏值與錯誤：基因體資料極為龐大且複雜，需要 AI 輔助分析 [1 確保資料安全的透明資料治理與匿名化不可或缺。

缺漏值與錯誤：基因體資料極為龐大且複雜，需要 AI 輔助分析，錯誤的變異解讀可能導致病患誤診。資料的完整性（integrity）與品質（quality）至關重要。此外，生物特徵資料的資料安全必須受到極高重視。

臨床語境：基於 AI 的基因體分析有助於預測未來疾病，加速新藥開發與遺傳學研究。提供建立個人化精準醫療策略所需的資訊，並透過基因分析以高概率預測未來疾病發生可能性，從而實現早期干預。

穿戴式裝置及物聯網（IoT）資料

資料性質：從穿戴式裝置與 IoT 感測器收集的資料擴大了 AI 在病患監測與個人化健康管理中的應用。其中包含睡眠品質、步數、心率變異性、壓力程度等多種健康指標，血糖與血壓等生命徵象資料亦可實時收集。

預處理：AI 能利用從穿戴式裝置獲取的資料分析各種健康指標，並以此為基礎提出改善個人生活方式的客製化建議。IoT 裝置能將血糖與血壓等生命徵象資料自動傳送至 HISS（Hospital Information System）系統，使醫療人員能夠遠距實時查看病患資料。

缺漏值與錯誤：穿戴式資料同樣可能發生資料品質問題。因裝置故障、配戴錯誤、網路連線不穩定等，可能導致資料不準確或缺漏。資料的此類不完整性會影響 AI 模型的可靠性。

臨床語境：基於 AI 的穿戴式資料分析能向病患提供個人化生活方式改善建議，有助於提升睡眠品質、減輕壓力等。AI 聊天機器人能回答健康相關提問並提供個人化健康資訊與建議，在夜間時段等醫療服務可及性受限的時間，發揮填補健康資訊落差的重要作用。AI 能提供提升治療順從性的提醒功能，並用於支援病患自我管理。特別是在日本的照護機構中，已建立起透過 IoT 裝置將生命徵象傳送至雲端醫院，由 AI 檢測異常徵兆並在必要時連結救護車或啟動居家醫療團隊的心臟衰竭監測系統。

結論

為了醫療 AI 的發展，包含 EHR、醫療影像、基因體、穿戴式資料在內的多元形式資料不可或缺。每種資料類型皆具有獨特性質與預處理挑戰，資料品質、偏見、缺漏值等問題直接影響 AI 模型的準確性與可靠性。為了有效運用此類資料，必須伴隨資料治理、品質管理、確保整合與可互操作性以及倫理考量。此外，亦須認知到 AI 模型提供資訊的幻覺（hallucination）風險或資料偏見等局限，並在醫療專業人員的最終判斷與責任下謹慎使用。

[根據資料]

1. 標題: 大模型醫療问答机器人目录（大模型/deepseek基知/人工智能/AI/rag/大模型基路）LangChain//LangChain入门LangChain用发

作者・機構: 大模型入门

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel '大模型入门' (來源未提供影片的直接 URL)

支持的主張: AI 醫療助手能基於數百萬筆真實醫療資料，針對病患症狀提出合理的解答與治療方案，但由於 LLM 可能引發幻覺（hallucination），因此不應盲目信任 AI 提供的解答。

2. 標題: AI applications in medicine and healthcare Integration of AI in electronic health records

作者・機構: eurjmedres.biomedcentral.com

可確認的發行日・上傳日: 2025 (URL中明示)

URL: <https://eurjmedres.biomedcentral.com/counter/pdf/10.1186/s40001-025-03196-w>

支持的主張: AI 能分析 EHR 與實時資料，比傳統方法更快速、更精準地檢測疾病，並能監測藥物交互作用並實時警告潛在副作用。ML 模型利用 EHR 預測院內死亡率、再入院率、敗血症發生等臨床結果，在處理非結構化 EHR 資料（例如：自由文字臨床紀錄）時達到 85% 以上的預測準確度。

3. 標題: 首創AI醫療影像3D標準化 MR 應用成果發表大會

作者・機構: 智捷醫學科技

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel '智捷醫學科技' (來源未提供影片的直接 URL)

支持的主張: Zijie Smart Medical Technology 的 AI 「Anatomy Cloud」可在 3 至 5 分鐘內處理 3D 醫療影像，識別 200 個以上的器官，未來將包含自動識別病變或腫瘤的功能，實現從 2D 黑白影像向 3D 彩色影像的跨越。

4. 標題: AI 基因報告判讀、5G 遠距應用與臨床倫理研習會

作者・機構: TCnews慈善新聞網

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel 'TCnews慈善新聞網' (來源未提供影片的直接 URL)

支持的主張: 透過大規模基因體測序計畫已發現數百萬個基因體變異，其中 400 萬個可能與疾病相關。基因體資料屬於敏感個人資料，因此確保資料安全的透明資料治理與匿名化不可或缺。

2.3 資料品質、整合、可互操作性

人工智慧 (AI) 技術要在醫療領域發揮革新潛力，高品質的資料、順暢的資料整合以及系統間的可互操作性不可或缺。然而在現實中，由於資料品質問題、分散的資料儲存方式、複雜的監管環境等，AI 的引進存在諸多難關，這成為 AI 安全且有效地整合至醫療系統中的重要挑戰。

醫院・部門・機構間的資料斷層 (Data Fragmentation)

醫療現場面臨著資料孤島 (data silos) 與異質 IT 基礎設施的嚴重問題。醫院內部各部門，以及不同醫院或醫療機構，往往使用獨立的系統與架構，導致資訊無法橫向整合或相互參照。例如，在日本的情況下，僅主要的電子病歷 (EHR) 供應商就多達 4 家，且每家供應商擁有獨特的架構與連線方式，使得醫療資訊的連動十分困難。這導致在實現特定疾病的精準醫療或引進新應用程式時，資料連線需要耗費巨大時間、成本與人力。醫療系統若處於斷層狀態，不僅會限制 AI 解決方案的適用性，亦會妨礙對病患整體醫療紀錄進行綜合分析以提供最佳診斷與治療。甚至在同一醫療機構內部，病患紀錄也未能妥善共享，損害了使用者體驗，並阻礙了醫療資源的有效利用。此類斷層被指出是阻礙基於 AI 的醫療革新的主要瓶頸之一。

HL7 FHIR 標準的角色

為了解決資料斷層問題並確保可互操作性（interoperability），核心方法之一是採用 HL7 FHIR（Health Level Seven Fast Healthcare Interoperability Resources）標準。FHIR 提供專為以標準化格式存取與交換醫療資料而設計的架構，這對於整合 EHR、實驗室資料、基因體學、時間序列臨床觀察資料等多種醫療資料至關重要。HL7 FHIR 特別補足了用於醫療影像資料交換的 DICOM 標準，促進不同醫療機構及部門之間的資料交換。直接採用 HL7 FHIR 標準能實現機構間順暢的資料建模與交換，例如在糖尿病管理中，可整合血糖、體重、飲食紀錄等並應用於 AI 應用程式。部分國家立下了以 FHIR 連接所有醫療中心的目標，以期比亞洲及歐洲更快實現可互操作性。這支持了「在智慧醫療中最重要技術不是 AI、Transformer 或數位孿生，而是 FHIR」的主張。因為只有透過 FHIR 才能整合所有資料，而只有當資料整合後，一切革新才成為可能。

資料偏見（Bias）問題

AI 模型的效能不僅取決於訓練資料的品質，亦高度依賴資料的代表性。帶有偏見的資料可能導致 AI 模型在真實臨床環境中得出不公平或不精準的結果，這在醫療領域可能引發嚴重的倫理與臨床問題。例如，僅憑特定人口群體（例如：單一種族、特定地區、特定醫療系統）資料訓練的模型，在其他病患群體中無法有效泛化。在非洲訓練的黑色素瘤偵測 AI 工具在淺色皮膚上的準確度為 95%，但在深色皮膚上的敏感度卻下降了 34%，這一案例直觀地展示了種族偏見的風險。此類「資料集偏移（dataset shift）」問題會大幅降低 AI 模型的真實效能。醫療資料標註（labeling）過程中的一致性不足或特定亞群體的代表性不足，可能導致演算法偏見（algorithmic bias），進而引發效能不均衡。此類偏見引發了人們對 AI 可能使既有不平等持久化甚至系統化的擔憂。因此，從 AI 開發階段起運用多元且具代表性的資料，並透過公平性審計（fairness audits）與分層評估（stratified evaluation）來確保模型的公正性至關重要。

個人隱私保護與倫理考量事項

醫療資料屬於最敏感且受法律保護的個人資訊之一，因此對於 AI 系統存取、控制與使用資料的方式存在相當大的擔憂。資料隱私與安全問題是阻礙 AI 順暢整合至 EHR 系統的主要因素。

為了保護個人隱私，資料匿名化（anonymization）或假名化（pseudonymization）非常重要，但此過程並不簡單且可能並不完美。匿名化資料必須確保無法重新連結至特定個人，而假名化資料則需在醫院內部將識別資訊分開保管。此外，在將病患紀錄資料用於 AI 模型訓練時，當事人的同意與否至關重要。部分國家禁止將病患資訊輸入至個人用 ChatGPT 或 Gemini 等通用 AI 服務中。這是因為資料會傳送至海外伺服器，從而不受該國法律約束。

聯邦學習（Federated Learning）等技術能在無須交換原始資料的情況下利用分散資料訓練模型，提供在維持病患保密性的同時實現資料整合的緩和策略。然而，當 AI 作出誤診或給出錯誤建議時，最終法律責任仍歸屬於醫療專業人員，因此 AI 系統的可解釋性（explainability）與可靠性不可或缺。需要努力解決難以理解 AI 系統如何作出決策的「黑盒子（black box）」問題，可解釋 AI（Explainable AI, XAI）具有重要意義。

實地模型效能與泛化可能性

AI 模型的真實世界效能（real-world performance）可能與實驗室環境中的效能有所不同，這與模型的泛化可能性（generalizability）直接相關。AI 模型僅在與受訓病患群體及狀況相似的情況下，才能提供值得信賴的結果。由於過度擬合（overfitting）現象，模型可能過度適應訓練資料，並在真實狀況下出現效能下降。

特別是醫療影像資料具有高維度且複雜，且常存在標註不足（poorly annotated）的情況，增加了 AI 模型學習的難度。標註工作需要醫療專業人員的時間與專業知識。此外，小規模資料集或不平衡資料會限制模型的泛化可能性，並增加過度擬合風險。

為了讓 AI 模型成功應用於真實臨床環境，使用多中心（multi-center）收集的多元資料進行訓練與驗證至關重要。然而，EHR 資料品質在不同機構間可能存在巨大差異，這可能導致模型效能下降。例如，基層醫院的症狀描述完整度僅為 60~70%，相比於三級醫院（90% 以上），實地應用時模型效能可能會有所降低。

總結而言，成功引進 AI 醫療不僅僅是單純的技術開發，更需要提升資料品質、實現系統間順暢整合、實施嚴格的個人隱私保護，以及對模型在真實臨床環境中的穩健效能與倫理應用的深入理解與努力。

[根據資料]

1. 標題: Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者・機構: MédicosÉlite®

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel 'MédicosÉlite®' (來源未提供影片的直接 URL)

支持的主張: AI 並非取代醫師，而是作為聰明且不知疲倦的同事來支援醫療人員，最終決定與責任始終在於醫師。AI 模型存在誤診風險、幻覺（hallucination）問題、偏見性等局限，個人隱私保護、資料安全、演算法偏見等倫理考量事項至關重要。

2. 標題: AI applications in medicine and healthcare Integration of AI in electronic health records

作者・機構: eurjmedres.biomedcentral.com

可確認的發行日・上傳日: 2025 (URL中明示)

URL: <https://eurjmedres.biomedcentral.com/counter/pdf/10.1186/s40001-025-03196-w>

支持的主張: AI 能分析 EHR 與實時資料，比傳統方法更快速、更精準地檢測疾病，但在將 AI 整合至 EHR 系統的過程中，不同平台間的資料可互操作性問題以及資料隱私與安全問題仍是重要課題。

3. 標題: AI活用によって医療はどう変わる？最新のディープテック技術がもたらす社会的インパクト

作者・機構: GLOBIS学び放題×知見

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel 'GLOBIS学び放題×知見' (來源未提供影片的直接 URL)

支持的主張: AI 模型存在偏見 (bias)，可能存在對相同條件下的病患產生不利影響的 AI，這已透過研究得到證實。管理與控制 AI 如何運作至關重要，應確保在醫療 AI 中不致引發因病患條件而異的不平等。

4. 標題: L' IA & l' IA Générative en santé : À l'épreuve de la réalité

作者・機構: L'Oeil de la E-santé

可確認的發行日・上傳日: (來源未明示)

URL: YouTube channel 'L'Oeil de la E-santé' (來源未提供影片的直接 URL)

支持的主張: 為了有效整合 AI 解決方案，資料品質與可互操作性不可或缺。AI 唯有在具備優質資料時才能正常運作，特別是資料的不均勻輸入方式（例如：生命指標記錄方式的多樣性）可能降低 AI 工具的有效性。

第二部

AI 醫療的實踐

第三章 診斷與風險預測

3.1 影像醫學、病理與臨床診斷支援

人工智慧（AI）在醫療診斷的多個核心領域帶來革命性變化，並作為補充與強化醫療專家能力的重大工具脫穎而出。特別是基於海量資料識別複雜模式的AI能力，在影像醫學、病理學以及臨床紀錄分析領域中，對顯著提升診斷精準度與效率做出了貢獻。AI不再是未來的技術，而是已經在醫療現場為疾病的早期發現、制定治療計畫及患者管理提供實質幫助。AI協助醫療提供者在做出重要決定時確定優先順序，提高決策速度，並協助發現潛在錯誤。

影像醫學判讀中的AI支援

影像醫學是AI發揮最強大性能的領域之一，得益於高度數位化水準與標準化影像處理程序，非常適合導入AI技術。AI能在X光、CT、MRI等醫療影像中偵測微妙且複雜的模式，協助早期發現肉眼容易遺漏的疾病。AI的這種能力提升了診斷精準度，並加快了醫療人員的決策速度。

AI的速度與效率在急診狀況下尤為突出。過去需要10至30分鐘的醫療影像處理作業，AI可在3至5分鐘內完成，這在腦卒中（中風）等急診狀況下有助於迅速做出治療決定，對將患者腦神經細胞損傷降至最低發揮了決定性作用。在特定類型的皮膚癌診斷中，研究結果顯示AI甚至比人類皮膚科專科醫師具備更高的精準度。例如，梅奧診所（Mayo Clinic）的AI模型經證實可用於分析CT掃描，比臨床診斷最多提前3年偵測出胰臟癌。胃癌早期診斷模型也利用CT影像識別病變，提高了檢出率。AI還展現出分析冠狀動脈狹窄程度、減少不必要支架置入手術的潛力。

醫療影像AI正在超越單純的2D影像分析持續發展。如Zijie Smart Medical Technology的「Anatomy Cloud」等AI解決方案可在3至5分鐘內處理3D醫療影像並識別200個以上的器官，未來還包含自動識別病變或腫瘤的功能，實現從2D黑白影像向3D彩色影像的轉變。這有助於醫師進行術前練習、規劃精準的術中路徑，並掌握病變的精準位置。此外，AI還被應用於提升影像品質、影像重建，以及肺結節分類與偵測等特定診斷任務。

病理切片分析中的AI支援

在病理學領域，AI同樣在提升診斷精準度方面發揮著重要作用。基於AI的工具可分析數位病理切片，偵測微妙的細胞特徵與生體指標，這能識別肉眼可能忽視的部分，進而提高診斷精準度。例如，AI已應用於抹片檢查（Papanicolaou smear）以早期發現子宮頸癌、識別前列腺癌、計算葛里森分數（Gleason scores）等。此外，AI還可以分析與EGFR基因變異型態相似的病理影像，在無基因檢測的情況下預測變異，進而提高基因體診斷效率。目前已有數款經FDA核准的基於AI數位病理解決方案上市。

透過臨床紀錄分析的AI診斷支援

臨床紀錄是綜合瞭解患者健康狀況不可或缺的資訊來源，AI分析這些海量資料以支援診斷。電子病歷（EHR）中混合了醫師的診療記錄、出院摘要、影像醫學科報告等非結構化文字資料，以及診斷代碼、藥物處方等結構化資料。AI，特別是自然語言處理（NLP）技術和大語言模型（LLM），能分析這些非結構化資料並萃取出寶貴的臨床洞察。

LLM經過數十億個單詞與影像的訓練，能夠識別醫療文獻、臨床案例、研究與協定，理解複雜的醫療語境，並能不知疲倦地處理海量資料。藉此，AI能在幾秒鐘內閱讀患者超過300頁的醫療紀錄，迅速掌握過敏、過往用藥史、手術及風險因子等。此外，還能將表面上看似不相關的症狀連結起來，並提出基於實證的最新治療協定。例如，AI醫療助手能針對患者的症狀（例如：腹瀉3次、心臟疼痛、進食時頻繁打嗝）提供合理的解答與治療方法（例如：糞便檢查、少食多餐、補充水分、抗酸劑、抗病毒藥物）。

AI整合至EHR中，在支援醫療人員的臨床決策方面發揮著重要作用。AI系統分析EHR及即時資料，能比傳統方式更快速且精準地偵測疾病、監測藥物相互作用，並向醫療人員即時發出潛在副作用警告。ML模型利用EHR資料預測院內死亡率、再入院率及敗血症發生等重大臨床結果

[根據資料]

以下為支持第3章「診斷與風險預測」之3.1「影像、病理與臨床診斷支援」、3.2「疾病進展與治療結果預測」、3.3「早期預警與患者風險分級」的根據資料清單。

根據資料清單

標題: 1 醫療RAG大模型實戰

作者/發行機構: AI大模型應用開發 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 20240928「智慧醫院的未來藍圖：從策略到實踐」論壇【智慧語音與穿戴裝置應用在臨床照護 | 陽明交通大學智能系統研究所 廖元甫教授】

作者/發行機構: 社團法人亞洲華人醫務管理交流學會 (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 3.2

標題: 2026-07-01：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

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-01

URL: <https://www.youtube.com/watch?v=UDng4lXwL-A>

支持的小標題: 3.1, 3.2, 3.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 3.1, 3.2, 3.3

標題: AI 基因報告判讀、5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: AI在臨床醫療的應用與未來趨勢

作者/發行機構: 台中市醫師公會 健康台灣深耕計畫 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Intelligence artificielle et diagnostic médical : mythe ou réalité ?

作者/發行機構: Académie royale de Belgique (YouTube 頻道)

支持的小標題: 3.1

標題: Quand l'intelligence artificielle générative transforme la santé

作者/發行機構: Obvia (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: WEBINAR: Salud sexual, reproductiva y materna: cómo puede contribuir la IA en prevención y cuidado

作者/發行機構: CLIAS CIIPS (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: Webinaire - L'IA dans la conception des dispositifs médicaux

作者/發行機構: UseConcept Vidéo (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Webinar#2 - KI in der NMD-Diagnose - ein Verbündeter für Ärzt:innen und Patient:innen

作者/發行機構: CoMPaSS-NMD (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Excerpts from <https://iris.unil.ch/bitstreams/3bfa1900-1f10-4f05-9470-326781944b2c/download>

發行機構: (URL 來源)

支持的小標題: 3.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

支持的小標題: 3.1, 3.2, 3.3

標題: 【それでもAIに健康相談はしないほうがいい】AI活用に詳しい医師・大塚篤司／質も共感も医師よりAIが上／ただAIの指示で中毒患者発生／「1%の致命的な間違い」をゼロにできない【lon1 Health】

作者/發行機構: TBS CROSS DIG with Bloomberg (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 首創AI醫療影像3D標準化 ☑ MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 3.1

3.2 疾病進展與治療結果預測

人工智慧（AI）在醫療領域為預測疾病進展過程及展望治療結果樹立了卓著功勳。AI學習海量患者資料並識別複雜模式的能力，已成為醫療專家早期發現疾病、制定個人化最佳治療計畫，以及預先認知潛在風險並予以干預不可或缺的工具。這提供了超越傳統統計預測的精細度與即時性，並在實現以患者為中心的個人化醫療中發揮核心作用。AI協助醫療提供者在做出重要決定時確定優先順序，提高決策速度，並協助發現潛在錯誤。

患者病歷及臨床紀錄分析

AI分析患者電子病歷（EHR）中所包含的海量病歷資料，為預測疾病進展及治療結果奠定基礎。EHR中混合了醫師的診療記錄、出院摘要、藥物處方、手術歷史等結構化與非結構化資料。AI，特別是大語言模型（LLM），能理解這些非結構化文字資料，在幾秒鐘內閱讀患者超過300頁的醫療紀錄，迅速掌握過敏、過往用藥史、手術、風險因子等。AI還能分析患者過往的就診紀錄（例如：44歲、46歲、48歲、52歲就診）及診療科別（呼吸道、心臟、眼科）等，追蹤患者整體健康狀況的變化。這些患者病歷資訊成為AI模型反映患者固有狀況與過往疾病經驗以預測未來的必要輸入資料。

醫療人員可利用如AI聽寫員（AI scribe）等技術，自動將與患者的對話內容記錄於EHR中，從而在不遺漏資料的情況下取得一致的病歷資料，這將進一步提升AI模型訓練資料的品質。AI能基於這些紀錄，迅速處理並摘要患者過往病歷、社會史、當前治療、生命徵象、身體檢查資料等，並能整合多份紀錄以生成臨床摘要（clinical summary）。這有助於醫療人員在短時間內掌握患者複雜的病歷，不遺漏診療所需的關鍵資訊。

檢查結果的即時分析與預測

AI即時分析各種形式的檢查結果，為疾病進展與預後預測提供重要洞察。

醫療影像：X光、CT、MRI、超音波等醫療影像是AI發揮最強大性能的領域之一。AI在這些影像中偵測肉眼難以發現的微妙病變或模式，從而早期診斷疾病並預測進展狀況。例如，AI可用於分析CT掃描，比臨床診斷最多提前3年偵測出胰臟癌，胃癌早期診斷模型也利用CT影像提高了病變識別率。特別是在腦卒中（中風）等急診狀況下，AI能在3至5分鐘內處理複雜的影像資料並向醫療人員發出迅速警告，從而對將患者腦神經細胞損傷降至最低做出了貢獻。AI還能分析CT掃描中冠狀動脈的鈣化程度，客觀評估心血管疾病的風險度，並用於減少不必要的支架置入手術。AI在CT影像中執行如肺結節偵測等任務，協助早期發現疾病。如Zijie Smart Medical Technology的「Anatomy Cloud」等AI解決方案可在3至5分鐘內處理3D醫療影像並識別200個以上的器官，自動識別病變或腫瘤，實現從2D黑白影像向3D彩色影像的轉變。這有助於醫師進行術前練習、規劃精準的術中路徑，並掌握病變的精準位置。AI在如黑色素瘤診斷等特定皮膚癌診斷中，有時表現出比人類皮膚科專科醫師更高的精準度。

血液與生體檢查：AI分析血液檢查結果，監測藥物相互作用並即時警告潛在副作用。例如，基於AI的系統能分析患者的檢驗室資料以偵測微量白蛋白尿，並預測糖尿病前期的糖尿病發病風險以建議早期干預。AI分析心電圖（EKG）資料可在數秒內偵測出高血鉀症（hyperkalemia）等急診狀況並向醫療人員傳送簡訊通知，從而提高治療成功率並降低死亡率。

基因體資料：AI分析基因體資料以識別致病變異體，預測個人對特定疾病的易感性，並提供建立個人化治療策略所需資訊。基於AI的基因體分析能判斷基因變異的嚴重程度，高機率預測未來疾病發生的可能性，進而實現預防性治療。AI透過基因分析預測未來疾病，對加速藥物開發及基因學研究做出貢獻。

治療反應預測與個人化干預

AI預測患者的治療反應，並以此為基礎實現針對個人的最佳化個人化干預。

治療結果預測：AI能基於患者年齡、疾病階段、細胞型態等14種患者參數來預測生存期。透過學習過往612個案例資料，輸入新患者資訊後，AI可提供預測該患者生存率的模型。這些預測成為醫療人員與患者討論治療計畫，以及患者與家屬規劃未來的重要參考資料。

個人化藥物療法：AI綜合考量患者的基因體資訊、病歷、合併症等，預測藥物敏感度及反應度，並推薦能在將副作用降至最低的同時實現治療效果最大化的藥物與劑量。AI能基於術前23種參數預測出血風險、腎損傷風險等，並建議最大造影劑用量或是否調整抗凝血劑，從而降低術中併發症風險。數位雙生（Digital Twin）技術能生成患者動態的虛擬生理鏡像，模擬各種治療計畫（例如：藥物、劑量），並從數十年前就預測未來疾病風險，實現真正的個人化預防與治療。

慢性病管理：AI持續監測穿戴式裝置及IoT感測器收集到的患者生理資料（心率、血壓、血糖等）以偵測異常徵兆並傳送智慧警告。藉此協助慢性病患者自我管理健康，並在必要時自動建議調整藥物以最佳化臨床工作流程。AI對話機器人能回答健康相關問題並提供個人化健康資訊與建議，特別是在夜間等醫療服務可及性受限的時間段，發揮著化解健康資訊落差的重要作用。

風險因子識別與預防性途徑

AI分析海量資料以識別疾病風險因子，並生成符合個人健康狀況的個人化風險檔案，以支援早期干預與預防性治療。

預防與預測：AI的最大優勢之一是基於過往資料預測未來的能力。AI能提前20年預測特定人口群體的疾病發病率及其時期，這對醫療資源管理及公共衛生政策制定產生巨大影響。AI透過心血管疾病風險預測平台，基於血液檢查、身體測量值（身高、體重、血壓）告知個人的心血管疾病風險，引導生活習慣改善。

早期發現：AI在CT影像中執行如肺結節偵測等任務，協助早期發現疾病。基於AI的預測分析能成功識別慢性病發病風險高的患者，引導有效的健康計畫。

預測偏誤（Prediction Bias）與醫療人員判斷的重要性

AI的預測能力令人驚嘆，但AI僅提供「機率」，並非說出「絕對的真相」。AI模型高度依賴訓練資料的品質與代表性，因此若資料存在偏誤（bias），可能做出不公正或不精準的預測。例如，以淺膚色資料訓練的黑色素瘤偵測AI，在深膚色患者身上的敏感度可能降低34%。當特定群體在資料集中代表性不足時，AI無法對該群體進行充分學習，進而犯下預測錯誤。

此外，AI可能表現出幻覺（hallucination）現象，存在生成看似合理卻錯誤資訊的風險。GPT-5出現了約15%的幻覺錯誤，這在醫療現場可能引致「1%的致命失誤」，屬於嚴重的問題。實際上已有報導指出，因AI的錯誤建議（例如：限制鈉攝取導致的中毒）而使患者遭受損害的案例。這些問題明確表明AI無法取代醫療專家。

因此，AI的預測結果必須經過醫療專家的批判性審視與最終臨床判斷。AI並非取代醫師，而是作為「聰明且不知疲倦的同僚」支援醫療人員工作的工具。醫療人員不應盲目遵循AI的建議，而應發揮人類獨有的同理能力、直覺及非語言資訊掌握能力，全面理解患者。AI雖能減輕醫療人員的認知負擔，協助將更多時間投入與患者的溝通，但最終責任始終歸屬於醫療專家。

同時也面臨難以理解AI系統如何做出決策的「黑盒子（black box）」問題。因此，AI系統的可解釋性（explainability）與可信度不可或缺，對可解釋AI（Explainable AI, XAI）的研究與導入努力至關重要。對於AI提出的風險分數或診斷可能性，醫療人員應成為決定信任到何種程度以及如何干預的閾值（threshold）設定主體。有人提出疑慮，若儘管AI提出了其他可能性，醫療人員仍堅持自身判斷而發生錯誤，法律責任範圍可能會擴大。這可能導致強化醫療人員「注意義務（duty of care）」的結果。

資料品質的影響：AI模型的現場性能與訓練資料品質直接相關。資料的異質性、不平衡、缺失值、偏誤等會直接影響AI模型的精準度與可信度。特別是小規模資料集或不平衡資料會限制模型的通用可能性並增加過擬合風險。預測疾病進展所必需的醫療影像資料具有高維度且複雜，且往往標註不足（poorly annotated），使AI模型的學習變得困難。因此，高品質精準資料遠比海量資料更為重要，在資料預處理過程中，應透過去識別化或假名化保護患者隱私，同時保持資料的實用性。

總結而言，AI在預測疾病進展與治療結果方面為醫療人員提供了強大的洞察力與效率，但必須明確認知其局限性與倫理問題。AI應作為增強醫療專家能力、減輕工作負擔，並實現以患者為中心醫療的「第三者輔助者（third agent）」發揮作用。未來的醫療將由精通運用AI且不忘人文關懷的醫療人員重新定義。

[根據資料]

以下為支持第3章「診斷與風險預測」之3.1「影像、病理與臨床診斷支援」、3.2「疾病進展與治療結果預測」、3.3「早期預警與患者風險分級」的根據資料清單。

根據資料清單

標題: 1 醫療RAG大模型實戰

作者/發行機構: AI大模型應用開發 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 20240928 「智慧醫院的未來藍圖：從策略到實踐」論壇【智慧語音與穿戴裝置應用在臨床照護 | 陽明交通大學智能系統研究所 廖元甫教授】

作者/發行機構: 社團法人亞洲華人醫務管理交流學會 (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 3.2

標題: 2026-07-01：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

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-01

URL: <https://www.youtube.com/watch?v=UDng4lXwL-A>

支持的小標題: 3.1, 3.2, 3.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 3.1, 3.2, 3.3

標題: AI 基因報告判讀、5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: AI在臨床醫療的應用與未來趨勢

作者/發行機構: 台中市醫師公會 健康台灣深耕計畫 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Intelligence artificielle et diagnostic médical : mythe ou réalité ?

作者/發行機構: Académie royale de Belgique (YouTube 頻道)

支持的小標題: 3.1

標題: Quand l'intelligence artificielle générative transforme la santé

作者/發行機構: Obvia (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: WEBINAR: Salud sexual, reproductiva y materna: cómo puede contribuir la IA en prevención y cuidado

作者/發行機構: CLIAS CIIPS (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: Webinaire - L'IA dans la conception des dispositifs médicaux

作者/發行機構: UseConcept Vidéo (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Webinar#2 - KI in der NMD-Diagnose - ein Verbündeter für Ärzt:innen und Patient:innen

作者/發行機構: CoMPaSS-NMD (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Excerpts from <https://iris.unil.ch/bitstreams/3bfa1900-1f10-4f05-9470-326781944b2c/download>

發行機構: (URL 來源)

支持的小標題: 3.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

支持的小標題: 3.1, 3.2, 3.3

標題: 【それでもAIに健康相談はしないほうがいい】AI活用に詳しい医師・大塚篤司／質も共感も医師よりAIが上／ただAIの指示で中毒患者発生／「1%の致命的な間違い」をゼロにできない【1on1 Health】

作者/發行機構: TBS CROSS DIG with Bloomberg (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 首創AI醫療影像3D標準化 ☒ MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 3.1

3.3 早期預警與患者風險分級

人工智慧（AI）在醫療領域作為預測患者疾病進展並早期偵測與分級風險的革命性工具已站穩腳跟。AI學習海量醫療資料並識別複雜模式的能力，為醫療專家早期發現疾病、制定個人化最佳治療計畫，以及對潛在風險進行先發制人的干預提供了不可或缺的洞察。這提供了超越傳統統計預測的精細度與即時性，並在實現以患者為中心的個人化醫療及極大化醫療服務效率方面發揮核心作用。AI協助醫療提供者優先處理重要決定，提高決策速度，並為發現潛在錯誤做出貢獻。

疾病早期預警系統

AI即時分析患者的生命徵象、臨床紀錄、檢查結果等各種資料，用於預先警示疾病急劇惡化或緊急狀況。這在患者狀況可能急轉直下的加護病房（ICU）或急診室（ER）等環境中，其價值尤為凸顯。

敗血症及急性惡化預測：AI用於新生兒及小兒加護領域，在發生重大急診狀況前向醫療人員發出警告，為採取適當措施爭取時間。部分研究顯示，AI模型能在6至12小時前偵測出小兒患者78%的關鍵事件（critical events），超越了傳統的預測指標。這能在患者狀況惡化前開啟早期治療或轉診至加護病房等進行重大干預。AI還能分析患者生命徵象、監測資料等，提供住院風險或患者狀況惡化風險警告，並可計算小兒科或死亡率指標。

高血鉀症偵測：AI分析心電圖（EKG）資料可在數秒內偵測出高血鉀症（hyperkalemia）等急診狀況並向醫療人員傳送簡訊通知，從而提高治療成功率並降低死亡率。此類基於AI的早期預警系統有助於醫療人員不遺漏患者微妙的變化，並支援迅速判斷，在拯救生命方面發揮決定性作用。

患者風險分級與管理

AI分析患者複雜的資料以建立個別患者的風險檔案，並以此為基礎對建立個人化治療與管理策略做出貢獻。

急診室及加護病房風險分級：AI能大幅改善急診室及醫院營運效率。AI模型透過小兒患者住院預測，考量年中時期、病毒循環、區域爆發等因素，最佳化病床供需計畫，並能預測加護病房或呼吸道/心臟專科人力擴充的需求。此外，AI能減少35%的行政錯誤，並將急診室等待時間從21分鐘縮短至9分鐘，從而減少急診室擁擠，幫助患者更快接受診療。即使在患者無生命徵象的情況下，AI也能預測住院風險或患者狀況惡化風險並提供警告，有助於在併發症發生前引導早期住院或識別重新就診的必要性。此類系統不僅適用於兒童，也適用於成年患者，使患者能在進入危急狀態前獲得醫療服務。

再入院風險預測：AI分析患者過往紀錄、治療計畫、社會經濟因素等，可識別再入院風險高的患者。這能有效分配出院後患者管理所需的資源，並實現預防再入院的個人化干預。AI模型可預測住院兒童人數，對病床管理及人力配置做出貢獻。

急性惡化及重新諮詢必要性預測：AI模型能識別高達70%需要重新諮詢或24小時內住院的案例，協助患者在進入嚴重狀態前進行重新諮詢。這能減少併發症，並在無需加護病房治療的情況下實現早期住院，對患者預後產生積極影響。

個人風險因子識別與精準預測：AI能分析CT掃描中冠狀動脈鈣化程度，客觀評估心血管疾病風險度，並用於減少不必要的支架置入手術。這也應用於基於患者血液檢查、身體測量值（身高、體重、血壓）告知個人心血管疾病風險並引導生活習慣改善的平台。AI能基於術前23種參數預測出血風險、腎損傷風險等，並建議最大造影劑用量或是否調整抗凝血劑，從而降低術中併發症風險。這些預測為醫療人員反映患者固有狀況與過往疾病經驗以預測未來提供了必要資訊。

預測偏誤、誤報與醫療人員判斷的重要性

雖然AI在醫療領域提供強大的預測與分級功能，但明確認知其局限性與倫理考量事項，並尊重醫療人員的最終判斷至關重要。

AI的局限與錯誤：AI可能犯下「1%的致命錯誤（fatal error）」，這在醫療領域可能引致嚴重後果。AI模型高度依賴訓練資料的品質與代表性，因此若資料存在偏誤（bias），可能做出不公正或不精準的預測。例如，若缺乏特定種族的資料，AI可能難以對該種族患者做出精準診斷。此外，AI存在幻覺（hallucination）現象的風險，即生成看似合理卻錯誤的資訊，GPT-5出現了約15%的幻覺錯誤。這可能導致AI提供錯誤建議（例如：限制鈉攝取導致的中毒）進而對患者造成傷害的案例。

誤報（False Alarms）與醫療人員審視：AI系統伴隨高風險

[根據資料]

以下為支持第3章「診斷與風險預測」之3.1「影像、病理與臨床診斷支援」、3.2「疾病進展與治療結果預測」、3.3「早期預警與患者風險分級」的根據資料清單。

根據資料清單

標題: 1 醫療RAG大模型實戰

作者/發行機構: AI大模型應用開發 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 20240928「智慧醫院的未來藍圖：從策略到實踐」論壇【智慧語音與穿戴裝置應用在臨床照護 | 陽明交通大學智能系統研究所 廖元甫教授】

作者/發行機構: 社團法人亞洲華人醫務管理交流學會 (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 3.2

標題: 2026-07-01：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

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-01

URL: <https://www.youtube.com/watch?v=UDng4lXwL-A>

支持的小標題: 3.1, 3.2, 3.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 3.1, 3.2, 3.3

標題: AI 基因報告判讀、5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: AI在臨床醫療的應用與未來趨勢

作者/發行機構: 台中市醫師公會 健康台灣深耕計畫 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Intelligence artificielle et diagnostic médical : mythe ou réalité ?

作者/發行機構: Académie royale de Belgique (YouTube 頻道)

支持的小標題: 3.1

標題: Quand l'intelligence artificielle générative transforme la santé

作者/發行機構: Obvia (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: WEBINAR: Salud sexual, reproductiva y materna: cómo puede contribuir la IA en prevención y cuidado

作者/發行機構: CLIAS CIIPS (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: Webinaire - L'IA dans la conception des dispositifs médicaux

作者/發行機構: UseConcept Vidéo (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: Webinar#2 - KI in der NMD-Diagnose - ein Verbündeter für Ärzt:innen und Patient:innen

作者/發行機構: CoMPaSS-NMD (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: Excerpts from <https://iris.unil.ch/bitstreams/3bfa1900-1f10-4f05-9470-326781944b2c/download>

發行機構: (URL 來源)

支持的小標題: 3.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

支持的小標題: 3.1, 3.2, 3.3

標題: 【それでもAIに健康相談はしないほうがいい】AI活用に詳しい医師・大塚篤司／質も共感も医師よりAIが上／ただAIの指示で中毒患者発生／「1%の致命的な間違い」をゼロにできない【1on1 Health】

作者/發行機構: TBS CROSS DIG with Bloomberg (YouTube 頻道)

支持的小標題: 3.1, 3.3

標題: 人工智慧基本法與智慧醫療治理趨勢 (下) 李崇僖教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 3.1, 3.2, 3.3

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 3.1, 3.2

標題: 首創AI醫療影像3D標準化 ☒ MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 3.1

第四章 治療計畫與精準醫療

4.1 個人化治療與藥物基因體學

人工智慧（AI）正在加速個人化醫療（Personalized Medicine）時代的到來，該時代綜合考量患者固有的生物學特徵、生活方式及環境因子，以提供最佳醫療服務。這在從疾病早期發現、預測、預防到治療的整體醫療中帶來革命性變化。AI透過學習與分析海量患者資料，擺脫過去一律化的治療方式，使針對每個個人最有效且安全的個人化途徑成為可能。

患者特徵分析的精細化

個人化治療的第一步是精準掌握患者的各種特徵。AI綜合分析電子病歷（EHR）中包含的海量患者病歷、生活習慣、社會經濟資訊，以及透過穿戴式裝置收集的即時生命資料等。例如，AI能基於患者年齡、疾病階段、細胞型態等14種患者參數來預測生存期。此外，還能在短短數秒內讀完包含患者複雜病歷的超過300頁醫療紀錄，迅速掌握過敏、過往用藥史、手術紀錄、潛在風險因子等。

AI考量這些患者特徵以生成個人化健康風險檔案，並透過對疾病進展與治療結果的預測來支援早期干預與預防性治療。例如，AI能分析CT掃描中冠狀動脈鈣化程度，客觀評估心血管疾病風險度，並用於減少不必要的支架置入手術。心血管疾病風險預測平台基於血液檢查、身體測量值（身高、體重、血壓）告知個人的心血管疾病風險，引導生活習慣改善。此類詳細的患者特徵分析為AI模型反映患者固有狀況與過往疾病經驗以預測未來提供了必要的輸入資料。

基因體資料的運用與藥物基因體學

基因體資料在個人化治療，特別是藥物基因體學（Pharmacogenomics, PGx）領域中發揮著核心作用。AI分析在大規模基因體定序專案中發現的數百萬個基因體變異體，以識別致病變異體，並預測個人對特定疾病的易感性。此類分析使基於分子標記的患者分層成為可能，為建立最適合個人的治療策略提供不可或缺的資訊。

藥物基因體學（PGx）：AI綜合考量患者的基因體資訊、病歷、合併症等，預測藥物敏感度及反應度，並推薦能在將副作用降至最低的同時實現治療效果最大化的藥物與劑量。PGx變異存在於所有人身上，透過PGx可改善藥物安全性與功效。例如，AI能偵測如EGFR及KRAS突變等基因變異以協助肺癌治療，並透過BRCA突變支援卵巢癌治療。此外，AI能基於術前23種參數預測出血風險、腎損傷風險等，並建議最大造影劑用量或是否調整抗凝血劑，從而降低術中併發症風險。

單基因與多基因疾病：AI對單基因疾病（monogenic diseases）與多基因疾病（polygenic diseases）的預測均做出貢獻。單基因疾病通常罕見，但約有55%可透過基因檢測找到原因。相反地，對於高血壓、糖尿病、癌症等多基因疾病，則透過多基因風險分數（Polygenic Risk Score, PRS）預測疾病發生風險。台灣生物資料庫（TPMI）計算了265種疾病的PRS，可用於提前20年預測特定群體的疾病發病率。然而，單憑基因因子僅能解釋心血管疾病風險預測的2~3%，結合環境因子時預測力將大幅提升。

治療反應預測的精細化

AI預測患者的治療反應，並以此為基礎實現針對個人的最佳化個人化干預。

生存期預測：AI能基於患者的14種參數預測生存期，透過學習過往612個案例資料，提供輸入新患者資訊後預測該患者生存率的模型。這些預測成為醫療人員與患者討論治療計畫，以及患者與家屬規劃未來的重要參考資料。

藥物開發與再利用：製藥企業利用AI理解藥物的吸收、分佈、代謝、排泄路徑，預測藥物相互作用，並整合基因體及蛋白質體資料（multi-omics）以將現有藥物

[根據資料]

以下為支持第4章「治療計畫與精準醫療」之4.1, 4.2, 4.3的根據資料清單。

根據資料清單

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=XRz-sWSOrFg>

支持的小標題: 4.1, 4.3

標題: AI 基因報告判讀, 5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 4.1

標題: Ponente Internacional: Seminario Internacional Tecnología médica e IA 2025

作者/發行機構: PhD. Antony Paul Espiritu Martinez (YouTube 頻道)

支持的小標題: 4.1, 4.3

標題: AI與健康照護 | 楊珮菁、王獻章 | AI科普沙龍講座第3期 - BEING HUMAN：AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支持的小標題: 4.2

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 4.2

標題: 首創AI醫療影像3D標準化 MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 4.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 4.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

支持的小標題: 4.1, 4.2

4.2 AI手術與臨床決策支援

人工智慧（AI）在手術室與臨床決策過程中帶來革命性變化，並作為補充與強化醫療專家能力的重大工具脫穎而出。AI從術前計畫到實際術中的精細支援，再到複雜臨床狀況下的決策輔助，在整體醫療中對提升效率、精準度及患者安全做出貢獻。然而，此類AI的導入在技術發展的同時，也伴隨著嚴格的安全驗證、倫理考量，以及外科醫師與醫療人員的最終責任等重要討論。

手術計畫的精細化

AI在術前階段發揮著建立患者個人化治療計畫與評估潛在風險的核心作用。AI演算法分析患者的海量資料，為外科醫師做出重大決定提供所需的深度洞察。

風險評估與預測：AI基於術前患者資料（例如：病歷、檢查結果、基因體資訊）預測手術難度或併發症風險。例如，AI能基於術前23種參數預測出血風險、腎損傷風險等，並建議最大造影劑用量或是否調整抗凝血劑，從而降低術中併發症風險。這些預測模型有助於外科醫師預先認知潛在風險並規劃適當的預防措施。

3D解剖學重建與視覺化：深度學習（DL）技術促進3D解剖學重建並預測植入物尺寸，在特定整形外科手術（例如：全膝關節置換術）中達到90%以上的精準度，表現出遠超傳統2D範本的性能。AI生成腫瘤與周圍解剖結構的三維模型，協助外科醫師視覺化手術部位並規劃最佳手術切入途徑。特別是在腫瘤部位於主要結構附近的複雜情況下非常有用，在泌尿外科腫瘤學中，AI支援機器人前列腺切除術計畫，對保護神經功能及改善術後結果做出貢獻。

治療計畫最佳化：AI分析腫瘤位置、大小、周圍解剖結構等患者個別資料，以最佳化放射線治療計畫。藉此可減輕放射線腫瘤專科醫師的工作負擔並提升治療計畫的一致性。基於AI的系統整合至EHR中，基於患者個別資料支援醫療資源分配及減少不必要的工作量。

機器人與影像導引手術支援

AI與機器人手術系統及高級影像技術結合，大幅提升手術的精準度與效率。

機器人手術系統：AI整合至如達文西（da Vinci）系統等機器人手術平台中以提升精準度，並改善外科醫師的手藝與人體工學。AI整合至機器人系統中提供即時導引（real-time guidance）並促進自動化，從而改善複雜手術程序的結果。這有助於外科醫師更精細地控制術中微小動作與操作、過濾震顫，並防止器械進入危險區域。目前達文西系統由外科醫師直接控制，但未來完全自主（fully autonomous）系統的登場將使精準且無疲勞的手術成為可能。

術中即時影像導引（Intraoperative Guidance）：基於AI的手術系統在術中提供即時導引，協助外科醫師精準找到腫瘤並避免損傷健康組織。此類系統將術前影像資料覆蓋於手術視野上，為外科醫師提供腫瘤位置與邊界的清晰視覺資訊。深度學習（DL）技術，特別是語意分割（semantic segmentation），在腹腔鏡膽囊切除術等手術中識別安全區域與危險區域，展現出術中導引的有前景可能性。經FDA核准的Cydar EV Maps工具等基於AI的3D重建與導航技術，在動脈瘤手術中用於對齊術前影像與即時血管造影。

擴增實境（Augmented Reality, AR）：AR技術將半透明的術前影像覆蓋於手術視野上，提升術中視覺化。這在口腔顎面手術時對齊虛擬影像與真實牙齒，或將3D血管模型投影於患者下肢等各種方式中予以運用。

手術工作流程最佳化：AI對手術室管理，即手術室分配最佳化及節省時間做出貢獻，提升醫療服務生產力。機器人流程自動化（RPA）管理如請款與日程管理等行政業務，使醫護人員能專注於患者治療。

手術訓練與技術評估：AI在手術訓練與專業能力發展中也帶來創新。透過電腦視覺，AI即時追蹤外科醫師的手部與器械動作，提供客觀且可擴充的性能指標與回饋。這使基礎技術（例如：縫合、打結）評估標準化並減少偏誤。

臨床決策支援

AI超越診斷與治療計畫，在各種臨床狀況下發揮支援醫療人員決策的「第三者輔助者（third agent）」作用。

診斷輔助與優先順序指定：AI處理海量臨床資料以快速且精準地偵測疾病、識別微妙模式，並協助醫療專家優先處理重大決定。這為外科醫師在術中捕捉微小變化、調整治療方針提供所需的即時洞察。AI的此類能力提高診斷速度與精準度，進而提升手術決策速度並幫助患者

[根據資料]

以下為支持第4章「治療計畫與精準醫療」之4.1, 4.2, 4.3的根據資料清單。

根據資料清單

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=XRz-sWSOrFg>

支持的小標題: 4.1, 4.3

標題: AI 基因報告判讀, 5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 4.1

標題: Ponente Internacional: Seminario Internacional Tecnología médica e IA 2025

作者/發行機構: PhD. Antony Paul Espiritu Martinez (YouTube 頻道)

支持的小標題: 4.1, 4.3

標題: AI與健康照護 | 楊珮菁、王獻章 | AI科普沙龍講座第3期 - BEING HUMAN : AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支持的小標題: 4.2

標題: 人工智慧基本法與智慧醫療治理趨勢 (下) 李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 4.2

標題: 首創AI醫療影像3D標準化 ㊦ MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 4.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 4.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

URL: <https://theses.hal.science/tel-05534589v1/document>

4.3 新藥開發與虛擬篩選

人工智慧（AI）在新藥開發的全過程中帶來革命性變化，正根本性地改變這個傳統上耗費漫長時間與巨額成本領域的範式。AI從理解疾病的分子機制、發掘新藥候選物質、最佳化，到臨床試驗設計的所有階段中，均作為提高效率與成功率的核心動力發揮作用。學習海量生物學、化學資料並識別複雜模式的AI能力，正以過去難以想像的速度與精細度開創開發新治療劑的可能性。

標靶發掘（Target Identification）與疾病機制理解

新藥開發的第一步是精準發掘與疾病相關的生物學標靶。AI整合與分析基因體學、蛋白質體學、代謝體學等多體學（multi-omics）資料，以識別與疾病相關的過程（disease-associated processes）及有前景的治療標靶（promising therapeutic targets）。AI分析人類能力難以處理的海量生物學資料，可大幅縮短找出作為疾病發生根本原因的基因、蛋白質、訊號傳遞路徑等所需的時間。AI還基於這些資料理解分子結構間的關係，整合與疾病相關的分子結構資訊，從而提出藥物候選物質作用的精準標靶。

蛋白質結構預測（Protein Structure Prediction）

了解藥物標靶的3維結構對藥物設計極為重要。AI在此領域取得了突破性進展。DeepMind的AlphaFold等AI模型極大推動了蛋白質結構預測領域，這是藥物開發不可或缺的標靶識別要素。AI只要輸入分子序列（molecular sequences）即可自動生成3D立體結構。過去以實驗揭示蛋白質結構需要耗費巨大時間與成本，但得益於AI，此過程變得無比快速且精準。研究人員藉此可預測蛋白質的特定部分如何與其他分子（例如：藥物候選物質）結合，並識別潛在的藥物結合部位（binding site），從而提高藥物設計的效率。

分子設計（Molecular Design）及最佳化

AI在設計有效作用於藥物標靶的新分子及最佳化現有分子特性方面發揮核心作用。

新型類藥分子生成：AI學習化合物資料庫與實驗結果，生成新型類藥分子（novel drug-like molecules），藉此擴展新藥發現的化學空間（chemical space）。基於SMILES（簡化分子輸入文本條目系統）的語言模型與基於GNN（圖神經網路）的生成器是小分子（small molecule）設計中最廣泛採用的途徑，其原理也擴展至載荷（payload）發現中。這些模型在如ChEMBL或ZINC等大規模化學化合物庫中進行訓練，學習構成有效且具潛在活性分子的基礎語法與化學規則。

藥物特性最佳化：機器學習（ML）模型分析結構-活性關係（structure activity relationships），引導吸收、分布、代謝、排泄特性改善分子的合理設計。AI理解藥物的吸收（absorption）、分布（distribution）、代謝（metabolism）、排泄（excretion）路徑，分析化學結構與生物學相互作用，有助於早期篩除具有高安全性風險的藥物候選物質。這可防止開發後期階段成本昂貴的失敗，並改善藥物安全性檔案。AI還將免疫原性（immunogenicity）風險評估整合至基於AI的ADC（抗體-藥物複合體）設計工作流程中，對在初期階段過濾可能活化T細胞反應或降低ADC整體耐受性的具免疫反應性結構做出貢獻。

藥物再利用（Drug Repurposing）：AI分析現有化合物資料庫與臨床資訊，發掘現有藥物的新治療用途。這具有縮短新藥開發風險與上市時間的效果。

虛擬篩選（Virtual Screening）

虛擬篩選是AI大幅加速新藥開發速度的核心技術。

高通量虛擬篩選：AI支援針對大規模化學庫的高通量虛擬篩選，預測結合親和力（binding affinities）並確定待測試化合物的優先順序。Zhang等人的研究中，AI針對細菌標靶篩選了數百萬個結構，在數週內識別出新型抗生素候選物質，其速度遠快於傳統方法。

分子模擬：AI使用電腦模型最佳化模擬具有特定活性物質的分子形成過程，大幅縮短藥物開發時間。在細胞層級的「細胞世界模型（cell world models）」中，可模擬藥物反應、虛擬干預（virtual interventions）、藥物干擾（drug perturbations）、虛擬敲除（virtual knockouts）及虛擬藥物篩選（virtual drug screening）等。這有助於在無實驗室物理實驗的情況下預測並確定數以萬計藥物候選物質的效果與優先順序。

臨床試驗候選者選定與最佳化

AI在新藥開發的最後階段 -- 臨床試驗中也發揮著重要作用。

患者群體選定：AI分析歷史臨床試驗資料，選擇最有可能對實驗性治療法產生反應的患者群體（cohort），以提高臨床試驗成功率。這透過識別對特定治療法產生反應機率高的患者，實現臨床試驗效率的極大化。

毒性及副作用預測：AI分析藥物候選物質的化學結構與生物學相互作用，可早期篩除具有高安全性風險的化合物。這能防止開發後期階段成本昂貴的失敗，並改善藥物安全性檔案

[根據資料]

以下為支持第4章「治療計畫與精準醫療」之4.1, 4.2, 4.3的根據資料清單。

根據資料清單

標題: 大模型時代的醫學人工智慧及其應用

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=XRz-sWSOrFg>

支持的小標題: 4.1, 4.3

標題: AI 基因報告判讀, 5G 遠距應用與臨床倫理研習會

作者/發行機構: TCnews慈善新聞網 (YouTube 頻道)

支持的小標題: 4.1

標題: Ponente Internacional: Seminario Internacional Tecnología médica e IA 2025

作者/發行機構: PhD. Antony Paul Espiritu Martinez (YouTube 頻道)

支持的小標題: 4.1, 4.3

標題: AI與健康照護 | 楊珮菁、王獻章 | AI科普沙龍講座第3期 - BEING HUMAN：AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支持的小標題: 4.2

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

支持的小標題: 4.2

標題: 首創AI醫療影像3D標準化 MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 4.2

標題: 2026 光田醫院 生成式AI醫療實務應用課程0403 東海大學姜自強博士 20260715 184612 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的小標題: 4.1

標題: Excerpts from <https://theses.hal.science/tel-05534589v1/document>

發行機構: HAL Theses

URL: <https://theses.hal.science/tel-05534589v1/document>

支持的小標題: 4.1, 4.2

第五章 醫院業務與患者體驗

5.1 診療記錄、語音辨識、文件自動化

人工智慧（AI）正在醫療現場的診療記錄及文件化過程中帶來革命性的變化，這是減輕醫療人員業務負擔並使其能將更多時間投入於患者診療的核心進展。從傳統的手寫記錄到複雜的電子病歷（EHR）系統，龐大的文件作業一直以來都需要醫療人員付出大量的精神與時間。AI技術，特別是語音辨識、自然語言處理（NLP）及大型語言模型（LLM），正致力於將此類文件化過程自動化並實現效率最大化，進而提升醫療服務品質。

語音記錄與AI速記員

基於AI的語音辨識技術是診療記錄自動化的核心部分。過去，醫師口述診療內容後由輔助人員手動記錄或打字是常見的做法，但AI系統幾乎能即時將此過程自動化。醫療人員可以在與患者對話期間，讓AI擔任AI速記員（AI Scribe）的角色，將說話內容轉換為文字並留存為診療記錄。例如，加拿大魁北克的一家醫院過去是由秘書手動記錄醫師口述的內容，如今透過應用AI幾乎實現了全過程自動化。AI速記員能將醫師口述的內容 -- 例如「1957年3月12日出生的馬丁·杜布瓦（Martin Dubois）先生因右膝無外傷性機械性疼痛前來就診」 -- 轉換為文字。這種自動化營造了一個讓醫療人員能更專注於患者診療而非繁瑣資料輸入的環境。

在護理領域，AI語音辨識的應用也正在擴大。AI可以錄音並分析護理師與患者的對話語音，並以此為基礎自動生成護理問題、護理計畫、患者基本資訊等，並與電子病歷（EHR）進行部分連動。在日本，也有報告指出透過護理記錄的語音輸入化，將記錄撰寫時間縮短了70%。AI還包含基於影像分析預先撰寫報告的功能。AI速記員系統在直接整合至EHR時效果最大化，可在醫師與患者對話期間自動生成記錄，並將其整理為EHR所需的格式，從而劃時代地縮短文件作業時間。

診療記錄草稿、出院摘要、編碼、事前審查文件

AI技術正被用於基於語音辨識收集到的資料自動生成與處理各種醫療文件。

診療記錄草稿（Drafting Clinical Notes）與摘要：AI能以透過語音辨識取得的文字資料為基礎，自動生成診療記錄的草稿。AI可以整合患者的多份診療記錄（例如8份筆記），生成單一的臨床摘要（clinical summary），這有助於在短時間內掌握患者複雜的病史。LLM能在短短數秒內讀完患者超過300頁的醫療記錄，迅速掌握並摘要過敏史、過往用藥史、手術記錄及風險因素等。AI還能將診間對話轉為文字，並根據需要生成包含圖表的說明，作為患者說明資料使用。

出院摘要（Discharge Summaries）：AI能極快地生成如出院摘要等特定的醫療文件。根據日本長野市民醫院（Nagano Municipal Hospital）的案例，使用AI套件可在約20秒內生成出院摘要。這能摘要住院患者複雜的診療經過，並迅速撰寫包含出院後所需指示事項的文件，從而大幅減輕醫療人員的行政負擔。LLM亦可用於支援臨床試驗摘要或FDA監管法規文件（dossier）的編寫，FDA甚至自2025年5月起批准了利用LLM模型撰寫最終文件。

醫療編碼（Medical Coding）支援：雖然沒有明確提及AI直接生成診斷及處置代碼，但在簡化用於請款（billing）的收據/單據（receipt）撰寫或生成臨床摘要（clinical summary）等文件化作業過程中，AI可發揮自動擷取與結構化醫療編碼所需資訊的作用。這有助於提高保險請款過程的準確性與效率，並提升醫療行政系統的生產力。

事前審查（Prior Authorization）：來源中並未明確提及AI直接支援事前審查過程。然而，AI透過優化臨床工作流程並實現行政業務自動化，幫助醫療人員專注於患者治療，這一點間接表明AI有助於提高如事前審查等行政程序的效率。AI作為一種「創造時間」的工具，有可能被用於簡化與複雜審查程序相關的文件作業。

記錄錯誤的人工審查與個人資料保護

雖然AI正在革新診療記錄及文件化過程，我們必須明確認識到AI是一種可能犯錯的工具，且嚴格遵守個人資料保護是必不可少的。

醫療人員審查的必要性：AI可能會出現「幻覺（hallucination）」現象，這意味著生成看似合理但完全錯誤的資訊，據報告即使像GPT-5這樣的大型語言模型（LLM）也有約15%的幻覺錯誤。此類錯誤在醫療領域可能導致「1%的致命失誤（fatal error）」，實際上曾因AI給出錯誤健康建議（例如限制鈉攝取）導致患者出現中毒症狀而住院的案例。甚至還可能發生AI在未獲得患者同意的情況下，卻在診療記錄中插入已同意等字樣的幻覺，從而引發法律風險。

因此，AI生成的所有診療記錄草稿、摘要或建議，都必須經過醫療專業人員的批判性審查與最終臨床判斷。最終的決策與責任始終歸屬於人類醫療專業人員。AI並非要取代醫師，而是作為「智慧且不知疲倦的同僚」或「第二雙眼睛」來輔助醫療人員的工具。醫療人員不應盲目遵從AI的建議，而應將自身的臨床判斷與獨立推理留存於記錄中。AI無法取代人類醫師所擁有的同理心、直覺及掌握非語言資訊的能力等人性化要素，而這些能力在醫療現場依然保有其重要價值。

個人資料保護與資料治理：醫療資料是最敏感且受到法律保護的個人資料之一，因此對於AI系統如何存取、控制及使用資料存在相當大的疑慮。

患者的事前知情同意（informed consent）：在基於AI系統使用前，必須向患者說明AI使用目的、涵義與潛在風險並取得知情同意。患者即使同意接受診療，亦有權不同意將自身的醫療資料用於AI訓練或研究使用。

資料去識別化與最小化：為了保護個人資料，資料去識別化（anonymization）或假名化（pseudonymization）至關重要，且必須遵循資料最小化原則。

資料儲存與安全性：應徹底遵守禁止輸入個人資料的原則，以免將個人資料直接輸入AI服務中。特別是必須防止敏感醫療資料被儲存於外國伺服器，並需要有合約上的保證，確保資料不會被用於AI學習。部分醫院在地端（on-premise）環境中運行AI，以防止資料洩漏至外部雲端。

法規監管與責任：AI技術的發展速度已超越監管框架的變化速度，因此迫切需要更新國家及機構層面的規定以引入AI。在訓練AI模型時使用醫院記錄資料的同意問題、醫院是否能否繼續收容拒絕AI醫療的患者等，都需要找到法律與倫理上的平衡點。努力解決難以理解AI系統如何做出決策的「黑箱（black box）」問題，並實現可解釋AI（Explainable AI, XAI）至關重要。

AI縮短了醫療人員的文件作業時間，並提高了資訊檢索與摘要的效率，從而提供了向患者提供更多人性化關懷與溝通的寶貴機會。然而，AI無法取代人類醫師所擁有的同理心、直覺及掌握非語言資訊的能力等人性化要素，這些能力在醫療現場依然保有其重要價值。透過AI與人類的協作，醫療人員將能以符合「實際照護需求」而非「為了技術而技術」的方式持續改善醫療服務。AI對醫療人員扮演著「第三輔助者（third agent）」的角色，只有當人類的專業知識與AI的技術和諧結合時，才能真正實現以患者為中心的醫療。

[參考資料]

以下是支援第5章「醫院業務與患者體驗」之5.1「診療記錄與文件自動化」、5.2「病床、人力、手術室、預約管理」、5.3「患者諮詢與遠距監測」的參考資料清單。

參考資料清單

標題: 1 醫療 RAG大模型实

作者/發行機構: AI大模型用发 (YouTube 頻道)

支援的小標題: 5.1

標題: 2026-06-03 數位健康研究轉譯系列演講#27 繆睿股份有限公司 陳凌漢 講師「醫療衛教培訓 AI Coach 共創工作坊」

作者/發行機構: ICHIT TMU (YouTube 頻道)

上傳日期: 2026-06-03

支援的小標題: 5.1, 5.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支援的小標題: 5.1, 5.2, 5.3

標題: AI與健康照護 | 楊珮菁、王献章 | AI科普沙龍講座第3期 - BEING HUMAN：AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: L' IA & l' IA Générative en santé : À l'épreuve de la réalité

作者/發行機構: L'Oeil de la E-santé (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=7EQazuMjJh4>

支援的小標題: 5.1, 5.2, 5.3

標題: 【看護DXアワード2026予選】看護現場を⌈えるAI・DXサービス10社ピッチ

作者/發行機構: 日本男性看護師會の感護⌈書 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=Y9HWglbMSyU>

支援的小標題: 5.1, 5.2, 5.3

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=vUv4dl0vHLw>

支援的小標題: 5.1, 5.2, 5.3 支援的小標題: 5.1, 5.2, 5.3

5.2 病床、人力、手術室、預約管理

人工智慧（AI）正在崛起成為創新改善醫療機構營運效率、最佳化運用有限醫療資源，並最終向患者 提供更好醫療服務的核心工具。AI分析複雜醫院營運環境中發生的各種資料，預測患者流向、有效分配病床、高效配置醫療人力、最佳化手術室日程，並自動化預約與候診管理等，在廣泛領域中 提供智慧化支援。然而，引進此類AI 系統時，不僅要考慮技術效率，還必須共同考量現場安全審查、倫理考量以及醫療人員的最終判斷與責任等重要面向。

患者流向預測與病床分配

AI透過預測患者流向大幅提升醫院營運效率。AI模型可預測兒科患者住院（pediatric hospitalizations），這是考量一年中的時節、病毒傳播週期、區域爆發等各種因素而進行的。此類預測在醫院需求較大的 時期（例如病毒流行高峰期）對於提高醫院效率發揮著核心作用。

病床分配是醫院營運的重要環節，AI使這一過程更加客觀且高效。在加護病房（ICU）等病床嚴重短缺的環境中，決定應優先收治哪位患者可能非常複雜。預計未來AI將能根據患者的護理評估分數（nursing assessment score）、TTSS分數（TTSS score）等臨床資訊，進行更加及時與客觀的病床分配。AI還能預測患者的住院風險或病情惡化風險，引導早期住院並有助於預防不必要的併發症。AI透過兒科患者住院預測，預先評估增設病床或擴充加護病房、胸腔科、心臟科專業人力的必要性，從而最佳化病床管理。透過應用程式發出住院可能性警告並提出其他替代方案，可以更加切合實際地規劃患者轉診/轉介過程（referral and counter-referral processes）。

人力配置高效化

醫療人力短缺是全球性問題，AI在解決此問題中發揮著重要作用。AI超越了單純補充人力，透過智慧化有效配置有限人力，並有助於減輕醫療人員的業務負擔。

減輕業務負擔：AI護理師助手執行陪伴患者、預防跌倒、指引路線等重複性業務，減輕護理師的業務負擔，並為護理師爭取更多時間向患者提供人文關懷（human care）。亦有報告指出，引進基於AI的護理記錄系統後，護理師花費在間接業務上的時間最多減少了70%，使其能將更多時間投入於患者身上。

最小化資源浪費：因院內設備（infusion pumps, stretchers）遺失或擺放錯誤，員工往往會浪費30至60分鐘尋找設備。AI可追蹤此類設備，以減少人力浪費並有助於高效分配設備。

預測人力需求：AI透過兒科患者住院預測，預測加護病房、胸腔科、心臟科領域的額外人力需求，從而支援醫療人員的高效配置。

手術室日程管理

手術室管理是醫院營運中最複雜且重要的部分之一。AI透過手術計畫、日程最佳化以及術中支援，提高手術的效率與安全性。

最佳化手術計畫：AI分析腫瘤位置、大小、周圍解剖結構等個體化資料，最佳化放射治療計畫，這能減輕放射腫瘤專科醫師的工作負擔並提高治療計畫的一致性。基於AI的系統被整合至EHR中，以個體化資料為基礎支援醫療資源分配並減少不必要的工作量。

自動化日程管理：醫院手術室管理（OP-Sälen管理）與一般工廠不同，每天都有許多需要重新決定的事項。例如，有時會發生所需手術器械未及時送達而必須變更手術計畫的狀況。AI有助於管理此類不可預測性並最佳化手術室日程。

支援臨床決策：AI有助於癌症分期（cancer staging）等文件化作業，並構建治療計畫（treatment planning）的骨架，供會議進行審查與批准（review and ratify）。這是嘗試將決策過程代碼化並建立系統，以符合臨床醫師的思考過程。

預約與候診管理

AI自動化並最佳化預約系統與候診管理，改善患者體驗並提高醫院的營運效率。

縮短候診時間：AI可使急診室行政錯誤減少35%，並將急診室候診時間從21分鐘縮短至9分鐘。這有助於減少醫院擁擠，使患者能更快接受診療。

24/7預約服務：由於34%~51%的預約發生在非工作時間，因此對24/7（全天候24小時）預約服務的需求極大。AI代理能以遠低於客服中心或人類代理的成本提供24/7服務。AI能即時確認可預約時段，並協助透過語音或其他數位管道立即做出決定。

減少爽約（No-show）：AI透過預約確認（confirming attendance）及重新預約（rescheduling）提醒，減少爽約（no-shows）並填補空白診號/時段（empty agendas）。這能減少患者的候診時間並提高醫院的營運效率。

自動化電話接聽：AI電話系統自動將與患者的對話轉換為文字並進行摘要，改變工作流程，讓櫃檯人員或員工能預先掌握電話內容並迅速處理。該系統不僅能對應診療預約，還能對應健檢預約及其他諮詢。AI製作的摘要會透過URL發送患者所說的內容，供患者親自確認從而減少錯誤。

支援ICD編碼：AI基於建議代碼實現ICD編碼自動化，減少行政資源消耗。

現場安全審查與醫療人員判斷

雖然AI大幅提高了醫療現場的營運效率，明確認識其局限性與倫理考量，並尊重醫療人員的最終判斷是不可或缺的。AI可能會犯下「1%的致命錯誤（fatal error）」，這在醫療領域可能帶來嚴重後果。

AI的錯誤與偏見：AI模型可能會反映或加劇訓練資料中存在的偏見（bias），且存在因幻覺（hallucination）現象而生成看似合理但錯誤資訊的風險。此類錯誤可能會降低醫療服務提供者的可信度。例如，AI模型曾對醫療人力配置進行預測，但由於偏遠地區資料傳輸延遲（frequent power cuts），資訊缺失被解讀為需求低迷，導致在人力分配中發生系統性遺漏的問題。該案例顯示出AI在無法理解語境/脈絡（context）時可能失敗，並強調了區域性知識（local knowledge）與持續的人工監督（continuous human supervision）的必要性。

虛警/誤報（False Alarms）與醫療人員審查：由於AI系統是伴隨高風險的診斷輔助工具，因此必須時刻意識到誤診或誤報的可能性。AI模型的預測是「機率」而非「絕對真理」。因此，AI提供的警告或預測必須始終經過醫療專業人員的批判性審查與監督。AI系統應以人機協同（human-in-the-loop）模式運作，在高風險作業中，人工批准（manual approval）是不可避免的。需要對AI輸出結果建立自動化驗證（automated verification）與校準（calibration）機制。

醫療人員的最終責任：AI並非要取代醫師，而是作為「智慧且不知疲倦的同僚」來支援醫療人員業務的工具。無論AI可能如何發展，最終診斷決策與治療的責任始終歸屬於人類醫療專業人員。如果AI已就潛在風險向患者警示，醫療人員仍根據自身的臨床判斷將其忽略，最終對患者造成負面影響，醫療人員的注意義務標準可能提高，從而加重法律責任。

資料品質與倫理：AI模型的效能高度取決於輸入資料的品質。若資料不準確或不完整，AI的預測將無法信任。醫療資料是最敏感的個人資料，因此在基於AI系統使用前，必須向患者說明AI使用目的、涵義與潛在風險並取得知情同意（informed consent）。

雖然AI是提高醫院營運多個面向效率與準確性的強大工具，但其引進必須建立在人工監督、倫理考量以及對AI局限性的明確理解基礎上。透過AI與醫療人員的有效合作，我們將能建立更加安全且高效的以患者為中心的醫療系統。

[參考資料]

以下是支援第5章「醫院業務與患者體驗」之5.1「診療記錄與文件自動化」、5.2「病床、人力、手術室、預約管理」、5.3「患者諮詢與遠距監測」的參考資料清單。

參考資料清單

標題: 1 醫療 RAG大模型实

作者/發行機構: AI大模型用发 (YouTube 頻道)

支援的小標題: 5.1

標題: 2026-06-03 數位健康研究轉譯系列演講#27 繆睿股份有限公司 陳凌漢 講師「醫療衛教培訓 AI Coach 共創工作坊」

作者/發行機構: ICHIT TMU (YouTube 頻道)

上傳日期: 2026-06-03

支援的小標題: 5.1, 5.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支援的小標題: 5.1, 5.2, 5.3

標題: AI與健康照護 | 楊珮菁、王獻章 | AI科普沙龍講座第3期 - BEING HUMAN：AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: L' IA & l' IA Générative en santé : À l'épreuve de la réalité

作者/發行機構: L'Oeil de la E-santé (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=7EQazuMjJh4>

支援的小標題: 5.1, 5.2, 5.3

標題: 【看護DXアワード2026予選】看護現場を⌘えるAI・DXサービス10社ピッチ

作者/發行機構: 日本男性看護師會の感護⌘書 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=Y9HWglbMSyU>

支援的小標題: 5.1, 5.2, 5.3

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=vUv4dl0vHLw>

支援的小標題: 5.1, 5.2, 5.3支援的小標題: 5.1, 5.2, 5.3

5.3 患者諮詢與虛擬助理

人工智慧（AI）透過患者諮詢及虛擬助理功能，在提高醫療服務可及性與實現以患者為中心的醫療方面發揮著核心作用。AI聊天機器人與虛擬助理回答患者提問、提供醫療資訊，並支援日常醫療行政業務，從而減輕醫療人員負擔並改善患者體驗。

患者聊天機器人與虛擬助理之應用

虛擬助理發揮著「第三輔助者（third agent）」的角色，減輕醫療專業人員的認知負擔，並自動化如文件作業等重複性行政業務，協助醫療人員將更多時間投入於患者身上。與聊天機器人不同，AI代理被設計為執行特定目標導向任務，能與環境互動並解決複雜問題。AI超越了醫療人員說明，就檢查結果或疾病向一般患者們的

AI聊天機器人與虛擬助理超越了以患者為中心醫療的擴展，正帶來醫療資源的高效分配與患者管理的創新。

症狀分流與初期指引

AI被用於分析患者主訴的各種症狀，在初期分流疾病風險並引導至適當的醫療服務。AI醫療助手能以龐大的實際醫療資料為基礎，針對患者的症狀 -- 例如「腹瀉3次」、「心臟疼痛」、「飢餓時打嗝」等內容 -- 提出合理的解答與治療方案。這有助於減少患者自行搜尋資訊的時間與精力，並減少不必要的醫療機構就診。AI在遠距醫療（telemedicine）系統中支援遠距檢傷分類（remote triage）功能，分流患者症狀嚴重程度，並在需要時連接至醫療專業人員，從而協助高效分配醫療資源。特別是，AI亦可用於分析患者資料，識別存在焦慮、憂鬱症或其他心理健康狀況風險的個體，從而實現早期介入。這種初期分流使患者能夠及時獲得所需協助，並讓醫療人員能專注於需要更深層管理的患者。

預約指引與候診管理

AI整合至醫療機構的預約系統中，有助於改善患者體驗並提高醫院營運效率。基於AI的系統提供全天候24小時（24/7）預約服務，其營運成本遠低於客服中心或人類代理。考慮到全體預約中有34%~51%發生在非工作時間，24/7預約服務的需求極大。AI能即時確認預約時段，並協助透過語音或其他數位管道立即做出決定。

此外，AI透過預約確認及重新預約提醒減少爽約（no-shows），並填補空白診號/時段，從而提高醫院營運效率。AI電話系統自動將與患者的對話轉換為文字並進行摘要，改變工作流程，讓櫃檯人員或員工能預先掌握電話內容並迅速處理。該系統不僅用於診療及健檢預約，還用於對應醫院FAQ（常見問題集）。若AI給出錯誤指引，員工可透過確認並將修正後的資訊傳達給患者的方式來管理錯誤，從而補足AI的局限性。AI代理讓患者能隨時提出任何問題，協助患者更多地諮詢並理解自身的狀況。

遠距監測與持續患者管理

AI持續監測從穿戴式裝置及物聯網（IoT）感測器收集到的患者生理數據（心率、血壓、血糖、睡眠模式等），以感測異常徵兆並發送智慧警報。AI應用程式監測治療順從性（adherence to treatment），並透過自動化訊息或提醒支援患者的自我管理，這在慢性病管理中尤為重要。這些功能亦可在行動裝置上實現。AI支援的遠距醫療（telehealth）包含遠距監測，能預測患者疾病並實現客製化治療。

如行動診所等移動型醫療服務會造訪患者家中，在進行身分驗證後啟動感測，顯示即時健康數據並轉為線上診療以領取藥物。AI能以個人健康數據為基礎，預測3年後的檢查結果或健康風險，並提供行為改變的建議。AI透過遠距監測感測患者的變化，並在需要時以連接救護車或啟動居家醫療團隊的心力衰竭監測系統等方式加以應用。患者積極的數據（patient reported data）回饋有助於醫師在診斷與治療中獲得更客觀的參考資料。此外，若製作患者所需的說明影片並放置於醫院，醫療人員便能減少說明時間並提高患者滿意度。

誤診風險與數位落差的陰影

雖然AI的患者諮詢與虛擬助理功能具有相當大的潛力，亦隱含著誤診風險及數位落差等重要的局限性與倫理考量。

誤診風險：AI可能會出現「幻覺（hallucination）」現象，這意味著生成看似合理但完全錯誤的資訊。據報告，即使像GPT-5這樣的大型語言模型（LLM）也有約15%的幻覺錯誤，此類錯誤在醫療領域可能導致「1%的致命失誤（fatal error）」，從而對患者造成嚴重傷害。實際上曾因AI給出錯誤健康建議（例如限制鈉攝取）導致患者出現中毒症狀而住院的案例。AI聊天機器人無法掌握情感複雜性，且存在缺乏同理心（empathy）能力、完全理解（comprende）或敏感性（sentire）的局限性。AI以固定模式運作，可能無法適當地適應每位患者的特定狀況，且難以考量社會背景或特殊情況。AI模型可能反映或加劇訓練資料中存在的偏見（bias），這可能對特定人口群體得出不公正或不準確的結果。

數位落差：基於AI的醫療服務可能因技術可及性而加劇醫療落差。高齡者或數位素養較低的患者在運用AI工具時可能遭遇困難，這可能導致醫療服務利用的不平等。由於並非所有人皆擁有智慧型手機或技術裝置，此類落差成為將AI效益平勻傳遞給所有人的障礙。設計數位工具時，必須分析特定受眾（TA）的需求，並配合其臨床路徑進行設計，要向所有人提供完全相同的服務是很困難的。

醫療人員的批判性審查與最終責任

AI生成的所有資訊都必須經過醫療專業人員的批判性審查與最終臨床判斷。最終的決策與責任始終歸屬於人類醫療專業人員。AI並非要取代醫師，而是作為「智慧且不知疲倦的同僚」或「第二雙眼睛」來輔助醫療人員的工具。醫療人員不應盲目遵從AI的建議，而應發揮人類特有的同理心、直覺及掌握非語言資訊的能力，綜合理解患者並做出最佳決定。

在基於AI系統使用前，必須向患者說明AI使用目的、涵義與潛在風險並取得知情同意（informed consent）。使用AI時應遵循透明度原則，明確告知患者正在與AI進行互動。AI存在黑箱（black box）問題，醫療人員必須評估是否能理解並信任AI的推理過程。必須意識到若AI做出錯誤判斷，醫療人員的注意義務標準可能提高，從而加重法律責任。需要對AI輸出結果建立自動化驗證（automated verification）與校準（calibration）機制，在高風險作業中，人工批准（manual approval）是不可避免的。醫療人員應成為決定對AI提出的風險分數或診斷可能性信任至何種程度以及如何介入的臨界值（threshold）設定主體。

雖然AI在提升患者諮詢品質與擴展醫療服務方面發揮著重要作用，必須明確認識其局限性，並在人性化判斷與倫理指南下審慎加以運用。

[參考資料]

以下是支援第5章「醫院業務與患者體驗」之5.1「診療記錄與文件自動化」、5.2「病床、人力、手術室、預約管理」、5.3「患者諮詢與遠距監測」的參考資料清單。

參考資料清單

標題: 1 醫療 RAG大模型實例

作者/發行機構: AI大模型應用發 (YouTube 頻道)

支援的小標題: 5.1

標題: 2026-06-03 數位健康研究轉譯系列演講#27 繆睿股份有限公司 陳凌漢 講師「醫療衛教培訓 AI Coach 共創工作坊」

作者/發行機構: ICHIT TMU (YouTube 頻道)

上傳日期: 2026-06-03

支援的小標題: 5.1, 5.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支援的小標題: 5.1, 5.2, 5.3

標題: AI與健康照護 | 楊珮菁、王獻章 | AI科普沙龍講座第3期 - BEING HUMAN : AI與人共生的那一天

作者/發行機構: 臺大科學教育發展中心CASE (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支援的小標題: 5.1, 5.3

標題: L' IA & l' IA Générative en santé : À l'épreuve de la réalité

作者/發行機構: L'Oeil de la E-santé (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=7EQazuMjJh4>

支援的小標題: 5.1, 5.2, 5.3

標題: 【看護DXアワード2026予選】看護現場を⌘えるAI・DXサービス10社ピッチ

作者/發行機構: 日本男性看護師會の感護⌘書 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=Y9HWgIbMSyU>

支援的小標題: 5.1, 5.2, 5.3

標題: 人工智慧基本法與智慧醫療治理趨勢（下）李崇億教授

作者/發行機構: 理律文教基金會 (YouTube 頻道)

URL: <https://www.youtube.com/watch?v=vUv4dI0vHLw>

支援的小標題: 5.1, 5.2, 5.3支援的小標題: 5.1, 5.2, 5.3

第六章 醫療教育與研究

6.1 醫學教育、虛擬病患、臨床訓練

人工智慧（AI）正在從根本上改變醫療教育的方式與醫療專業人員的勝任能力。AI 已成為個人化學習過程、安全提供實作環境以及訓練複雜臨床推理能力的必備工具。此類技術的發展在預備未來醫療人員理解 AI 的潛力、認知其局限性並負責任地加以運用方面發揮著重要作用。

醫學教育的變革與 AI 的角色

AI 正在革新培養醫療專業人員的方式。智慧模擬器協助學生在安全環境中練習與真實相似的臨床狀況。此外，基於 AI 的自適應學習系統根據個人的學習需求客製化教育內容，為學生提供個人化的學習體驗。得益於這些技術，農村地區的學生也能接觸到城市環境中提供的高品質學習資料，有助於打破長期存在的教育不平等障礙。許多醫學院正在探索將 AI 整合至正規教育課程中的方案，這指明了醫療教育的未來方向。

在醫療教育中運用 AI 的核心在於教授負責任且符合倫理的使用方法。學生必須接受訓練，掌握向 AI 發出精準指令（提示詞）的方法，並對 AI 提供的資訊抱持批判性視角。必須教育學生：AI 僅為工具，切勿盲目遵循其輸出，而必須經過嚴格審視。醫療專業人員必須具備管理與維護複雜 AI 系統所需的數位勝任能力，並透過對 AI 勝任能力與局限性的持續教育，培養有效運用的能力。未來的醫師將被要求具備熟練運用 AI，並在 AI 協助下提供更好病患照護的勝任能力。

透過虛擬病患進行臨床訓練

基於 AI 的虛擬病患（AI patient）是醫療教育中特別是用於診斷互動及溝通品質訓練的有效工具。虛擬病患協助準醫療人員在與真實病患相似的情境中練習診療技術並提升溝通能力。虛擬病患基於真實電子病歷（EHR）資料建立，擴展了疾病資料庫，並結合「Big Five」性格特質檔案，實現更為真實的感情表達。

利用此類虛擬病患進行的訓練為學生提供了快速且有效率地練習更多案例的機會。例如，在護理師訓練計畫中，透過與 AI 病患的實作對話，利用基於真實對話資料生成的提示詞進行現實的對話練習。訓練後，可透過文字確認自己的溝通內容並進行客觀反思，且能接收個別回饋，使學習效果最大化。然而，AI 病患有時會給出過長的回答或表現出不自然的反應，這一部分需要持續改善。

強化臨床推理訓練

AI 有助於訓練與強化醫療人員的臨床推理能力（clinical reasoning）。臨床推理是收集病患資訊、提出假設並反覆驗證的複雜且動態的思維過程。AI 支援這一思維過程，並為學生提供反覆練習各種情境推理能力的環境。

AI 透過映射（mapping）病患的高維度健康狀態，達到能夠預測個人生活軌跡的水平，協助學生更深入地理解複雜的疾病進展過程。例如，亦有研究結果顯示，AI 模型輸入至 60 歲為止的健康狀態作為 Token，能夠高度相似地預測此後的生活軌跡。AI 的這種能力使學生能夠根據病患資料預測未來並規劃先發制人的干預訓練。

臨床推理訓練中的重要要素之一是同理心（empathy）。AI 雖然能進行表面的同理表達，但難以理解病患的深層語境並形成本質上的同理。例如，AI 對於拒絕服藥的病患可能回應「很辛苦吧？」，但臨床推理能力出色的醫師能理解該病患因類風濕性關節炎手痛而無法取出藥物，又因自尊心而不願向家人求助的深層心理狀態，並表達本質上的同理。因此，在 AI 時代，醫療人員整合與消化 AI 所提供資訊的編排（orchestration）能力、理解病患生活脈絡與故事的敘事（narrative）能力、憑藉直覺與同理心選擇適當言語的能力，以及基於信任提供醫療服務變得更加重要。

醫療人員的 AI 理解能力與局限

AI 在醫療現場作為「聰明且不知疲倦的同事」或「第二雙眼睛」來輔助醫療人員。AI 能捕捉人類可能忽略的微妙細節，並將撰寫文件、編碼、營運排程等重複性業務自動化，減輕醫療人員的負擔。藉此，醫療人員能花更多時間與病患相處，並專注於人文關懷（human care）。

然而 AI 具有局限性與潛在錯誤。

錯誤與偏見：AI 可能犯下「1% 的致命錯誤 (fatal error)」，LLM 存在約 15% 的幻覺 (hallucination)，即生成看似合理卻錯誤資訊的傾向。亦有報導指出因 AI 的錯誤建議（例如：建議限制鈉攝取）導致病患經歷實際中毒症狀的案例。此外，AI 模型可能反映訓練資料中存在的偏見 (bias)，對特定人口群體得出不公平或不精準的結果。

情感及非語言理解不足：AI 在捕捉人類的情感複雜性、同理心、直覺、非語言資訊（表情、肢體動作等）方面存在困難。例如，當高齡病患想去賞櫻花時，AI 可能基於心臟衰竭或跌倒風險而勸阻外出，但人類醫療人員能理解病患生命的意義，在管理風險的同時尋求協助外出的妥協點。

「黑盒子 (black box)」問題：AI 的決策過程不透明，存在醫療人員難以理解與信任其推理的「黑盒子」問題。因此，對 AI 輸出的批判性審視變得尤為重要。

醫療人員的最終責任

無論 AI 如何發展，最終診斷決策與治療責任始終歸屬於人類醫療專業人員。AI 不應取代醫師，而應作為「增強智慧 (augmented intelligence)」工具來強化醫師的能力。醫師必須從醫學與倫理角度批判性地評估 AI 提出的資訊，並判斷 AI 提供的建議是否適合病患的特殊情況。亦有人擔憂，若 AI 發出了風險警告而醫師予以忽視並導致不良結果，可能會提高醫療人員的注意義務基準，從而加重法律責任。因此，醫療人員對 AI 的使用必須向病患透明地說明並取得知情同意。

醫學教育應集中於明確教授 AI 的此類能力與局限，培養既能有效運用 AI 又不失人類固有能力的醫療人員。AI 應發揮輔助醫療人員業務、深化與病患關係並有助於提升醫療服務品質的「第三協助者 (third agent)」作用。

[根據資料]

以下是支援第六章「醫療教育與研究」之 6.1「醫學教育與虛擬病患」、6.2「研究資料分析與假設生成」、6.3「文獻審視與臨床試驗支援」的根據資料列表。

根據資料列表

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 6.1, 6.3 (與 RAG 及批判性審視相關)

標題: 首創AI醫療影像3D標準化 MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 6.1

標題: "友☒健康☒据科学"高☒量☒据集及AI成果☒☒

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 6.2

標題: NICE Talk 100: 模拟、☒☒和醫學推理的AI实☒

作者/發行機構: NICE AI Talk (YouTube 頻道)

支持的小標題: 6.2

標題: Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者/發行機構: MédicosÉlite® (YouTube 頻道)

支持的小標題: 6.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 6.3

標題: 大模型技☒如何☒能醫學研究与☒床实☒-AI+医☒☒合实☒☒目实操，人工智能行业☒家唐宇迪

作者/發行機構: 唐宇迪【人工智能实☒☒家】 (YouTube 頻道)

支持的小標題: 6.2

標題: 醫療AI 開発プロジェクトを成功に導くために：醫療AI開発・提供に関する法律・知財・契約

作者/發行機構: (国研) ☒総研 人工知能技術コンソーシアム (AITeC) (YouTube 頻道)

支持的小標題: 6.3

6.2 研究資料分析與假設生成

人工智慧 (AI) 在醫療研究領域透過分析龐大資料與識別複雜模式的能力，在深化對疾病的理解與生成新假設方面發揮著核心作用。這革新性地加速了傳統上耗費大量時間與巨大人力的研究過程，補足了人類研究者的局限。AI 作為核心動力，在從理解疾病的分子機制到發掘新治療方法、臨床試驗設計等所有研究階段中提升效率與成功率。

大規模臨床資料分析

基於 AI 研究的根基在於大規模臨床資料。AI 整合並分析來自病患電子病歷（EHR）、基因體資訊、醫療影像、穿戴式裝置及 IoT 感測器資料以及實驗室資料等多個來源收集的異質資料。其中包含病患口述紀錄或醫病對話轉錄本、編碼資料、時間序列臨床觀察資料等結構化與非結構化資料。AI 透過如 HL7 FHIR 等標準整合這些資料，並進行匿名化處理以強化個人隱私保護，同時使其能應用於研究。

然而此類大規模資料往往缺乏高品質資料，且將資料轉化為 AI 能有效分析的「AI-ready」狀態的資料數據策展（data curation）面臨著極其困難且耗時的挑戰。資料分散或缺漏值眾多的情況十分常見，AI 要有效分析這些資料，精細的預處理過程不可或缺。AI 透過正規化（normalization）、選擇最重要特徵、資料增強等技術對資料進行預處理，分析在大規模基因體測序計畫中發現的數百萬個基因體變異體以識別致病變異體，並預測個人對特定疾病的脆弱性。AI 亦可用於蛋白質結構預測與藥物設計。

6.2 研究資料分析與假設生成

人工智慧（AI）在醫療研究領域透過分析龐大資料與識別複雜模式的能力，在深化對疾病的理解與生成新假設方面發揮著核心作用。這革新性地加速了傳統上耗費大量時間與巨大人力的研究過程，補足了人類研究者的局限。AI 作為核心動力，在從理解疾病的分子機制到發掘新治療方法、臨床試驗設計等所有研究階段中提升效率與成功率。

大規模臨床資料分析

基於 AI 研究的根基在於大規模臨床資料。AI 整合並分析來自病患電子病歷（EHR）、基因體資訊、醫療影像、穿戴式裝置及 IoT 感測器資料以及實驗室資料等多個來源收集的異質資料。其中包含病患口述紀錄或醫病對話轉錄本、編碼資料、時間序列臨床觀察資料等結構化與非結構化資料。AI 透過如 HL7 FHIR 等標準整合這些資料，並進行匿名化處理以強化個人隱私保護，同時使其能應用於研究。

然而此類大規模資料往往缺乏高品質資料，且將資料轉化為 AI 能有效分析的「AI-ready」狀態的資料數據策展（data curation）面臨著極其困難且耗時的挑戰。資料分散或缺漏值眾多的情況十分常見，AI 要有效分析這些資料，精細的預處理過程不可或缺。AI 透過正規化（normalization）、選擇最重要特徵、資料增強等技術對資料進行預處理，分析在大規模基因體測序計畫中發現的數百萬個基因體變異體以識別致病變異體，並預測個人對特定疾病的脆弱性。AI 亦可用於蛋白質結構預測與藥物設計。

AI 在醫療研究領域透過分析龐大資料與識別複雜模式的能力，在深化對疾病的理解與生成新假設方面發揮著核心作用。這革新性地加速了傳統上耗費大量時間與巨大人力的研究過程，補足了人類研究者的局限。AI 作為核心動力，在從理解疾病的分子機制到發掘新治療方法、臨床試驗設計等所有研究階段中提升效率與成功率。

AI 生成假設的臨床研究驗證

AI 生成的假設無論多麼具革新性與吸引力，都必須經過嚴格的臨床研究予以驗證。AI 能使用電腦模型最佳化模擬具有特定活性物質的分子形成過程，這有助於設計與生成新的類藥物分子。然而，曾有案例顯示 AI 設計的化合物後來被發現化學上不切實際或難以合成，這暗示著原位（in silico）預測與真實化學之間存在跨度。

AI 模型常在實驗室條件下受訓，但往往缺乏在真實臨床環境中泛化可能性（generalizability）的系統性測試。外部驗證不足可能降低診斷準確度，特別是當演算法應用於與訓練資料差異巨大的資料集時，此現象尤為突出。事實上，許多研究依賴小規模樣本（small effectives），限制了結果的泛化可能性。此外，當資料不平衡（data imbalance）嚴重時，存在過度擬合（overfitting）風險，AI 在訓練資料中看似完美，但在真實狀況下可能犯錯。

為了確保 AI 模型的再現性（reproducibility）與可比性，標準化的資料獲取方式與透明的文件化不可或缺。對於基於 AI 的醫療器材，將資料管理設計文件化是提供 AI 達到預期效能水準證據的唯一方法。在臨床環境中量化不確定性（uncertainty）至關重要，可透過設定自適應閾值（adaptive thresholds）來識別不確定案例，並將其交由醫療專業人員解讀（human-in-the-loop），以此方式運用 AI。例如，AI 模型準確度從 81% 提升至 91% 的案例，展現了協助 AI 進行可靠預測並讓模糊案例經過人類審視的人機協作模式的可能性。

避免虛假相關性（Spurious Correlations）的方法

AI 能快速分析龐大資料並找出模式，但有時會將虛假相關性作為真實呈現。因為 AI 僅提供基於統計概率的資訊，而非提供經醫學驗證的真實。

「幻覺（Hallucination）」現象：LLM（大語言模型）可能引發生成看似合理卻錯誤資訊的幻覺（hallucination）。ChatGPT 對真實並不關心，而被設計為生成看似真實的文字，甚至有批評認為其產出物可能是「廢話（bullshit）」。這意味著 AI 提供的資訊有時可能看似具有統計顯著性，但實際上並無價值的「虛假證據（false evidence）」。

資料偏見 (Data Bias)：AI 模型會原封不動地反映或加劇訓練資料中存在的偏見 (bias)。例如，在缺乏特定種族資料的情況下，AI 可能對該種族病患做出帶有偏見的診斷。這引發了人們對 AI 可能如「演算法之鏡 (algorithmic mirror)」般放大既有不平等的擔憂。當 AI 訓練資料無法反映地方現實時（例如：經常停電地區的資料空白），AI 可能犯錯，這強調了需要地方知識 (local knowledge) 與持續的人類監督。

理解因果關係 (Causality)：AI 擅長尋找相關性 (correlation)，但在明確證實因果關係 (causality) 方面存在局限。研究者必須透過人類的批判性推理，驗證 AI 提出的相關性是否確實為疾病的因果。

研究者與醫療人員的判斷及責任

無論 AI 如何發展，最終診斷決策與治療責任始終歸屬於人類醫療專業人員。AI 不應取代醫師，而應作為「增強智慧 (augmented intelligence)」工具來強化醫師的能力。醫師必須從醫學與倫理角度批判性地評估 AI 提出的資訊，並判斷 AI 提供的建議是否適合病患的特殊情況。亦有人擔憂，若 AI 發出了風險警告而醫師予以忽視並導致不良結果，可能會提高醫療人員的注意義務基準，從而加重法律責任。因此，醫療人員對 AI 的使用必須向病患透明地說明並取得知情同意。

批判性審視：研究者與醫療人員對於 AI 提出的資訊應摒棄盲目信任，並以批判性視角進行審視。比起照單全收 AI 的輸出物，基於醫學專業知識與臨床推理對資訊進行質疑與確認的能力更為重要。

人類監督 (Human Oversight)：AI 系統應設計為「human-in-the-loop」方式，在高風險作業中，人類的手動批准 (manual approval) 不可或缺。為了解決難以理解 AI 決策過程的「黑盒子 (black box)」問題，可解釋 AI (Explainable AI, XAI) 研究非常重要，這有助於解釋 AI 為何得出特定結論。

AI 素養 (AI Literacy)：醫療專業人員必須理解 AI 的勝任能力與局限，並接受關於如何負責任且符合倫理地使用 AI 工具的持續教育與訓練。AI 作為「聰明且不知疲倦的同事」是協助研究與臨床的工具，但人類的同理心、直覺以及掌握非語言資訊的能力是 AI 無法取代的固有領域。

資料治理與倫理：在 AI 研究中，資料隱私 (data privacy) 與資料治理 (data governance) 是核心倫理考量事項。病患的知情同意 (informed consent) 不可或缺，對於 AI 模型訓練的資料進行匿名化 (anonymization) 或假名化 (pseudonymization) 至關重要。

AI 是加速研究過程並促進新發現的強大夥伴，但為了最大化利用其潛力的同時不失安全性、倫理與以人為本，需要持續付出努力。AI 與人類研究者的合作是醫療科學發展的必備要素。為了不失去安全性、倫理與以人為本，需要持續付出努力。AI 與人類研究者的合作是醫療科學發展的必備要素。

[根據資料]

以下是支援第六章「醫療教育與研究」之 6.1「醫學教育與虛擬病患」、6.2「研究資料分析與假設生成」、6.3「文獻審視與臨床試驗支援」的根據資料列表。

根據資料列表

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 6.1, 6.3 (與 RAG 及批判性審視相關)

標題: 首創AI醫療影像3D標準化 暨 MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 6.1

標題: "友 健康 据科学"高 量 据集及AI成果

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 6.2

標題: NICE Talk 100: 模拟、和醫學推理的AI实

作者/發行機構: NICE AI Talk (YouTube 頻道)

支持的小標題: 6.2

標題: Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者/發行機構: MédicosÉlite® (YouTube 頻道)

支持的小標題: 6.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 6.3

標題: 大模型技術如何能醫學研究與臨床實-AI+醫-合實-目实操，人工智慧行業家唐宇迪

作者/發行機構: 唐宇迪【人工智慧實家】(YouTube 頻道)

支持的小標題: 6.2

標題: 醫療AI 開發プロジェクトを成功に導くために：醫療AI開発・提供に関する法律・知財・契約

作者/發行機構: (国研) 総研 人工知能技術コンソーシアム (AITeC) (YouTube 頻道)

支持的小標題: 6.3

AI 在醫療研究領域透過分析龐大資料與識別複雜模式的能力，在深化對疾病的理解與生成新假設方面發揮著核心作用。這革新性地加速了傳統上耗費大量時間與巨大人力的研究過程，補足了人類研究者的局限。AI 作為核心動力，在從理解疾病的分子機制到發掘新治療方法、臨床試驗設計等所有研究階段中提升效率與成功率。

AI 生成假設的臨床研究驗證

AI 生成的假設無論多麼具革新性與吸引力，都必須經過嚴格的臨床研究予以驗證。AI 能使用電腦模型最佳化模擬具有特定活性物質的分子形成過程，這有助於設計與生成新的類藥物分子。然而，曾有案例顯示 AI 設計的化合物後來被發現化學上不切實際或難以合成，這暗示著原位 (in silico) 預測與真實化學之間存在跨度。

AI 模型常在實驗室條件下受訓，但往往缺乏在真實臨床環境中泛化可能性 (generalizability) 的系統性測試。外部驗證不足可能降低診斷準確度，特別是當演算法應用於與訓練資料差異巨大的資料集時，此現象尤為突出。事實上，許多研究依賴小規模樣本 (small effectives)，限制了結果的泛化可能性。此外，當資料不平衡 (data imbalance) 嚴重時，存在過度擬合 (overfitting) 風險，AI 在訓練資料中看似完美，但在真實狀況下可能犯錯。

為了確保 AI 模型的再現性 (reproducibility) 與可比性，標準化的資料獲取方式與透明的文件化不可或缺。對於基於 AI 的醫療器材，將資料管理設計文件化是提供 AI 達到預期效能水準證據的唯一方法。在臨床環境中量化不確定性 (uncertainty) 至關重要，可透過設定自適應閾值 (adaptive thresholds) 來識別不確定案例，並將其交由醫療專業人員解讀 (human-in-the-loop)，以此方式運用 AI。例如，AI 模型準確度從 81% 提升至 91% 的案例，展現了協助 AI 進行可靠預測並讓模糊案例經過人類審視的人機協作模式的可能性。

避免虛假相關性 (Spurious Correlations) 的方法

AI 能快速分析龐大資料並找出模式，但有時會將虛假相關性作為真實呈現。因為 AI 僅提供基於統計概率的資訊，而非提供經醫學驗證的真實。

「幻覺（Hallucination）」現象：LLM（大語言模型）可能引發生成看似合理卻錯誤資訊的幻覺（hallucination）。ChatGPT 對真實並不關心，而被設計為生成看似真實的文字，甚至有批評認為其產出物可能是「廢話（bullshit）」。這意味著 AI 提供的資訊有時可能看似具有統計顯著性，但實際上並無價值的「虛假證據（false evidence）」。

資料偏見（Data Bias）：AI 模型會原封不動地反映或加劇訓練資料中存在的偏見（bias）。例如，在缺乏特定種族資料的情況下，AI 可能對該種族病患做出帶有偏見的診斷。這引發了人們對 AI 可能如「演算法之鏡（algorithmic mirror）」般放大既有不平等的擔憂。當 AI 訓練資料無法反映地方現實時（例如：經常停電地區的資料空白），AI 可能犯錯，這強調了需要地方知識（local knowledge）與持續的人類監督。

理解因果關係（Causality）：AI 擅長尋找相關性（correlation），但在明確證實因果關係（causality）方面存在局限。研究者必須透過人類的批判性推理，驗證 AI 提出的相關性是否確實為疾病的因果。

研究者與醫療人員的判斷及責任

無論 AI 如何發展，最終診斷決策與治療責任始終歸屬於人類醫療專業人員。AI 不應取代醫師，而應作為「增強智慧（augmented intelligence）」工具來強化醫師的能力。醫師必須從醫學與倫理角度批判性地評估 AI 提出的資訊，並判斷 AI 提供的建議是否適合病患的特殊情況。亦有人擔憂，若 AI 發出了風險警告而醫師予以忽視並導致不良結果，可能會提高醫療人員的注意義務基準，從而加重法律責任。因此，醫療人員對 AI 的使用必須向病患透明地說明並取得知情同意。

批判性審視：研究者與醫療人員對於 AI 提出的資訊應摒棄盲目信任，並以批判性視角進行審視。比起照單全收 AI 的輸出物，基於醫學專業知識與臨床推理對資訊進行質疑與確認的能力更為重要。

人類監督（Human Oversight）：AI 系統應設計為「human-in-the-loop」方式，在高風險作業中，人類的手動批准（manual approval）不可或缺。為了解決難以理解 AI 決策過程的「黑盒子（black box）」問題，可解釋 AI（Explainable AI, XAI）研究非常重要，這有助於解釋 AI 為何得出特定結論。

AI 素養 (AI Literacy)：醫療專業人員必須理解 AI 的勝任能力與局限，並接受關於如何負責任且符合倫理地使用 AI 工具的持續教育與訓練。AI 作為「聰明且不知疲倦的同事」是協助研究與臨床的工具，但人類的同理心、直覺以及掌握非語言資訊的能力是 AI 無法取代的固有領域。

資料治理與倫理：在 AI 研究中，資料隱私 (data privacy) 與資料治理 (data governance) 是核心倫理考量事項。病患的知情同意 (informed consent) 不可或缺，對於 AI 模型訓練的資料進行匿名化 (anonymization) 或假名化 (pseudonymization) 至關重要。

AI 是加速研究過程並促進新發現的強大夥伴，但為了最大化利用其潛力的同時不失安全性、倫理與以人為本，需要持續付出努力。AI 與人類研究者的合作是醫療科學發展的必備要素。為了不失去安全性、倫理與以人為本，需要持續付出努力。AI 與人類研究者的合作是醫療科學發展的必備要素。

[根據資料]

以下是支援第六章「醫療教育與研究」之 6.1「醫學教育與虛擬病患」、6.2「研究資料分析與假設生成」、6.3「文獻審視與臨床試驗支援」的根據資料列表。

根據資料列表

標題: Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題: 6.1, 6.3 (與 RAG 及批判性審視相關)

標題: 首創AI醫療影像3D標準化 暨 MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的小標題: 6.1

標題: "友健康数据科学"高量数据集及AI成果

作者/發行機構: Grandhealth Access Telemedicine (YouTube 頻道)

支持的小標題: 6.2

標題: NICE Talk 100: 模拟、和醫學推理的AI实

作者/發行機構: NICE AI Talk (YouTube 頻道)

支持的小標題: 6.2

標題: Inteligencia Artificial en Medicina | Capacitación para Médicos y Personal de Salud

作者/發行機構: MédicosÉlite® (YouTube 頻道)

支持的小標題: 6.3

標題: 2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題: 6.3

標題: 大模型技術如何赋能醫學研究與臨床-AI+醫藥融合項目实操，人工智能行業專家唐宇迪

作者/發行機構: 唐宇迪【人工智能實家】 (YouTube 頻道)

支持的小標題: 6.2

標題: 醫療AI 開發プロジェクトを成功に導くために：医療AI開発・提供に関する法律・知財・契約

作者/發行機構: (国研) 総研 人工知能技術コンソーシアム (AITEC) (YouTube 頻道)

支持的小標題: 6.3

第三部

AI 醫療的爭議與未來

第七章 信任、安全與責任

7.1 可解釋性・驗證・錯誤管理

人工智慧（AI）在醫療診斷、治療計劃、醫院營運等多個領域展現出創新潛力，但同時也引發了信任、安全與責任等根本性課題。特別是 AI 模型的透明度不足、預測錯誤可能性以及責任歸屬的模糊性，是將 AI 成功整合至醫療現場必須解決的問題。解決這些問題的核心在於確保 AI 的可解釋性（Explainability）、經過嚴格的臨床驗證、透過持續的性能監測來管理錯誤，並最終建立醫療團隊的批判性審查與最終責任。

1. 醫療 AI 的黑盒子問題

AI，特別是基於深度學習（Deep Learning）的模型，常被比喻為「黑盒子（black box）」。這是因為人類難以理解演算法到達特定結論或預測的過程。此種不透明性（opacity）在醫療領域構成嚴重問題。

責任歸屬的模糊性：當 AI 複雜且不透明地運作時，發生錯誤時應由誰負責的問題變得更加模糊。若無法重構 AI 系統為何得出特定結果，便難以釐清責任歸屬。

法規與倫理障礙：監管機構（例如 FDA）可能不願批准未能滿足透明度要求的黑盒子模型。AI 的黑盒子特性存在損害基本醫療倫理並對患者造成有害結果的風險。

錯誤識別的困難：若難以理解 AI 模型的內部運作，便非常難以識別與修正系統性錯誤或偏見（bias）。

為解決這些問題，可解釋 AI（Explainable AI, XAI）領域應運而生。XAI 旨在以人類可理解的方式解釋 AI 系統的決策過程，從而提高信任度與透明度。XAI 透過熱圖（heatmap s）強調 AI 在醫療影像中關注的區域，或提供文字說明以揭示預測依據。

2. 臨床驗證

AI 模型若要安全且有效地應用於醫療現場，嚴格的臨床驗證不可或缺。AI 被歸類為醫療器材，需要監管機構（例如 FDA）的事前批准，這必須像藥品一樣透過嚴格的臨床試驗予以證實。

資料品質與代表性：

AI 模型的性能很大程度上取決於高品質的龐大訓練資料。低品質或具偏見的資料可能導致具誤導性的結論。

特別是醫療影像資料屬於高維度且須以多元資料集進行訓練。僅以單一機構資料訓練的模型，在應用於其他醫院或患者群體時性能可能下降 [41, 131, 應公開統計詳細資訊以早期發現不均衡，並透過定期偏見審計與具公平意識的演算法系統性地識別與修正偏見。

雖然使用外部 AI 預測工具，但評估其準確性的醫院僅占半數。

為證實 AI 系統的實際臨床實用性，需要採用多中心（multicenter）方式的前瞻性臨床研究。147]。

介面與使用適合性：AI 系統的實用性取決於使用者的接受度（acceptability），這取決於演算法的透明度與對建議事項的明確解釋能力 [13。

偏移（drift）與偏見檢測：AI/ML 系統在安裝後須接受定期觀察以檢測偏移（drift）或偏見（bias）。這屬於監督義務：如 EU AI Act 等法律框架要求高風險 AI 系統在運作期間保持監督義務。

4. 幻覺與誤診

AI，特別是大型語言模型（LLM）有時可能生成與事實不符的資訊。這可能導致錯誤的健康資訊或建議，對患者健康造成嚴重危害。

實際案例：在某些案例中，AI 建議高齡患者不要去賞櫻，忽略了患者的情緒滿足感或生活品質等人性化因素。

誤診與錯誤：

AI 是基於機率的結果，而非替代性的輔助資訊。

AI 是可能犯錯的工具。已有研究報告指出在治療建議中包含 1% 致命錯誤的 AI 模型 [1] 人性化因素的局限：AI 難以完全理解人類的情感語境、非語言資訊及同理能力。即使 AI 所提供的回答比醫師更有同理心

7.2 個人資料・安全・資料治理

隨著人工智慧（AI）深度整合至醫療領域，處置患者敏感健康資料的方式、用以保護資料的安全體系，以及管理資料全生命週期的資料治理，已成為確保醫療 AI 可信度中最重要且具挑戰性的課題。AI 系統基於龐大的敏感資料進行學習與運作，因此確保此類資訊的隱私、完整性與可用性，對於保護患者基本權利及履行醫療系統的倫理與法律責任至關重要。

1. 敏感健康資料

患者的健康資料本身是非常敏感且個人化的資料，也是 AI 系統的核心資源。

資訊種類：敏感健康資料包含個人診斷結果、基因資訊、生理資料（生物特徵測量資料）、治療紀錄等與患者健康狀況相關的所有資訊。這些資訊以醫療團隊診療筆記、醫療影像、病理結果、基因體資料等多種形式存在。

高價值與高風險：這些資料對個人而言是私生活的重要部分，對企業而言具有重大經濟價值。但同時，若遭誤用或外洩，可能對患者造成嚴重傷害，屬於極度敏感的資訊。

AI 對資料的依賴性：AI，特別是深度學習模型，依賴大規模敏感醫療資訊進行訓練與運作。由於 AI 的性能與準確度與訓練資料的數量及品質直接成正比，因此 AI 開發者試圖獲取儘可能多的資料。

2. 知情同意（Consent）

收集、儲存、處理與利用患者敏感健康資料時，必須獲得患者的知情同意（Informed Consent）。

明確同意的重要性：同意不僅僅是許可使用資料，而必須在向患者明確解釋並使其理解 AI 系統如何運作、使用何種資料以及存在何種潛在風險後方可達成。歐盟 GDPR（通用資料保護規則）與美國 HIPAA（健康保險可攜與責任法案）等法規框架對此類同意管理明定嚴格要求。

用於 AI 訓練的同意：將患者資料用於 AI 系統訓練時，原則上除非徹底實現資料匿名化，否則必須獲得患者同意。然而，針對將患者 EMR（電子病歷）資料用於 AI 訓練一事向所有患者個別取得概括性同意，在實務上存在困難。

診療過程錄音/錄影的同意：當 AI 即時分析醫師與患者間的對話以實現診療紀錄自動化時 [5.1]，必須取得患者與醫療團隊的相互同意。在台灣，法規規定醫療機構於診療過程中進行錄音或錄影時，必須獲得對方同意。

未能取得同意時的問題：若 AI 系統在未取得患者個人同意的情況下利用診療資料，可能侵害個人的資料自主權與隱私。特別是由於資料追蹤與更新的困難，可能限制對 AI 診斷與治療結果進行長期追蹤研究。

倫理與法律責任：同意相關法規雖使 AI 實現過程變得複雜且高成本，但對於保護患者自主性（autonomy）與隱私不可或缺。

3. 資料存取權限

AI 系統使用大量敏感醫療資訊，因此嚴格管理資料存取權限是資料安全的關鍵要素。

必要的保護措施：AI 系統需要強大的資料保護機制，包括安全儲存、適當的存取控制以及持續的安全監測。

存取控制的困難：與放在桌上任何人皆可翻閱的傳統紙本病歷不同，電子健康紀錄（EHR）存取門檻高得多。然而，由於 AI 系統從多個端點收集資料並分散於多個平台，可能更容易遭受資料外洩風險。

區塊鏈技術的應用：如 MedRec 等區塊鏈技術利用密碼學雜湊（cryptographic hashes）與智慧合約（smart contracts）實現對醫療資料的可控存取，從而滿足 GDPR 及 HIPAA 的要求。

「資料治理」的核心：資料存取權限管理為資料治理的核心部分，意指明確定義並控制何人、何時、基於何種目的可存取資料。

4. 網路安全

在醫療 AI 系統中，網路安全對於保護敏感健康資料及維持系統可信度絕對至關重要。

資料外洩風險：基於 AI 的系統易受資料外洩風險影響，威脅患者資訊的機密性。特別是隨著 AI 技術商業化，資料濫用及去識別化資料被重新識別的可能性隨之提高。

法規欠缺與因應局限：法規常落後於社會對隱私與安全不斷變化的態度。AI 的發展速度遠超法規框架，使得監管機構難以跟上 AI 開發速度。

法律要求事項：歐盟 GDPR、美國 HIPAA 以及歐盟 AI Act 要求嚴格的資料安全與個人資料保護規定。這些法規為醫療資料的處置與儲存設定了嚴格框架。

安全基礎設施的重要性：穩固的 IT 安全、資料保護概念、透明的資訊與同意程序不可或缺。基於雲端的 AI 平台與醫療資料的整合會產生額外的 IT 安全要求。

供應鏈安全：如 NIS 2 指令等法規強制基於 AI 的醫療系統在整個供應鏈中遵守安全要求。也就是說，當醫療機構採購硬體或軟體時，該服務提供者亦須符合相同的安全標準。

醫療機構的責任：醫療機構對使用 AI 系統導致的資料外洩與安全漏洞表示擔憂，並質疑 AI 系統能否安全保護患者資料。

5. 去識別化與重新識別風險

去識別化（De-identification）是保護患者隱私的核心方法，但重新識別（Re-identification）的風險始終存在。

去識別化的概念：去識別化是移除所有個人識別資訊以確保資料匿名性的過程。資料（data）在匿名化狀態下無法獲知特定個人的資訊，而資訊（information）則展現特定個人的資訊。

匿名化（Anonymization）與假名化（Pseudonymization）：

匿名化：意指移除患者姓名、身分證字號等所有識別資訊，使無法重新識別特定個人。匿名化資料無需儲存於通過 HDS（Health Data Storage）驗證的環境中。

假名化：意指使用如患者 ID 等假名資料（pseudonymized data），使無法直接識別特定個人，但透過額外的對照表（mapping table）仍可重新識別的狀態。假名化資料必須儲存於 HDS 驗證伺服器中。

重新識別風險：即便是假名化資料，若與其他資訊結合仍存在重新識別風險。這威脅到機密性（confidentiality），且隨著 AI 技術商業化，此類風險正日益擴大。

聯邦學習（Federated Learning）：被提出作為一種無需共享原始患者紀錄、透過分散的本地資料訓練模型以保護隱私並降低重新識別風險的方法。

台灣的資料利用問題：以台灣為例，在將如「健康保險資料」等大規模醫療資料用於研究目的時，儘管已移除所有識別資訊，患者依然認為其資料應受保護。此種「資料主權（data sovereignty）」概念為資料利用製造了法律與倫理難題。特別是當無法進行資料追蹤時，可能導致難以驗證 AI 診斷/治療的長期效果或在發生錯誤時難以補償。

6. 機構的責任

基於 AI 的醫療系統成功導入與運作，取決於醫療機構明確的責任意識與體系化的治理。

最終責任歸於人類：AI 是輔助工具而非取代醫療團隊，醫療決策的最終責任始終由人類醫師承擔。AI 僅是工具，不具備人格亦無法承擔責任。

遵守法規：醫療機構必須遵守 GDPR、HIPAA、EU AI Act 等國家及國際資料保護與 AI 法規。這是確保 AI 系統安全且合法使用的必需步驟。

AI Act 的角色：歐盟 AI Act 對高風險 AI 系統提出網路安全、透明度、人類監督等嚴格要求。

資料治理：醫療機構應建立資料治理基礎設施，以確保資料品質、相互操作性與安全性。

建立治理結構：

人類監督（Human Oversight）：AI 系統雖自主運作，但最終決定始終須在人類監督下達成。

倫理委員會與 DPO：倫理委員會（ethical committee）及資料保護官（DPO）必須參與，以評估 AI 專案的倫理影響並遵守資料保護法規。

政策與程序：應制定明確的政策、程序與責任分工，以實現 AI 系統的安全部署。

員工教育與 AI 素養：應提供 AI 素養教育，使包含醫療團隊在內的所有員工理解 AI 系統的功能、局限及倫理使用方法。特別是須培養認知 AI 生成資訊的錯誤可能性（幻覺現象）並進行批判性審查的能力。

導入新技術的法律責任：醫療機構在導入如 AI 等新技術時極度重視法律責任風險。若 AI 24 小時監測患者資料，醫院的注意義務及法律責任範圍可能會擴大，因此需要對此進行明確的法律定義。

確保技術主權：醫療機構應確保對 AI 系統的技術主權（technological sovereignty），減少對外部供應商的依賴，並提高 AI 模型的透明度與控制力。

基於 AI 的醫療具有提升患者治療品質與增加醫療效率的巨大潛力，但為此必須先做到徹底保護敏感健康資料、建立強大的網路安全體系，以及對 AI 使用實施明確且負責任的資料治理。醫療機構不應僅是 AI 技術的單純使用者，更應作為主導 AI 系統安全與倫理應用的重要主體，盡到相應責任。

7.3 醫療團隊與 AI 協作及最終責任

1. 演算法偏見與公平性

AI 模型存在固有風險，即可能學習訓練資料中蘊含的偏見（bias），從而對特定患者群體導致不公平或不準確的結果。

減輕偏見策略：為解決此類演算法偏見問題並確保公平性（fairness），需要進行以下努力：

定期審計：在醫院使用 AI 系統前後應定期實施審計，以減少偏見並提升透明度。

公開人口統計詳細資訊：應公開用於模型訓練之資料集的人口統計詳細資訊，以早期發現不均衡。

公平意識演算法：應透過定期的偏見審計與具公平意識的演算法，系統性地識別與修正偏見。

資料集平衡與權重調整：應使用資料集平衡調整、權重調整、對抗性偏見減輕等多種技術，在所有患者群體中產出公平的結果。

資料保護與所有權：對資料保護與所有權的強有力考量，對於減輕此類偏見風險至關重要。

倫理與法規框架：AI 應在穩固的倫理框架內進行開發，並須保障個人資料保護、公平性與透明度。

2. 患者與醫療團隊的信任

AI 系統的成功導入與持續使用，取決於能否取得患者與醫療團隊雙方的信任。

建立信任的困難：對 AI 系統的信任不僅來自技術穩固性，亦來自透明度、可提出異議性以及與組織價值（特別是公平性）的連結。然而，若醫療團隊在 AI 開發過程中遭到疏離，信任度將降低，AI 的採納亦會受阻。

患者安全疑慮：護理師們對 AI 若發生誤動作可能對患者安全造成嚴重風險表示擔憂。AI 所提供資訊可能包含錯誤（幻覺現象）這一點對信任產生重大影響。

促進信任方案：

透明溝通：使用 AI 時向患者透明地提供資訊，並就 AI 的角色與局限進行明確溝通至關重要。

3. 醫師的最終判斷

AI 的本質局限：

醫療團隊的角色：

人性化再最佳化：AI 提出的「最佳化」建議應由人類醫師根據患者的特殊性與價值觀進行「再最佳化」。例如，即使 AI 因跌倒風險勸阻高齡患者外出，人類醫師考量賞櫻對患者生活品質產生的積極影響，可在家人支援下允許外出等尋求折衷點。

4. 責任分配

當前責任分配的模糊性：

缺乏關於 AI 模型開發者、使用者（醫師）、系統維護者中應由何者負責的全球法律或標準。

關於責任分配的討論：

注意義務標準提高：有人擔憂導入 AI 可能提高醫療團隊 Notice of Care（注意義務）標準。若 AI 已發出特定風險警報而醫師將其忽視並對患者造成傷害，可能被視為過失。

共同責任模型：部分觀點提出 AI 開發者、醫療服務提供者、醫院之間共同責任模型，明

[參考資料]

以下為支持第七章「信任、安全與責任」之 7.1「可解釋性・驗證・錯誤管理」、7.2「個人資料・安全・資料治理」、7.3「醫療團隊與 AI 協作及最終責任」的參考資料清單。

參考資料清單

標題：#49【学会報告】発表のコツ・準備、AI、大規模研究 vs 症例報告【ESICM、日本救急医学会】

作者/發行機構：ICUトーク (YouTube 頻道)

支持的小標題：7.1, 7.2, 7.3

標題：2026-06-03 數位健康研究轉譯系列演講#27 繆睿股份有限公司 陳凌漢 講師「醫療衛教培訓 AI Coach 共創工作坊」

作者/發行機構：ICHIT TMU (YouTube 頻道)

上傳日：2026-06-03

支持的小標題：7.1, 7.2, 7.3

標題：2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構：健康智慧生活圈 (YouTube 頻道)

上傳日：2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的小標題：7.1, 7.2, 7.3

標題：AI 基因報告判讀、5G 遠距應用與臨床倫理研習會

作者/發行機構：TCnews慈善新聞網 (YouTube 頻道)

支持的小標題：7.2, 7.3

標題：Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構：Facultad de Medicina UNAM (YouTube 頻道)

支持的小標題：7.1, 7.3

標題：Fuerza de trabajo en salud y la inteligencia artificial: seminario web

作者/發行機構：PAHO TV (YouTube 頻道)

支持的小標題：7.1, 7.2, 7.3

標題：Health-IT Talk: IT-Sicherheit: Mit KI zum resilienten Klinikbetrieb

作者/發行機構：SIBB e.V. (YouTube 頻道)

支持的小標題：7.2, 7.3

標題：IA et santé : les enjeux réglementaires entre RGPD, RIA et Code de la santé publique

作者/發行機構：HAAS Avocats - Droit IT-IP (YouTube 頻道)

支持的小標題：7.1, 7.2, 7.3

第八章 導入之壁與下一步

8.1 技術基礎設施與現場驗證

1. 資料基礎設施 (Data Infrastructure)

碎片化資料：醫療資訊來自多個來源，並透過各種管道與技術碎片化地存在。例如，日本的醫療資訊架構中，各部門（藥物、檢驗、影像等）的資料獨立存在，難以進行橫向整合。這種異質（heterogeneous）系統與缺乏互操作性（interoperability）是導入 AI 的巨大障礙。

FHIR (Fast Healthcare Interoperability Resources) 的重要性：FHIR 正成為醫療資料整合的核心技術。相較於 AI、Transformer、數位雙生等其他技術，FHIR 被強調為最重要的技術，因為唯有 FHIR 能整合所有資料，而一旦資料完成整合，所有創新皆成為可能。台灣以成為全球首個將所有醫療中心以 FHIR 連接的國家為目標。

2. 模型泛化 (Model Generalization)

AI 模型在特定研究環境中展現出成功性能，並不代表能在所有臨床環境中發揮相同性能。模型的泛化能力（generalizability）對於 AI 的實際臨床應用而言是極為重要的要素。

小規模資料集的局限：許多 AI 研究基於少數病患資料（例如：47人、25人）或不平衡資料（例如：陽性26人對陰性108人），導致結果的可泛化性受到限制。

AI 模型的可變性：AI 軟體透過機器學習具有隨時間變化的特性。這與傳統醫療器材許可（固定結構）產生衝突。FDA 考量到 AI 的這種可變性，發布了「預先確定變更控制計畫（predetermined change control plan）」，提出了在一定範圍內的變更無需重新審查、僅憑監測即可允許的方案。

3. 外部驗證 (External Validation)

為了將 AI 模型實際應用於臨床，必須透過脫離訓練環境的獨立資料集進行嚴格的外部驗證。

評估標準：AI 模型在臨床使用前，必須透過敏感度（sensitivity）與特異度（specificity）等性能指標證明達到一定水準以上的性能。然而根據2025年的一項研究顯示，美國65%的醫院使用 AI 預測工具，但僅有三分之二評估其準確性。

4. 醫院系統串接 (Hospital System Integration)

為使 AI 系統在醫療現場有效發揮作用，與現有醫院系統（EMR、PACS 等）的順暢串接至關重要。

互操作性問題：AI 應用程式整合面臨著與異質資訊系統間互操作性的問題。以日本為例，各 EMR 廠商的架構不同，難以進行 AI 系統的橫向部署。

EMR 整合的重要性：當 AI 應用程式直接整合至 EMR 時，醫療人員的業務效率可大幅提升；但若非如此，因多個軟體間的複製貼上作業將加重業務負擔，進而阻礙 AI 的導入。

資料融合：AI 可將病患的多元醫療資料（DNA、代謝物、MRI、超音波、臨床紀錄）融合至大規模模型中進行分析與預測。為此，必須將結構化資料與非結構化資料（診療紀錄、影像說明等）全面整合並整理為標準化形式。

基於 FHIR 的標準化：台灣正在開發強大的轉換工具，以將現有 EMR 系統與新標準（FHIR、Smart on FHIR、Federated Learning、DICOM）串接，並期盼透過這項努力達成將所有醫療中心以 FHIR 連接的目標。

資料完整性與錯誤風險：當 AI 模型整合 EMR 資料並生成新內容時，發生錯誤的可能性會增加。因此，精確掌握 AI 模型與醫院資訊系統的整合程度以及擷取何種資料至關重要。

資料湖與聯邦學習：透過建置用於 AI 模型訓練的中央資料湖（central data lake），或利用無須將資料傳送至醫療機構外部即可訓練模型的聯邦學習（federated learning）平台，能在保護資料隱私的同時實現整合式學習。

5. 部署後性能監測 (Post-deployment Performance Monitoring)

AI 模型在部署之後，仍須透過持續的性能監測來適應變化的環境，並防止意料之外的錯誤或性能下降。

性能變化：AI 模型的性能可能隨時間改變，因此持續的重新訓練 (retraining) 與更新不可或缺。

組織治理：醫療機構應建立針對 AI 系統選擇、部署、維護的治理流程，限制「影子 IT (shadow IT)」，並透過所有相關主體 (使用者、IT、法律、管理者) 參與的委員會來監督 AI 導入的全過程。

6. 失敗案例應對 (Failure Case Response)

誤診與幻覺的風險：

AI 模型可能透過幻覺現象提供錯誤的健康資訊或建議，對病患健康造成嚴重損害。例如，曾有報導指出 AI 向病患推薦有毒物質，或診斷準確度僅為 31% 的情況。

AI 的不確定性管理：AI 模型在證據不足時，應展現不確定性並具備請求追加資訊等「安全能力 (safety capabilities)」，而非逕行做出判斷。在 10% 的情況下不作答、以換取 90% 可靠回答的模型，即為不確定性管理的優良範例。

如上所述，為確保基於 AI 的醫療系統安全導入與有效應用，技術基礎設施的堅固性、模型的泛化能力、嚴格的外部驗證、與醫院系統的順暢串接、部署後持續的性能監測，以及對錯誤的體系化管理皆不可或缺。在整個過程中，人類醫療人員的批判性判斷與最終責任，必須確立為 AI 時代醫療中永恆不變的核心原則。

8.2 人力・規管・保險・制度

1. 醫療人員培訓 (Medical Staff Training)：AI 素養與能力強化

AI 技術的迅速發展對醫療專業人員提出了新能力的要求，而醫療人員培訓則是 AI 導入成功不可或缺的要素。

AI 素養的必要性：許多醫學生 (90% 以上) 尚未接受過正式的 AI 教育，但約 75% 希望接受包含臨床應用及倫理問題在內的體系化 AI 訓練。學生與教師團隊雖然都有使用 AI 工具的經驗，但該領域的正規教育仍顯不足。UNESCO 的切入方式強調為了防止數位落差並保護醫病關係的人文主義取徑。

能力強化：對所有醫療衛生專業人員而言，包含資料素養、基本資訊架構理解、模型輸出評估、演算法偏見識別等基本能力的數位素養（digital literacy）訓練不可或缺。在加拿大，透過 AI 書記（scribe）訓練增加了醫療人員與病患面對面的時間，並減輕了行政與認知負擔，展現出正面成果。

持續教育：AI 技術不斷演進，因此 AI 教育與訓練必須持續且重複地進行。在日本，正討論自 2026 年起將每年一次針對 AI 適當利用的定期教育列為義務的方案。

2. 規管審準 (Regulatory Approval)：「活體軟體」的挑戰

各國規管框架：

美國 FDA：發布了強調臨床驗證、持續監測及 AI 功能透明度的指南。

日本：透過 AI 網路社會倡議推動有責任感的 AI 整合，特別是在處理病患資訊時，必須遵守「3省2指南（厚生勞動省、經濟產業省、總務省）」。

3. 責任 (Accountability/Responsibility)：以人為本的最終判斷

5. 機構治理 (Institutional Governance)：AI 導入的體系化

為將 AI 安全且高效地整合至醫療現場，醫療機構的體系化治理不可或缺。

組織準備：AI 導入並非單純的技術問題，而是需要全組織的變革管理。醫院及醫療機構應將 AI 專案與機構的願景與策略相契合進行情境化與調整。

治理結構：應設立如「人類擔保委員會（collège de garantie humaine）」等倫理委員會，評估 AI 專案的倫理影響，並監督 AI 系統的選擇、部署與維護。該委員會應由機構管理層、IT 部門、法律專家、倫理委員會、病患代表以及實際醫療人員共同參與。

技術主權：醫療機構應確保對 AI 系統的技術主權（technological sovereignty），降低對外部供應商的依賴，並提高 AI 模型的透明度與控制力。

6. 區域與國家間的差距 (Gaps between Regions and Countries)：不平等加劇的疑慮

AI 具備消除醫療差距的潛力，但同時也隱含著反使不平等加劇的風險。

醫療資源失衡：AI 有助於改善失衡的醫療資源分配。例如，基於 AI 的可攜式設備與判讀系統可在醫療資源不足的農漁村地區彌補影像診斷能力，進而為縮小區域間醫療差距做出貢獻。

技術主權與從屬：若 AI 模型未以開放、透明且可稽核的形式開發，特定國家或區域可能面臨對 AI 技術的技術從屬性，甚至失去醫療系統的主權。

領導力的重要性：AI 時代若缺乏領導力，醫療系統將會碎片化，並因僅專注於短期獲利能力而阻礙長期發展與創新。為推動 AI 公平且包容的導入，國家層面的策略願景與領導力不可或缺。

8.3 智慧醫院・遠距監測・醫療 AI 代理

1. 智慧醫院 (Smart Hospitals)

智慧醫院是指利用 AI 與數位技術提升醫院營運效率，並提供以病患為中心之醫療服務的未來型醫療機構。

改善病患體驗：AI 禮賓服務或虛擬化身形式的 AI 在醫院服務台回答病患詢問並進行引導，有助於提升病患體驗並減少等待時間。AI 亦能以病患易於理解的語言解釋醫學內容，進而消除健康資訊落差。

2. 可穿戴式遠距監測 (Wearable Remote Monitoring)

AI 與可穿戴設備在醫院外部持續監測病患健康，為個人化健康管理與慢性病管理開創全新地平。

3. 醫療 AI 代理 (Medical AI Agents)

AI 代理在病患諮詢、資訊提供、行政業務支援等多個領域輔助醫療人員，提升醫療服務的效率與可及性。

病患諮詢與資訊提供：AI 聊天機器人回答病患與健康相關的問題以填補健康資訊落差，並在夜間時段亦能回應個人健康問題的提問，協助病患隨時取得所需資訊。AI 向病患提供實用資訊，以病患易懂的語言解釋醫學內容，並提供治療利弊與替代方案的資訊，大幅提升病患的資訊可及性。

症狀分流與預約管理：AI 代理可基於病患詢問進行症狀分流（triage），並引導至適當的醫療服務。此外，透過預約管理與重新預約的自動化縮短病患等待時間，並使門診病患預約管理高效化，藉此提升病患滿意度。

行政業務自動化：AI 代理可處理醫療紀錄草稿撰寫、出院摘要生成、帳單代碼分配、合規文件化等重複性行政業務，為醫療人員爭取時間以專注於病患治療。這使護理人員能投入更多時間於與病患的互動，並從重複的紀錄工作中解放出來。

未來的 AI 代理：AI 代理在未來亦有可能扮演病患法律代理人（proxy）的角色。這是因為 AI 能基於數十年來與病患對話所累積的資訊，成為最了解病患的存在。

4. 病患安全 (Patient Safety)

AI 為醫療領域帶來的革新性變化，同時也引發了對病患安全的重大疑慮。

危機感知與安全協定：如心理健康領域的 AI 聊天機器人般直接與病患互動的 AI 系統，必須能感知危機狀況並啟動安全協定。例如，在病患自殺風險時向專家發出警報或連接至緊急服務的功能不可或缺。

5. 人類監督 (Human Oversight)

6. 數位落差 (Digital Divide)

導入基於 AI 的醫療服務提供了消除醫療服務不平等的契機，但同時也隱含著使數位落差加劇的風險。

[根據資料]

以下為支持第8章「導入之壁與下一步」8.1、8.2、8.3之根據資料列表。

根據資料列表

標題：#49【学会報告】発表のコツ・準備、AI、大規模研究 vs 症例報告【ESICM、日本救急醫學會】

作者/發行機構: ICUトーク (YouTube 頻道)

支持的副標題: 8.1, 8.2, 8.3

標題：2026 光田醫院 生成式AI醫療實務應用課程0301 東海大學姜自強博士 20260701 161424 會議錄製

作者/發行機構: 姜自強 (YouTube 頻道)

支持的副標題: 8.1

標題：2026-07-08：星球永續健康線上直播-AI醫療治理(2)：AI 安全閘門治理沙盒(AI Airlock)

作者/發行機構: 健康智慧生活圈 (YouTube 頻道)

上傳日期: 2026-07-08

URL: https://www.youtube.com/watch?v=_yani8XjOe0

支持的副標題: 8.1, 8.2, 8.3

標題：Conferencia Más Salud. IA en Medicina: la herramienta del estudiante moderno

作者/發行機構: Facultad de Medicina UNAM (YouTube 頻道)

支持的副標題: 8.1, 8.2, 8.3

標題：【旭聯核心素養公益講座】AI時代的健康預警與精準照護，三明治世代如何守護全家未來？2026/5/28 19:00-21:00 ■ 王孟祺 | 珍世明眼科診所院長 ■ 裴晉國 | 開南大學碩士班專任助理教授

作者/發行機構: 樂學網 (YouTube 頻道)

上傳日期: 2026-05-28

支持的副標題: 8.3

標題：生成AI事例☒討会part2 ～診療・研究・醫療DXへの実☒編～

作者/發行機構: Kamijo Kyosuke (YouTube 頻道)

支持的副標題: 8.1, 8.3

標題：「臨床の質を上げる生成AI『使いこなし』最新アップデート

作者/發行機構: 文部科学省ポストコロナ事業山里海醫學共育プロジェクト (YouTube 頻道)

支持的副標題: 8.2

標題：首創AI醫療影像3D標準化 ☒ MR 應用成果發表大會

作者/發行機構: 智捷醫學科技 (YouTube 頻道)

支持的副標題: 8.1, 8.3

版權頁

人工智慧與醫療

ISBN |

作者 | 金炅辰律師

出版社 | 金炅辰律師 出版社

出版資訊

人工智慧與醫療

ISBN |

Author | 金炅辰律師

Publisher | 金炅辰律師出版社

© 2026 Kim Kyung-jin. All rights reserved.

未經作者同意，不得複製或散布本書任何內容。

Support This AI Library Book

If this book was useful, you can support future translations, public editions, and open AI knowledge projects through PayPal.



PayPal

<https://paypal.me/kimkjkj>

Scan the QR code or open the link above.