



統計科學研究所  
INSTITUTE OF  
STATISTICAL SCIENCE

中央研究院  
生命科學圖書館  
LIFE SCIENCE LIBRARY

# AI 在生物醫學的應用

中央研究院統計所/智慧健康計畫  
李易儒 博士

Sep. 8, 2026

# 為什麼是現在？



## 2024

諾貝爾化學獎

頒給 AlphaFold 團隊，肯定 AI 解決逾50年蛋白質折疊難題的貢獻；原始演算法 AlphaFold2 於2021年發表於 Nature，準確度達到與冷凍電顯、X光繞射等實驗方法相當的水準。

意義：這是AI首次因為在基礎科學上的具體貢獻獲得諾貝爾獎，象徵AI已具備可驗證、可複製的科學研究能力，而非僅止於工具輔助。

Jumper et al., Nature, 2021, 596:583–589



## 2023

LLM 進入臨床視野

以ChatGPT為代表的大型語言模型快速普及，Nature Medicine 發表系統性回顧，整理LLM在醫療上的應用現況與限制，涵蓋臨床決策、病歷摘要、病人衛教等多個場景。

意義：導入速度遠快於傳統醫療科技，也讓「怎麼負責任地使用」的討論，變得比過去更迫切。這正是稍後倫理段落要談的重點。

Thirunavukarasu et al., Nature Medicine, 2023, 29:1930–1940

# 為什麼是現在？

## **Areas where AI is having impact**

- Imaging, pathology and diagnostic support
- Clinical decision support, triage and risk prediction
- Documentation, discharge summaries and patient communication
- Precision medicine, drug discovery and public health surveillance

# 為什麼是現在？

## Ethical Concerns

- The potential for AI to make erroneous decisions; for algorithmic bias and hallucinations
- Determining who is responsible when AI is used to support decision-making;
- Data privacy and security;
- Maintaining **public trust** in the development and use of AI technologies;
- Effects on the roles and skill-requirements of healthcare professionals;
- Implications for regulation as AI algorithms may undergo continuous updates and improvements, requiring ongoing validation to ensure that their performance remains consistent over time

# 為什麼是現在？

## Public's Concerns

- Survey conducted by King's Health Partners, Responsible AI UK, and the Policy Institute at King's College London May 2026  
<https://www.kcl.ac.uk/policy-institute/assets/use-of-ai-in-healthcare.pdf>
- 2,093 UK adults aged 18+, online survey March 2026
- 15% of the public have used AI chatbots for health advice instead of contacting a GP
- 37% support and 38% oppose the use of AI in clinical decision-making
- Anxiety about safety and accuracy is main emotion (39%)
- Across a range of clinical scenarios, the public consistently place **more trust in doctors than in AI** –with trust highest for psychological therapy, where 78% say they trust a doctor more, and just 5% trust AI

# 為什麼是現在？

## Public's Concerns

- 58-63% of respondents say they should be told in advance and given the option to opt-out of out of AI use.
- If an NHS-approved AI disagreed with a doctor's diagnosis, a majority (55%) say a **second doctor should review** before any decision is made, with just 7% willing to follow the AI's advice on the basis that it may be more accurate.
- Three-quarters (76%) say AI tools used in patient care should be **officially approved and regulated**, even if this slows down their adoption.
- Only **17%** take the opposing view that doctors should be able to choose AI tools freely without official approval.

# 為什麼是現在？

## Recommendations

- Focus on developing AI tools that save GPs time in their routine work, not replacing clinical judgement.
- Immediately clarify professional liability and safe AI use and regulatory and governance frameworks.
- Develop interim local AI practice guidance until clearer national guidance and regulatory and governance frameworks are available.
- **Test, learn and share collectively.**
- Help to educate patients about the use of AI in practice.
- Develop comprehensive and structured training and education programmes.

# 為什麼是現在？

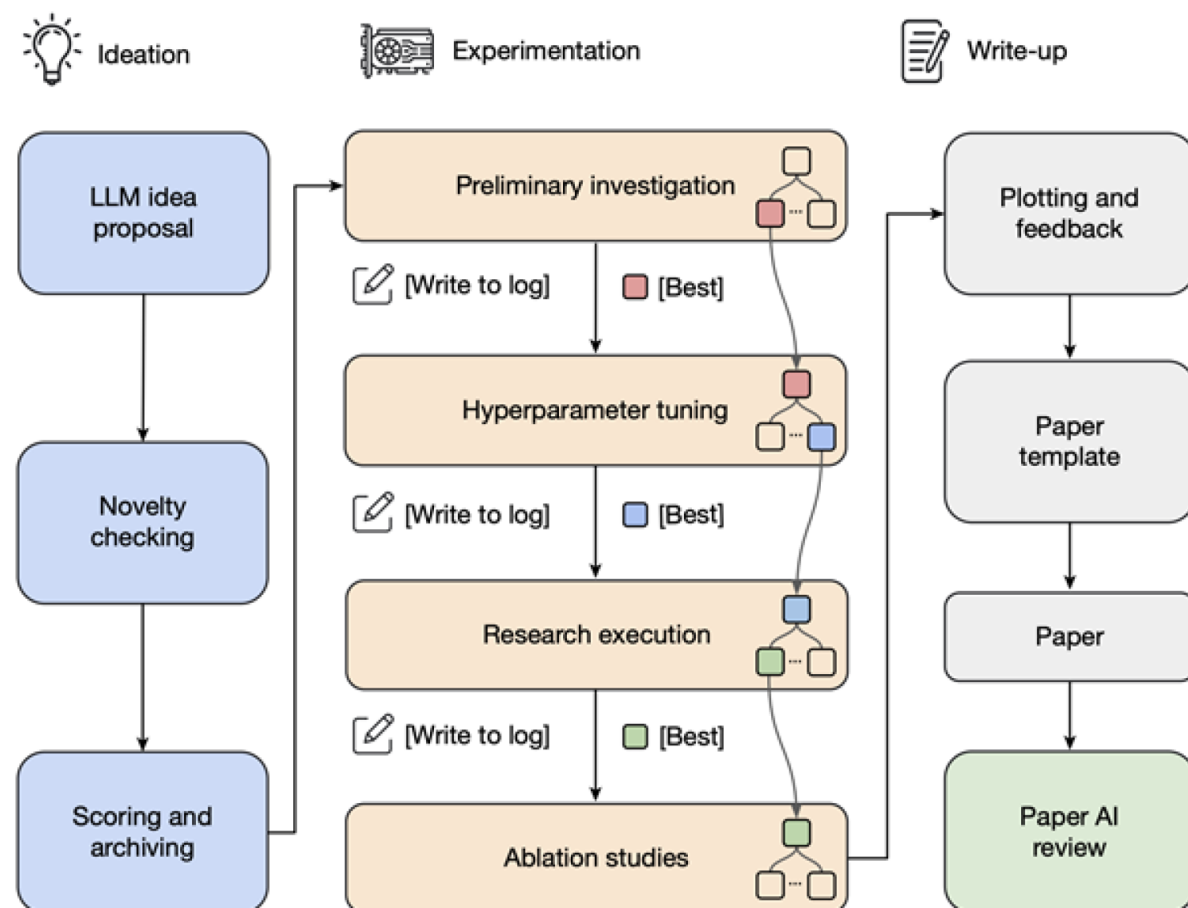
## Looking beyond the technology

- Focus has been on the technical aspects of AI, when successful implementation requires an understanding of the socio-technical context
- AI in healthcare is not just a technical innovation
- It reshapes clinical decisions, institutional responsibilities and patient trust
- The key question is not “Can we use AI?” but “Under what conditions should we trust it
- Ethical and regulatory safeguards must be built in from the start

# Towards end-to-end automation of AI research

[Chris Lu](#), [Cong Lu](#), [Robert Tjarko Lange](#), [Yutaro Yamada](#) ✉, [Shengran Hu](#), [Jakob Foerster](#), [David Ha](#) ✉ & [Jeff Clune](#) ✉

[Nature](#) **651**, 914–919 (2026) | [Cite this article](#)



# 生醫 AI 四象限

接下來的案例導覽，依此架構逐一展開



## 影像 / 病理判讀

放射科、病理科AI輔助診斷



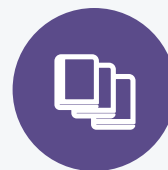
## 分子設計 / 藥物研發

蛋白質結構預測、藥物篩選



## 基因體 / 精準醫療

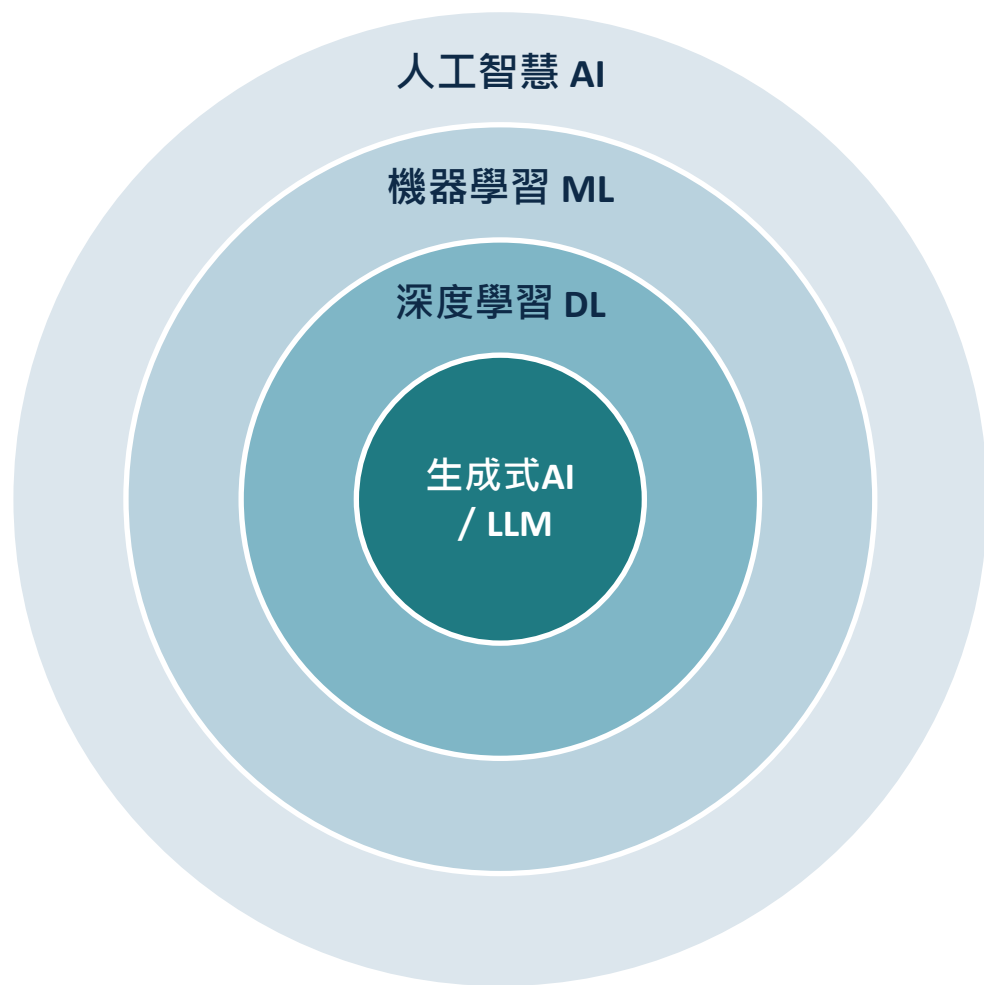
多基因風險評分、多體學整合



## 文獻知識 / 臨床決策

文獻回顧、證據整合、知識管理

# AI / 機器學習 / 深度學習 / 生成式AI



- 人工智慧（AI）：最大範疇，泛指讓機器展現智慧行為的技術
- 機器學習（ML）：從資料中學習規則，而非人工寫死規則
- 深度學習（DL）：以多層神經網路實現的機器學習方法
- 生成式AI / 大型語言模型（LLM）：從大量文本學會「語言與知識合成」，近年深度學習中最受矚目的分支



重點提醒：LLM 擅長生成「看起來合理」的內容，但這不等於「正確」的內容，這一點會在稍後的倫理討論以真實數據說明。

# Generative learning

including graph networks

## Metalearning

AI to train AI

## Reinforcement Learning

including adversarial debiasing

## Bayesian inference

Gaussian processes, Neural processes

## Privacy

dataset condensation, quantum encoding

## NLP

including Large Language Models

## Foundation models

fusing different data modes

# Hospital Care

NHS Trusts, International

## Primary Care

CPRD, RCGP, OpenSafely

## Cohort Studies

UK Biobank, China Kadoorie Biobank

## LMICs

ISARIC, OUCRU Vietnam

## Medicines / Molecules

PSI Vaccines, GSK, Chemistry, AZ, BMS, Pfizer

## Animal Health

Royal Veterinary College

## Social Care

Oxford County



案例導覽

# 醫學影像與病理 AI

從FDA核准現況到具體診斷應用

---

- FDA核准AI醫材的現況地圖
- 糖尿病視網膜病變、IDx-DR、病理切片三個實證案例
- 監管透明度的限制與挑戰

# FDA核准AI醫材：現況地圖

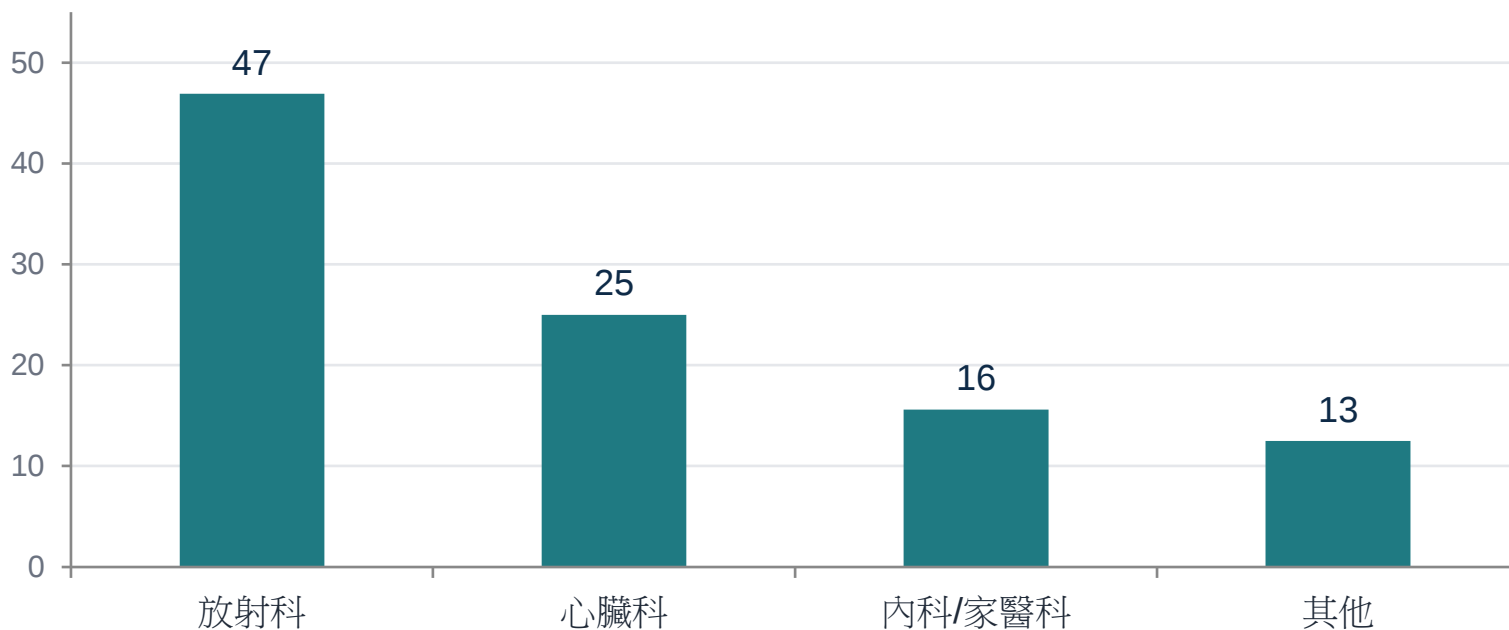
64

項 FDA 核准的  
AI/ML 醫材

( 截至2020年資料庫整理研究 )

僅 45% 的官方公告  
明確使用「AI/ML」字樣

各科別占比 (%)



 Benjamins, Dhunoo & Meskó, npj Digital Medicine, 2020, 3:118

## 三個具體應用案例



### 糖尿病視網膜病變篩檢

深度學習演算法對可轉診病變的判讀敏感度與特異度，達到與眼科專家相當的水準（訓練集逾12萬張眼底影像）。

 [\*Gulshan et al., JAMA, 2016, 316\(22\):2402–2410\*](#)



### IDx-DR：首個自主判讀AI醫材

樞紐試驗證實可在不需醫師覆核下自主判讀，成為FDA首個核准的自主判讀AI診斷系統。

 [\*Abràmoff et al., npj Digital Medicine, 2018, 1:39\*](#)



### 病理切片AI輔助判讀

弱監督式深度學習模型在攝護腺癌、基底細胞癌與乳癌淋巴結轉移判讀，AUC均超過0.98，可排除65–75%免人工複閱切片。

 [\*Campanella et al., Nature Medicine, 2019, 25\(8\):1301–1309\*](#)

## 限制與挑戰：透明度不足

45%

的FDA官方公告  
明確使用「AI/ML」字樣

45%

明確標示 AI/ML

55%

未明確標示

其餘超過一半的核准案例，官方公告並未清楚標示這是AI / 機器學習醫材

- 監管透明度仍有進步空間：多數核准公告未明確揭露AI/ML技術屬性
- 「輔助判讀」與「取代判讀」的角色分際，仍是持續討論的議題
- 對研究單位而言，這正是推動AI素養教育可以著力之處
- 使用者（醫師、病人）有權清楚知道自己正在接觸的是AI輔助的判讀結果

延伸思考：如果連監管公告都不夠透明，臨床端的醫師與病人又該如何被清楚告知？

📖 *Benjamins, Dhunoo & Meskó, npj Digital Medicine, 2020, 3:118*

案例導覽

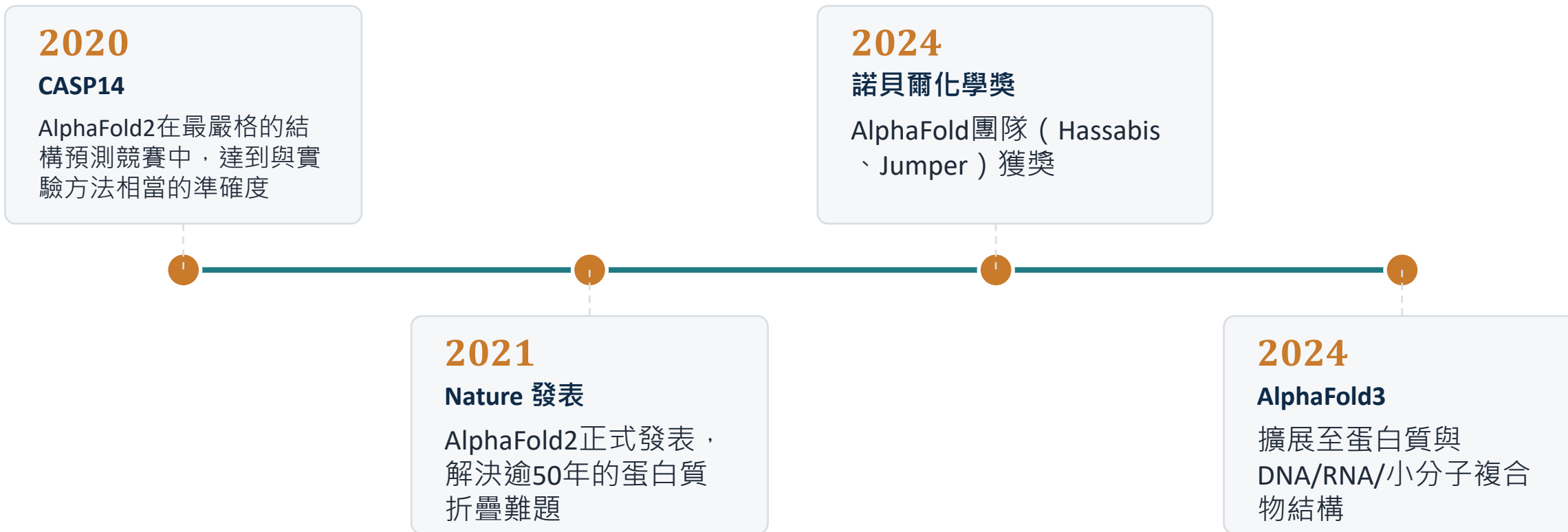
# 蛋白質結構與藥物研發

從AlphaFold2到AlphaFold3的技術演進

---

- AlphaFold的故事：從CASP14到諾貝爾化學獎
- AlphaFold3如何擴展到藥物研發應用

# AlphaFold 的故事



📖 [\*Jumper et al., Nature, 2021, 596:583–589\*](#)

📖 [\*Abramson et al., Nature, 2024, 630:493–500\*](#)

# 從單鏈結構到藥物研發



## AlphaFold2

預測單一蛋白質鏈的立體結構，準確度媲美冷凍電顯、X光繞射等實驗方法。

 [\*Jumper et al., Nature, 2021, 596:583–589\*](#)



## AlphaFold3

改用擴散模型架構，將預測擴展到蛋白質與DNA/RNA/小分子配體的複合物結構。DeepMind衍生公司 Isomorphic Labs 運用此技術加速藥物標靶篩選流程。

 [\*Abramson et al., Nature, 2024, 630:493–500\*](#)



案例導覽

# 精準醫療與基因體學

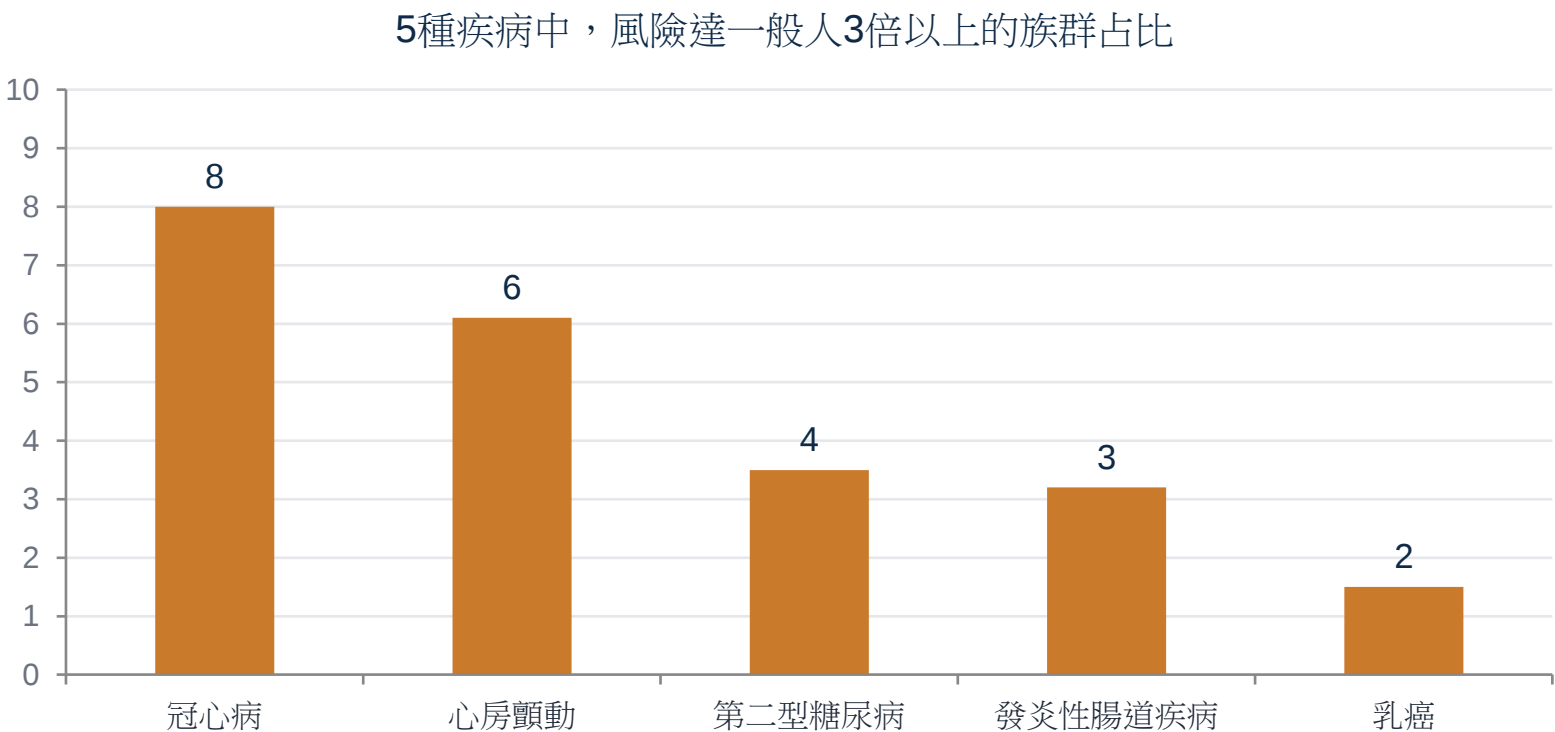
多基因風險評分與族群公平性議題

---

- 多基因風險評分 ( PRS ) : 辨識堪比單基因突變的風險族群
- 限制 : 資料代表性不足恐加劇健康不平等

# 多基因風險評分 ( PRS )

2018 年 *Nature Genetics* 研究：全基因體多基因風險評分，可辨識出風險程度堪比單基因突變帶原者的族群



20×

冠心病高風險族群  
人數，是傳統單基  
因帶原者的20倍規  
模

 *Khera et al., Nature Genetics, 2018, 50:1219–1224*

## 限制：資料代表性與健康不平等



現行PRS在歐洲血統族群的準確度，明顯高於其他血統族群。若不積極納入多元族群的基因體資料，臨床導入PRS恐加劇既有的健康不平等。

- 根源：現有GWAS研究樣本仍以歐洲血統族群為主
- 後果：PRS預測準確度隨基因分歧程度而下降，非歐裔族群受益較少
- 呼籲：擴大多元族群基因體研究，是負責任導入精準醫療AI工具的前提

 Martin, Kanai, Kamatani, Okada, Neale & Daly, Nature Genetics, 2019, 51:584–591

案例導覽

# 生醫文獻與知識管理

圖書館最切身相關的AI應用場域

---

- 文獻爆炸問題，與LLM如何協助文獻回顧
- 案例工具scite：引用語境的支持 / 反駁 / 僅提及分析

## TRIPOD+AI

**TRIPOD+AI statement: updated guidance for reporting clinical prediction models that use regression or machine learning methods**

*BMJ* 2024 ; 385 doi: <https://doi.org/10.1136/bmj-2023-078378> (Published 16 April 2024)

Cite this as: *BMJ* 2024;385:e078378

## STARD-AI

Consensus Statement | Published: 15 September 2025

**The STARD-AI reporting guideline for diagnostic accuracy studies using artificial intelligence**

*Nature Medicine* (2025) | [Cite this article](#)

## DECIDE-AI

**Reporting guideline for the early stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI**

*BMJ* 2022 ; 377 doi: <https://doi.org/10.1136/bmj-2022-070904> (Published 18 May 2022)

Cite this as: *BMJ* 2022;377:e070904

## CONSORT-AI

Consensus Statement | [Open access](#) | Published: 09 September 2020

**Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: the CONSORT-AI extension**

*Nature Medicine* 26, 1364–1374 (2020) | [Cite this article](#)

## Checklists for health AI evaluation

## OPTICA

CASE STUDY

f X in

**Evaluation of AI Solutions in Health Care Organizations — The OPTICA Tool**

Published August 14, 2024 | *NEJM AI* 2024;1(9) | DOI: 10.1056/AIcs2300269 | VOL. 1 NO. 9 | Copyright © 2024

## FURM

ARTICLE

f X in

**Standing on FURM Ground: A Framework for Evaluating Fair, Useful, and Reliable AI Models in Health Care Systems**

Published September 18, 2024 | *NEJM Catal Innov Care Deliv* 2024;5(10) | DOI: 10.1056/CAT.24.0131

VOL. 5 NO. 10 | Copyright © 2024

## FUTURE-AI

**FUTURE-AI: international consensus guideline for trustworthy and deployable artificial intelligence in healthcare**

*BMJ* 2025 ; 388 doi: <https://doi.org/10.1136/bmj-2024-081554> (Published 05 February 2025)

Cite this as: *BMJ* 2025;388:e081554

## POLARIS-GM

Correspondence | Published: 30 June 2025

**International partnership for governing generative artificial intelligence models in medicine**

*Nature Medicine* 31, 2836–2839 (2025) | [Cite this article](#)

# 文獻爆炸，AI如何協助？



## 文獻量持續指數成長

科學文獻的產出量長期呈現指數成長，人工檢索與閱讀已難以窮盡特定領域的所有文獻，這正是身於研究者的我們，感受最深的問題。



## LLM輔助文獻回顧

LLM可執行語意搜尋、跨文獻證據整合、系統性文獻回顧的自動化資料萃取。  
2023年Nature Medicine回顧文章，將文獻檢索與研究支援列為LLM在醫療領域的重要應用場景之一。

 *Thirunavukarasu et al., Nature Medicine, 2023, 29:1930–1940*

# 案例工具：scite 的引用語境分析

- 傳統引用次數只告訴我們「被引用了幾次」，需要再深入搜索「怎麼被引用的」
- scite以深度學習模型分析引用語境，自動判讀每筆引用是「支持」「反駁」還是「僅提及」原始研究
- 方法論正式發表於2021年 Quantitative Science Studies 期刊，並非黑盒行銷話術
- 對研究者的實際幫助：下判斷前，先看這篇論文是否被後續研究挑戰過

對圖書館服務的啟發：這類「引用脈絡分析」工具，可補足傳統書目資料庫只看引用數的盲點。

## 支持 Supporting

例：「本研究結果與Smith et al. 一致，進一步證實...」

## 反駁 Contrasting

例：「與先前報告不同，我們並未觀察到...」

## 僅提及 Mentioning

例：「如Smith et al.所述，該方法...」（背景引用）

 Nicholson, Mordaunt, Lopez, Uppala, Rosati, Rodrigues, Grabitz & Rife, Quantitative Science Studies, 2021, 2(3):882–898



# 中場休息

10 分鐘

接下來要動手操作：AlphaFold Server、Consensus、Gemini Notebook  
請先用手機打開網頁備用

---

已完成：開場概念 · 四象限案例導覽      接下來：實作示範 · 倫理討論

# Demo 1：蛋白質結構預測



## AlphaFold Server

DeepMind官方免費平台，輸入序列即可預測蛋白質與DNA/RNA/小分子複合物結構，非商業用途免費使用。背後即AlphaFold3技術。

[alphafoldserver.com](https://alphafoldserver.com)



## ESMFold

貼上胺基酸序列，數秒內出結果，速度快，適合搶時間的live demo；僅支援單鏈、序列長度上限約400aa。

[esmatlas.com/resources](https://esmatlas.com/resources)

## Demo 2：證據導向文獻問答



### Consensus

輸入研究問題，AI直接彙整逾2億篇同行評審文獻，給出證據導向答案並附引用來源，具有免費額度。

[consensus.app](https://consensus.app)



### Elicit

展示AI自動化文獻資料萃取、系統性文獻回顧，可將多篇論文的方法、樣本數、結果自動整理成表格，適合深入示範。

[elicit.com](https://elicit.com)

# Demo 3：文件對話與語音摘要



## Gemini Notebook

(原 NotebookLM，Google 於 2026 年 7 月 16 日更名)

上傳論文後直接對話問答，並可生成語音摘要（Audio Overview / 類 Podcast 形式）。

- 可上傳 PDF、Google 文件、網址、YouTube 影片等多種來源
- 所有回答皆附引用來源，可回頭核對原文段落

[notebook.google.com](https://notebook.google.com)

## 完整工具連結一覽

AlphaFold Server

[alphafoldserver.com](https://alphafoldserver.com)

分子結構

ESMFold

[esmatlas.com/resources?action=fold](https://esmatlas.com/resources?action=fold)

分子結構

Consensus

[consensus.app](https://consensus.app)

文獻證據

Elicit

[elicit.com](https://elicit.com)

文獻證據

Gemini Notebook

[notebook.google.com](https://notebook.google.com)

知識管理

提醒：工具的免費額度與功能可能隨時間調整，建議使用前再次確認。

## 六、挑戰與倫理討論

# 幻覺、偏誤、法規與公平性

---

- 幻覺：ChatGPT生成醫學文章中47%參考文獻是虛構的
- 偏誤：真實演算法種族偏誤案例 ( Science, 2019 )
- 法規與可重現性：510(k)核准比例、族群代表性

# 幻覺與可信度：真實的量化數據

2023 年 Cureus 研究：讓 ChatGPT-3.5 生成 30 篇短篇醫學文章，逐一核對每一筆參考文獻

47%

參考文獻  
完全虛構、不存在

93%

文獻至少存在  
一項資訊錯誤

錯誤類型包括：PMID 錯誤（93%）、卷期／頁碼錯誤（64%）、出版年份錯誤（60%）

 [Bhattacharyya, Miller, Bhattacharyya & Miller, Cureus, 2023, 15\(5\):e39238](#)

## 資料偏誤與代表性：真實案例

- 2019年 Science 研究，美國一項廣泛使用（影響逾2億病患）的醫療資源分配演算法
- 存在系統性種族偏誤：演算法以「醫療支出」作為疾病嚴重度的代理變項
- 但黑人病患因長期就醫可近性較低，醫療支出系統性偏低
- 結果：  
同樣風險評分下，黑人病患實際上病得更重

**17.7% → 46.5%**

修正偏誤後，獲得額外照護資源的黑人病患比例

這個數字反映的影響規模，涉及數以百萬計的病患照護資源分配

 Obermeyer, Powers, Vogeli & Mullainathan, Science, 2019, 366:447–453

## 法規與可重現性：兩個提醒

# 85.9%

的FDA核准AI醫材，是透過相對簡化的510(k)路徑核准；僅1.6%經過最嚴謹的PMA（上市前核准）審查。

 *Benjamins, Dhunoo & Meskó, npj Digital Medicine, 2020, 3:118*



### 呼應3-3：訓練資料代表性

不管是影像、文字還是基因體資料，訓練資料的代表性，始終是AI模型能不能可重現、能不能公平應用在不同族群身上的關鍵前提。

 *Martin et al., Nature Genetics, 2019, 51:584–591*

## 八、延伸主題一

# 政策法規動態

監理框架如何跟上AI醫材的演進速度

---

- FDA的「預先變更控制計畫」(PCCP)：讓AI醫材可以持續學習
- EU AI Act：醫材如何被畫入「高風險」類別
- 台灣TFDA的AI/ML醫材查驗登記技術指引

# FDA PCCP：讓AI醫材「持續學習」的新框架



## 問題出在哪？

傳統醫材審查假設產品定型不變；但AI/ML模型會隨著新資料不斷更新，每次更新若都要重新送件審查，將嚴重拖慢AI醫材的疊代速度。

2022年FDORA法案第515C條，賦予FDA明確法源；2024年12月發布正式定案指引



## PCCP怎麼解決？

製造商可在初次上市申請時，「預先」說明未來會如何修改模型、如何驗證、影響評估方法，經FDA核准後，之後的疊代更新可依此計畫執行，不必每次重新送件。

2024年適用範圍由「機器學習」擴大到所有「AI賦能」的醫材軟體功能 (AI-DSF)

 [FDA. Marketing Submission Recommendations for a PCCP for AI-Enabled Device Software Functions \(Final Guidance, Dec 2024\)](#)

# PCCP的三大要素

1

## 修改內容說明

具體列出未來預計對AI模型進行哪些類型的修改（如：擴增訓練資料、調整判讀閾值）

2

## 驗證方法

說明每次修改後，將如何測試驗證模型的安全性與有效性維持不變

3

## 影響評估

評估修改對整體裝置效能、風險與使用族群可能造成的影響

五項跨國指導原則（2023，FDA / Health Canada / MHRA 共同發布）：聚焦明確、風險導向、實證支持、透明溝通、全生命週期管理

## FDA核准現況：一組數字看規模

95,147

項醫材經510(k)或  
De Novo路徑核准  
(截至2024年底)

1,678

項醫材經最嚴謹的  
PMA (上市前核准)  
路徑核准

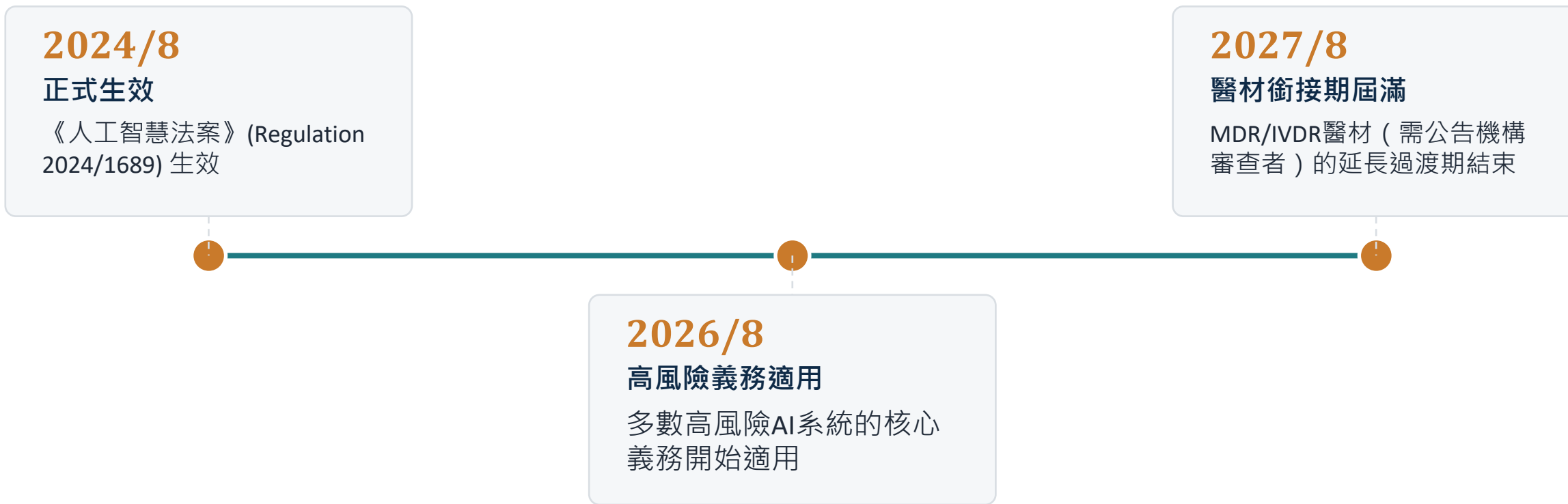
53,315

件PMA補充申請，  
平均每項裝置  
3,177件補充案

這代表：AI醫材的疊代速度極快，也解釋了為什麼FDA需要PCCP這類「持續學習」的審查框架，來因應大量的後續修改申請。

[📖 \*Predetermined Change Control Plans: Guiding Principles for Advancing Safe, Effective, and High-Quality AI-ML Technologies \(PMC, 2025\)\*](#)

# EU AI Act：一張圖搞懂醫材時程



重點：MDR/IVDR分類為Class IIa以上的醫材（約占放射科AI商用醫材的75%），幾乎都會被自動畫入AI Act的「高風險」類別。

 [Tandem Health; ReedSmith, EU AI Act and Medical Devices](#)

# EU AI Act：高風險AI醫材的五大要求



## 資料治理

訓練、驗證、測試資料須符合品質與代表性要求



## 技術文件

完整記錄系統設計、開發與驗證過程



## 人為監督

確保人類可以有效理解、監督與介入AI系統決策



## 透明溝通

向使用者清楚說明系統能力與限制



## 全生命週期管理

上市後持續監測、更新與風險管理

 [ReedSmith, The EU AI Act and Medical Devices: Navigating High-Risk Compliance](#)

# 台灣TFDA動態：本土的AI醫材指引

- 食藥署已公告《人工智慧/機器學習技術之醫療器材軟體查驗登記技術指引》，明訂臨床性能驗證應考慮的重點
- 指引要求製造業者說明：研究對象、病患族群（含年齡、族裔、人種）、參與驗證之醫事人員資格、臨床資料取得方式與統計分析方法
- 也參考國際醫療器材法規論壇（IMDRF）2017年發布的SaMD臨床評估指引
- 食藥署設有「智慧醫療器材資訊暨媒合平台」，提供查驗登記諮詢輔導與產業媒合



## 與國際接軌

台灣的技術指引架構參考FDA與IMDRF的既有框架，同時要求驗證族群的代表性。呼應了我們稍早在精準醫療段落討論過的族群公平性議題。

 [衛生福利部食品藥物管理署，智慧醫療器材資訊暨媒合平台](#)

# 三方法規比較一覽

|       | 美國 FDA         | 歐盟 EU AI Act     | 台灣 TFDA      |
|-------|----------------|------------------|--------------|
| 核心工具  | PCCP（預先變更控制計畫） | 風險分級（高風險強制合規）    | 查驗登記技術指引     |
| 最新進展  | 2024/12 定案指引   | 2026/8 高風險義務適用   | 已公告AI/ML技術指引 |
| 核心精神  | 允許模型持續疊代更新     | 資料治理 + 人為監督 + 透明 | 族群代表性 + 臨床驗證 |
| 對業者意義 | 免每次修改重新送件      | 須通過公告機構合規審查      | 需揭露驗證族群組成    |

## 政策趨勢小結：三個共同方向

**01**

從「一次核准」走向「持續管理」

無論FDA的PCCP或EU AI Act的上市後監測要求，都反映監理機關認知到AI模型會持續演進，審查機制也必須跟著動態化。

**02**

資料代表性成為硬性要求

從台灣TFDA的驗證族群揭露，到EU AI Act的資料治理義務，「訓練資料夠不夠具代表性」已從倫理呼籲變成合規要求。

**03**

透明與可追溯是共同底線

三個法規體系都要求業者能向監管機關與使用者說明系統的能力邊界與決策依據，這與我們稍早討論的倫理原則高度呼應。

延伸主題二

# 產業故事與趣聞

從實驗室到臨床、從突破到教訓—四個真實案例

---

- Isomorphic Labs：AlphaFold走出實驗室的三十億美元交易
- Insilico Medicine：全球第一個AI發現且完成陽性臨床試驗的藥物
- IBM Watson：一堂6,200萬美元的教訓
- Google AI Co-Scientist：AI自己提出的科學假設

# Isomorphic Labs：AlphaFold走出實驗室

DeepMind 衍生公司

還記得前面介紹的AlphaFold3嗎？2021年DeepMind成立了獨立子公司Isomorphic Labs，專注把蛋白質結構預測技術應用在真實的藥物研發上。2024年1月，這家公司在同一天內，宣布了兩筆合計近30億美元的重量級合作案。

## 公司背景

- 2021年從Google DeepMind獨立成立
- 由DeepMind共同創辦人 Demis Hassabis 擔任執行長
- 核心技術：新一代AlphaFold + 自研生成式化學模型

## 合作對象

- 2024年1月同日宣布與Eli Lilly、Novartis簽約
- 鎖定過去被視為「難成藥」的標靶
- 2026年初已從標靶辨識推進到臨床前候選藥物

## 數字看合作規模：近30億美元

Eli Lilly

**\$45M**

簽約時立即支付

**最高可達 \$1.7B**

依里程碑分階段給付

Novartis

**\$37.5M**

簽約時立即支付

**最高可達 \$1.2B**

依里程碑分階段給付

 *Fierce Biotech; Inside Precision Medicine, Isomorphic Labs deals with Lilly and Novartis (2024)*

# 這代表什麼？

- 從學術突破到商業驗證：AlphaFold從一篇論文，走到能讓製藥巨頭掏出真金白銀的技術平台，只花了不到三年
- 「生成式」藥物設計：科學家設定期望的藥物特性（如高結合力、低毒性），讓AI直接提出候選分子結構
- 2024年諾貝爾化學獎，成為這類合作案的重要信任背書。Hassabis與Jumper正是Isomorphic Labs的核心人物
- 產業觀察：傳統藥物開發平均耗時逾10年、成本超過20億美元，AI能否真正縮短這個週期，仍待後續臨床試驗結果驗證



延伸提問：如果AI設計的分子真的縮短了研發週期，藥價會因此下降，還是因為授權金而維持高檔？這是產業界持續辯論的問題。

 [MedPath; Wedbush, The \\$3 Billion Bet: How Isomorphic Labs is Rewriting the Rules of Drug Discovery \(2026\)](#)

# Insilico Medicine：全球第一個AI發現的藥物

## 里程碑

如果說Isomorphic Labs的故事還在進行中，Insilico Medicine已經跑到了終點線的前一哩路：2025年，他們的候選藥物**rentosertib**，成為全球第一個「AI發現標靶 + AI設計分子」並完成陽性臨床試驗的藥物，結果發表在Nature Medicine。

### AI怎麼做到的

- PandaOmics：AI引擎辨識出新藥物標靶 TNIK
- Chemistry42：生成式化學引擎設計全新分子結構
- 從標靶辨識到候選藥物僅費時約18個月

### 治療對象

- 適應症：特發性肺纖維化（IPF），一種缺乏有效藥物的慢性肺病
- 2023年獲FDA孤兒藥認定
- 2025年6月完成第二期臨床試驗數據發表

## 臨床數據：肺功能顯著改善

**+98.4 mL**

60mg劑量組患者的  
用力肺活量 (FVC) 平均改善

12週治療期間

**-20.3 mL**

安慰劑組患者的  
用力肺活量同期變化

疾病自然惡化的對照

這是全球第一次，一款「AI發現標靶+AI設計分子」的藥物，在隨機雙盲對照臨床試驗中，展現出具統計意義的療效訊號。

📖 [Xu, Ren, Wang et al., Nature Medicine, 2025, 31:2602–2610](#)

# 這代表什麼？

- 驗證了「AI原生」藥物研發的完整鏈路：從文獻探勘找標靶，到生成化學設計分子，全程都有AI深度參與
- 研究團隊也坦承：這只是第二期早期試驗（71位病人），距離最終上市，還需要更大規模、更長期的第三期試驗驗證
- 對研究人員的啟發：這類「AI原生」研究的方法學論文，往往同時發表在化學、生技與臨床期刊上，跨領域文獻檢索能力更形重要
- 產業觀察：Insilico已有18項臨床前候選藥物、7項獲得IND核准，顯示這不是單一個案，而是一套可複製的研發流程



時間對比：從標靶辨識到候選藥物僅約18個月，相較傳統藥物開發前期探索階段動輒3到6年，速度相當驚人。

 [Insilico Medicine, Phase IIa Results Announcement \(2025\)](#)

# IBM Watson：一堂6,200萬美元的教訓

## 警世案例

不是所有AI醫療故事都是成功案例。2011年IBM Watson在益智節目《Jeopardy!》擊敗人類冠軍後，IBM隨即宣布要讓它成為「AI醫師」。2013年，MD Anderson癌症中心與IBM合作開發腫瘤治療顧問系統，但這個雄心勃勃的計畫，最終在2016年黯然收場。

### 計畫緣起

- 2013年MD Anderson與IBM簽約合作
- 目標：開發「Oncology Expert Advisor」腫瘤治療顧問系統
- 運用自然語言處理彙整病歷、搜尋文獻資料庫

### 為何失敗

- 始終未能與醫院既有的Epic電子病歷系統整合
- 範疇不斷擴張、時程一再延宕
- 2016年德州大學系統稽核揭露問題，計畫隨即終止

## 失敗的代價：數字說明一切

**\$62M**

MD Anderson投入  
的總花費

**4年**

從簽約到  
終止合作

**0**

最終實際導入  
臨床使用的病患數

IBM內部文件甚至在2016年9月坦承，這套系統「尚未準備好用於人體研究或臨床用途」  
但對外的行銷口徑，仍持續保持樂觀。

 [IEEE Spectrum, How IBM Watson Overpromised and Underdelivered on AI Health Care](#)

## 三個教訓

### 1 系統整合比演算法更難

Watson從未真正與醫院既有的電子病歷系統相容，工程師被迫手動輸入資料。再厲害的模型，也敵不過現實世界的資料管線問題。

### 2 對外樂觀、對內保留，是危險的落差

當內部評估已認定系統「尚未準備好」，對外行銷卻持續加碼宣傳，這種落差終將反噬機構信譽。

### 3 「準確度媲美專家」≠ 「臨床上真的有用」

IBM曾主張Watson建議與專家意見相符率達90%，但這只證明它學會了模仿現有做法，未必代表它能改善治療結果。

# Google AI Co-Scientist：AI自己提出的科學假設

## 研究輔助新典範

如果AI不只能回答問題，還能主動「提出」值得驗證的科學假設呢？2025年一篇發表於Advanced Science的研究，讓一套以Gemini為基礎的多代理人AI系統（AI co-scientist），針對「肝纖維化」這個治療選擇有限的疾病，提出藥物再利用的候選建議，並實際在人類肝臟類器官中驗證。

### AI怎麼運作

- 多代理人架構，模擬「假設生成→實驗規劃」的科學推理流程
- 分析大量文獻，找出「隱藏」的分子關聯
- 同時提出可執行的實驗驗證步驟，不只是給答案

### 實際驗證方式

- 在多譜系人類肝臟類器官（microHO）中測試
- 評估25種藥物的抗纖維化效果與毒性
- 由科學家全程審核每一項AI建議

## 實驗結果：舊藥新用的91%

### Vorinostat

一款已獲FDA核准上市多年的抗癌藥物

91%

降低TGF $\beta$ 誘發的  
染色質結構變化

- AI co-scientist是由AI建議、由科學家實驗驗證的組合。這款藥物是系統推薦的候選藥物之一，並非唯一結果
- 同時促進了肝臟實質細胞的再生，顯示除了抗纖維化，還有輔助修復的潛力
- 研究團隊強調：AI生成的假設仍需經過嚴謹的實驗驗證，AI的角色是「增強」而非「取代」科學家的判斷

 [\*Guan et al., Advanced Science, 2025, e08751\*](#)

# 這代表什麼？

- 從「回答問題」到「提出假設」：這類系統代表LLM應用的下一步，不只是文獻摘要工具，而是主動參與研究設計的夥伴
- 「藥物再利用」（drug repurposing）特別適合AI輔助：已上市藥物的安全性數據齊全，AI只需要找出「原本沒想到的新用途」
- 研究團隊也做過對照：同一套方法先前已用於基因表現型分析，協助發現新的小鼠基因功能與人類遺傳診斷
- 對研究者的提醒：這仍是「人在迴圈中」（human-in-the-loop）的模式。AI負責擴大搜尋範圍，人類負責最終把關



與今天四象限的呼應：這個案例同時涉及「分子藥物」與「文獻知識」兩個象限。AI從文獻中提出假設，再回到濕實驗室驗證。

 [\*Guan et al., Advanced Science, 2025, e08751\*](#)

# 四個故事，一次回顧

| 案例                     | AI做了什麼               | 結果                | 一句話啟示           |
|------------------------|----------------------|-------------------|-----------------|
| Isomorphic Labs        | 以AlphaFold技術設計候選藥物分子 | 近\$30億美元藥廠合作案     | 學術突破可以快速商業化     |
| Insilico Medicine      | AI發現標靶 + 設計分子        | 全球首個完成陽性二期試驗的AI藥物 | AI原生藥物研發已走進臨床   |
| IBM Watson             | 彙整病歷、提供治療建議          | \$62M投入後專案終止      | 系統整合與誠實溝通同樣重要   |
| Google AI Co-Scientist | 提出藥物再利用假設            | 人類肝臟類器官中驗證有效      | AI可以是提問者，不只是答題者 |

# 政策與產業連結彙整

以下為政策法規與產業新聞來源，屬新聞／官方文件，不同於前述學術文獻，故另列於此

|                                     |                   |                      |
|-------------------------------------|-------------------|----------------------|
| FDA PCCP最終指引 ( 2024/12 )            | fda.gov           | <a href="#">開啟連結</a> |
| EU AI Act醫材合規解析                     | tandemhealth.ai   | <a href="#">開啟連結</a> |
| 台灣TFDA智慧醫材資訊平台                      | aimd.fda.gov.tw   | <a href="#">開啟連結</a> |
| Isomorphic Labs與Lilly / Novartis合作案 | fiercebiotech.com | <a href="#">開啟連結</a> |
| Insilico Medicine Phase IIa結果公告     | insilico.com      | <a href="#">開啟連結</a> |
| IBM Watson健康照護案例分析                  | spectrum.ieee.org | <a href="#">開啟連結</a> |

生醫領域專家研習課程

# 研究分析與論文寫作的實用工具 兼談來源正確性的把關

01

# 全球現況掃描

AI 在生醫領域，目前究竟走到哪裡？

# 藥物開發：全球案例地圖

## 1 AlphaFold 3

英國 | DeepMind × Isomorphic Labs

由蛋白質結構預測，擴展到蛋白質—DNA、小分子結合等生物分子複合體預測，已成為結構生物學的「生產級」研究工具。

## 2 Insilico Medicine

香港 / 中國

旗艦藥物 rentosertib ( TNIK 抑制劑，治療特發性肺纖維化 )：從標的確認到 Phase 1 僅約 30 個月，Phase IIa 結果已刊登於《Nature Medicine》。

## 3 Exscientia × 住友製藥

英國 / 日本

DSP-1181 ( 強迫症候群藥物 ) 曾被視為首個由 AI 從頭設計、進入人體臨床試驗的分子。

## 4 Eli Lilly × NVIDIA

美國

打造藥廠專屬「AI 工廠」超級運算平台，並與 Purdue 大學合作 2.5 億美元研究計畫。

# 醫學影像與病理：台灣在地案例

## 1 aetherAI

台灣 | 數位病理

亞洲數位病理 AI 領導廠商，2026 年 5 月於台灣創新板掛牌，協助病理科醫師以 AI 輔助判讀切片、加速診斷流程。

## 2 Deep01

台灣 | 急重症影像

研發腦出血 AI 偵測系統，已進入台灣 10 多家醫院使用，並取得 20 多個國家的獨家代理合約，是台灣智慧醫療走向國際的代表案例。

## 現實查核：發現層快，臨床層仍是瓶頸

117

AI 相關藥物資產  
已進入人體臨床試驗  
( 63 家公司，ASCO  
2026 分析 )

51.3%

完成 Phase 1 的  
候選藥物比例

6.8%

完成 Phase 2 的  
候選藥物比例

*重點：AI 大幅加速了「發現」與「設計」階段，但臨床試驗仍受限於生物學本質的高失敗率。  
這也是我們今天聚焦「研究與寫作」層工具的原因：這是每位研究者現在就能立即受益的一層。*

## 小結：AI 在生醫領域的三個成熟層次

- 已成熟：文獻挖掘與知識探勘、醫學影像輔助判讀（如病理、腦影像）
- 快速發展中：藥物標的發現、分子從頭設計、蛋白質結構預測
- 仍屬早期：AI 輔助臨床試驗設計、真實世界證據整合、AI 直接臨床決策
- 每位研究者「現在就能用」的一層：研究文獻探索與論文寫作

# 02

## 研究文獻與證據整合工具箱

從「大海撈針」到「分鐘級」的文獻回顧

## 傳統文獻回顧的三個痛點

- 文獻量爆炸性成長：每年新增的生醫論文以百萬篇計，人工全面掌握已不可能
- 人工篩選極度耗時：一篇系統性文獻回顧，篩選與資料抽取常需耗費數週人力
- 跨文獻證據整合困難：同一問題常有數十篇甚至上百篇論文，結論彼此矛盾，難以快速判斷「文獻共識」
- **AI 工具的價值，就在把「探索—篩選—整合」這三步，從數週壓縮到數小時甚至數分鐘**

# 文獻探索層：找到對的論文

## 1 Semantic Scholar 免費 | AI2 出品

免費的學術發現引擎，提供 AI 摘要（TL;DR）與引用網絡，是多數研究者「畫地圖」的起點。

## 2 Connected Papers 視覺化引用圖

以視覺化圖譜呈現一篇論文與相關文獻的關聯，快速掌握某個子領域的關鍵論文與發展脈絡。

## 3 ResearchRabbit 推薦引擎

類似「文獻界的 Spotify」，可訂閱關注的論文集合，持續推薦新發表的相關研究。

## 4 Perplexity 通用 AI 搜尋

跨網路與學術文獻的即時搜尋，附帶逐句引用來源，適合快速掌握某議題的整體輪廓。

# 系統性文獻回顧層：結構化擷取數據

## 1 Elicit 1.38 億篇論文

目前最成熟的結構化文獻回顧工具：可批次篩選上千篇摘要，並以「欄位」自動擷取樣本數、族群、介入、結果等資料，直接產出比較表。

## 2 Rayyan 雙盲篩選

系統性文獻回顧常用的協作篩選平台，支援多人雙盲審查與衝突標記。

## 3 Covidence 全流程管理

涵蓋篩選、資料抽取到品質評估的系統性回顧全流程工具，可將篩選時間縮短達50%。

# Consensus：這題文獻站在哪一邊？

- 輸入一個是非型研究問題（例如：「維生素 D 補充是否能降低骨折風險？」）
- Consensus 會掃描相關論文，將每篇的立場歸類為「支持 / 反對 / 不確定」，並以視覺化比例呈現「文獻共識」
- 適合用在：寫 Introduction 需要快速確認某個論點的文獻風向、或審查一篇投稿時快速做「合理性檢查」
- 提醒：共識視覺化是起點，不是終點。仍需點開關鍵論文確認方法學品質

## 案例

### CASE STUDY

# scite.ai：這篇論文，後來被推翻了嗎？

- 全球資料庫規模超過 12 億筆「已分類」的引用語句，標示每一次引用是「支持（supporting）」「反對（contrasting）」還是「僅提及（mentioning）」
- 真實情境：準備在稽核報告中引用一篇 2019 年抗凝血劑試驗，用 scite 一查，發現後續兩篇更大型研究已經推翻原始結論。結論因此修正
- 這是目前唯一能回答「一篇論文的說法，後來的文獻怎麼看待它」的工具類型
- 建議：論文中每一個「決定性」的關鍵引用，投稿前都用 scite 查一次

## 案例

### CASE STUDY

# 生醫專用挖礦與團隊筆記整合

## 1 PubTator 3.0

NER F1 = 0.93

NIH 開發的生醫實體辨識工具，可自動標記文獻中的基因、疾病、化合物名稱，加速知識挖掘。

## 2 知識圖譜 (iKraph 等)

藥物 - 疾病關聯挖掘

結合 PubMed 全文與公開資料庫，建構知識圖譜，用以推論尚未被直接研究過的藥物 - 疾病潛在關聯。

## 3 NotebookLM

團隊文獻筆記

上傳自己的 PDF 文獻庫，建立「有溯源依據」的問答與摘要，適合研究團隊共用、彼此可回頭核對出處。

# 2026 新趨勢：Agentic AI 科學研究代理人

- FutureHouse 的多代理人系統（Robin、Crow、Falcon、Finch，架構於 PaperQA2）：可自動完成文獻搜尋、閱讀與分析的完整流程
- Google Co-Scientist：多代理人協作，模擬「提出假說—檢索文獻—評估證據」的科研思考流程
- 現況：仍屬早期階段，速度快但仍需嚴格人工查核，不建議直接採用其結論作為投稿依據

# 建議工作流程：四步驟組合拳

- 1 畫地圖** 用 Semantic Scholar 或 ResearchRabbit 掌握領域關鍵論文與發展脈絡
- 2 看共識** 用 Consensus 快速掌握文獻對關鍵問題的整體立場
- 3 建資料表** 用 Elicit 針對篩選出的論文，結構化擷取樣本數、方法、結果等資料
- 4 查引用** 用 scite 核實關鍵引用是否仍站得住腳，避免引用已被推翻的結論

# 03 論文寫作與編修工具箱

# 語言潤飾與文法工具

## 1 Paperpal Elsevier 系 | 投稿把關

Manuscript Checker 涵蓋 30 多項期刊編審常見檢查項目，並內建 AI Detector，適合投稿前最後一道語言與格式把關。

## 2 Trinko 技術文法 | 隱私導向

針對科學與技術文本訓練，誤判率較低，並提供「發表就緒度檢查（Publication Readiness Check）」。

## 3 Writefull LaTeX / Overleaf 整合

唯一原生整合 Overleaf 的工具，提供依段落分類的「句型庫（Sentence Palette）」，協助組織 Abstract、Methods 等段落。

## 4 Grammarly 通用型 | 最後防線

非學術專屬，但涵蓋範圍最廣，適合作為潤飾流程的最後一道通用檢查。

# 草稿撰寫與通用 LLM 的角色

- Jenni AI：逐段自動完成 + 串接引用資料庫，適合快速產生初稿骨架
- Claude / ChatGPT / Gemini：整理摘要、重組段落邏輯、協助資料視覺化描述、搭配 Excel / 程式碼協助資料分析與圖表製作
- 這類通用型工具最適合的角色是「表達與整理」。把你已確認的事實與數據，寫得更清楚、更有邏輯
- 關鍵原則：讓 AI 幫你「說得更好」，但事實、數據與引用永遠要由你自己查證

## 原創性與 AI 使用偵測

- iThenticate：業界標準的抄襲比對工具，多數期刊投稿系統已內建
- Paperpal AI Detector：偵測文本中可能由 AI 生成的段落，投稿前可先自我檢查
- 現況：越來越多期刊在審稿流程中「同時」查核抄襲比例與 AI 生成痕跡，兩者缺一不可

# 全球期刊 AI 使用規範現況

| 組織 / 期刊 | 最新規範重點                                                                    |
|---------|---------------------------------------------------------------------------|
| ICMJE   | 2026 年 1 月更新新增「Section V：AI 於出版之使用」。AI 不得列為作者，須揭露用於草擬、編輯、翻譯、圖像生成、資料分析等各環節 |
| COPE    | 2025 年準則：明訂捏造引用可構成研究不端行為，出版社須建立 AI 使用揭露與查核機制                              |
| JAMA    | 延續 ICMJE 精神，並要求作者對 AI 協助生成之內容承擔完全責任                                       |
| WAME    | 全球醫學編輯協會早於 2023 年即倡議「禁止 AI 列名作者、強制揭露」，是目前規範的重要源頭                          |

## 小結：寫作工具箱使用三原則

- 潤飾類工具（ Paperpal / Trinko / Writefull ）：可以放心使用，風險低、效益高
- 生成草稿類工具：務必人工核對每一個事實陳述與引用來源，不可直接採信
- 揭露 AI 使用：這已經是國際期刊的「趨勢」而非「選項」，投稿前務必查閱該期刊最新政策

# 04

## 準確性風險與品質控管

AI 幫你找到方向，但你要對內容負責

## 觸目驚心的數字：捏造引用正在氾濫

1/277

2026 年受審查的  
PubMed 論文中，  
每 277 篇就有 1 篇  
含捏造參考文獻

12×

兩年內，生醫文獻中  
捏造引用數量  
成長了 12 倍

14.69 萬

2025 單年，從 250 萬篇  
論文稽核樣本中  
偵測到的 AI 生成假引用筆數

資料來源：Retraction Watch (2026)；Nature「Hallucinated citations are polluting the scientific literature」(2026)。這不是假設性風險，而是正在發生的現實。

## 真實案例：一本學術書的撤回

- 2025 年 7 月，Springer Nature 撤回一本學術書籍。原因是書中 46 篇參考文獻，三分之二「不存在或有嚴重錯誤」
- 諷刺的是，這本書還特別有一節討論「ChatGPT 的倫理議題」，作者卻始終未承認使用 AI 撰寫
- 另一起案例：法國 Toulouse 大學學者 Cabanac 收到 Google Scholar 通知，發現自己被引用在一篇牙科期刊論文中。但他的研究領域與牙科完全無關，那筆引用根本是被 AI 捏造出來的
- 這類案例正快速增加，且不限於小型期刊。連知名出版社的正式出版品都未能倖免

## 案例

### CASE STUDY

## 為什麼會發生：理解 LLM 的「生成」本質

- 語言模型的本質是「根據統計模式預測下一個字」，不是「查詢資料庫」
- 沒有連接文獻資料庫、純粹依賴訓練記憶回答的聊天機器人，即使產生的引用格式完美無瑕，內容也可能完全是虛構的
- 早期研究者與小型團隊的論文最容易「中鏢」。往往缺乏資深研究者把關、審查流程也較寬鬆
- **格式正確 ≠ 內容真實**

# 品質控管 SOP：五個步驟

- 1 選對工具** 優先使用「有溯源機制」的工具（Elicit / Consensus / scite / Semantic Scholar），避免只靠純聊天機器人搜尋文獻
- 2 逐筆核對** 每一筆關鍵引用都點進 DOI 或 PubMed 連結，親自確認論文確實存在、內容確實如此
- 3 引用審計** 對論文中「決定性」的關鍵論點，用 scite 等工具查核後續文獻是否已推翻該結論
- 4 主動揭露** 依循 ICMJE / COPE 規範，清楚說明 AI 用於哪個環節（潤飾 / 草擬 / 資料分析等）
- 5 交叉查核** 投稿前由團隊或同儕再過一次查證流程，尤其是早期研究者與小型團隊的文稿

「AI 能幫你更快找到方向，  
但『這篇文獻真的存在、真的這樣說』，  
永遠是研究者自己的責任。」

# 05

## 工具選擇框架

你的下一步：從哪一個工具開始？

# 依任務對照工具地圖

| 研究任務           | 推薦工具                              |
|----------------|-----------------------------------|
| 文獻探索、掌握領域全貌    | Semantic Scholar / ResearchRabbit |
| 系統性文獻回顧、結構化擷取  | Elicit / Rayyan / Covidence       |
| 快速判斷文獻共識方向     | Consensus                         |
| 核實關鍵引用是否仍成立    | scite.ai                          |
| 團隊文獻筆記與知識庫     | NotebookLM                        |
| 語言潤飾與文法檢查      | Trinka / Writefull                |
| 投稿前總體把關        | Paperpal                          |
| 資料整理、初稿撰寫、圖表描述 | Claude / ChatGPT + Excel / 程式碼協作  |

# 三種入門組合建議

## 1 精簡免費版

個人研究者起步

Semantic Scholar ( 免費 ) + Consensus 免費層 + Claude / ChatGPT 協助整理與潤飾。零成本即可大幅提升效率。

## 2 標準研究室

小型實驗室

在免費版基礎上，加入 Elicit ( 結構化擷取 ) + scite ( 引用查核 ) + Trinka ( 語言潤飾 )，涵蓋完整研究到投稿流程。

## 3 系統性回顧團隊

多人協作 / 大型計畫

再加入 Rayyan 或 Covidence ( 多人篩選 ) + NotebookLM ( 團隊知識庫 ) + Paperpal ( 投稿總把關 )，適合正式系統性回顧或統合分析計畫。

## Q&A 與討論

- 你目前的文獻回顧流程中，最耗時的步驟是哪一個？今天介紹的哪個工具可能解決它？
- 如果要在你的實驗室建立「AI 使用揭露」規範，你會怎麼寫？由誰把關？
- 面對「每 277 篇就有 1 篇假引用」的現實，你所屬的期刊或機構，目前有配套的查核機制嗎？



謝謝聆聽 敬請指教

yijulee@stat.sinica.edu.tw

## 參考文獻（1／4）

- [1] Jumper, J., Evans, R., Pritzel, A., et al. (2021). Highly accurate protein structure prediction with AlphaFold.  
*Nature*, 596(7873), 583–589. <https://doi.org/10.1038/s41586-021-03819-2>
- [2] Abramson, J., Adler, J., Dunger, J., et al. (2024). Accurate structure prediction of biomolecular interactions with AlphaFold3.  
*Nature*, 630(8016), 493–500. <https://doi.org/10.1038/s41586-024-07487-w>
- [3] Gulshan, V., Peng, L., Coram, M., et al. (2016). Development and validation of a deep learning algorithm for detection of diabetic retinopathy in retinal fundus photographs.  
*JAMA*, 316(22), 2402–2410. <https://doi.org/10.1001/jama.2016.17216>
- [4] Abràmoff, M. D., Lavin, P. T., Birch, M., Shah, N., & Folk, J. C. (2018). Pivotal trial of an autonomous AI-based diagnostic system for detection of diabetic retinopathy in primary care offices.  
*npj Digital Medicine*, 1, 39. <https://doi.org/10.1038/s41746-018-0040-6>

## 參考文獻（2／4）

- [5] Campanella, G., Hanna, M. G., Geneslaw, L., et al. (2019). Clinical-grade computational pathology using weakly supervised deep learning on whole slide images.

*Nature Medicine*, 25(8), 1301–1309. <https://doi.org/10.1038/s41591-019-0508-1>

- [6] Benjamins, S., Dhunoo, P., & Meskó, B. (2020). The state of artificial intelligence-based FDA-approved medical devices and algorithms: an online database.

*npj Digital Medicine*, 3, 118. <https://doi.org/10.1038/s41746-020-00324-0>

- [7] Khera, A. V., Chaffin, M., Aragam, K. G., et al. (2018). Genome-wide polygenic scores for common diseases identify individuals with risk equivalent to monogenic mutations.

*Nature Genetics*, 50(9), 1219–1224. <https://doi.org/10.1038/s41588-018-0183-z>

- [8] Martin, A. R., Kanai, M., Kamatani, Y., Okada, Y., Neale, B. M., & Daly, M. J. (2019). Clinical use of current polygenic risk scores may exacerbate health disparities.

*Nature Genetics*, 51(4), 584–591. <https://doi.org/10.1038/s41588-019-0379-x>

## 參考文獻（3／4）

- [9] Thirunavukarasu, A. J., Ting, D. S. J., Elangovan, K., Gutierrez, L., Tan, T. F., & Ting, D. S. W. (2023). Large language models in medicine.

*Nature Medicine*, 29(8), 1930–1940. <https://doi.org/10.1038/s41591-023-02448-8>

- [10] Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations.

*Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>

- [11] Nicholson, J. M., Mordaunt, M., Lopez, P., Uppala, A., Rosati, D., Rodrigues, N. P., Grabitz, P., & Rife, S. C. (2021). scite: A smart citation index that displays the context of citations and classifies their intent using deep learning.

*Quantitative Science Studies*, 2(3), 882–898. [https://doi.org/10.1162/qss\\_a\\_00146](https://doi.org/10.1162/qss_a_00146)

- [12] Bhattacharyya, M., Miller, V. M., Bhattacharyya, D., & Miller, L. E. (2023). High rates of fabricated and inaccurate references in ChatGPT-generated medical content.

*Cureus*, 15(5), e39238. <https://doi.org/10.7759/cureus.39238>

## 參考文獻（4／4）

- [13] Xu, Z., Ren, F., Wang, P., et al. (2025). A generative AI-discovered TNIK inhibitor for idiopathic pulmonary fibrosis: a randomized phase 2a trial.

*Nature Medicine*, 31, 2602–2610. <https://doi.org/10.1038/s41591-025-03743-2>

- [14] Guan, Y., et al. (2025). AI-Assisted Drug Re-Purposing for Human Liver Fibrosis.

*Advanced Science*, e08751. <https://doi.org/10.1002/advs.202508751>

### 謝謝聆聽！

本講義 / 投影片中所有學術文獻引用與工具連結，均於課前逐一查證。如發現任何連結失效或內容有誤，歡迎與講者或中研院生醫圖書館聯繫更新。