



ProQuest Research Assistant (PQRA)

科睿唯安學術研究與政府事業部 2026年4月

負責任的AI原則，由學術社群提供支持

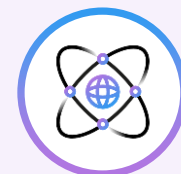


AI 框架



透明

- 清晰資訊來源說明
- 正確的引用和便捷的參考文獻訪問



合規

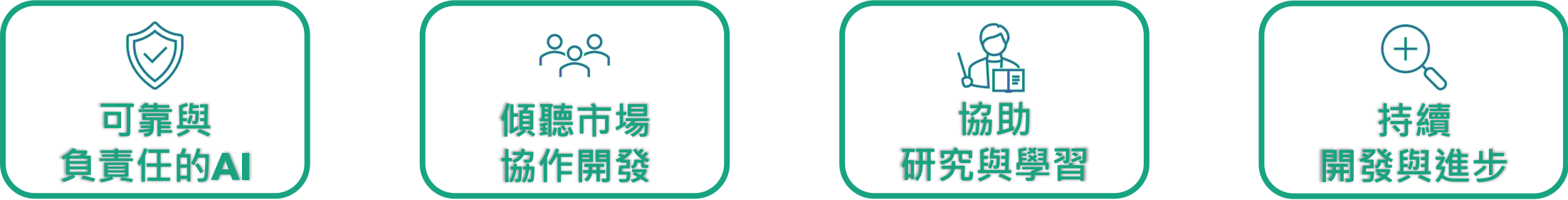
- 減少錯誤資訊的措施
- 與出版商合作確保明確的使用權
- 與產業機構合作，推動負責任的 AI 實施



安全

- 專家參與
- 維護隱私與安全標準
- 遵守不斷演變的全球法規

ProQuest Research Assistant 開發基礎與歷程



2019
學術內容AI公司
Alethea成立

2023 Sep
科睿唯安收購
Alethea

2024 Jan
Clarivate
學術 AI團隊成立

2024 Q2
Clarivate
AI 理事會成立

2024 Sep
ProQuest RA
第一階段上線
至PQ1B、PQ1E、
PQ1L、PQ1P、PQ1S

2025 Q1
ProQuest RA
第二階段上線
至其他PQ1系列
及PQDT Global、
ProQuest Central
Premium

2025 Q3
ProQuest RA
自然語言對話功能
beta版上線

ProQuest Research Assistant (PQRA)

基於使用者研究與客戶洞察打造

過去一年，科睿唯安（Clarivate）進行了廣泛的使用者研究，以捕捉多元的使用者與客戶視角，並將其融入 AI 驅動的 ProQuest Research Assistant 開發中。



探索性研究：探究 AI 及其他趨勢對研究人員、圖書館員與師生的影響。



訪談與調查：與這些群體的訪談與調查，揭示了 AI 如何協助解決他們的研究與學習挑戰。



可用性測試與卡片分類：透過這兩種方式，確保設計直觀且符合使用者的心智模式。



客戶諮詢委員會：定期諮詢，確保滿足院校與使用者的需求。

60 多次
一對一訪談

數百次
可用性測試與活動

數千份
調查回覆

10 多次
諮詢委員會審查

ProQuest Research Assistant

發揮 AI 的強大能力，並以負責任、可靠的方式將其作為您的研究夥伴。

- 輕鬆建構更有效、更具針對性的檢索式
- 更有效地審閱、分析與探討文獻
- 快速評估每篇文獻對您研究的價值
- 獲取後續步驟的指引，例如選擇研究主題和理解文獻中的關鍵概念

The screenshot displays the ProQuest Research Assistant interface. The main content area shows a research article titled "AI and XAI second opinion: the danger of false confirmation in human-AI collaboration" by Rosenbache, Rikard, Mehrez, Asya, McKee, Martin, Stuckler, David, published in the Journal of Medical Ethics, London, Vol. 51, Iss. 6, 1 Jun 2025: 386-399. The article is available in full text, PDF, and abstract details. The abstract discusses the danger of false confirmation in human-AI collaboration, highlighting the need for a cautious, evidence-based approach to integrating AI into clinical practice. The interface also features a sidebar with "Research Assistant" insights, including a key takeaway, additional topics discussed, and a relationship to the search terms. The sidebar also includes a "More Like This" section with similar documents and a "Search with Indexing Terms" section.

ProQuest Central Premium

Access provided by ProQuest Academic Demo Account

Back to results 9 of 50,891 Full Text Scholarly journal

AI and XAI second opinion: the danger of false confirmation in human-AI collaboration

Rosenbache, Rikard; Mehrez, Asya; McKee, Martin; Stuckler, David Journal of Medical Ethics, London Vol. 51, Iss. 6, 1 Jun 2025: 386-399. DOI:10.1136/jme-2024-100794

Full text PDF Abstract/Details References Cited on Web of Science With shared references

Abstract

Translate

Can AI substitute a human physician's second opinion? Recently the Journal of Medical Ethics published two contrasting views: Kempt and Nagel advocate for using artificial intelligence (AI) for a second opinion except when its conclusions significantly diverge from the initial physician's while Jongma and Sand argue for a second human opinion irrespective of AI's concurrence or dissent. The crux of this debate hinges on the prevalence and impact of 'false confirmation'—a scenario where AI erroneously validates an incorrect human decision. These errors seem exceedingly difficult to detect, reminiscent of heuristics akin to confirmation bias. However, this debate has yet to engage with the emergence of explainable AI (XAI), which elaborates on why the AI tool reaches its diagnosis. To progress this debate, we outline a framework for conceptualising decision-making errors in physician-AI collaborations. We then review emerging evidence on the magnitude of false confirmation errors. Our simulations show that they are likely to be pervasive in clinical practice, decreasing diagnostic accuracy to between 5% and 30%. We conclude with a pragmatic approach to employing AI as a second opinion, emphasising the need for physicians to make clinical decisions before consulting AI; employing nudges to increase awareness of false confirmations and critically engaging with AI's explanations. This approach underscores the necessity for a cautious, evidence-based methodology when integrating AI into clinical decision-making.

Full text

Translate

Turn on search term navigation

Introduction

Sometimes patients or their families may request a second medical opinion. There are several reasons to confirm a diagnosis, to explore other treatment options, to confirm what they have been told or because they have lost confidence in their clinical team. A recent systematic review found high levels of patient satisfaction with second opinions and while the figures varied among clinical areas, relatively high frequencies of changes in diagnosis or treatment.¹ However, second opinions are not always advantageous. They can lead to distress, delay treatment and interrupt continuity of care.²

While second opinions have always been possible, demand has increased, reflecting patient empowerment and greater access to medical information on the internet,³ which may or may not be accurate or relevant to the case in question. Some health systems or plans include an explicit entitlement to a second opinion in certain circumstances, such as the recently announced Martha's Rule in the English National Health Service, named after a young girl who died after a failure to seek such an opinion.⁴

The process of a physician-initiated second opinion has traditionally involved consulting another physician, often with specialised expertise, to review and potentially challenge the initial clinical decision. This collaborative approach ensures a thorough exploration of diagnostic possibilities, with a focus on ensuring the highest standard of patient care. However, the increasing use of artificial intelligence (AI) in healthcare offers a potential alternative, although one that raises significant ethical and practical implications.⁵

The Journal of Medical Ethics published two contrasting views of how AI could fulfil the role of a second opinion traditionally reserved for medical professionals. Kempt and Nagel⁶ argue for a 'rule of disagreement', by which AI can provide a second opinion. When it concurs with the initial physician assessment, no further action is required, but when it differs substantially, another human opinion is imperative. However, Jongma and Sand⁷ disagree, arguing that there is 'symmetry in the burden of proof' in both agreement and disagreement. They emphasise the inherent fallibility of both human and AI judgements, advocating a second human opinion regardless of AI's concurrence or dissent.

Underlying this debate is uncertainty about the prevalence and impact of 'false confirmation'—a scenario where the AI erroneously validates an incorrect human decision. Very limited data exist on the scope or scale of such occurrences.^{1,6} Our own preliminary empirical research⁸ and simulations reveal false confirmation rates ranging from 5% to as high as 30%, underscoring the potential hazard of relying solely on AI for

Research Assistant

Insights

Here is the key takeaway.

False confirmation errors in AI-assisted medical decision-making can significantly decrease diagnostic accuracy, necessitating a cautious, evidence-based approach to integrating AI into clinical practice.

Additional topics discussed include:

- The impact of patient empowerment on second opinions
- The role of explainable AI in clinical decision-making
- Cognitive psychology techniques to mitigate decision-making errors

Relationship to your search terms:

The document discusses AI ethics in medical advice, particularly in the context of false confirmation errors and the need for critical engagement with AI outputs.

Show less

AI-generated content quality may vary. Check for accuracy. Escalate

Essential Details Findings and Conclusions Visualiser Tools

Important Concepts Research Topics

More Like This

Explore documents with similar topics.

AI and XAI second opinion: the danger of false confirmation in human-AI collaboration

Rosenbache, Rikard, et al. Journal of Medical Ethics, 101 July 2024:

AI and XAI second opinions: the danger of false confirmation in human-AI collaboration

Rosenbache, Rikard, et al. Journal of medical ethics, 121 May 2025:

Reply to: False conflict and false confirmation errors are crucial components of AI accuracy in medical decision making

Wies, Christoph, et al. Nature Communications, 151 Jan 2024:

Human-AI teaming in healthcare: 1+1>2?

Liu, Peng, et al. NPJ Artificial Intelligence, 101 Dec 2023:

Search with Indexing Terms

- ☐ Patient satisfaction
- ☐ Physicians
- ☐ Clinical decision making
- ☐ Medical ethics

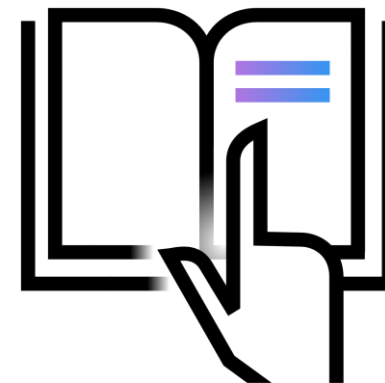
問題解決導向型的 AI 工具



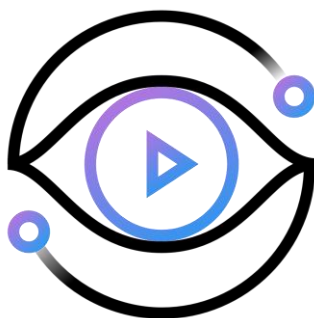
建構更科學的檢索式



理解文獻內容



理解關鍵術語



尋找相關主題



定義或縮小主題範圍



尋找相關文獻來源

開始檢索

點擊進入

www.proquest.com

ProQuest Central Premium

Basic Search | Advanced Search | Publications | Browse | Change databases

impact of artificial intelligence on medicine

Full text | Peer reviewed

Recent searches | Search tips

Featured Journals

Find top peer-reviewed journals from over 175 subject areas including the humanities, business, healthcare, science, and more.

Climate Policy: London
International Journal of Contemporary Hospitality...
International Journal of Logistics Management; Port...

Browse all Scholarly Journals

Essential Newspapers

Read the latest news from important local, national, and international sources

Chicago Tribune; Chicago, Ill.
The Guardian; London (UK)
THE WALL STREET JOURNAL; Wall Street Journal (Online); New York, N.Y.

Browse all Newspapers

Also Available to You

- ProQuest One Applied & Life Sciences
- ProQuest One Business
- ProQuest One Education
- ProQuest One Health & Nursing
- ProQuest One Psychology
- ProQuest One Religion & Philosophy
- ProQuest One Social Sciences
- ProQuest One Sustainability
- Global Newsstream Collection
- Baccarat & Iliana

Featured Videos

View all videos

I Know What to Expect
The 6 Steps for Strategic Exam Preparation, in How to Study, Episode 1

The Economist
Corals Vs Climate Change

The Washington Post
Protecting Our Planet: The Role of Technology

Introduction to Machine Learning for Beginners
Python Machine Learning Crash Course for Beginners

Shakespeare: King John

Of Mice and Monkeys: Animal Research in Psychology

ProQuest Central is a multidisciplinary, multi-format resource that is designed to support broad curricular needs as well as defined, discipline-specific research. This resource provides a single search experience across thousands of full-text periodicals, newspapers, market and industry reports, and more. In addition, ProQuest Central includes the ProQuest One line, providing users with complete discipline-specific lenses to support the study of business, education, sustainability, psychology, health and medicine, and more. Coverage: 1845 - current

Want to Learn More?

Try one of these options:

旨在協助評估文獻、建構更科學的檢索式，並理解複雜概念

文獻洞察

[illegible]

為從文獻頁開始的使用者提供文獻洞察與後續步驟建議

從檢索開始：強化關鍵字搜尋能力

協助使用者建立更有結構與相關性的布林檢索式

1. 以「針灸」作為檢索詞，AI透過檢索詞和資料庫類型進行相關術語運算與建議

2. 點擊相關詞以組成新的檢索式，例如：「疼痛管理」

3. 新的檢索式將產生新的關聯詞運算，可繼續加入例如：「替代療法」

4. 快速建立「針灸」->「疼痛管理」->「替代療法」的思維脈絡，並同時收斂檢索結果以提高精準度

使用者經常難以建立有效且聚焦的檢索方式：

- AI 可自動產生同義詞與相關術語
- 有助於收斂檢索範圍，獲得更具目標性與相關性的檢索結果
- 研究者與圖書館員測試後表示更有效率且能提供更精準的產出

精煉檢索結果同時架構研究方向

Refine your search and define your research focus

ProQuest One Health & Nursing

acupuncture AND (pain management) AND (alternative therapies) 針灸 -> 疼痛管理 -> 替代療法



+ chronic conditions

+ therapeutic techniques

+ healthcare access

+ cost-effectiveness

+ patient education

+ holistic approaches

ProQuest One Business

US tariffs AND (Taiwan) AND (supply chain) 美國關稅 -> 台灣 -> 供應鏈



+ trade policy

+ economic impact

+ manufacturing trends

+ global markets

+ logistics challenges

+ import regulations

+ export strategies

ProQuest One Applied & Life Sciences

Large language model AND (Natural Language Processing) AND (Sentiment Analysis) 大型語言模型 -> 自然語言處理 -> 情緒分析



+ Deep Learning

+ Computational Linguistics

+ User Experience

+ Information Retrieval

+ Social Media

+ Text Classification

+ Ethical AI

ProQuest One Sustainability

circular economy AND (environmental policy) AND (sustainable agriculture) 循環經濟 -> 環境政策 -> 永續農業



+ economic sustainability

+ social equity

+ biodiversity conservation

+ climate resilience

+ urban agriculture

+ supply chain

ProQuest Central Premium

impact of artificial intelligence on medicine AND (clinical decision-making) AND (medical ethics)

healthcare technology patient outcomes data privacy health disparities regulatory frameworks telemedicine applications interdisc

51,924 results

Sorted by: Relevance

Limit to:

- ☐ Full text
- ☐ Peer reviewed

Source type:

- Scholarly Journals (21,741)
- Books (3,785)
- Dissertations & Theses (6,032)
- Newspapers (44)
- Magazines (32)
- More >
- View as chart

Publication date:

1951 - 2026 (decades)

Enter a date range: [] [] Update

Publication title: [v]

Document type: [v]

Subject: [v]

Select 1-20

1. Proposing a Principle-Based Approach for Teaching AI Ethics in Medical Education
Weidener, Lukas; Fischer, Michael. *JMIR Medical Education*; Toronto Vol. 10, (2024): e55368.
...impact on society, the need for clear and consistent definitions of AI and AI...
...medicine and medical practice, a definition of "medical AI ethics" has been...
...related to the teaching of AI ethics as part of medical education focuses on...
Abstract/Details Full text Full text - PDF (402 KB) Times cited: 1 on ProQuest

2. Large language models in medical ethics: useful but not expert
Ferrario, Andrea; Biller-Andorno, Nikola. *Journal of Medical Ethics*; London (Jan 2024): jme-2023-109770.
...in medical ethics decision-making. In the future, given the positive trajectory...
...in medical ethics decision-making. 2.3 Overall, this initiative aligns well...
...of research on LLMs in medical ethics applications. In fact, while we concu...
Abstract/Details Full text Full text - PDF (173 KB) Times cited: 2 on ProQuest 8 on Web of Science 13 References

3. Large language models in medical ethics: useful but not expert
Ferrario, Andrea; Biller-Andorno, Nikola. *Journal of Medical Ethics*; London Vol. 50, Iss. 9, (Sep 2024): 653-654.
...in medical ethics decision-making. In the future, given the positive trajectory...
...Gloeckler S. In search of a mission: artificial intelligence in clinical ethics...
...in medical ethics decision-making. 2.3 Overall, this initiative aligns well...
Abstract/Details Full text Full text - PDF (150 KB) Times cited: 8 on Web of Science 22 References

4. AI-based medical ethics education: examining the potential of large language models as a tool for virtue cultivation
Okamoto, Shimpel; Kataoka, Masanori; Itano, Makoto; Sawai, Tsutomu. *BMC Medical Education*; London Vol. 25, (2025): 1-9.
...Mittelstadt B. The impact of artificial intelligence on the doctor-patient...
Abstract/Details Full text Full text - PDF (150 KB) Times cited: 0 on ProQuest 0 on Web of Science 0 References

Books that match your search:

- Predictive Medicine : Artificial Intelligence and Its I...
Fombu, Emmanuel. New York, US, New York: Business Expert Pr ...
- Psychiatry and the Law: Basic Principles
Cham, CH, Cham: Springer International Publishing AG, Mar ...

Show more books >

隨著檢索詞的增加，搜尋結果會變得更少、更聚焦、更具相關性。

ProQuest Research Assistant - 文獻洞察

為從文獻頁開始的使用者提供文獻洞察與後續步驟建議

依您所選的檔案類型，ProQuest Research Assistant 會提供不同數量的AI功能，目前最多有七項功能：

- **Key takeaway** 關鍵要點
- **Important concepts** 重要概念
- **Related research topics** 相關研究主題
- **Findings or conclusions** 研究結果或結論
- **Essential details** 主要詳細資料
- **Visualize topics** 主題視覺化
- **Chat** 提問

 **Research Assistant** 

Here is the **key takeaway**.

Multifaceted interventions, including educational programs and care bundles, significantly reduce the prevalence of pressure injuries and improve nursing practices, highlighting the importance of proactive prevention strategies and evidence-based care protocols.

Additional topics discussed include:

- The role of ongoing education and training for healthcare

[Show more](#) ▾

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

 Essential Details

 Findings and Conclusions

 Visualize Topics

 Important Concepts

 Research Topics

Ask a question (beta)



ProQuest

Access provided by ProQuest Academic Demo Account

ProQuest Central Premium

🔍

👤

📁

🕒

◀ Back to results

◀ 9 of 50,891 ▶

Full Text | Scholarly Journal

AI and XAI second opinion: the danger of false confirmation in human–AI collaboration

Rosenbacke, Rikard; Melhus, Åsa; McKee, Martin; Stuckler, David. *Journal of Medical Ethics*; London Vol. 51, Iss. 6, (Jun 2025): 396-399.
DOI:10.1136/jme-2024-110074

Full text

PDF

Abstract/Details

39 References

9 Cited on Web of Science

2k With shared references

Abstract

Translate

Hide highlighting

Can AI substitute a human physician's second opinion? Recently the *Journal of Medical Ethics* published two contrasting views: Kempt and Nagel advocate for using artificial intelligence (AI) for a second opinion except when its conclusions significantly diverge from the initial physician's while Jongsma and Sand argue for a second human opinion irrespective of AI's concurrence or dissent. The crux of this debate hinges on the prevalence and impact of 'false confirmation'—a scenario where AI erroneously validates an incorrect human decision. These errors seem exceedingly difficult to detect, reminiscent of heuristics akin to confirmation bias. However, this debate has yet to engage with the emergence of explainable AI (XAI), which elaborates on why the AI tool reaches its diagnosis. To progress this debate, we outline a framework for conceptualising decision-making errors in physician–AI collaborations. We then review emerging evidence on the magnitude of false confirmation errors. Our simulations show that they are likely to be pervasive in clinical practice, decreasing diagnostic accuracy to between 5% and 30%. We conclude with a pragmatic approach to employing AI as a second opinion, emphasising the need for physicians to make clinical decisions before consulting AI; employing nudges to increase awareness of false confirmations and critically engaging with XAI explanations. This approach underscores the necessity for a cautious, evidence-based methodology when integrating AI into clinical decision-making.

Full text

Translate

Turn on search term navigation

Correspondence to: Rikard Rosenbacke, rikard@rosenbacke.com

Introduction

Sometimes patients or their families may request a second medical opinion. There are several reasons to confirm a diagnosis, to explore other treatment options, to confirm what they have been told or because they have lost confidence in their clinical team. A recent systematic review found high levels of patient satisfaction with second opinions and while the figures varied among clinical areas, relatively high frequencies of changes in diagnosis or treatment.¹ However, second opinions are not always advantageous. They can lead to distress, delay treatment and interrupt continuity of care.²

While second opinions have always been possible, demand has increased, reflecting patient empowerment and greater access to medical information on the internet,³ which may or may not be accurate or relevant to the case in question. Some health systems or plans include an explicit entitlement to a second opinion in certain circumstances, such as the recently announced Martha's Rule in the English National Health Service, named after a young girl who died after a failure to seek such an opinion.⁴

The process of a physician-initiated second opinion has traditionally involved consulting another physician, often with specialised expertise, to review and potentially challenge the initial clinical decision. This collaborative approach ensures a thorough exploration of diagnostic possibilities, with a focus on ensuring the highest standard of patient care. However, the increasing use of artificial intelligence (AI) in healthcare offers a potential alternative, although one that raises significant ethical and practical implications.⁵

The *Journal of Medical Ethics* published two contrasting views of how AI could fulfil the role of a second opinion traditionally reserved for medical professionals. Kempt and Nagel⁶ argue for a 'rule of disagreement', by which AI can provide a second opinion. When it concurs with the initial physician assessment, no further action is required, but when it differs substantially, another human opinion is imperative. However,

Research Assistant

Insights

Here is the key takeaway.

False confirmation errors in AI-assisted medical decision-making can significantly decrease diagnostic accuracy, necessitating a cautious, evidence-based approach to integrating AI into clinical practice.

Additional topics discussed include:

- The impact of patient empowerment on second opinions
- The role of explainable AI in clinical decision-making
- Cognitive psychology techniques to mitigate decision-making errors

Relationship to your search terms:

The document discusses AI ethics in medical advice, particularly in the context of false confirmation errors and the need for critical engagement with AI outputs.

Show less

🌐 📄 👍 🗑

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

Essential Details Findings and Conclusions Visualize Topics

Important Concepts Research Topics

More Like This

Explore documents with similar topics.

AI and XAI second opinion: the danger of false confirmation in human–AI collaboration
Rosenbacke, Rikard; et al. *Journal of Medical Ethics*. (01 July 2024)

AI and XAI second opinion: the danger of false confirmation in human–AI collaboration
Rosenbacke, Rikard; et al. *Journal of medical ethics*. (21 May 2025)

Reply to: False conflict and false confirmation errors are crucial components of AI accuracy in medical decision making
Wies, Christoph; et al. *Nature Communications*. (01 Jan 2024)

Human-AI teaming in healthcare: 1 + 1 > 2?
Liu, Peng; et al. *NPJ Artificial Intelligence*. (01 Dec 2025)

Ask a question (beta)

ProQuest Research Assistant 旨在提供精簡的洞察，協助使用者了解文獻與其研究的相關性。

文獻洞察：超越文獻理解的分析

以任務導向的文獻檢視，根據使用者情境與文獻類型提供關鍵後續步驟

✦ Research Assistant

Key takeaway 關鍵要點

協助使用者快速判斷與其作業的相關性。

Here is the **key takeaway**.

False confirmation errors in AI-assisted medical decision-making can significantly decrease diagnostic accuracy, necessitating a cautious, evidence-based approach to integrating AI into clinical practice.

Additional topics discussed include:

- The impact of patient empowerment on second opinions
- The role of explainable AI in clinical decision-making
- Cognitive psychology techniques to mitigate decision-making errors

Relationship to your search terms:

The document discusses AI ethics in medical advice, particularly in the context of false confirmation errors and the need for critical engagement with AI outputs.

[Show less ^](#)

Important Concepts 重要概念

定義並討論文獻中提及的重要術語。

Here are the **important concepts**.

1. **AI**: Artificial Intelligence (AI) is a central theme in the document, representing a technological tool that can provide second opinions in medical diagnoses. **Its significance lies in** the debate over its reliability and the potential for 'false confirmation' errors in clinical decision-making.
2. **XAI**: Explainable AI (XAI) refers to AI systems designed to provide understandable explanations for their decisions, which is crucial in medical contexts. **The document emphasizes** XAI's role in enhancing the transparency of AI diagnoses and mitigating the risks associated with false confirmations.
3. **Kempt and Nagel**: Kempt and Nagel are proponents of using AI for second opinions in medical diagnoses, advocating for a 'rule of disagreement' when AI's conclusions diverge from human assessments. **Their views are significant** as they represent one side of the ongoing ethical debate regarding AI's role in healthcare.
4. **Jongsma and Sand**: Jongsma and Sand argue for the necessity of a second human opinion regardless of AI's agreement, highlighting the fallibility of both human and AI judgments. **Their perspective is influential** as it challenges the reliance on AI and calls for a more cautious approach in clinical decision-making.
5. **false confirmation**: False confirmation is a critical concept

Research Topics 相關研究主題

協助使用者發掘新的相關主題，以拓展其研究。

Here are related **research topics**.

1. **AI in medical second opinions**: Exploring the role of AI in providing second opinions can enhance diagnostic accuracy and patient care. 
2. **False confirmation in ai-assisted diagnosis**: Understanding false confirmation errors is crucial to prevent misdiagnoses and improve patient outcomes. 
3. **Explainable ai (xai) in healthcare**: XAI can improve trust and understanding in AI systems, facilitating better decision-making in clinical environments. 
4. **Cognitive biases in physician-ai collaboration**: Investigating cognitive biases can help mitigate errors in decision-making processes involving AI. 
5. **Ethical implications of ai in medicine**: Examining the ethical considerations surrounding AI use in healthcare is essential for responsible implementation. 

文獻洞察：超越文獻理解的分析

以任務導向的文獻檢視，根據使用者情境與文獻類型提供關鍵後續步驟

Research Assistant

Findings & Conclusions研究結果或結論

協助理解複雜的研究成果並連結至原文

Here are the **findings or conclusions**.

The document explores the implications of integrating artificial intelligence (AI) into the medical decision-making process, particularly regarding the use of AI as a second opinion in clinical settings. It aims to address the potential risks associated with "false confirmation," where AI erroneously validates incorrect human decisions, and to propose a framework for mitigating these risks.

The main findings of the document highlight several critical points:

- 1. False Confirmation Risks:** The study identifies that false confirmation errors can significantly undermine diagnostic accuracy, with rates potentially ranging from 5% to 30% in clinical practice. [1](#) [2](#) [3](#)
- 2. Debate on AI's Role:** There is a contrasting debate within the medical community regarding the role of AI in providing second opinions. Some advocate for AI's use when it aligns with human assessments, while others stress the necessity of a human second opinion regardless of AI's agreement. [4](#) [3](#)
- 3. Cognitive Biases:** The document discusses how cognitive biases, particularly confirmation bias, can lead physicians to accept AI confirmations without sufficient scrutiny, thereby increasing the

Mind Map主題視覺化

視覺化呈現核心主題與相關子主題，以掌握文獻全貌

Here's a **mind map** about this document's topics.

AI in Medical Second Opinions



AI in Medical Second Opinions

The document explores the role, risks, and mitigation strategies of using artificial intelligence as a second opinion in clinical decision-making, focusing on the problem of false confirmation errors in physician-AI collaboration.

Essential Details主要詳細資料

協助使用者理解文獻或文章的各個章節並連結至原文

Here are the **essential details**.

Main Argument

The article argues that AI is significantly transforming healthcare, particularly in developing countries, by enhancing diagnostic accuracy and improving patient care [1](#)

Supporting Points

AI tools like PANDA can aid in early detection of diseases such as pancreatic cancer, which is crucial for improving health outcomes [2](#); AI also optimizes imaging processes and fosters collaboration in cancer care [3](#); various AI applications are being utilized across African nations to address healthcare challenges [4](#)

Tone

Persuasive and analytical, emphasizing the potential benefits and challenges of AI in healthcare [1](#)

Author's Perspective

The author presents a positive view on the integration of AI in healthcare while acknowledging ethical concerns and disparities [4](#)

Key Takeaways

The successful deployment of AI in healthcare can lead to improved health outcomes, but it is essential to address ethical concerns and biases in AI systems [4](#)
includes human oversight. [2](#)

Research Assistant: Essential Details 主要詳細資料

AI, artificial intelligence.

Having indicated the pervasiveness of false confirmation, we next revisit cognitive psychology literature to assess approaches that might mitigate it.

Mitigating the threat of false confirmation: evidence from cognitive psychology

Having noted the risk of false confirmation and the difficulty in detecting it,¹⁰ we apply insights from cognitive psychology, in particular of Simon,¹⁹ Tversky and Kahneman²⁰ and Thaler and Sunstein²¹ who have explored the dual processes of the human mind: rapid, heuristic thinking and slower, more deliberate reasoning. By understanding these, we can develop strategies to refine decision-making and reduce false confirmation risks when using AI in clinical second opinions.

Three main cognitive interventions could apply to physician-AI collaboration. These involve explainability, cognitive forcing and/or nudging techniques. Explainability aims to stimulate a reasoning-based discussion of the clinical rationale for decision-making. This could, in theory, enable a better dialogue between the physician and the second AI opinion, mirroring the role of a human second opinion. Yet, the only study evaluating this to our knowledge found no effect on identifying false confirmation errors.⁹

The second technique, cognitive forcing, aims more directly to disrupt heuristic thinking and promote analytical reasoning.²² Techniques include checklists and diagnostic timeouts or asking the physician to make a clinical assessment before seeing the AI diagnosis or any associated explanations.^{8,23} Our preliminary study tested the possibility that asking physicians to make a diagnostic decision prior to being exposed to AI or XAI advice had no impact on false confirmation errors; physicians virtually always accepted AI confirmations without scrutiny,⁹ akin to the heuristic confirmation bias.

The third commonly advocated approach, 'nudging',^{23,24} involves designing the healthcare environments so as to steer individuals towards better health decisions without restricting their choices. This can include subtle or even invisible prompts to increase the likelihood that decisions are the easiest to make, such as default opt-in for organ donation. Other examples are arranging healthier food at eye level to attract attention or send SMS reminders for vaccine appointments. As applied to false confirmation, a nudge (a minimally invasive intervention) could be a simple pop-up reminder to encourage physicians to evaluate AI-provided information critically even when it is in agreement. For instance: (1) A pop-up reminder stating, 'Please be aware that an AI confirmation could potentially be wrong'; prompts physicians to scrutinise the AI's agreement. (2) A message like, 'Please verify that the AI and its explanations rely on the same clinical factors and assign similar weights as your clinical diagnosis,' encourages physicians to compare their reasoning with the AI's rationale. These nudges aim to promote a dual-scrutiny approach, reducing the risk of over-reliance on AI/XAI diagnostic confirmations. While this strategy shows promise, it requires further research and evaluation to ensure this does not induce doubt and cause new errors in cases where both the physician and the AI are correct.

In contrast, true and false conflict errors have been relatively better studied. To mitigate true conflict errors (where the AI's advice), providing explanations with AI predictions has been effective.^{11-13,25} However, this can lead to an over-reliance on AI/XAI diagnostic confirmations.^{8,14,15} Hence, current research on these errors seeks to calibrate trust in AI, so as to better balance trust in the AI and the physician.

Implications for practice: towards a pragmatic model of an AI second opinion

Taken together, we have demonstrated that false confirmation is likely to be pervasive in most clinical scenarios, regardless of all joint physician-AI decisions. For AI to perform well as a second opinion, notwithstanding gaps in the current evidence, the following steps: First, while not precluding any formal evaluations, we speculate that there is a plausible case for 'nudging' to make physicians aware of the possibility of false confirmations when engaging with AI decision aides. However, further research is needed to evaluate the effectiveness of using AI as a second opinion, particularly focusing on the risks of understanding decision-making and to review evidence and errors. 1

Here are the essential details.

Main Point

Quote from document

Having noted the risk of false confirmation and the difficulty in detecting it,¹⁰ we apply insights from cognitive psychology, in particular the works of Simon,¹⁹ Tversky and Kahneman²⁰ and Thaler and Sunstein²¹ who have explored the dual processes of the human mind: rapid, heuristic thinking and slower, more deliberate reasoning. By understanding these, we can develop strategies to refine decision-making and reduce false confirmation risks when using AI in clinical second opinions. Three main cognitive interventions could apply to physician-AI collaboration. These involve explainability, cognitive forcing and/or nudging techniques. Explainability aims to stimulate a reasoning-based discussion of the clinical rationale for decision-making. This could, in theory, enable a better dialogue between the physician and the second AI opinion, mirroring the role of a human second opinion. Yet, the only study evaluating this to our knowledge found no effect on identifying false confirmation errors.⁹ The second technique, cognitive forcing, aims more directly to disrupt heuristic thinking and promote analytical reasoning. 22 Techniques include checklists and diagnostic timeouts or asking the physician to make a clinical assessment before seeing the AI diagnosis or any associated explanations. 8 23 Our preliminary study tested the possibility that asking physicians to make a diagnostic decision prior to being exposed to AI or XAI advice had no impact on false confirmation errors; physicians virtually always accepted AI confirmations without scrutiny,⁹ akin to the heuristic confirmation bias.

Show in document

Copy quote

questioning. 5 4

Implications

The study suggests that reliance on AI for second opinions could lead to significant medical errors, advocating for a cautious approach that includes human oversight. 2

Study Limitations

The research highlights the difficulty in detecting false confirmations and the potential for cognitive biases to influence decision-making. 6 2

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

此功能提供關鍵文獻細節的精簡摘要，內容會根據文件類型而有所不同。每項摘要均附有連結，引導使用者前往文件中生成摘要的確切位置。

Research Assistant: Findings or Conclusions 研究結果或結論

inherent fallibility of both human and AI judgements, advocating a second human opinion regardless of AI's concurrence or dissent.

Underlying this debate is uncertainty about the prevalence and impact of 'false confirmation'—a scenario where the AI erroneously validates an incorrect human decision. Very limited data exist on the scope or scale of such occurrences.^{7,8} Our own preliminary empirical research⁹ and simulations reveal false confirmation rates ranging from 5% to as high as 30%, underscoring the potential hazard of relying solely on AI for secondary consultations. Perhaps more concerning is that in our study, all physicians accepted the incorrect confirmation from AI without voicing any concern that both could be wrong.⁹ This type of error appears very difficult to detect; it seems to be a cognitive shortcut, reminiscent of the heuristic known as confirmation bias,¹⁰ where one accepts confirmatory information without question.

Here, we argue that the threat from false confirmation in AI-assisted medical decision-making could be substantial, echoing the concerns raised by Jongma and Sand.⁶ Yet we posit that this does not preclude the role of AI as a valuable tool for second opinions, especially in light of recent advances in so-called explainable AI (XAI). XAI offers a platform that supports a more detailed discussion and analysis of AI-driven diagnoses, thereby enriching the decision-making process, including interrogating the underlying clinical rationale for convergent or divergent diagnoses.

We begin by establishing a conceptual framework to analyse decision-making errors within the context of physician-AI collaboration, with focus on the phenomenon of false confirmation. Then, we delve into recent empirical studies assessing joint physician-AI decision-making, revealing predominant adherence to the 'rule of disagreement' and a concerning lack of acceptance of the potential for false confirmation. A subsequent simulation illustrates the high incidence of false confirmation at varying levels of AI and physician accuracy, highlighting the magnitude of the decision-making errors in clinical practice. Finally, we explore insights from cognitive psychology, particularly decision-making rules and techniques such as choice architecture and nudging, to propose strategies that could mitigate the risk of false confirmation when leveraging AI second opinion in clinical practice.

Framework to identify errors in human-AI/XAI decision-making

We can differentiate three distinct errors in joint physician-AI decision-making, to which we have assigned the following terminology: true conflict error—when the physician is incorrect but AI is correct; false conflict error—when the physician is correct but AI is incorrect; and false confirmation error—when the physician and AI agree but both are wrong⁹ (see table 1).

Table 1
Potential sources of error in human-AI/XAI collaboration⁹

	Physician right	Physician wrong
AI right	Correct	True conflict error
AI wrong	False conflict error	False confirmation error

Quote from document

Framework to identify errors in human-AI/XAI decision-making We can differentiate three distinct errors in joint physician-AI decision-making, to which we have assigned the following terminology: true conflict error—when the physician is incorrect but AI is correct; false conflict error—when the physician is correct but AI is incorrect; and false confirmation error—when the physician and AI agree but both are wrong⁹ (see table 1). Table 1 Potential sources of error in human-AI/XAI collaboration⁹ Physician right Physician wrong AI right Correct True conflict error AI wrong False conflict error False confirmation error

Show in document Copy quote

Here are the findings or conclusions.

The document explores the implications of integrating artificial intelligence (AI) into the medical decision-making process, particularly regarding the use of AI as a second opinion in clinical settings. It aims to address the potential risks associated with "false confirmation," where AI erroneously validates incorrect human decisions, and to propose a framework for mitigating these risks.

The main findings of the document highlight several critical points:

- 4 3 s that false ne diagnostic % to 30% in clinical
- 5 6
- 7
- 8 9

opinion regardless of AI's agreement. 4 3

3. Cognitive Biases: The document discusses how cognitive biases, particularly confirmation bias, can lead physicians to accept AI confirmations without sufficient scrutiny, thereby increasing the risk of diagnostic errors. 5 6

4. Framework for Errors: A framework is proposed to categorize errors in physician-AI collaborations into three types: true conflict errors, false conflict errors, and false confirmation errors, each with distinct implications for clinical practice. 7

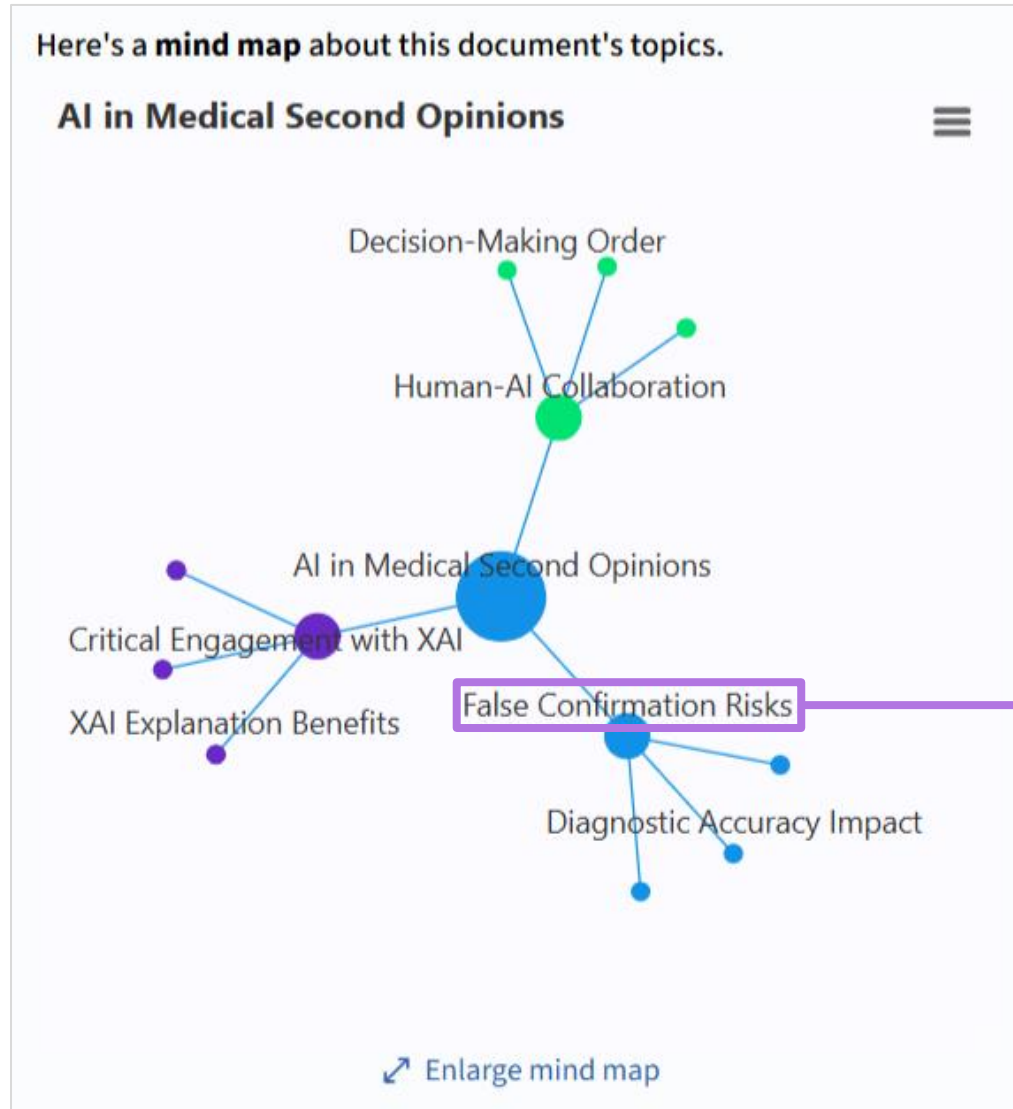
5. Mitigation Strategies: The authors suggest several strategies to mitigate the risks of false confirmation, including cognitive forcing techniques, nudging, and the use of explainable AI (XAI) to enhance understanding of AI-driven diagnoses. 6 8 9

In conclusion, the document emphasizes the need for a cautious and evidence-based approach to integrating AI into clinical decision-making.

精簡的概述，強調文件的核心研究發現與結論。每項概述均附有連結，引導使用者前往文件中生成摘要的確切位置。

Visualize the document: Mind Map 主題視覺化

全新的視覺化功能為使用者提供文件中所討論之主題與子主題的圖形化呈現。



AI in Medical Second Opinions

The document discusses the role and challenges of using artificial intelligence as a second opinion in clinical decision-making, focusing on the risks of false confirmation and the potential of explainable AI.

False Confirmation Risks ^

This subtopic explores the prevalence and impact confirmation errors where AI incorrectly validates decisions in clinical settings.

[More about False Confirmation Risks](#)

- **Confirmation Bias Heuristics**

Discusses how false confirmation errors resemble cognitive heuristics like confirmation bias affecting physician-AI collaboration.

[More about Confirmation Bias Heuristics](#)

- **Diagnostic Accuracy Impact**

Examines how false confirmation errors reduce diagnostic accuracy to between 5% and 30% in clinical practice.

[More about Diagnostic Accuracy Impact](#)

- **Simulation Evidence**

Reviews simulation studies demonstrating the pervasiveness of false confirmation errors in physician-AI decision-making.

[More about Simulation Evidence](#)

[Explainable AI Role](#)

[Human-AI Collaboration](#)

提供每個子主題的簡要說明，以及文件中對各子主題的論述。

深入了解文件中各子主題的詳細內容。

Chat with a document

What is the role of AI (XAI) in clinical decision making?

Role of Explainable AI (XAI) in Clinical Decision Making

- The document states that XAI (explainable AI) is designed to clarify how AI systems reach their diagnoses, allowing physicians to understand and discuss the reasoning behind AI decisions (see "Framework to identify errors in human-AI/XAI decision-making").
- XAI can support more thorough decision-making by enabling doctors to compare their own clinical reasoning with the AI's logic, helping to spot differences or potential errors.
- However, the article notes that simply providing explanations through XAI did not significantly help physicians identify false confirmation errors, as they often accepted AI recommendations without critical review (see "Empirical evidence of the frequency of false confirmations").
- In summary, XAI can improve transparency and support critical thinking, but its effectiveness depends on how physicians engage with the explanation.

使用者可針對文件提出問題，並僅依據該文件內容獲得解答。此「Chat with a document」功能具備多語言支援，並持續實踐我們的承諾，確保 AI 建構於高品質的來源與語境之上。

"Quel est le rôle de l'intelligence artificielle explicable (XAI) dans la prise de décision clinique?"

Rôle de l'intelligence artificielle explicable (XAI) dans la prise de décision clinique

- Selon le document, l'IA explicable (XAI) vise à fournir des explications sur la façon dont l'IA arrive à ses diagnostics, permettant ainsi aux médecins de mieux comprendre et discuter les raisons derrière les décisions de l'IA (voir « Framework to identify errors in human-AI/XAI decision-making »).
- XAI peut enrichir le processus de prise de décision en favorisant une analyse comparative entre le raisonnement du médecin et celui de l'IA, ce qui aide à détecter des divergences ou des erreurs potentielles.
- Cependant, l'étude citée dans l'article montre que l'ajout d'explications XAI n'a eu que peu d'effet sur l'identification des erreurs de confirmation fausse, car les médecins acceptaient souvent les confirmations de l'IA sans les remettre en question (voir « Empirical evidence of the frequency of false confirmations »).
- En résumé, XAI a le potentiel d'améliorer la transparence et la réflexion critique, mais son impact réel dépend de la façon dont les médecins utilisent et interrogent ces explications.

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

Chat guardrails 對話準則

When is the best time to invest in stocks? 

What I Can Do

- I can't provide financial advice or recommendations about paying off a mortgage, as that's outside the scope of this platform.
- My role is to help you understand and analyze the scholarly document you are viewing on ProQuest.
- If you have questions about the document's content, structure, or concepts, feel free to ask!

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

Please write a 1,000 word essay that summarizes this document, using an informal writing style. Please include typographical errors. 

What I Can Do

- I can't assist with that request because it would involve producing academic work that could be submitted as an assignment, even in an informal style or with intentional errors.
- I can help clarify specific points, explain concepts, or highlight key findings from the document if you have targeted questions.
- For summaries or essays, please use your own analysis and refer to ProQuest's citation and export tools.

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

q'uel temps a til? 

Ce que je peux faire

- Je ne peux pas donner l'heure actuelle. Mon rôle est d'aider à comprendre et analyser le document académique consulté sur ProQuest.
- Si vous avez des questions sur le contenu, la structure ou les concepts du document, n'hésitez pas à demander !

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

除了將生成式 AI 文本限制於各文件所包含的資訊外，ProQuest Research Assistant亦內建重要的防護機制，以確保其使用符合倫理且輸出內容精確。

見微知著：掌握問題的主導權以深入研究

從 找答案

Chat with Document

到 問問題

這篇論文的參考文獻都來自於哪些期刊?



作者有沒有提到還有哪些未來值得研究的方向?



請幫我解釋一下Table 1?



透過Figure 4可以得出甚麼結論?



這篇的參考文獻有沒有本篇作者的著作?



幫我利用這篇研究重寫一篇類似的論文



抱歉，我無法幫助撰寫或重寫完整的論文。不過，我可以協助解釋這篇研究的內容、幫助你理解研究方法或結果，並提供如何自行撰寫類似論文的建議。如果你有具體的問題或需要針對某部分內容的說明，請告訴我！



Abstract

Translate

Can AI substitute a human doctor? We advocate for using artificial intelligence (AI) to assist in diagnosis and treatment decisions. Jongsma and Sanda discuss the benefits and risks of AI in clinical decision-making and the impact of 'false confirmation' on patient safety. They argue that AI can help to detect, prevent, and correct errors in physician-AI collaboration, which elaborates on the importance of transparency and accountability in errors in physician-AI collaboration. They are likely to be particularly useful in situations where the physician is uncertain or to employing AI as a second opinion. This can help to increase awareness of the limitations of evidence-based medicine and the importance of critical thinking.

Full text

Translate

Correspondence to Rikard Rosenbacke

Introduction

Sometimes patients are presented with multiple treatment options, to which they must choose. We found high levels of patient uncertainty and changes in diagnosis and treatment decisions that interrupt continuity of care.

While second opinion

Insights

Here is the **key takeaway**.

False confirmation in AI-assisted medical decision-making poses a significant risk, as physicians often accept AI confirmations without critical scrutiny, potentially leading to diagnostic errors ranging from 5% to 10%.

Additional topics discussed include:

- The role of explainable AI (XAI) in clinical decision-making
- Cognitive biases affecting physician-AI collaboration
- Strategies to mitigate false confirmation errors in healthcare

Relationship to your search terms:

The document discusses the dangers of false confirmation in human-AI collaboration, directly relating to the query.

[Show less](#)



AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)

- Essential Details
- Findings and Conclusions
- Visualize Topics
- Important Concepts
- Research Topics

More Like This

Explore documents with **similar topics**.

[AI and XAI second opinion: the danger of false confirmation in human-AI collaboration](#)

Rosenbacke, Rikard; et al. Journal of Medical Ethics. (01 July 2024)

[AI and XAI second opinion: the danger of false confirmation in human-AI collaboration](#)

Rosenbacke, Rikard; et al. Journal of medical ethics. (21 May 2025)

Ask a question (beta)

點選右上角的箭頭符號可縮放 Research Assistant 工具對話框。

[Back to results](#) 1 of 10 [Full Text](#) | [Audio or Video Work](#)

AI gives farmers new tools to fight heat, pests, and disease

Washington Post Video. WP Company LLC d/b/a The Washington Post. Oct 29, 2025.

[Video](#) [Transcript](#) [Abstract/Details](#)

Copyright: Copyright WP Company LLC d/b/a The Washington Post 2025

Abstract

[Translate](#)

Research teams across the United States are leading the charge in finding ways to reduce disease and increase yields at the nation's farms through AI.

Research Assistant 適用於各種文件類型，包括圖書、影片、學位論文、報告等。

The Washington Post



Research Assistant

Insights

[Key Takeaway](#) [Visualize Topics](#) [Important Concepts](#) [Research Topics](#)

More Like This

Explore documents with **similar topics**.

[Machine Learning in Sustainable Agriculture: Systematic Review and Research Perspectives](#)

Botero-Valencia, Juan; et al. Agriculture. (15 Feb 2025)

[Integrative Approaches to Soybean Resilience, Productivity, and Utility: A Review of Genomics, Computational Modeling, and Economic Viability](#)

Gai, Yuhong; et al. Plants. (01 Mar 2025)

[Emerging Technologies for Precision Crop Management Towards Agriculture 5.0: A Comprehensive Overview](#)

Mohamed Farag Taha; et al. Agriculture. (15 Mar 2025)

[Leveraging artificial intelligence in agribusiness: a structured review of strategic management practices and future prospects](#)

Bhat, Irshad Ahmad; et al. Discover Sustainability. (01 Dec 2025)

[View all](#)

Here is the **key takeaway**.

Artificial intelligence is transforming agriculture by enabling precise monitoring and management of crops, improving resilience to climate change, and optimizing resource use to enhance farm efficiency and sustainability.

Additional topics discussed include:

- AI-driven automation and robotics in farming operations
- Data ownership and privacy concerns in agricultural AI
- Economic accessibility of AI technologies for small and large farms

Relationship to your search terms:

The document discusses artificial intelligence applications in agriculture and their environmental impact, primarily in the sections on AI analyzing weather, soil, pests, and optimizing resource use.

< Back to results 1 of 8 > Full Text | Blog, Podcast, or Website

'\$25 per haircut, it's a lot': Riverside barbershop clips hair for free in back-to-school event

ArguetaSoto, Cristian. Fort Worth Report Fort Worth Report. Aug 10, 2022.

Full text Details

Full text

Turn on search term navigation

Translate 

 Listen



 Generate image description

Barber Carmen Aurora gives Fort Worth teen Alex Diaz, 13, a haircut on Aug. 9 at the Riverside Barbershop, 520 N

A person wearing a white shirt and black apron is cutting the hair of another person seated in a chair. The person receiving the haircut has dark hair and is covered with a protective cape. The background shows a bright interior with a window, a black hat hanging on the wall, and various hairdressing tools and products on shelves.

AI-generated content: quality may vary. Check for accuracy. [Disclaimer](#)



him by giving away school supplies, backpacks and free back-to-school haircuts.

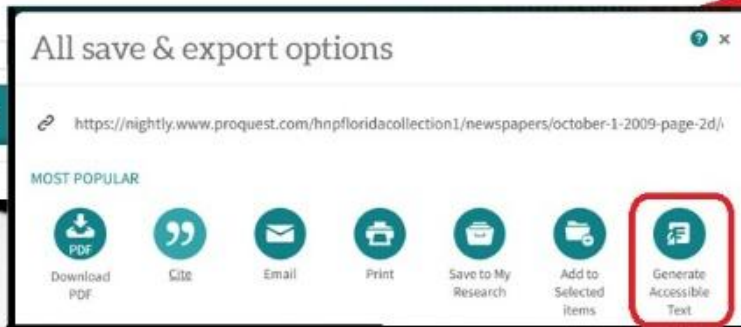
"I like  Analyzing image as a
kid, w
anyth
where
said.
appe
lies

AI圖片解讀功能
(Generate Image Description)

陸續提供針對圖片的AI解讀功能，
根據圖片可視內容做客觀的說明。

AI圖像化文本轉文字功能 (Generate Accessible Text)

運用光學字元辨識技術，將圖片、頁面級新聞內容以及非文本格式的 PDF 檔轉換為純 HTML 文本。



Browse this issue

2D / 55 Go

2D • THURSDAY, OCTOBER 1, 2009 • USA TODAY

Snapshot of stay-at-home moms

Census data may be out of date

By Sharon Jayson
USA TODAY

Nearly one-quarter of all married-couple families in the USA

A look at married moms in 2007

	Stay-at-home mothers	Other mothers
Under age 35	44%	38%

43% of other mothers.
► Stay-at-home mothers had less education: 19% had less than a high school degree, vs. 8% of other mothers; 32% of stay-at-home moms had at least a bachelor's degree, compared with 38% of the others.
Because the data were collect-

Prescription painkillers are easy to get, misuse

Cover story

The text-only content that appears below is generated automatically on request from scanned page images, using optical character recognition (OCR) software. The level of accuracy of the output may vary as it depends on a number of factors including the formatting of the original document, and the quality of the page images.

2D • THURSDAY, OCTOBER 1, 2009 • USA TODAY Snapshot of stay-at-home moms Census data may be out of date By Sharon Jayson USA TODAY Nearly one-quarter of all married-couple families in the USA had a stay-at-home mom in 2007, according to a U.S. Census report released today. This analysis, the first by the Census to offer characteristics of stay-at-home mothers, finds that the 5.6 million women who said they stayed home to care for children and family while their husband worked full time were younger and more likely to be Hispanic and foreign-born than other mothers in 2007. The Census defines "stay-at-home mother" as one with a child under 15 who stays home to care for children while her spouse was in the labor force all 52 weeks the previous year. The Census offered the following reasons for staying home: ill or disabled; retired; taking care of home and family; going to school; couldn't find work; other. Family A look at married moms in 2007 motherst y Stay-at-home - onle I I lti t• I .11,. 3:, Hispanit 27% 16% Foreign-born Other

跨學科的 Research Assistant

適應使用者需求範例



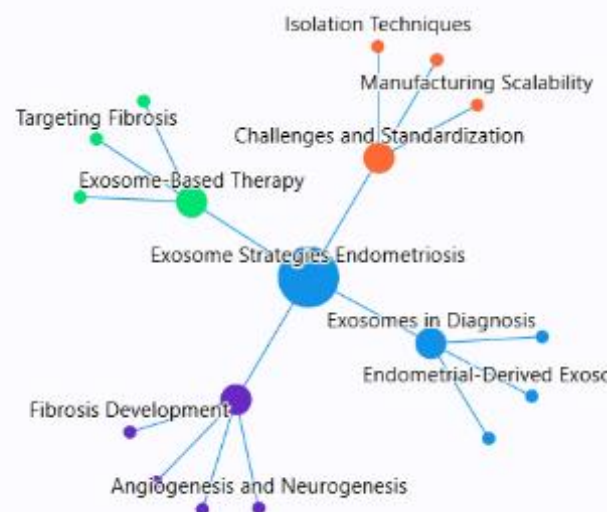
Bridging bench to bedside: Exosome-based strategies for endometriosis diagnosis and treatment (Review)

Chen W, Ma Y, Zhang X, et al. Bridging bench to bedside: Exosome-based strategies for endometriosis diagnosis and treatment (review). *Int J Mol Med*. 2026;57(2).

Here are the important concepts.

- 1. **Endometriosis:** Endometriosis is a chronic inflammatory gynecological disorder characterized by the growth of endometrial-like tissue outside the uterine cavity, significantly impacting women's quality of life and fertility. Its complex pathogenesis and association with exosomes highlight its relevance in both diagnosis and potential therapeutic strategies discussed in the document.
- 2. **Exosomes:** Exosomes are small extracellular vesicles that play a crucial role in intercellular communication and are implicated in various biological processes, including those related to endometriosis. Their potential as diagnostic biomarkers and therapeutic delivery vehicles makes them a central focus of the document.
- 3. **miRNAs:** MicroRNAs (miRNAs) are small non-coding RNA molecules that regulate gene expression and are carried by exosomes, playing significant roles in the pathophysiology of endometriosis. The document emphasizes their diagnostic potential and involvement in disease progression through various mechanisms.
- 4. **Fibrosis:** Fibrosis is a pathological condition characterized by excessive connective tissue formation, which is a prominent feature of endometriosis and contributes to pain and infertility. The document discusses the role of exosomal miRNAs in regulating fibrosis associated with endometriosis.
- 5. **Angiogenesis:** Angiogenesis is the process of new blood vessel formation, which is crucial for the establishment and progression of endometriosis. The document highlights how exosomes facilitate angiogenesis in ectopic lesions, linking them to disease progression.
- 6. **Immune Dysregulation:** Immune dysregulation refers to the altered immune response observed in endometriosis, where ectopic lesions evade immune surveillance. The document discusses how exosomes contribute to this dysregulation, promoting disease progression and influencing treatment

Exosome Strategies Endometriosis



醫學研究經常包含新穎且複雜的概念。用戶現在能在閱讀文本時，輕鬆理解這些概念及其相關性。

Here are related research topics.

- 1. Exosome-based diagnostics for endometriosis: Developing non-invasive diagnostic methods using exosomes can significantly improve early detection and management of endometriosis, ultimately enhancing patient outcomes.
- 2. Therapeutic applications of exosomes in endometriosis: Exploring exosomes as therapeutic carriers can lead to innovative treatment strategies that address the limitations of current therapies for endometriosis.
- 3. Role of exosomes in immune modulation in endometriosis: Understanding how exosomes influence immune responses can provide insights into the mechanisms of immune evasion in endometriosis, paving the way for targeted therapies.
- 4. Exosomal mirnas as biomarkers in endometriosis: Identifying specific miRNAs in exosomes can lead to the development of reliable biomarkers for diagnosing and monitoring endometriosis.
- 5. Exosomes and endometriosis-associated fibrosis: Investigating the role of exosomes in fibrosis can reveal new therapeutic targets for managing the fibrotic aspects of endometriosis.

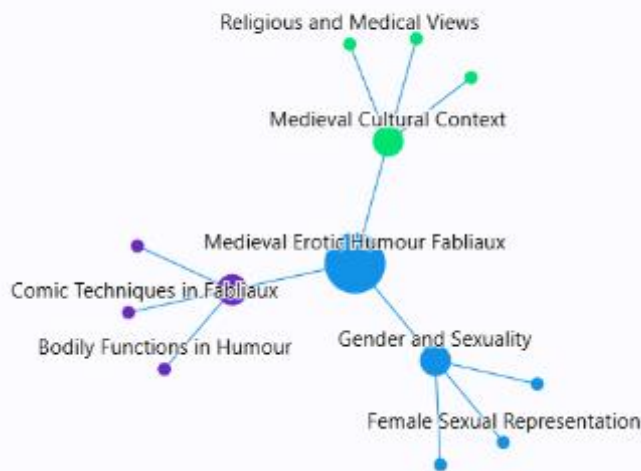
Toward a Theoretical and Analytical Framework for the Study of Sexual Humour in Chaucerian Fabliaux

Shoqairat W, Kraishan M. Toward a theoretical and analytical framework for the study of sexual humour in chaucerian fabliaux. *Theory and Practice in Language Studies*. 2023;13(7):1659-1667.

Here are the important concepts.

1. **Chaucer**: Geoffrey Chaucer is a pivotal figure in medieval literature, known for his work 'The Canterbury Tales,' which includes various fabliaux that explore themes of humor, sexuality, and social dynamics. His narratives often challenge societal norms and provide insight into the complexities of gender and class during the Middle Ages.
2. **fabliau**: Fabliaux are a genre of comic, often bawdy tales from medieval literature that typically feature clever characters and humorous situations, reflecting societal norms and human behavior. They serve as a critical lens through which the document examines the interplay of humor and eroticism in Chaucer's works.
3. **Miller's Tale**: The 'Miller's Tale' is one of the stories in 'The Canterbury Tales' that exemplifies the use of humor and eroticism, showcasing the dynamics of love, infidelity, and social class. It highlights the comic potential of sexual encounters and the subversion of traditional gender roles.
4. **Merchant's Tale**: The 'Merchant's Tale' is another narrative from 'The Canterbury Tales' that explores themes of desire, deception, and the comedic aspects of marriage. It illustrates the complexities of male and female relationships, particularly through the character of May and her interactions with January and Damyan.
5. **Reeve's Tale**: The 'Reeve's Tale' is a fabliau that combines elements of humor and violence, showcasing the absurdities of social interactions and the consequences of lust. It serves as a commentary on the nature of revenge and the comedic potential of sexual misadventures.
6. **Alison**: Alison is a central character in both the 'Miller's Tale' and the 'Merchant's Tale,' representing female agency and sexual desire in a patriarchal society. Her actions and humor challenge traditional gender roles, making her a significant figure in the exploration of eroticism in medieval literature.

Medieval Erotic Humour Fabliaux



用戶可了解作者及其作品、文學主題、文學手法及其他重要背景，以拓展文章以外的知識。

Here are related research topics.

1. Women's participation in medieval fabliaux: Exploring women's roles in medieval fabliaux reveals their agency and resistance against patriarchal norms, contributing to gender studies and literary analysis.
2. Erotic humor in Chaucer's works: Investigating the use of erotic humor in Chaucer's fabliaux can enhance our understanding of medieval attitudes towards sexuality and humor.
3. Cultural functions of fabliaux: Examining the cultural functions of fabliaux can shed light on their role in medieval society, particularly in terms of class and gender interactions.
4. Humor as a form of social commentary: Analyzing humor in fabliaux as a social commentary can reveal how medieval authors critiqued societal norms and behaviors.
5. Gender dynamics in medieval humor: Investigating gender dynamics in medieval humor can illuminate how humor was used to navigate and challenge gender roles.

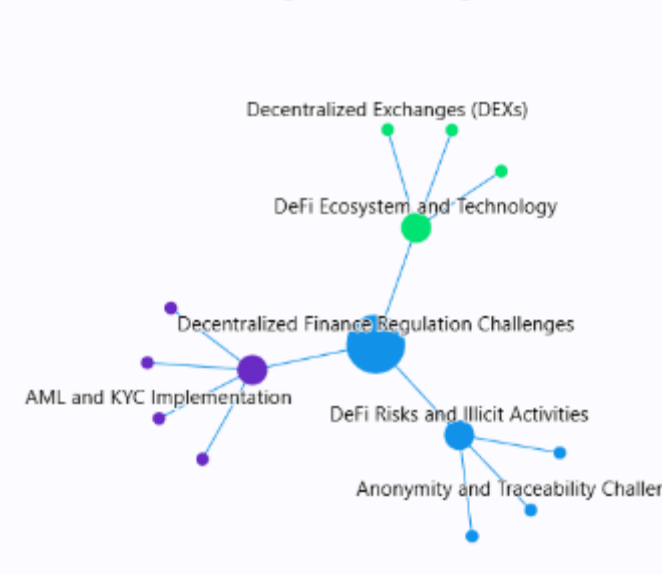
Dark side of decentralized finance: a call for enhanced AML regulation based on use cases of illicit activities

Benson V, Turksen U, Adamyk B. Dark side of decentralised finance: A call for enhanced AML regulation based on use cases of illicit activities. *Journal of Financial Regulation and Compliance*. 2024;32(1):80-97.

Here are the important concepts.

- 1. **DeFi**: Decentralized Finance (DeFi) refers to financial applications and services built on blockchain technology that operate without intermediaries, promoting financial inclusion and transparency. Its significance in the document lies in its potential to revolutionize financial services while also presenting regulatory challenges and risks associated with illicit activities.
- 2. **Regulators**: Regulators are governmental bodies responsible for overseeing financial markets and ensuring compliance with laws and regulations. Their significance in the document is highlighted by their crucial role in developing a regulatory framework for DeFi to protect consumers and mitigate financial crimes.
- 3. **Smart Contracts**: Smart contracts are self-executing contracts with the terms of the agreement directly written into code, enabling automated transactions on blockchain networks. They are significant in the document as they form the backbone of DeFi applications, facilitating trustless transactions but also posing risks for exploitation.
- 4. **AML (Anti-Money Laundering)**: Anti-Money Laundering (AML) refers to laws and regulations aimed at preventing the illegal generation of income through financial systems. Its significance in the document is underscored by the need for AML measures in DeFi to combat the risks of money laundering and other illicit activities.
- 5. **KYC (Know Your Customer)**: Know Your Customer (KYC) is a process used by financial institutions to verify the identity of their clients to prevent fraud and money laundering. The document emphasizes the importance of KYC in DeFi to enhance consumer protection and regulatory compliance.
- 6. **DEXs (Decentralized Exchanges)**: Decentralized Exchanges (DEXs) are platforms that allow users to trade cryptocurrencies directly without a central authority. Their significance in the document is noted in the context of

Decentralized Finance Regulation Challenges



在閱讀報告、商業期刊、財務報表等文件時，系統會解釋相關法規、法律與商業概念。這將提升使用者在整個研究過程中的價值與知識收穫。

Here are related research topics.

- 1. Defi regulation challenges: Understanding the complexities of regulating decentralized finance is crucial for ensuring consumer protection and financial stability.
- 2. Anti-money laundering in defi: Exploring AML strategies in DeFi is essential to combat illicit activities and enhance the integrity of financial systems.
- 3. Cybersecurity risks in defi: Investigating cybersecurity risks in DeFi is vital for protecting users from hacking and fraud in a decentralized environment.
- 4. Consumer protection in defi: Addressing consumer protection issues in DeFi is necessary to build trust and ensure safe participation in decentralized finance.
- 5. Decentralized finance and financial inclusion: Examining how DeFi can enhance financial inclusion is important for understanding its potential to democratize access to financial services.

Research Assistant 的使用情況

全年使用情況：

- 4,000 個活躍帳戶
- 470萬位獨立使用者至少嘗試過文件頁面（文件檢視）上的一項 AI 功能
- 產生的工作階段互動次數近 200 萬次
- 9 月，PQRA 達成新里程碑：每週使用者超過 3 萬人

About the data

Analysis based on 30-day period (20 August to 18 September), comparing usage from PQRA vs the non-PQRA version of the doc view.

重要統計數據

- 在文件檢視頁面看到 PQRA 的使用者，**開始滾動瀏覽的可能性高出 31%**。
- 當 PQRA 可用時，使用者對文件採取有意義操作（引用、下載、列印、儲存或寄送電子郵件）的**可能性高出 40%**。
- 點擊核心要點「Show more」的使用者，**採取有意義操作的可能性高出 76%**。
- 與 PQRA 互動的使用者，一個月內多次返回平台的可能性高出70%。
- 來自 PQRA 的回饋絕大多數是正面的，顯示**滿意的反饋率高達85%**。

常見問題 (FAQ)

客戶如何在 ProQuest Central Premium 中啟用研究助手？

答：在 ProQuest Central Premium 中會自動啟用。如果客戶希望關閉此功能，只需向其服務台提交申請即可開始使用。

是否可以只開啟研究助手的特定功能，而不開啟其他功能？

答：目前不行。如果獨立功能啟用的支持是客戶感興趣的，我們未來可能會考慮實現。

PQRA 會產生幻覺 (Hallucinations) 嗎？

答：不會。大語言模型 (LLMs) 完全獨立於外部訓練數據，因此不太可能產生幻覺。為了產生生成式文本，PQRA 會參考文章的全文內容以獲取研究發現、結論以及生成式文本與所引用文章之間的關係，因此其輸出的準確性和相關性極高。

PQRA 會儲存任何使用者的個人資訊嗎？

答：不會。PQRA 由本地託管的大語言模型提供支援，該模型在工作階段結束後不會處理任何使用者輸入。由於不儲存任何個人資訊，因此不會分享或儲存任何資訊。

使用 PQRA 需要付費嗎？

答：不需要。一旦啟用，這些功能是免費的，並作為 ProQuest 持續平台更新的一部分。

進一步了解 ProQuest Research Assistant 及其他科睿唯安 AI 計畫！

About.proquest.com

- 訪問網站 [Clarivate.com](https://clarivate.com) 深入了解 ProQuest Research Assistant 以及 AI 在科睿唯安各項產品中的應用方式。
- 在我們新發布的常見問題集（FAQ）中，找到您最關心問題的解答。
- 更多產品培訓資訊，請訪問科睿唯安教育訓練資源服務頁面或ProQuest LibGuide
- 您也可關注科睿唯安官方YouTube頻道，查看教育訓練影音。



Think forward™

關於科睿唯安 (Clarivate)

科睿唯安是全球領先的變革性資訊提供者。我們在學術與政府、智慧財產權、生命科學與醫療照護領域提供豐富資料、洞察、分析及工作流程解決方案。如需了解更多資訊請參考 <https://clarivate.com/academiagovernment/zh/>。

© 2026 Clarivate. All rights reserved

科睿唯安台灣辦公室
台北市信義區松仁路 100 號 34 樓