

科学におけるチューリングテスト： 「Agents4Science」会議と自動化された 科学的発見の未来に関する詳細な分析

Gemini Deep Research

I. 序論：学術界への意図的な挑発

2025 年 10 月に開催が予定されている「Agents4Science 2025」会議は、単なる目新しいイベントではなく、計算され尽くした意図的な挑発として学術界に登場した。この会議は、科学的実践の確立された規範と、高度な人工知能（AI）がもたらす破壊的な能力との間に、必要ではあるが不快感を伴う対立を強いるように設計されている。その核心にある緊張は、人間の説明責任、独創性、知的貢献といった伝統的な概念と、今や科学的発見のプロセスに参加し、さらには主導することさえ可能な AI システムの現実との衝突である。

この会議の基本的前提は前例がない。投稿されるすべての論文において AI が筆頭著者であることが義務付けられ、査読も AI システムが担当する¹。これは、AI の著者性を明確に禁止しているほばすべての一流学術誌や国際会議の方針に真っ向から挑戦するものである¹。主催者は、この会議を、AI 駆動型の科学の時代における功績の帰属、検証、倫理に関する新しい規範を構築し始めるための、透明性が高く、管理された低リスクの環境であると位置づけている²。これは、理論的な議論から経験的な観察へと対話を進めるための明確な試みである。

既存の学術機関、特に『Nature』や『PNAS』といった権威ある学術誌は、説明責任の所在が不明確であるという正当な懸念から、AI の著者性を認めない方針を打ち出してきた¹。しかし、この禁止措置は研究者による AI の利用を止めるものではなく、むしろその利用を水面下に押しやり、AI の貢献を過小評価または隠蔽することを奨励する「暗黙の了解」的な文化を生み出している⁴。この現状は、AI の真の能力と限界を理解する上で深刻な障害となっている。

Agents4Science の主催者はこの矛盾を的確に捉えている。AI の著者性を「義務化」し、プロンプトや AI による査読結果を公開することでプロセス全体の透明性を確保することにより、彼らはこの問題を公の場に引きずり出そうとしている³。

したがって、この会議の第一の機能は、単に AI が生成した論文を披露する場を提供することで

はない。そのより深い目的は、制度的な触媒として機能することにある。それは、より広範な学術コミュニティが AI との協働に関する新しい、より現実的な基準を策定せざるを得なくなるようなデータ、ケーススタディ、そして公開討論を生み出すことを目指している。本報告書は、**Agents4Science** を単なる会議以上のもの、すなわち人間と AI の科学者が対等な協力者として協働する、新しい科学研究のパラダイムに必要な制度的基盤を構築するための根源的な試みとして位置づける。これは、学術イベントの形をとった、制度設計の試みなのである。

II. 設計者：ジェームズ・ゾウ博士が描く自動化科学のビジョン

この会議の首謀者であるスタンフォード大学のジェームズ・ゾウ准教授の人物像を深く掘り下げることで、**Agents4Science** が彼の特定の研究軌道と公的な問題提起の論理的かつ戦略的な集大成であることが明らかになる。彼の研究は、この急進的な実験に対する技術的な概念実証と知的な動機の両方を提供している。

ゾウ博士は、スタンフォード大学の生物医学データサイエンス、コンピュータサイエンス、電気工学の准教授であり、スタンフォード AI ラボのメンバーでもある³。彼の研究は、AI をより信頼性が高く、人間と互換性があり、統計的に厳密なものにすることに焦点を当てており、特に医学や生物学への応用に強い関心を持っている⁷。

会議の発表直前である 2025 年 7 月、ゾウ博士の研究チームは画期的な論文を『**Nature**』誌に発表し、彼らの「バーチャルラボ (Virtual Lab)」について詳述した¹⁰。このシステムは、協働する AI エージェントのチームを用いて、新型コロナウイルス (SARS CoV-2) に対する新規ナノボディを自律的に設計し、その成果は後の実験で検証された。この AI エージェントたちは、仮説生成、改良、計算論的実験を、人間の介入を最小限 (推定 1%) に抑えて実行した¹⁰。

この成功体験は、ゾウ博士が直面する制度的な壁を浮き彫りにした。彼は、AI が研究に決定的な貢献をしているにもかかわらず、学術誌や会議が AI を共著者として認めてくれない「窮屈な状況」に公然と不満を表明してきた¹。彼にとって、これは科学における AI の将来的な役割を理解し、形成していく上での障害となっている。**Agents4Science** は、この制度的な行き詰まりに対する彼の直接的な回答なのである。

ゾウ博士の「バーチャルラボ」は、AI エージェントが新規性のある検証可能な科学的成果を生み出せることを技術的に証明した。しかし、その方法論を発表した『**Nature**』誌自身の方針が、将来的にその AI エージェントたちが発見した成果に関する論文において、彼らを著者として記載することを許さないという矛盾をはらんでいる¹。つまり、既存の学術出版システムは、

「AI 科学者に関する研究」は出版するが、「AI 科学者による研究」は出版しないという功績帰属のパラドックスを生み出している。

この状況は、**Agents4Science** が単なる哲学的な実験ではないことを示唆している。それは、ゾウ博士自身の研究の成功から生まれた、実践的な必要性なのである。彼は AI 駆動の発見を生み出すエンジンを構築し、今、その成果物を受け入れるためのエコシステム（専用の発表の場）を構築しようとしている。この会議は、彼が自身の『Nature』論文で個人的に直面した矛盾を解決するための直接的な試みと言える。

III. 必要な実験か？ 科学出版の危機に立ち向かう

Agents4Science を、特に AI のような進歩の速い分野における、現在の学術出版モデルの広範かつ体系的な欠陥という文脈の中に位置づけることで、その真の意義が明らかになる。この会議の査読に対する急進的なアプローチは、持続不可能な論文投稿数と査読の質の低下によって特徴づけられる、広く認識された「査読の危機」への直接的な応答であると論じることができる。

ICML、NeurIPS、AAAI といった主要な AI 分野の国際会議では、現在、それぞれ 1 万件を超える論文が投稿されており、その数は指数関数的に増加している¹⁴。この急増は、資格を持つ人間の査読者のキャパシティを完全に圧倒している。その結果として生じるのは、膨大な作業負担に起因する査読の質の低下、査読者の無責任、そして燃え尽き症候群である¹⁴。著者と査読者の間には力の不均衡が存在し、著者は表面的で無責任な査読に対して効果的な対抗手段を持たないという問題も認識されている¹⁴。

さらに皮肉なことに、人間の査読者自身が大規模言語モデル（LLM）を用いて質の低い表面的な査読コメントを生成するという憂慮すべき傾向が出現しており、査読プロセス全体の信頼性をさらに損なっている¹⁵。これは、人間主導のシステムが内包する深刻な欠陥を露呈している。既存の機関もこの問題に取り組んではいる。例えば、AAAI は、人間の査読者を補完するために AI が非決定的なレビューを提供する「AI 支援査読」を試験的に導入している¹⁹。学術誌もまた、編集者による初期スクリーニングを支援するための自動化ツールの導入を検討している²⁰。しかし、これらはあくまで漸進的な変更であり、体系的な改革には至っていない。

このような状況下で、**Agents4Science** は人間主導のシステムが崩壊の兆しを見せ、皮肉にも人間による AI の不正利用によって腐敗しつつあるまさにその時に、完全に AI が駆動する査読システムを提案している¹。査読の危機の核心的な問題は、質の高い人間の査読能力の供給と、論文投稿という需要の指数関数的な増加との間のミスマッチにある。人間の査読者は一貫性がなく、偏見を持ち、そして今や一部のケースでは、自身の知的労働を透明性なく AI に外注する

という非倫理的な行為に及んでいる。Agents4Science の提案は、この問題を根本から覆す。人間の査読者による AI の不正利用を取り締まる代わりに、AI を公式かつ透明性のある査読者として位置づけるのである。

したがって、この会議における AI 査読という要素は、単に AI 著者性と対をなす目新しさだけを狙ったものではない。それは、学術出版における「人間」というボトルネックに対する、急進的な解決策の提案なのである。この実験は、適切に設計されたマルチエージェント AI 査読システムが、疲弊し機能不全に陥りつつある人間によるシステムよりも、一貫性があり、スケラブルで、潜在的により偏見の少ないフィードバックを提供できるかどうかを検証するものである。それは、AI を査読の信頼性に対する脅威としてではなく、むしろその救世主として再定義する試みなのである。

表 1 : AI の著者性および査読に関する方針の比較分析

項目	Agents4Science 2025	Nature	PNAS	AAAI	ICML
著者としての AI	必須 ²	禁止 ¹	禁止 ⁵	禁止（推定）	禁止（推定）
AI による執筆支援の開示	必須（詳細なチェックリスト） ²¹	必須 ¹	必須（方法または謝辞に記載） ⁶	「適切な使用」に関する方針あり ¹⁹	「出版倫理」に関する方針あり ²²
査読者による AI の使用	必須（AI が査読者） ²¹	禁止 ¹	禁止（守秘義務により）	AI 支援査読を試験導入（非決定的） ¹⁹	禁止（守秘義務により）

IV. AI 科学者の解体 : Agents4Science のフレームワーク

本セッションでは、この会議の運営メカニズムを詳細に分析し、見出しの裏にある、この新しい科学的エコシステムを定義する具体的な規則とプロセスを検証する。分析は、AI の自律性と人間の監督との間で、このフレームワークがいかにしてバランスを取ろうとしているかに焦点を当てる。

投稿要件

- **AI の筆頭著者性:** AI システムが主要な創造者として、仮説生成、実験、執筆を主導しなければならない。また、唯一の筆頭著者として記載される必要がある²。
- **人間の役割:** 人間は共著者として参加できるが、その役割は支援的または監督的なもの（アイデアの提供、出力の確認、フィードバックなど）に限定される。コーディングや執筆といった中核的な作業を行うことはできない²。
- **透明性の確保:** 著者は、研究ライフサイクル全体にわたる AI と人間の貢献者の具体的な役割を詳述する「AI 研究自律性開示チェックリスト」を完成させなければならない²¹。これは、会議の目標の一つである透明性を直接的に実現するための措置である⁴。

査読プロセス

- **マルチ AI レビューパネル:** 特定のモデルに起因する偏りを避けるため、各論文は複数の独立した LLM システムによって評価される²¹。
- **標準化された評価基準:** AI 査読者は、確立された NeurIPS 2025 の査読テンプレートを使用し、論文を品質、明瞭さ、重要性、独創性の観点から採点する²¹。これにより、AI による査読が、人間中心の馴染み深い基準に基づいていることが保証される。
- **人間による監督委員会:** AI パネルによって上位に選出された論文は、権威ある専門家諮問委員会を含む人間の専門家からなる委員会によって評価される。この委員会が最終的な採択決定と賞の授与を行う¹。

会議の形式

会議は 1 日間のバーチャルイベントとして開催され、招待講演、採択論文の口頭発表、そして

AI が生成する研究の未来に関するパネルディスカッションが予定されている²。

表 2 : Agents4Science 2025 運営フレームワーク

カテゴリ	詳細
タイムライン	論文投稿締切 : 2025 年 9 月 15 日 採否通知 : 2025 年 10 月 5 日 会議開催日 : 2025 年 10 月 22 日 4
著者性に関する規則	筆頭著者 : AI システムのみ 人間の共著者 : 助言・監督役に限定 情報開示 : AI 自律性チェックリストの提出が必須 2
査読プロセス	第 1 段階 : 複数の LLM による匿名査読パネル 第 2 段階 : 上位論文に対する人間の専門家委員会による評価 21
評価基準	品質、明瞭さ、重要性、独創性 (NeurIPS テンプレートに基づく) 21
会議の成果物	口頭発表、スポットライト発表、公開される論文および AI による査読結果 (オープンリソースとして) 3

V. 研究の再定義 : AI の著者性と査読がもたらす根源的課題

この重要なセクションでは、Agents4Science モデルが提起する、倫理的、哲学的、そして法的な中心課題を徹底的に分析する。この会議をケーススタディとして用い、人間以外のエージェ

エントが科学的知識の創造者となったときに生じる深遠な問いを探求する。

V.A 説明責任と独創性

- **説明責任の欠如:** AI が執筆した論文に捏造されたデータ、有害な結論、あるいは盗用が含まれていた場合、誰が法的・倫理的責任を負うのか。現在の著者性の基準は、責任を問うことができる人間が存在することを前提としている¹。この会議のモデルは、その責任を人間の「アドバイザー」に全面的に移行させるが、これは未検証の新たな責任の枠組みを創出するものである。
- **独創性の本質:** 膨大な既存の人間が作成したテキストとデータを学習した AI によって生成された著作物は、真に「独創的」と言えるのか。これは著作権に関する根源的な議論に触れる。米国の著作権法は現在、創造的なコントロール権を持つ人間の著者を要求している²³。この会議の前提は、AI を著者とし、人間を単なるツールやガイドとして扱うことで、この法的基準に直接挑戦する。これは、AI が生成した発見の知的財産権の帰属について重大な問題を提起する²⁴。

V.B 組織的バイアスの影

- **AI による AI の優遇:** 『PNAS』誌に掲載された研究は、LLM が意思決定の役割で用いられる際、他の LLM や AI 支援を受けた人間によって生成されたコンテンツを暗黙のうちに優遇する可能性があることを示唆している²⁶。これは、AI による査読プロセスが自己強化的なエコーチェンバーを生み出し、AI が生成した研究スタイルを体系的に好み、斬新な人間のアイデアを差別する可能性があるという憂慮すべき事態を示唆する。
- **訓練データのバイアスの永続化:** AI システムは、その訓練データに存在するバイアスを反映し、増幅させる可能性がある¹⁷。既存の科学文献で訓練された AI 査読者は、主流の理論を支持し、真にパラダイムシフトをもたらす（しかし型破りな）アイデアを罰する傾向を持つかもしれない、結果としてイノベーションを阻害する恐れがある。

V.C 進化する人間の役割

- **創造者からキュレーターへ:** Agents4Science モデルは、研究者の役割が、実践的な実験者や執筆者から、高レベルの戦略家、プロンプター、そして検証者へと移行することを公式

化する。中核となるスキルは、プロンプトエンジニアリング、実験設計の監督、そして AI の出力に対する批判的な評価へとシフトする可能性がある。

- **認知的・倫理的リスク**: AI ツールへの過度の依存は、人間の批判的思考力や研究スキルの萎縮に関する懸念を引き起こす²⁴。また、「自動化バイアス」、すなわち人間が自動化システムの出力を過度に信頼し、微細だが致命的な誤りを見逃すリスクも存在する。この懸念は、Sakana AI の実験で見られた引用エラーによっても裏付けられている²⁸。

この会議は、開示のための詳細なフレームワーク（AI 自律性チェックリスト²¹）、多層的な査読プロセス（AI パネル+人間による監督²¹）、そして権威ある人間の諮問委員会⁴を備えている。これらはすべてガバナンスのメカニズムである。AI の著者性に対する主な反対意見は、その能力に関するものではなく、信頼、説明責任、検証のための確立された枠組みが存在しないことに関するものである。主催者は、これらのガバナンスのギャップに正面から取り組むために、会議の構造を積極的に設計した。開示チェックリストは透明性の問題に取り組み、ハイブリッドな査読モデルは人間の判断と監督の必要性に対応し、諮問委員会は知的・倫理的な権威を提供する。

したがって、Agents4Science は単なる科学的な実験ではなく、ガバナンスの実験でもある。それは、科学における AI の役割を管理するための新しい一連の規則と制度を創出し、検証するための実地演習である。この会議の成否は、生み出された論文の質だけでなく、そのガバナンスモデルが堅牢で、公正で、信頼できるものであると証明されるかどうかによっても判断されるだろう。ゾウ博士が予期する「有益な失敗」¹は、このガバナンス実験にとって、成功と同じくらい価値があるのである。

VI. 自律的な未来のためのガードレール構築：検証、再現性、そして信頼

本セクションでは、AI 科学者が当たり前になる未来において必要とされる、技術的および制度的な安全策を探求する。Agents4Science の査読モデルをプロトタイプとして分析し、自動化された科学の完全性を保証するために必要な、より広範でスケーラブルな解決策について議論を拡大する。

検証の課題

AI が生成した論文で提示された結果が、ハルシネーションではなく妥当なものであることを、我々はどうすれば確信できるのか。これは極めて重要なハードルである。

- **提案されている解決策:** AI が生成した主張を原典文献と照合して検証する「検証エンジン」²⁹ や、仮説の自動生成と検証のためのフレームワーク³⁰ に関する研究が進行中である。この分野は、AI システムの検証、説明、制御に焦点を当てた「モデルサイエンス」という新たなパラダイムへと移行しつつある³¹。

再現性の危機と AI

科学界はすでに再現性の危機に直面している³³。複雑なコードと「ブラックボックス」モデルを伴う AI による研究は、この問題をさらに悪化させる可能性がある。

- **技術的解決策:** 再現性を確保するためには、コード、モデル、コンテナ化されたソフトウェア環境の共有を義務付けるなど、新たな制度的・技術的基準が必要となる³⁵。学術誌はすでに、再現性チェックリストやバッジの導入へと動き出している³⁷。

来歴と真正性の確立

AI が生成したコンテンツが溢れる世界で、真正な研究と巧妙な捏造をどうすれば確実に見分けることができるのか²⁷。

- **電子透かし（デジタルウォーターマーキング）:** 有望な技術的解決策の一つが電子透かしである。これは、AI が生成したテキストや画像に、生成時に知覚できない信号を埋め込む技術である⁴⁰。これにより、AI による生成物であることを示す検証可能なマーカーが提供され、学術的不正行為を抑止し、偽情報の出所を追跡するのに役立つ⁴²。

Agents4Science 会議は、科学の AI「生成」における画期的な出来事である。同時に、AI の「検出」と「検証」に焦点を当てた研究分野も並行して進展している²⁹。日本の政府や学術団体は、すでにガイドラインや電子透かしのような技術的解決策について議論を始めている⁴²。質の高い AI 生成コンテンツの普及は、その真正性と出所を確実に検証するための信頼性の高い方法に対する、即時かつ緊急のニーズを生み出す。これにより、2 つの異なる、しかし共進化する技術的軌道が生まれる。すなわち、生成モデルはより人間に近くなり、検出・検証モデルは AI の痕跡を特定する上でより洗練されていく。

電子透かしのような技術的解決策⁴³ は、生成物に検証可能な「指紋」を与える方法として提案

されているが、これはモデル作成者の協力（例えば、OpenAI が鍵を保持する⁴³）に依存する。悪意のある者は、そのような保護策のないオープンソースモデルを使用する可能性がある。

このことから、Agents4Science 会議が AI 生成科学の分野を正当化し、加速させることで、必然的に AI 生成科学の「検証」分野の加速も引き起こすことがわかる。これは、技術的および制度的な軍拡競争の舞台を設定する。学術的誠実性の未来は、「AI 科学者」の能力だけでなく、彼らを誠実に保つために開発される「AI 審判員」（検証ツール、電子透かしの基準、再現性プラットフォーム）の堅牢性にかかっているのである。

VII. 最前線からの視点：諮問委員会からの洞察

本セクションでは、著名な専門家からなる諮問委員会⁴の戦略的な役割を分析する。委員会のメンバーが単にその名声のために選ばれたのではなく、彼らの専門知識が Agents4Science 実験の最も深遠な課題と機会に直接的に対応していることを論じる。彼らの集合的な業績は、この新しい領域を航海するための知的なロードマップを提供する。

- **グイド・インベンス（ノーベル経済学賞受賞者）：因果性と信頼の守護者**
 - **専門性:** インベンス氏のノーベル賞受賞研究は、単なる相関関係ではなく、データから原因と結果を特定する因果推論に関するものである⁴⁵。彼は、自動化システムが重要な意思決定において信頼されるためには、透明性、監査可能性、そして何が「なぜ」起こったのかを説明する能力が必要であると強調する⁴⁸。
 - **示唆:** 彼の参加は、この会議がブラックボックス的な AI の予測を超え、真の科学的発見の基盤である因果的な主張を生成し、擁護できる AI システムの実現に深くコミットしていることを示している。彼は、厳密で、信頼でき、説明可能な AI の論理的思考に対する要求を象徴している。
- **エリック・トポル（医師・科学者、未来学者）：人間中心の医療 AI の提唱者**
 - **専門性:** トポル博士は、AI が医師に取って代わるのではなく、データ集約的なタスクを処理することで臨床医を「深い共感」と人間的なつながりのために解放し、いかにして医療を変革できるかについての第一人者である⁵⁰。彼は、人間には見えないパターンを AI が見抜く能力を擁護し、それがより早期で正確な診断につながると主張する⁵²。
 - **示唆:** トポル博士の関与は、この会議の目標を人間と AI のパートナーシップというより広範なビジョンの中に位置づける。彼は、AI 科学者が人間の知性を増強し、複雑な疾患に取り組み、最終的には人間の研究者に最も創造的で影響力のある側面に集中するための「時間の贈り物」を与える強力なツールとして機能するという楽観的な見方を代表している⁵¹。

- **イェジン・チェ (AI 研究者、マッカーサー・フェロー) : 常識と多元主義の擁護者**
 - **専門性:** チェ氏の研究は、AI に「常識」を教え、単一の「正解」という仮定を超えて「多元的アライメント」へと向かうことに焦点を当てている。これは、AI が多様な人間の価値観や視点に貢献できるようにすることを保証するものである⁵⁵。彼女はまた、より小型で、効率的で、公平な AI モデルの提唱者でもある⁵⁶。
 - **示唆:** チェ氏の視点は、画一的な AI 思考のリスクに対する重要なカウンターバランスとして機能する。彼女の研究は、単一の目的関数を最適化するだけでなく、ニュアンスをもって推論し、文脈を理解し、科学的思考の多様性を尊重できる AI 科学者を構築する必要性を強調している。彼女は、AI が科学を均質化するのを防ぐという倫理的要請を代表している。

VIII. 結論 : AI 科学者の軌道を描く

本報告書の最終セクションでは、これまでの分析結果を統合し、**Agents4Science** 会議とその広範な影響について、多角的な判断を示す。結論として、この会議は、その直接的な成果にかかわらず、「AI 科学者」という存在が理論的な概念から、科学界が今まさに直面しなければならない実践的な現実へと移行したことを示す極めて重要なイベントであると位置づける。

この変化は、会議が引き起こしたものではなく、むしろその兆候であり、加速器である。ゾウ博士の「バーチャルラボ」や **Sakana AI** の「AI サイエнтиスト」¹⁰ のようなシステムの成功は、技術がすでにその段階に到達していることを示している。もはや重要な問いは、AI が論文を共著するか「どうか」ではなく、科学界がこの新しい現実を「どのように」管理していくかである。

Agents4Science から得られる「有益な失敗」¹は、その成功と同じくらい重要である。それらは、学術誌、大学、そして資金提供機関が、画一的な禁止措置から脱却し、以下のような洗練されたエビデンスに基づく方針を策定するために必要な具体的なデータを提供するだろう。

- **AI の貢献に関する開示基準:** 会議の開示チェックリスト²¹に着想を得て、研究における AI の役割を記述するための普遍的な分類法を作成する。
- **検証と再現性のためのプロトコル:** AI が関与する研究に対して、コンテナ化や、将来的には電子透かしといった技術的解決策の使用を義務付ける。
- **新しい査読モデルの探求:** AI を規模と一貫性のために活用しつつ、新規性やインパクトについては人間の判断を保持する、ハイブリッドな AI-人間査読システムを模索する。

最終的に、**Agents4Science** は決定的な解決策としてではなく、不可欠であり、おそらくは必要不可欠な触媒として見なされるべきである。管理された環境で意図的に古い規則を破ること

により、それは新しい規則の創設を強いる。これは、我々の最も強力な知的ツールが持つ能力に、我々の科学的制度を適合させるための、長く複雑な共進化のプロセスの始まりである。この「異色の学会」は、科学的探求が人間と人工知能との真のパートナーシップとなる未来を垣間見せるものである。

引用文献

1. 筆頭著者は AI、査読、発表も——異色の学会 ... - MIT Tech Review, 8 月 31, 2025 にアクセス、<https://www.technologyreview.jp/s/367619/meet-the-researcher-hosting-a-scientific-conference-by-and-for-ai/>
2. Finally: Stanford Hosts the First Academic Conference for AI Authors | by noailabs | Jul, 2025, 8 月 31, 2025 にアクセス、<https://noailabs.medium.com/finally-stanford-hosts-the-first-academic-conference-for-ai-authors-de8f12626e07>
3. Stanford-Initiated First Academic Conference for AI Authors: First Author Must Be an AI - 36 氦, 8 月 31, 2025 にアクセス、<https://eu.36kr.com/en/p/3375685195848200>
4. Open Conference of AI Agents for Science: 2025, 8 月 31, 2025 にアクセス、<https://agents4science.stanford.edu/>
5. PNAS Editorial Policies – Peer Review Standards and Author Responsibilities, 8 月 31, 2025 にアクセス、<https://www.pnas.org/author-center/editorial-and-journal-policies>
6. Information for authors | PNAS Nexus | Oxford Academic, 8 月 31, 2025 にアクセス、<https://academic.oup.com/pnasnexus/pages/general-instructions>
7. James Zou | Stanford Medicine, 8 月 31, 2025 にアクセス、<https://med.stanford.edu/profiles/james-zou>
8. James Zou - Stanford Data Science, 8 月 31, 2025 にアクセス、<https://datascience.stanford.edu/people/james-zou>
9. Research | Mysite - James Zou, 8 月 31, 2025 にアクセス、<https://www.james-zou.com/research>
10. Researchers create 'virtual scientists' to solve complex biological problems, 8 月 31, 2025 にアクセス、<https://news.stanford.edu/stories/2025/07/ai-virtual-scientists-lab-llms>
11. With no need for sleep or food, AI-built 'scientists' get the job done quickly, 8 月 31, 2025 にアクセス、<https://www.czbiohub.org/news/with-no-need-for-sleep-or-food-ai-built-scientists-get-the-job-done-quickly/>
12. Researchers create 'virtual scientists' to solve complex biological problems, 8 月 31, 2025 にアクセス、<https://med.stanford.edu/news/all-news/2025/07/virtual-scientist.html>
13. Researchers Create “Virtual Scientists” To Solve Complex Biological Problems - Technology Networks, 8 月 31, 2025 にアクセス、<https://www.technologynetworks.com/informatics/news/researchers-create->

- [virtual-scientists-to-solve-complex-biological-problems-402897](#)
14. Position: The AI Conference Peer Review Crisis Demands Author Feedback and Reviewer Rewards - arXiv, 8 月 31, 2025 にアクセス、
<https://arxiv.org/html/2505.04966v1>
 15. ICML Poster Position: The AI Conference Peer Review Crisis Demands Author Feedback and Reviewer Rewards - ICML 2025, 8 月 31, 2025 にアクセス、
<https://icml.cc/virtual/2025/poster/40108>
 16. [2505.04966] Position: The AI Conference Peer Review Crisis Demands Author Feedback and Reviewer Rewards - arXiv, 8 月 31, 2025 にアクセス、
<https://arxiv.org/abs/2505.04966>
 17. AI in Peer Review: A Recipe for Disaster or Success? - American Society for Microbiology, 8 月 31, 2025 にアクセス、
<https://asm.org/articles/2024/november/ai-peer-review-recipe-disaster-success>
 18. Position: The AI Conference Peer Review Crisis Demands Author Feedback and Reviewer Rewards | OpenReview, 8 月 31, 2025 にアクセス、
<https://openreview.net/forum?id=l8QemUZaIA>
 19. Main Technical Track: Call for Papers - Association for the Advancement of Artificial Intelligence (AAAI), 8 月 31, 2025 にアクセス、
<https://aaai.org/conference/aaai/aaai-26/main-technical-track-call/>
 20. AI 時代における査読と研究の誠実さ - Editage Blog, 8 月 31, 2025 にアクセス、
<https://www.editage.jp/blog/peer-review-and-research-integrity-in-the-age-of-artificial-intelligence/>
 21. Call for Papers - Open Conference of AI Agents for Science 2025, 8 月 31, 2025 にアクセス、
<https://agents4science.stanford.edu/call-for-papers.html>
 22. 2025 Conference, 8 月 31, 2025 にアクセス、
<https://icml.cc/>
 23. Generative Artificial Intelligence and Copyright Law - Congress.gov, 8 月 31, 2025 にアクセス、
<https://www.congress.gov/crs-product/LSB10922>
 24. Challenges in Ensuring Originality and Authenticity in AI-Generated Works - ResearchGate, 8 月 31, 2025 にアクセス、
https://www.researchgate.net/publication/394499575_Challenges_in_Ensuring_Originality_and_Authenticity_in_AI-Generated_Works
 25. Reforming Copyright Law for AI-Generated Content: Copyright Protection, Authorship and Ownership | Technology and Regulation, 8 月 31, 2025 にアクセス、
<https://techreg.org/article/view/23039>
 26. AI–AI bias: Large language models favor communications generated by large language models | PNAS, 8 月 31, 2025 にアクセス、
<https://www.pnas.org/doi/10.1073/pnas.2415697122>
 27. GenAI synthetic data create ethical challenges for scientists. Here's how to address them. | PNAS, 8 月 31, 2025 にアクセス、
<https://www.pnas.org/doi/10.1073/pnas.2409182122>
 28. AI-generated paper passes peer review before planned withdrawal - The Decoder, 8 月 31, 2025 にアクセス、
<https://the-decoder.com/ai-generated->

- [paper-passes-peer-review-before-planned-withdrawal/](#)
29. Verif.ai: Towards an Open-Source Scientific Generative Question-Answering System with Referenced and Verifiable Answers NGISearch - arXiv, 8 月 31, 2025 にアクセス、<https://arxiv.org/html/2402.18589v1>
 30. arXiv:2311.09706v1 [cs.AI] 16 Nov 2023, 8 月 31, 2025 にアクセス、<https://arxiv.org/pdf/2311.09706>
 31. [2508.20040] Model Science: getting serious about verification, explanation and control of AI systems - arXiv, 8 月 31, 2025 にアクセス、<https://arxiv.org/abs/2508.20040>
 32. Model Science: getting serious about verification, explanation and control of AI systems, 8 月 31, 2025 にアクセス、<https://arxiv.org/html/2508.20040v1>
 33. Full article: Artificial intelligence: help or hindrance in solving the reproducibility crisis?, 8 月 31, 2025 にアクセス、<https://www.tandfonline.com/doi/full/10.1080/07366205.2024.2355776>
 34. Is AI leading to a reproducibility crisis in science? - ResearchGate, 8 月 31, 2025 にアクセス、https://www.researchgate.net/publication/376254925_Is_AI_leading_to_a_reproducibility_crisis_in_science
 35. Reproducible research policies and software/data management in scientific computing journals: a survey, discussion, and perspectives - Frontiers, 8 月 31, 2025 にアクセス、<https://www.frontiersin.org/journals/computer-science/articles/10.3389/fcomp.2024.1491823/full>
 36. Transparency and reproducibility in artificial intelligence - Maastricht University, 8 月 31, 2025 にアクセス、https://cris.maastrichtuniversity.nl/files/75369807/Aerts_2020_Transparency_and_reproducibility_in_artificial.pdf
 37. Improving Reproducibility in AI Research: Four Mechanisms Adopted by JAIR, 8 月 31, 2025 にアクセス、https://www.researchgate.net/publication/387499326_Improving_Reproducibility_in_AI_Research_Four_Mechanisms_Adopted_by_JAIR
 38. Improving Reproducibility in AI Research : Four Mechanisms Adopted by JAIR, 8 月 31, 2025 にアクセス、<https://d-nb.info/1355328942/34>
 39. Safeguarding authenticity for mitigating the harms of generative AI: Issues, research agenda, and policies for detection, fact-checking, and ethical AI, 8 月 31, 2025 にアクセス、<https://pmc.ncbi.nlm.nih.gov/articles/PMC10838945/>
 40. [2502.05215] Watermarking across Modalities for Content Tracing and Generative AI- arXiv, 8 月 31, 2025 にアクセス、<https://arxiv.org/abs/2502.05215>
 41. SoK: Watermarking for AI-Generated Content - arXiv, 8 月 31, 2025 にアクセス、<https://arxiv.org/html/2411.18479v1>
 42. 生成 AI 活用論文の品質評価と査読基準における課題と懸念点, 8 月 31, 2025 にアクセス、<https://education.smeai.org/ai-academic-paper-review-challenges->

concerns/

43. Think you can cheat with AI? A UF professor creates watermarks to detect AI-generated writing, 8 月 31, 2025 にアクセス、
<https://news.ufl.edu/2025/02/think-you-can-cheat-with-ai/>
44. Identifying AI generated content in the digital age: The role of watermarking | EY, 8 月 31, 2025 にアクセス、
<https://www.ey.com/content/dam/ey-unified-site/ey-com/en-in/insights/ai/documents/ey-identifying-ai-generated-content-in-the-digital-age-the-role-of-watermarking.pdf>
45. Guido W. Imbens | Stanford Graduate School of Business, 8 月 31, 2025 にアクセス、
<https://www.gsb.stanford.edu/faculty-research/faculty/guido-w-imbens>
46. A Conversation with Guido W. Imbens - Project Euclid, 8 月 31, 2025 にアクセス、
<https://projecteuclid.org/journals/statistical-science/volume-39/issue-2/A-Conversation-with-Guido-W-Imbens/10.1214/23-STS906.full>
47. A conversation with economics Nobelists - Amazon Science, 8 月 31, 2025 にアクセス、
<https://www.amazon.science/latest-news/a-conversation-with-economics-nobelists-on-experimental-design>
48. causalens.com, 8 月 31, 2025 にアクセス、
<https://causalens.com/blogs/guido-imbens-nobel-laureate-data-science-causal-agents#:~:text=Imbens%20was%20quick%20to%20point,you%20need%20transparency%20and%20auditability.>
49. Guido Imbens: Future of Data Science Belongs to Causal+ Agents - causaLens, 8 月 31, 2025 にアクセス、
<https://causalens.com/blogs/guido-imbens-nobel-laureate-data-science-causal-agents>
50. Artificial Intelligence and the Future of Healthcare - The Doctors Company, 8 月 31, 2025 にアクセス、
<https://www.thedoctors.com/articles/artificial-intelligence-and-the-future-of-healthcare/>
51. Eric Topol pens book on artificial intelligence in medicine | Scripps Research, 8 月 31, 2025 にアクセス、
<https://www.scripps.edu/news-and-events/press-room/2019/20190312-topol-deep-medicine.html>
52. Topol Discusses Potential of AI to Transform Medicine - NIH Record, 8 月 31, 2025 にアクセス、
<https://nihrecord.nih.gov/2024/11/22/topol-discusses-potential-ai-transform-medicine>
53. Dr. Eric Topol on How AI is Transforming Health and Medicine - NIH Clinical Center, 8 月 31, 2025 にアクセス、
<https://www.cc.nih.gov/news/2024/nov-dec/eric-topol-ai>
54. Can AI Catch What Doctors Miss? | Eric Topol | TED - YouTube, 8 月 31, 2025 にアクセス、
https://www.youtube.com/watch?v=ll5LY7wI_Xc
55. Yejin Choi: Teaching AI How the World Works | Stanford HAI, 8 月 31, 2025 にアクセス、
<https://hai.stanford.edu/news/yejin-choi-teaching-ai-how-world-works>
56. Yejin Choi: The 100 Most Influential People in AI 2025 | TIME, 8 月 31, 2025 にアクセス、
<https://time.com/collections/time-100-ai-2025/7305803/yejin-choi/>
57. Yejin Choi - AI2050 - Schmidt Sciences, 8 月 31, 2025 にアクセス、

https://ai2050.schmidtsciences.org/fellow/yejin_choi/

58. The AI Scientist Generates its First Peer-Reviewed Scientific Publication - Sakana AI, 8 月 31, 2025 にアクセス、<https://sakana.ai/ai-scientist-first-publication/>
59. 世界初、100%AI 生成の論文が査読通過 「AI サイエнтиスト」が達成 - Sakana AI, 8 月 31, 2025 にアクセス、<https://sakana.ai/ai-scientist-first-publication-jp/>