

ScientistTwoが提示する自律型研究の到達点と閉ループ自動化の構造的限界

Gemini 3.8 Flash

Google Cloud AI Researchおよびウオーターラー大学の研究チームが2026年9月に公開した『ScientistTwo: Pioneering the Human Knowledge Frontier with Autonomous AI』は、機械学習分野における科学的発見の全サイクルを完全自律化したマルチエージェント・フレームワークである¹。人間が定式化した未解決課題と既存のコードベースを入力とし、ベースラインの弱点特定から仮説立案、計算資源を節約する二段階実験、因果関係を検証する自動アブレーション、論文執筆、そして追加実験を伴う模擬査読・反論(レバタール)に至る一連の研究活動を人間の介入なしで遂行する¹。ICLR、ICML、NeurIPSに採択された論文に基づく107のベンチマーク課題において、86課題(成功率80.4%)で人間の最先端(SOTA)ベースラインを上回り、平均25.2%の相対性能向上を達成したと報告されている²。

本フレームワークの工学的価値は、先行システムが抱えていた架空引用や再現不能なコード生成といった致命的欠陥を排除し、完全実行可能なコードベースと学術論文の双方を同時に出力する「Chain-of-Evidence(CoE)監査機構」を確立した点にある¹。生成された1,814件の文献参照におけるハルシネーション(虚偽生成)発生率はゼロであり、49件の無作為抽出監査において実行再現率は100%を記録した²。自動査読エージェント(Stanford Agentic Reviewer)によるブラインド評価でも、トップ会議に採択された人間論文と同等の採択率(75.8%対76.3%)を獲得するなど、ソフトウェア工学としての自律エージェントの統合度において現行最高水準に到達している⁴。

しかし、本作が標榜する「知識のフロンティアの開拓」という主張には重大な留保が必要である。相対改善率の分布は平均25.2%に対して中央値が7.7%にとどまり、大半の成功は既存パラダイム内部における局所的なハイパーパラメータ調整やモジュール付加の域を出ていない⁴。人間専門家による精査では「手法の妥当性・理論的深み」が人間論文を下回り(2.9/5)、評価スコアそのものもLLM査読器に対する敵対的最適化のバイアスを強く内包している⁴。課題1件あたり約3,765ドルの計算コストや、人間による課題定式化への絶対的依存を鑑みれば、ScientistTwoは新たな科学概念を創出する「科学者」ではなく、既存ベンチマークの漸進的改善を工業化した「自動化された通常科学の実装エンジン」と評価するのが実態に即している²。

課題抽出から実証的反論に至る閉ループ実行パイプラインの機構

ScientistTwoの中核を成すのは、研究課題目標 G を受け取り、査読基準を満たす論文 P^+ と完

全に検証された実行可能コードベース C^+ の組を自律的に導出する変換プロセス

$(P^+, C^+) = \mathcal{A}(G)$ である⁵。Sakana AIが発表した「The AI Scientist」などの先行システムが単一のプロンプトチェーンによる場当たり的なコード改変に依存していたのに対し、ScientistTwoは各研究工程を専門のエージェントへ分解し、動的なフィードバック結合を持つシーケンシャルなパイプラインとして設計されている⁴。

パイプラインの始点において、人間研究者は解決すべき具体的課題と基盤となるSOTA論文およびコードベースを提示する²。まず「Limitation Extractor」が原著論文の論理展開と実験結果を走査して構造的ボトルネックを特定し、「Limitation Verifier」がその弱点が実験的に反証・検証可能であるかを判定する⁴。この二重チェックを通過した課題に対し、「Novelty Checker」が外部の学術インデックスと照合して新規性を確認した上で、「Idea Evolver」が具体的な改善仮説を生成する⁴。この仮説生成は無作為な探索ではなく、抽出された弱点の解消に直接結びつけられた「錨付けされた生成 (Anchored Generation)」として機能し、過去の失敗実験のログを負例として再利用することで探索空間を効率的に絞り込む³。

実験段階では、無制限な計算資源の消費を抑止するために「サブセット先行評価 (Subset-First Evaluation)」が徹底される²。提案された仮説群は、まずベンチマークデータの縮約サブセット上でベースラインと比較され、性能改善の統計的シグナルを示さない仮説は早期に棄却される²。このスクリーニングを通過した有望な仮説のみがフルスケールの本実験環境へと昇格し、多様なデータセットと評価指標で検証される³。

フルスケール実験で改善が確認された手法に対しては、「Analyzer」が自動アブレーション (構成要素消去) 実験を実行する⁴。追加されたモジュールを体系的に除去・置換し、観測された性能向上が提案アルゴリズムの本質的寄与によるものか、偶発的なハイパーパラメータの変動によるものかを厳密に分離する⁴。性能向上に寄与しない余剰なコード改変はこの段階で自動的に削ぎ落とされ、手法の簡潔性と解釈可能性が担保される⁴。

執筆段階では「Writer Agent」が実験データ、図表、数式を統合して論文初稿を起草し、それを「ScholarPeer」などの模擬査読エージェントが10点満点で審査する²。評点が合格基準 (8点) に満たない場合、「Rebuttal Agent」が起動する⁴。従来の自動執筆ツールがテキストの修正のみで反論していたのに対し、ScientistTwoの反論機構は査読者の指摘や批判的疑念を解消するための追加実験スクリプトを自律設計し、コードを再実行して得られた新データを論文に組み込む²。「Meta-Review Agent」はこの改訂プロセスを統括し、基準を満たすまで仮説の微修正や追加実験のループを反復させる²。

最終成果物は「Chain-of-Evidence (CoE) 監査エンジン」に引き渡され、外部ネットワークから完全に隔離された環境でコードの再実行が行われる²。論文内に記載された数値やグラフとの100%一致 (Score Verification)、報酬ハッキングやテストデータ漏洩の不在 (Specification Compliance)、引用DOIの実在検索照合 (Reference Verification)、数式論理とPythonコードの一致 (Method-Code Alignment) が機械的に検証され、すべての監査を通過した成果物のみが出力される²。

パイプライン実行段階	トリガーおよび入力データ	担当エージェントと自律処理機構	到達状態および検証基準
課題特定・仮説定式化	人間が定義した研究課題、既存SOTA論文、動作可能なコードベース ² 。	Limitation Extractor / Verifier が既存研究の理論的・実験的欠陥を抽出。Idea Evolver が過去の失敗履歴を参照しながら	ボトルネックの具象性判定および学術検索による新規性確認の通過 ⁴ 。

		ら解決策を立案 ⁴ 。	
二段階実験検証	立案された複数の新規アルゴリズム仮説 ² 。	Evaluator がデータセットの代表縮約版で全仮説を初期評価。改善を示した上位候補のみを本番フルスケールデータセットで実行 ² 。	サブセット上での改善シグナル検出、フルデータセットでのベースライン凌駕 ⁴ 。
因果的アブレーション	フルスケール実験で性能向上を示した新規実装コード ⁴ 。	Analyzer が各構成コンポーネントを体系的に除去する実験群を自動実行。性能向上への寄与度が低い冗長コードを自動削除 ⁴ 。	性能向上をもたらした因果コンポーネントの特定と手法の極小化 ⁴ 。
論文執筆・査読・反論	検証済みアルゴリズム、実験ログ、図表アーティファクト ² 。	Writer Agent が初稿を執筆。ScholarPeer が10点満点で採点し、8点未満の場合はRebuttal Agent が補足実験コードを自律実行して改訂 ² 。	Meta-Review Agent による採録基準判定の達成(査読反論ループの収束) ² 。
証拠連鎖監査	生成論文ドラフトおよび最終実行コードベース ² 。	CoE Audit Engine が隔離環境で完全再実行。引用文献の実在検索照合、テスト漏洩検査、本文数式とコード実装の論理整合性を走査 ² 。	数値再現率100%、仕様違反0件、架空引用0件、コード・テキスト整合率100% ⁴ 。

主要国際会議107課題の実証成果と機械審査における非対称性

ScientistTwoの実証評価は、機械学習のトップカンファレンス(ICLR、ICML、NeurIPS)に直近で採択された論文から抽出された107件の研究課題を対象として実施された²。対象領域は、大規模言語モデル、ロボティクス、強化学習、神経科学、音声処理、ロバストネス、ゲーム理論、差分プライバシー、数理最適化、時系列解析など、広範なサブフィールドに及ぶ²。

全107課題中、システムは86課題において人間が設計したSOTAベースラインの性能を上回り、全体成功率は80.4%を記録した²。成功した86課題におけるベースラインに対する相対改善率は、算術平均で25.2%に達した²。しかし、この数値の背景には顕著な分布の歪みが存在する。相対改善率の中央値(Median)は7.7%にとどまり、平均値との間に3倍以上の乖離が生じている⁴。この乖離は、ごく一部の課題で極端に大きな性能向上が得られた一方で、大多数の課題におけるゲインは1桁台前半から中盤の漸進的なスコア向上であったことを実証している⁴。

論文品質の客観的測定のため、内部査読のScholarPeerとは別に、検証時のデータ汚染を排除した「Stanford Agentic Reviewer」を用いた独立ブラインド審査が行われた⁴。全生成論文に対するStanford Agentic Reviewerの平均スコアは5.7/10、自動採択率は72.1%に達した²。特にNeurIPS由来の33課題における比較では、ScientistTwoの論文が平均5.6/10の評点と75.8%の採択率を獲得し、人間著者が執筆した原著採択論文(平均5.5/10、採択率76.3%)と統計的にほぼ同等の水準を示した⁴。

実証実験では、反論ループが論文品質の向上に与える因果的影響もアブレーションによって定量化されている⁴。査読・反論プロセスを一切介さず初稿のまま評価された場合、Stanford Agentic Reviewerによる採択率は49.0%に低迷していた⁴。これに対し、1ラウンドの模擬査読と追加実験を伴う反論を実行した結果、採択率は73.5%へと大幅に向上した⁴。さらにScholarPeerを用いた検証では、反論なしの46.9%から2ラウンドの反論を経て93.9%まで採択率が上昇しており、批判に対する追加実験の自動執行が機械審査における評価を決定的に押し上げている実態が確認された⁴。一方、人間専門家9名による33本の論文に対するブラインド評価(1～5スケール、3.0が人間論文と同等水準)では、機械審査とは異なる評価の非対称性が浮き彫りとなった⁴。人間専門家による単独評価スコアは平均3.7/5、人間原著論文との相対比較スコアは総合で3.0と判定され、形式的な体裁は人間研究者と同等と認められた⁴。しかし項目別の内訳を見ると、「実験の網羅性(3.5)」や「アブレーション設計(3.3)」といった計算資源を投入して網羅的に実行可能な検証作業では人間を上回った反面、「手法の妥当性・理論的深み(Method Soundness)」は2.9と人間を下回る評価にとどまった⁴。

評価指標および測定項目	人間著者の採択論文(基準値)	ScientistTwo(実測値)	測定対象および統計的性質
研究課題の改善成功率	該当なし(ベースライン)	80.4%(86 / 107課題) ²	ICLR/ICML/NeurIPS採択論文の課題 ²
相対性能改善率(平均値)	0.0%(基準ベースライン)	25.2% 向上 ²	成功した86課題における算術平均 ²
相対性能改善率(中央値)	0.0%(基準ベースライン)	7.7% 向上 ⁴	平均値との乖離(偏重分布の存在) ⁴
ScholarPeer 平均	該当なし	7.5 / 10(採択率	システム内部の査読

評点 / 採択率		91.9%) ²	エージェント ²
Stanford Reviewer 評点 / 採択率	5.5 / 10(採択率 76.3%) ⁴	5.6 / 10(採択率 75.8%) ⁴	NeurIPS由来33課題 の独立外部審査 ⁴
反論ループ未実施 時の採択率	該当なし	49.0%(Stanford Reviewer) ⁴	レバタタル機構のア ブレーション結果 ⁴
人間専門家評価:総 合スコア	3.0 / 5(比較同等水 準) ⁴	3.0 / 5(単独評価 3.7 / 5) ⁴	専門家9名による33 本のブラインド審査 ⁴
人間専門家評価:手 法の妥当性	3.0 / 5(基準水準) ⁴	2.9 / 5(人間優位) ⁴	アルゴリズムの動機 付けと理論的厳密性 ⁴
人間専門家評価:実 験の網羅性	3.0 / 5(基準水準) ⁴	3.5 / 5(エージェント 優位) ⁴	比較対象の広さとベ ンチマークの規模 ⁴
人間専門家評価:消 去実験設計	3.0 / 5(基準水準) ⁴	3.3 / 5(エージェント 優位) ⁴	アブレーションの徹 底度と因果分離 ⁴
架空引用文献(ハル シネーション)	0.0%	0.0%(0 / 1,814件) ²	検索拡張照合による 全件実在性監査 ²
コード完全再実行に よる再現率	不明(人間論文では ばらつき大)	100.0%(49 / 49課 題) ⁴	クリーン環境での数 値完全一致監査 ⁴
課題あたりの平均消 費リソース	数週間～数ヶ月の研 究工数	2.5日 / 3,765米ドル ⁴	NeurIPS 33課題の計 算・API合算コスト ⁴

先行自律研究システムとの系譜的比較と工学的差別化

自律型AIサイエンティストの系譜において、初期のパラダイムを切り拓いたのはSakana AIの「The AI Scientist」(2024年)であった⁴。同システムはアイデア生成からLaTeX論文のコンパイルまでを完全自動で完結させたことで大きな注目を集めたが、実質的には事前定義されたテンプレートコードに対する小規模なテキスト置換とスクリプト実行に過ぎず、実験の失敗に対する動的な修復能力や多角的なアブレーション機構を欠いていた⁴。その結果、出力される論文には存在しない学術論文の引用(幻覚文献)が頻発し、生成コードが意図せぬ無限ループやリソース浪費を引き起こすなど、学術出版の水準には程遠い脆弱性を露呈していた⁴。

続く「The AI Scientist-v2」(2025年)では、エージェント型の木探索(Tree Search)を導入してワークショップ級の自動化を目指したものの、因果的な寄与の分離や体系的な反論機構の実装には至ら

なかった⁴。2026年5月に発表された「ScientistOne」は、「Chain-of-Evidence(証拠連鎖)」の概念を導入し、実験証拠と記述内容の追跡可能性を担保することで再現性の改善に道筋をつけたが、探索空間の最適化や査読指摘に応じた自律的な追加実験の実行といった高度な双方向ループは確立できていなかった⁴。

ScientistTwoは、これらの先行システムの限界を工学的に解消している⁴。AutoSOTAとの直接比較(共有されたICLR 5課題)において、AutoSOTAが既存コードのハイパーパラメータ調整にとどまり新規アルゴリズム要素を導入できなかったのに対し、ScientistTwoは成功した4課題すべてにおいて明確に新規なアルゴリズム構成要素を組み込み、完全なベンチマーク実行と査読可能な論文ドラフトの生成を完遂した⁴。さらに、前述の「VD-STrans」から「BXT-Transducer」への発展に見られるように、自らが自律生成した先行成果を次の研究のベースラインとして再入力し、さらなる性能向上を達成する「知識の複利化(Compounding Discovery)」の実証に成功した点は、先行システムには見られない独自の進境である⁵。

学術界の受容とAI査読を巡る方法論的論争

ScientistTwoの論文公開に対し、AI研究コミュニティおよび計算科学分野の反応は、そのソフトウェア工学的な卓越性を称賛する声と、研究の認識論的妥当性を懸念する声へと二極化している⁴。肯定的な評価として最も広く共有されているのは、科学研究の自動化において長年の課題であった「再現性の担保」と「ハルシネーションの排除」に対し、確実な技術的解決策を提示した点である²。1,814件の引用文献を網羅的に検証して架空文献を完全に排除し、49件のコード監査において隔離環境での100%再現を達成した実績は、従来の「体裁は整っているが中身が破綻しているLLM生成論文」の壁を打ち破るマイルストーンとして高く評価された²。また、査読エージェントからの批判に対し、単なるテキストの言い換えではなく「新規コードを実装して追加実験を実行する」という科学的レバットの本質をエージェントシステムとして実体化させた設計は、自律型ソフトウェア開発の極致として受け止められている²。

一方で、コミュニティの議論(Redditのr/aicuriosityやalphaXivの議論フォーラム等)では、評価方法の正当性に関する深刻な疑義が提起されている⁴。最大の批判は、論文の可否判定が「ScholarPeer」や「Stanford Agentic Reviewer」といった自動査読エージェントに依存している点に向けられている²。研究者らからは、「LLMが好む論理展開、定型的な章立て、整然とした実験フォーマットに対してLLM査読器が高得点を付与しているに過ぎず、自己言及的な共鳴室(Echo Chamber)を形成しているのではないか」という指摘が相次いでいる²。

さらに、システムに入力された107の課題がすべて「既に人間がコードベース、前処理パイプライン、評価メトリクスを整備し終えた問題」である点も論争の的となっている²。研究活動において最も本質的かつ困難とされる「未知の現象から問いを抽出する作業」を人間に外注している以上、これを「自律的な科学的先駆者」と呼ぶのは過大広告であり、実態は「極めて洗練された自動ベンチマーク・オプティマイザ」に過ぎないという冷淡な視線も根強く存在する¹。

認識論的パズル解きと合成ペーパーミル化への批評的視座

ScientistTwoが示した圧倒的な工学的達成度の背後には、計算論的アプローチが原理的に直面せざるを得ない5つの構造的・制度的限界が存在する。

1. 査読の自己言及性とGoodhartの法則による評価の歪み

本研究における最大の方法論的アキレス腱は、論文の価値基準をAI査読器に置いたことによる「評価の循環論法」である²。大規模言語モデルを基盤とする査読器は、文脈の一貫性、網羅的な消去

実験の記述、礼貌な反論トーンに対して高い評点を与える強いバイアスを有する。ScientistTwoのWriter AgentとRebuttal Agentは、これらの査読器の受容パターンに対して実質的な強化学習・プロンプト最適化を行っている状態に近い²。

「指標が目標になると、それは良い指標ではなくなる」というGoodhartの法則の通り、AI査読器による高いスコア (ScholarPeer 7.5/10、Stanford Reviewer 5.7/10) は、真の科学的ブレークスルーを証明するものではなく、査読モデルの選好に対する過剰適合の結果である可能性が高い²。人間専門家によるブラインド審査で「手法の妥当性・理論的深み」が2.9と人間論文を下回っていた事実は、AI査読器が論理構造の深層にある概念的浅薄さを見抜けず、表層的な実験ボリュームと形式美に欺瞞されていた現実を物語っている⁴。

2. クーンの「通常科学」の局所探索とパラダイム転換の不在

科学哲学者トーマス・クーンの枠組みを用いれば、ScientistTwoの動作機序は完全に「通常科学 (Normal Science)」のパズル解きに閉じ込められている。既存手法の弱点を抽出し、モジュールを追加・置換して既存ベンチマークのスコアを更新するプロセスは、定義された探索空間における勾配降下法的な局所ヒルクライミングに過ぎない⁴。

相対改善率の中央値が7.7%にとどまるというデータは、システムが生成した成果の大半が「既存技術に対する漸進的なパッチ当て」であることを証明している⁴。アインシュタインの相対性理論、量子力学の定式化、あるいは近年のディープラーニングにおけるトランスフォーマー構造の提唱のように、評価指標や問題の前提そのものを破壊・再構築する「パラダイム転換 (Paradigm Shift)」を、弱点補正型の局所探索アルゴリズムから生み出すことは原理的に不可能である³。

3. 計算資本の寡占と研究の経済的非対称性

1課題の完遂に「約2.5日」の実行時間と「約3,765ドル」の合算コストを要するという経済的仕様は、自律研究システムの利用可能性を激しく偏らせる⁴。人間研究者の人件費と比較して安価であるという見解も存在するが、80.4%の成功率の裏で約2割の試行は成果を出せずに計算資源を浪費しており、100件の論文を機械生成するためには数十万ドルのコンピュータ予算が必要となる⁴。

この構造は、最先端の自律研究を主導できる主体を、巨大な計算インフラと資金力を持つ一握りのハイパースケーラーやエリート機関に固定化させる⁴。結果として、研究の優先順位が「計算資源を大量投入してベンチマークの数値を力づくで押し上げられる領域」に偏重し、潤沢な計算資源を持たない理論的研究やニッチな学術領域が周縁化されるという学術の構造的歪みを加速させるリスクがある。

4. 砂場環境への依存と課題設定能力の欠落

ScientistTwoが発揮する自律性は、高度に人工的な「砂場 (サンドボックス)」の内部に限定されている¹。システムが正常に動作するためには、以下の前提条件が人間によって完全に整備されていなければならない²。

- 反証可能で明確に言語化された研究課題
- 依存関係が解決され、バグなく動作する既存SOTAのコードベース
- スカラー値として客観的に評価可能なベンチマークデータセット

自然科学の歴史において最も創造性を要し、知的能力が試されてきたのは、「何が解くべき真の課題であるかを発見すること」であり、「測定不能だった事象を測定可能にする実験パラダイムを創出すること」である。ScientistTwoは、問題設定と測定系が完全に固定された後の「実験ルーティンワーク」を高速代行しているに過ぎず、科学的発見の最も本質的な部分を依然として人間に依存している²。

5. 学術エコシステムの機能不全と合成ペーパーミルの脅威

ScientistTwoがもたらす最も破壊的な社会的リスクは、学術出版制度そのものの崩壊である⁴。再現性監査を通過し、文献捏造がなく、アブレーションが完備された論文が、1本あたり数千ドル・数日のコストで無制限に自動生成される世界が到来したとき、arXivや国際会議の査読システムは機械生成論文の洪水に直面することになる⁴。

人間の査読リソースは既に危機的飽和状態にあり、大量に流入する「工学的に欠陥はないが、概念的洞察に乏しい論文」の波を精査することは不可能である。防衛のために査読プロセス側もAIエージェントへの依存を深めれば、学術界は「AIが論文を書き、AIが査読し、AIが引用し合う」という、人間の理解と批判的対話から完全に乖離した自己目的的シミュレーション空間へと変質する危険を孕んでいる。

ScientistTwoが成し遂げた真の功績は、科学的思考の自動化というよりも、仮説検証型エンジニアリングの徹底した工業化とパイプライン統合にある。小規模実験による早期枝刈り、因果的寄与を分離するアブレーションの自動化、そしてコードと論文の完全一致を保証する監査機構は、計算機科学における実験の再現性を底上げする基盤技術として極めて実用的な価値を持つ²。しかし、このシステムを「人間の知識のフロンティアを開拓する自律的先駆者」として無批判に受容することは、科学の進歩を単なるベンチマーク指針の漸進的最適化と混同する過ちを招きかねない。自律型AIが真に科学の相転移を担うためには、与えられた盤面で最適解を探すエージェントの洗練にとどまらず、盤面そのものを疑い、新たな評価規範や理論体系を創出するための認識論的枠組みの再定義が不可欠である。

引用文献

1. Pioneering the Human Knowledge Frontier with Autonomous AI - arXiv, <https://arxiv.org/abs/2609.19644>
2. ScientistTwo Automates Autonomous AI Research - SuperGok, <https://supergok.com/scientisttwo/>
3. ScientistTwo: Pioneering the Human Knowledge Frontier with, <https://academy.dair.ai/papers/scientisttwo-pioneering-the-human-knowledge-frontier-with-autonomous-ai-2609.19644>
4. ScientistTwo: Pioneering the Human Knowledge Frontier ... - alphaXiv, <https://www.alphaxiv.org/abs/2609.19644>
5. Pioneering the Human Knowledge Frontier with Autonomous AI, <https://www.themoonlight.io/fr/review/scientisttwo-pioneering-the-human-knowledge-frontier-with-autonomous-ai>
6. ScientistTwo: 자율 AI가 인간 지식의 최전선을 개척하는 방법 - YouTube, https://m.youtube.com/watch?v=foi8pMg_Ulw
7. ScientistTwo Google AI Agent Outperforms Human Research at Top, https://www.reddit.com/r/aicuriosity/comments/1wkti0v/scientisttwo_google_ai_agent_outperforms_human/