

AI for Science の進化と知財業務への影響

— 科学研究の閉ループ化が特許制度と企業知財実務にもたらす変容 —

Claude Fable 5.1

2026 年 9 月 17 日

I. はじめに——問題設定と検証方針

本稿の出発点は、AI for Science (AI4Science) を「AI を科学へ応用すること」ではなく、「科学研究のループ全体——文献探索→仮説生成→理論計算・シミュレーション→実験計画→実験→結果解析→次の仮説——を閉ループ化し、科学的方法そのものの一部を AI が実行する動き」と捉える視座である。この枠組みに立てば、注目すべき指標は「AI が人間より速く計算できるか」ではなく、「AI がどこまで科学的方法そのものを閉ループ化できるか」となる。

本稿は、この枠組みを 2026 年 9 月時点の一次ソースに照らして検証したうえで、その進化が知財制度および企業知財部門の実務に与える影響を論じる。記述に当たっては、各節で【確認済み事実】と【分析・推論】を明示的に区別し、事実には番号引用を付した。法制度・審査実務が未確定の論点については断定を避け、複数の見解と不確実性を併記している。

1. 本稿の結論（要旨）

出発点の枠組み——AI が仮説生成と探索を担い、科学者が「探索を実行する人」から「探索空間と評価条件を設計する人」へ移行する——は、方向性としては 2026 年時点の一次ソースによって裏付けられる。ただし「科学的方法の完全な閉ループ化」は、材料科学・創薬・数学の一部において部分的に実現しているにとどまり、再現性の確保、第三者による検証、開かれた研究課題における判断力、自己評価バイアスという四つの壁が依然として大きい。

知財面では、「発明者は自然人に限る」という原則が主要各国で確定した一方、AI による大量発明・大量公開が進歩性判断・先行技術・出願戦略を根底から揺さぶり始めている。競争優位の源泉は、公開されやすい「発見結果」から、データ・モデル・自律実験室ノウハウという「探索装置」側へと移りつつある。これが本稿の中心的な見立てである。

II. AI for Science とは何か——定義・系譜・出発点の検証

1. 語の由来と定着（【確認済み事実】）

「AI for Science」という語は、米国エネルギー省（DOE）が2019年7月から10月にかけてアルゴンヌ国立研究所、オークリッジ国立研究所、ローレンス・バークレー国立研究所で開催した4回のタウンホール会合を踏まえ、2020年に公表した報告書AI for Scienceによって定着した¹。同報告書には、4回の会合に1,000名を超える米国の科学者・技術者が参加した旨が記されている¹。共同議長はRick Stevens（アルゴンヌ）、Kathy Yelick（バークレー）、Jeff Nichols（オークリッジ）である¹。DOEはその後、対象を安全保障・エネルギーに広げたAI for Science, Energy, and Security Report（2023年）を公表している。

日本では、科学技術振興機構研究開発戦略センター（JST-CRDS）が短報シリーズ「AI for Science Spotlights」を展開している。Vol.1「AI for Scienceの現在地——研究プロセスを変えるAIの広がり——」は、AIによる文献調査やデータ解析の支援もAI for Scienceであり、AIが仮説を生成し実験計画を立ててロボット実験を制御し、その結果をもとに次の実験条件を提案することもAI for Scienceであるとしたうえで、さらにAIが問いの設定、仮説生成、証拠の評価、研究方針の判断に関与するようになれば、科学者の知的活動そのものにAIが関与することになる、と整理する²。同Vol.1は、AI for Scienceを「支援系のAI for Scienceから自律系のAI for Scienceまで」の連続体として捉える枠組みを示している²。

Vol.2「AI エージェントの科学研究への展開」（2026年8月7日）は、AI エージェントの科学研究への適用が「研究者の問いに答えるような支援ツールから、文献調査、仮説生成、計算、データ解析、実験計画、結果評価などを連続して進めるシステムへと発展している」と述べ、代表例としてGoogleのCo-Scientist、および米国の非営利研究機関FutureHouseが開発したマルチエージェントシステム「Robin」を挙げる³。Robinは、文献探索を担うCrowやFalcon、データ解析を担うFinchなど異なる役割を持つエージェントを一つの研究ワークフローとして連携させ、仮説生成・実験計画・データ解析・結果を踏まえた次の仮説の提示までを反復的に進める仕組みであり、乾性加齢黄斑変性を対象とした研究で実際に運用された³。

2. 出発点の枠組みの検証（【分析・推論】）

出発点論考が示す「支援ツール→科学的方法の実行主体」という段階論は、JST-CRDSの「支援

系→自律系」という連続体の把握と実質的に同型である^{2,3}。したがって本稿では、AI for Science を「閉ループ化の程度」という単一の軸で位置づけ、段階ごとに知財業務への影響を検討する方法を採る。

なお用語の整理として、「AI for Science（科学のための AI）」、「Science of AI（AI そのものの科学的解明）」、「AI Scientist（自律的に研究を遂行するシステム）」、「サイエンスエージェント」、「自律科学」は相互に重なりつつも異なる概念であり、論者によって指示範囲が揺れている点には留意を要する。

III. 技術と実装の現状

1. 主要マイルストーンの時系列（【確認済み事実】）

以下に、AI for Science の展開を示す主要な出来事を、一次ソースと査読・第三者検証の状況とともに整理する。

時期	事項	一次ソース／検証状況
2020 年	DOE 報告書 AI for Science 公表。語の定着	OSTI、DOI:10.2172/1604756 ¹
2023 年 11 月 29 日	GNoME：新規結晶構造 220 万件を発見、うち約 38 万件を安定と予測し Materials Project に提供。DeepMind は、世界各地の研究室の外部研究者が並行研究のなかでこれらの新構造のうち 736 件を実験的に作製したと述べる	Nature（2023）、DeepMind 公式 ⁴ 。後述の重複指摘・訂正要求あり
2023 年 11 月 29 日	A-Lab（バークレー研究所）：17 日間の連続稼働による自律合成	Nature、DOI:10.1038/s41586-023-06734-w ⁵ 。後に訂正、再現性論争あり ⁶
2024 年 5 月	AlphaFold 3：タンパク質に加え DNA・RNA・リガンド・イオンとの複合体構造を予測	Nature（2024） ⁷ 。査読済
2024 年 7 月 25 日	AlphaProof／AlphaGeometry 2：国際数学オリンピック問題で銀メダル相当の得点	DeepMind 公式 ⁸ 、Nature（2025 年 11 月 12 日）DOI:10.1038/s41586-025-09833-y ⁹ 。査読済
2024 年 10 月	ノーベル化学賞（Hassabis、Jumper／AlphaFold）	ノーベル財団

時期	事項	一次ソース／検証状況
2025 年 5 月 14 日	AlphaEvolve：進化的コーディングエージェント。4×4 複素行列積を 48 回の乗算で実行する手法などを発見	DeepMind 白書、arXiv:2506.13131 ¹⁰ 。査読前テクニカルレポート
2025 年 6 月 3 日	レントセルチブ (ISM001-055) 第 IIa 相試験結果を Nature Medicine に掲載。生成 AI 創薬として初の臨床 POC	Nature Medicine ¹¹ 。査読済
2025 年 11 月 24 日	米大統領令「Launching the Genesis Mission」	ホワイトハウス ¹²
2026 年 3 月	Sakana AI「The AI Scientist-v2」生成論文が ICLR ワークショップで査読通過	arXiv:2504.08066 ¹³ 。ワークショップ査読、人間による事前選別あり
2026 年前半	Google「AI co-scientist」：マルチエージェントによる仮説生成、cf-PICI 仮説が独立実験と一致	Nature (2026) ¹⁴ 、JST-CRDS Vol.2 ³ 。査読済
2026 年 8 月 6 日	Evo 1/2 によるバクテリオファージ全ゲノム生成と実験検証	Science、DOI:10.1126/science.aec2657 ¹⁵ 。査読済
2026 年 8 月 18 日	Anthropic：Claude によるタンパク質設計・分析化学の加速を報告	Anthropic 公式 ¹⁶ 。自己申告値、湿式検証は外部機関
2026 年 9 月 9 日	レントセルチブ第 III 相試験 (GENESIS-IPF-3) 初回投与。生成 AI 創薬として世界初の第 III 相	Insilico 公式 ¹⁷
2026 年 9 月 6～8 日	OpenAI「Astra」によるナビエーストークス関連の証明を公表	OpenAI 公式、Nature 報道 ¹⁸ 。重大な論争あり（後述）

2. ナビエーストークス主張の検証（【確認済み事実】）

2026 年 9 月上旬、OpenAI は内部モデル「Astra」を用い、多数のエージェントを長時間並列稼働させてナビエーストークス問題に関する結果を得たと公表した。しかし、これを「未解決問題の解決」と表現することには重大な留保が必要である。

第一に、示されたのは強制項を伴う（forced）設定における有限時間爆発であり、クレイ数学研究所のミレニアム懸賞問題が主として問うている非強制系の存在と滑らかさとは異なる^{18,19}。第二に、OpenAI 自身が懸賞金を請求しない旨を明言し、クレイ数学研究所は当該問題を未解決として扱い続

けている¹⁹。第三に、先行研究者のデータ利用をめぐり、ニューヨーク大学の数学者から公然たる批判が提起された²⁰。

【分析】出発点論考の「AI が巨大な理論空間を探索する装置になる」という把握は、方向性としては本件によって支持される。しかし「AI が未解決問題を解いた」という表現は現時点では過大であり、査読・第三者検証は未完了である。知財実務の観点からは、AI による「発見の主張」と「検証済みの技術的裏付け」との間に構造的なギャップが存在するという事実こそが重要であり、この点は後述の記載要件論に直結する。

3. 開かれた研究課題における限界（【確認済み事実】）

AI エージェントの「研究遂行能力」を最も厳密に測った一次証拠が、CRUX による論文「Can AI agents conduct open-ended AI research?」（arXiv:2607.27191）である²¹。同研究は、未公開の NeurIPS 投稿論文の中心的な研究課題をフロンティアエージェントに与え、数日間と相応の計算資源を割り当てて自律的に解かせる「シャドー評価（shadow evaluation）」という手法を採った²¹。

エージェントは文献調査、GPU 環境のデバッグ、数百回にのぼる実験、投稿可能な体裁の原稿作成までを人手の介在なしに完遂した。しかし原著者による評価は明確に不合格であり、失敗の主因は創造性と判断の欠如、および初期アプローチへの早期固着にあった²¹。また同論文は、AI 生成論文をブラインド査読で評価する手法が過負荷かつ確率的で質が低いことを指摘し、代替評価としてシャドー評価を提案している²¹。

【分析】この結果は、「サイエンスエージェント化」が工学的実行力の面では顕著な水準に達している一方、開かれた研究課題における判断力は未成熟であることを示している。出発点論考が述べる「科学者は探索空間と評価条件を設計する人へ移る」という命題は、裏を返せば「評価条件の設計こそが AI に委ねられない領域として残っている」ことを意味する。これは知財実務における「何を守るか」の判断が人間に残る理由と同型である。

4. 自律実験室の実像（【確認済み事実】）

化学・材料科学における閉ループ化の代表例が A-Lab である。2023 年 11 月の Nature 論文は、17 日間の連続稼働により多数の標的化合物を自律合成したと報告した⁵。しかし、その後、実際には目的物が合成できていない事例が相当数含まれること、X線回折データの解釈に問題があることを指摘する反論が Palgrave、Schoop らから提起され、同論文は訂正された。訂正版では成功件数が下方修正されている⁶。姉妹関係にある GNoME 論文についても、既知物質との重複を指摘する声が上がっ

た⁶。

【分析】出発点論考の「化学・材料科学で自律実験室が先行している」という認識は事実として正しい。しかし「閉ループが安定稼働している」段階ではなく、成功率・再現性・第三者検証の三点において概念実証の域を出ていない。知財実務にとってこの差は決定的である。自律実験室が生み出すデータをそのまま出願の裏付けとして用いれば、記載要件で足をすくわれる危険がある。

5. 生命科学・創薬における進捗（【確認済み事実】）

生成 AI 創薬の臨床的到達点を示すのがレントセルチブ（ISM001-055）である。第 IIa 相 GENESIS-IPF 試験は中国 22 施設で特発性肺線維症患者 71 名を対象とした二重盲検プラセボ対照試験であり、主任研究者は北京協和医学院の Zuojun Xu 教授である¹¹。60mg1 日 1 回投与群では 12 週時点の努力肺活量（FVC）の平均変化が+98.4mL であったのに対し、プラセボ群は-20.3mL であった¹¹。ただし各群の症例数は限られており、より大規模なコホートでの検証を要する旨が研究者自身により明示されている¹¹。

2026 年 9 月 9 日、同薬是北京協和医院および上海市肺科医院で第 III 相試験（GENESIS-IPF-3、NCT07687459）の初回投与に至った。52 週投与、中国 47 施設、320 例を予定し、主要評価項目は 52 週にわたる FVC 年間低下率である¹⁷。生成 AI が標的同定と分子設計の双方を担った医薬品が第 III 相に進んだ初の事例とされる¹⁷。

ゲノム設計の領域では、Evo 1/2 を用いてバクテリオファージの全ゲノムを生成し、合成・実験検証に至った成果が Science に掲載された¹⁵。デュアルユース懸念への配慮から、ヒト・動植物の病原体データを学習から除外し、対象を ΦX174 に限定する措置が採られている¹⁵。

Anthropic は 2026 年 8 月 18 日、Claude を用いた de novo タンパク質バインダー設計および分析化学の加速について報告した¹⁶。ただし公表値は自己申告であり、湿式検証は外部機関が担当している点に留意が必要である¹⁶。

IV. 政策と各国の動向（【確認済み事実】）

1. 米国——ジェネシス・ミッション

2025 年 11 月 24 日の大統領令「Launching the Genesis Mission」は、DOE を主導機関とし、データ・計算資源・AI モデル・研究設備を共通基盤上で接続する「米国科学安全保障プラットフォーム

(American Science and Security Platform)」の構築を掲げた¹²。JST-CRDSによれば、2026年7月22日、DOEは同ミッションの最初の研究開発テーマ群として277件のプロジェクトを選定した²²。その内訳は、国家科学技術課題に直接対応する17領域（先進製造、バイオテクノロジー、重要鉱物、原子力、核融合、量子、マイクロエレクトロニクス、材料、電力網、地下資源など）と、プラットフォームの横断的ニーズに対応する4領域（HPCコード開発、科学的推論のためのAI、AI駆動型科学ワークフローのサイバーセキュリティなど）から構成される²²。

政府横断の展開も進み、保健福祉省と国立衛生研究所は生命・医学分野を担う「Bio Genesis Mission」を開始し、生体システムの予測、バイオ製造、生物学的脅威、小児がん、創薬・臨床応用、慢性疾患の6課題を対象としている²³。NASAは宇宙システムの設計・運用へのAI利用と150ペタバイトを超える蓄積データへのAI適用を進め、NISTはDOE科学局と協力覚書を締結してAIエージェントとロボティクスによる製造高度化等に着手した²³。

【分析】ジェネシス・ミッションの構造上の特徴は、個別重点領域へAIを適用する「縦」の研究と、複数分野で汎用利用できるAIモデル・研究ワークフロー・自律実験・セキュリティを整備する「横」の研究を同一ポートフォリオで並行して進める点にある²²。この「横」の基盤こそが、後述する知財戦略上の「探索装置」に相当する。国家がこの層に集中投資しているという事実は、企業が同じ層をどう押さえるかという問いを突きつけている。

2. 日本——科学技術政策と知財政策

文部科学省は2026年5月21日、総合科学技術・イノベーション会議の場で「AI for Scienceの推進に向けた基本的な戦略方針」の推進状況を示した²⁴。

知財政策については、2026年6月12日、知的財産戦略本部が「知的財産推進計画2026——成長戦略を支える知財戦略の推進——」を決定した²⁵。同計画は、日本成長戦略17分野の戦略的不可欠性・競争優位性の確保のため、IPランドスケープを活用した勝ち筋の特定、集中的な知財投資、知財・無形資産ガバナンスを進めるとともに、他国の知財覇権戦略や知財侵害リスクへの対抗措置としての知財の積極的活用、国際標準戦略の成長戦略との一体的推進を図るとする²⁶。また、令和8年度中に知財・無形資産ガバナンスガイドラインの改訂を行う旨も明記されている²⁶。

2026年8月25日、知的財産戦略本部は「生成AIの適切な利活用等に向けた知的財産の保護及び透明性に関するプリンシプル・コード」を公表した²⁷。同コードは、生成AI事業者が行うべき透明性の確保および知的財産権保護のための措置の原則を定め、生成AI技術の進歩の促進と知的財産権

の適切な保護の両立に向け、権利者や利用者にとって安全・安心な利用環境を確保することを目的とする²⁸。適用対象は「生成 AI 開発者」および「生成 AI 提供者」である²⁸。法的拘束力はないが、コンプライ・オア・エクスプレインの手法を採用しているため、事業者が受入れを表明した場合には各原則について実施するか、実施しない理由を合理的に説明する必要がある²⁹。権利者による個別開示請求は、今後 AI 学習データの内容を把握する実務上の手段として活用される可能性があり、知財紛争の場面での活用も想定されている²⁹。

さらに 2026 年 8 月 3 日、内閣府は「価値創造を加速する知財・無形資産投資・活用のガイダンス」を公表した³⁰。

3. 欧州・中国

欧州連合では AI Act が施行され、その後のオムニバス・パッケージにより一部規制の適用時期・範囲が調整された。中国国家知識産権局（CNIPA）は 2023 年 12 月 21 日に専利審査指南を改正し（2024 年 1 月 20 日施行）、AI を発明者として記載できない旨を明確化している³¹。

【注記】依頼者が言及された中国「知的財産権保護・運用『第 15 次 5 カ年規画』（国発〔2026〕30 号）」については、本調査の範囲では一次ソースを確認できなかった。真偽未確認として扱う。

V. 批判・限界・リスク

1. 再現性と第三者検証（【確認済み事実】）

前述の A-Lab 論文の訂正および GNoME 論文への重複指摘は、AI 主導の「発見」が第三者検証と再現性の面で脆弱でありうることを示す代表例である^{5,6}。この問題は、単なる学術上の論争にとどまらず、当該データを根拠に出願された特許の有効性に直結する。

2. 自己評価バイアスと評価の困難（【確認済み事実】）

Sakana AI のシステムに含まれる AI 査読機能を第三者が検証した研究では、OpenReview で人間により採択された論文を含む大半を却下する強い保守的バイアスと、人間の判断との乖離が報告されている³²。また AI 生成論文の査読通過事例は、ワークショップ・トラックにおけるものであり、かつ人間による事前選別を経ている¹³。

3. デュアルユースと資源の寡占（【確認済み事実】）

生成ゲノム設計については、安全措置が講じられている一方で、原理的にはより危険な病原体設計へ転用しうるとの警告がバイオセキュリティ研究機関から発せられている¹⁵。また、フロンティア級の科学的探索には巨額の計算資源を要することが報じられており、科学的発見の能力が一部の大規模事業者に集中するリスクが露呈している¹⁸。

【分析】出発点論考の「巨大な理論・探索空間を AI が並列探索する」という把握は正しい。しかし「探索できること」と「正しく検証・再現できること」の間には大きなギャップがある。知財実務上、このギャップこそが記載要件論・進歩性論の核心となる。

VI. 知財制度への影響

1. 発明者要件——各国の帰結（【確認済み事実】）

DABUS 事件を通じて、実体審査を伴う主要法域はいずれも AI を発明者と認めないという結論に収斂した。各法域の帰結は次のとおりである。

法域	結論	根拠・出典
米国	否定	Thaler v. Vidal, 43 F.4th 1207 (Fed. Cir. 2022)、最高裁は上告受理せず ³³
米国（審査運用）	2024 年 2 月ガイダンス→2025 年 11 月 28 日改訂	89 FR 10043（2024 年 2 月 13 日） ³⁴ 、Revised Inventorship Guidance for AI-Assisted Inventions（2025 年 11 月 28 日） ³⁵
欧州（EPO）	否定	法律審判部 J 8/20・J 9/20（2021 年 12 月 21 日） ³⁶
英国	否定	Thaler v Comptroller-General [2023] UKSC 49（2023 年 12 月 20 日） ³⁷
ドイツ	AI 自体は否定、自然人を記載すれば可	BGH X ZB 5/22（2024 年 6 月 11 日） ³⁸
豪州	否定（確定）	Commissioner of Patents v Thaler [2022] FCAFC 62
中国	否定	CNIPA 專利審査指南改正（2023 年 12 月 21 日公布、2024 年 1 月 20 日施行） ³¹
韓国	否定	KIPO 拒絶、ソウル行政裁判所が支持（2023 年）
日本	否定（確定）	東京地判令和 6 年 5 月 16 日（令和 5 年（行ウ）第 5001 号） ³⁹ 、知

法域	結論	根拠・出典
		財高判令和7年1月30日（令和6年(行コ)第10006号） ⁴⁰ 、最高裁第二小法廷決定令和8年3月4日により確定 ⁴¹

日本の DABUS 事件では、出願人が特許法 184 条の 5 第 1 項所定の国内書面の「発明者の氏名」欄に「ダバス（DABUS）、本発明を自律的に発明した人工知能」と記載したのに対し、特許庁長官が自然人の氏名を記載するよう補正命令を発し、応じなかったため出願を却下した⁴⁰。知財高裁は控訴を棄却したが、その理由付けは第一審と異なり、AI 発明の問題は立法により解決されるべきものとして、より踏み込んだ見解を示した⁴⁰。すなわち、特許権は天与の自然権ではなく、「発明を奨励し、もって産業の発達に寄与する」ことを目的とする特許法に基づいて付与されるものであり、その制度設計は国際協調の側面を含め一国の産業政策の観点から議論されるべき問題である、との判示である⁴⁰。報道によれば、最高裁第二小法廷は令和8年3月4日付の決定で上告を退け、「発明者は人間に限られる」との判断が確定した⁴¹。

米国の審査運用については、2025 年 11 月 28 日の改訂ガイダンスが 2024 年 2 月 13 日ガイダンスを全面的に撤回した³⁵。改訂ガイダンスは、（i）発明者たり得るのは自然人のみであることを再確認し、（ii）AI システムを実験機器・ソフトウェア・研究データベースなど発明過程を補助する他の道具と同視し、（iii）AI 支援の有無にかかわらず伝統的な「conception（着想）」テストを一律に適用し、（iv）自然人が一人のみ関与する場合には Pannu 要件の適用を排除し、これを複数の自然人の共同発明者性判断に限定する、という枠組みを採った⁴²。同改訂は、大統領令 14179（Removing Barriers to American Leadership in Artificial Intelligence）を実施するものと位置づけられている⁴³。

【分析】各国は AI を発明者と認めない点で収斂したが、「人間の貢献をどの水準で要求するか」は依然として流動的である。米国は Pannu 要件からの離脱により実務上の運用を簡素化した一方、着想の有無という事実認定に回帰したことで、むしろ発明者性が訴訟で争われる余地が広がったとの指摘もある⁴²。「Human over the loop」型の研究体制、すなわち人間が探索空間と評価条件のみを設計する体制では、どの請求項に誰が着想を寄与したかの立証が実務上の最大の課題となる。

【分析】日本法においては、職務発明制度（特許法 35 条）が自然人発明者の存在を前提とする以上、AI 活用研究における発明者特定の失敗は、相当の利益の支払対象の不明確化、ひいては権利帰属の不安定化に直結する。対応策としては、①請求項ごとの着想寄与者の記録、②AI プラットフォーム

利用規約および共同研究契約における権利帰属条項の点検、③プロンプト・AI 出力・実験データの来歴管理、の三点を制度化する必要がある。

2. 進捗性と当業者水準（【確認済み事実】および【分析】）

米国特許商標庁は 2024 年 2 月 27 日、進捗性判断に関する更新ガイダンスを公表した⁴⁴。同ガイダンスは KSR 判決の柔軟なアプローチと Graham 要素を再確認するものであるが、AI の利用が当業者の技術水準を引き上げるか否かという論点には踏み込んでいない。この論点については、別途 2024 年 4 月 30 日の意見募集（Request for Comments、89 FR 25417）において、AI と先行技術・当業者の関係を含む複数の設問が提示されたが⁴⁵、これに基づく方針採択には至っていない。

【分析】学界では、AI 利用により当業者の水準が上昇し技術が予測可能になるのだから、非自明性・進捗性のハードルを引き上げるべきだとする議論が有力に主張されている。しかしこれは立法論・政策提言であって、現行の審査基準ではない。「AI で容易に探索できる」ことが進捗性否定の論拠となる可能性は理論的には存在するが、判例・審査基準として確立していない。実務上の含意は、材料組成・分子探索の領域における選択発明・パラメータ発明において、「AI が大量に探索した中から選んだ」という事実だけでは技術的意義の主張が構造的に弱くなる、という点にある。予期せぬ効果や技術的意義の実証を従来以上に厚く準備すべきである。

3. 記載要件——AI 予測・シミュレーションデータの扱い（【確認済み事実】および【分析】）

AI 予測データが実施可能要件を満たすかを正面から扱った米国判例は、現時点では存在しないと指摘されている⁴⁶。その背景には、広範な機能的クレームについて過度の実験を要せず全範囲を実施可能とする指針が明細書に必要であったとした Amgen v. Sanofi 判決（598 U.S. 594 (2023)）がある⁴⁶。欧州実務については、技術的效果は in vitro または in vivo のデータで裏付けるべきであり、in-silico データのみに依拠することは避けるべきであるとの指摘がある⁴⁷。これは拡大審判部 G 2/21 が示したプラウジビリティをめぐる枠組みに連なる論点である。

日本については、特許庁が「AI 関連技術に関する特許審査事例」を公表し、2024 年 3 月 13 日に事例を追加している⁴⁸。もっとも、AI 予測データが湿式実験データを代替してサポート要件・実施可能要件を満たすかという問題を正面から解決した公表指針は確認できない。

【分析】前述の A-Lab 論争が示す自律実験室データの再現性の弱さ⁶は、AI 予測由来データの記載要件上の取扱いを厳格化する方向に働く蓋然性が高い。企業としては、当面「AI 予測＋実測による

裏付け」をセットで押さえる実務を標準とすべきである。in-silico データのみで広い権利範囲を取りにいく戦略は、各国の実務が固まるまでは高リスクと評価せざるを得ない。

4. 先行技術と新規性（【確認済み事実】および【分析】）

知的財産推進計画 2026 は、生成 AI による先回り大量生成の問題を課題として認識している²⁵。また同計画は、AI と知財をめぐる論点について、利活用の推進から「ルール整備・保護」へと重心を移している。

【分析】生成 AI による大量生成と大量の防御的公開は、先行技術の母集団を爆発的に増大させ、調査の網羅性という前提そのものを揺るがす。さらに GNoME のような大規模予測データベースの公開は⁴、新規性・進歩性判断における「公知」の外延を実質的に拡大しうる。自社の自律実験室が生成するログ・データについて、いつ・どの粒度で公開するかという管理は、そのまま自社出願のタイミング戦略の一部となる。従来「研究データ管理」と呼ばれてきた領域が、知財戦略の中核に繰り上がるということである。

VII. 知財戦略・知財部門業務への影響と対応策

1. 研究現場の変化（【確認済み事実】）

キリンホールディングスと GenerativeX は 2026 年 8 月 6 日、AI エージェントを研究活動そのものに組み込んだ新たな研究スタイルを構築し、2026 年よりキリンの一部の研究部門で運用を開始したと発表した⁴⁹。研究者が都度 AI を呼び出すのではなく、研究活動の流れの中に AI が常時存在する設計とし、仮説形成、情報探索、壁打ち、ドキュメント化、知見共有を一体で支援するものである⁴⁹。仮説や議論、探索履歴は蓄積され、組織全体で活用できる仕組みが取り入れられている⁴⁹。同社が挙げる背景には、情報探索・過去資料検索・仮説整理・共有業務によって思考が何度も中断される「思考の分断」という課題がある⁴⁹。

知財部門側では、NTT データが 2026 年 8 月 19 日、AI を活用した発明発掘を支援するサービスを先行提供開始すると発表した⁵⁰。同サービスは、AI により発明を自動生成するのではなく、研究者・開発者の発明発掘のプロセスを AI が支援するアプローチを採用し、AI が研究開発部門と知財部門をつなぐパートナーとして機能して、発明相談から発明提案書作成までのプロセスを支援するものである⁵⁰。提供開始は 2026 年 10 月以降とされる⁵⁰。

2. 発明発掘の閉ループ化（【分析・推論】）

研究現場に AI エージェントが常駐し、仮説が大量に生成される体制が広がれば、知財部門はその流れの下流で待つ従来型の運用を維持できない。麒麟の事例が示す「探索履歴の組織的蓄積」⁴⁹と、NTT データのサービスが示す「研究開発部門と知財部門をつなぐ AI」⁵⁰は、いずれも同じ方向、すなわち発明発掘プロセス自体の閉ループ化を指している。

ここで筆者が提案したい枠組みが、知財業務の「Human over the loop」化である。出発点論考は、科学者が「探索を実行する人」から「探索空間と評価条件を設計する人」へ移ると論じた。同じ転換が知財部門にも起こる。先行技術調査、IP ランドスケープ分析、明細書の一次ドラフトは AI が担い、知財担当者は「何を守るか（探索空間）」と「どの基準で権利化・秘匿・公開を振り分けるか（評価条件）」の設計者へ移行する。これは分析であり提案であって、確立した実務ではない。

3. 出願戦略の再設計（【分析・推論】）

発見が大量並列化すれば、全件出願は費用面でも実務面でも成立しない。必要となるのは、出願・防御的公開・営業秘密の三択を価値評価に基づいて高速に振り分けるトライアージの仕組みである。

オープン・クローズ戦略の観点からは、自律実験室の運用ノウハウ、学習データ、モデルの重みを営業秘密として保持し、最終的に到達した組成物・用途を特許で押さえる二層構造が有力である。組成物クレーム、マーカッシュ形式、用途発明の設計も、AI 探索を前提に組み直す必要がある。また AI による発見・改良サイクルの高速化は、特許の経済的寿命を短縮し、IP ランドスケープの更新頻度を引き上げる。

さらに上位の論点として、AI→科学→ハードウェア→AI という自己強化ループの中で知財が果たす役割がある。ジェネシス・ミッションが「縦」の応用研究と「横」の基盤整備を同一ポートフォリオで進めている構造²²が示唆するとおり、競争優位の源泉は、公開されやすい「発見結果」よりも、データ・モデル・自律実験室ノウハウという「探索装置」の側に移りつつある。この層を営業秘密と限定的な特許の組合せで押さえられる企業が、長期的な優位を得ると考えられる。これは筆者の見立てであり、実証された命題ではない。

4. 発明記録とガバナンス（【分析・推論】）

AI 生成発明の「発明記録」管理は、発明者適格の立証と職務発明対価算定の双方の前提となる。具体的には、着想を寄与した自然人の特定、プロンプト・AI 出力・実験データの来歴記録を、研究の進行と同時に残す仕組みが必要である。生成 AI 利用ガイドラインについては、プリンシプル・コード²⁸が示す透明性確保の枠組みと整合させておくことが望ましい。

5. 進化段階と知財業務への影響の対応関係（【分析】）

進化段階	科学の実態（代表例）	知財業務への主な影響	主な対応
支援系	文献探索・データ解析の支援。AlphaFold 等の道具的 AI	先行技術調査・明細書作成の効率化	AI ツール導入と品質管理
仮説生成系	Co-Scientist、Robin、キリンの AI ネイティブ研究環境	発明発掘の大量化、着想寄与者の特定困難化	発明記録の制度化、発明者性の立証設計
自律実験系	A-Lab 等の自律実験室	データ来歴と公知時期の管理、記載要件の厳格化	ログ管理、営業秘密と特許の二層構造
自律研究系	AI Scientist、フロンティアモデルによる探索	進歩性論・インセンティブ論の再構築、探索装置の秘匿	ポートフォリオ・トライアージ、制度動向の監視

VIII. 結論と提言

1. 総括

AI for Science は、科学研究の各段階を部分的に閉ループ化しつつある。しかし、①再現性と第三者検証、②開かれた課題における判断力、③自己評価バイアス、④デュアルユース、⑤計算資源の寡占、という五つの壁が残っており、「科学的方法の完全な閉ループ化」にはなお距離がある。

知財の側では、「発明者は自然人」という原則が主要各国で確定した結果、争点は制度の入口から実務の内部へ移った。すなわち、人間の貢献をいかに立証するか、AI による大量発明の下で進歩性・記載要件・先行技術をどう扱うか、そして探索装置そのものをどう秘匿するか、である。

2. 段階的提言

即時（概ね 3 か月以内）——AI 利用発明の発明記録テンプレート（請求項ごとの着想寄与者、プロンプト、AI 出力、実験データ来歴）を整備する。生成 AI 利用ガイドラインをプリンシプル・コードの枠組みと整合させる。共同研究契約および AI プラットフォーム利用規約の権利帰属条項を点検する。

短期（概ね 1 年以内）——発明発掘、先行技術調査、IP ランドスケープに AI エージェントを試験導入する。出願・防御的公開・営業秘密のトリアージ基準を策定する。材料・創薬分野において「AI 予測＋実測」による記載要件対応を標準化する。

中期（1～3 年）——知財部門を「Human over the loop」型の体制へ再設計する。自律実験室のノウハウ、データ、モデルの営業秘密管理体制を構築する。進歩性・記載要件に関する各国審査実務の変化を継続的に監視する。

3. 判断を変えるべき指標（トリガー）

次の事象が生じた場合には、上記の方針を速やかに見直す必要がある。第一に、米国特許商標庁が 2024 年 4 月の意見募集⁴⁵を踏まえ、AI が当業者水準を引き上げる旨の方針を採択した場合。第二に、日本特許庁または欧州特許庁が、in-silico データによるサポート要件充足について明示的な指針を公表した場合。第三に、ナビエーストックス等の AI 主導の主張が査読および第三者検証を通過した場合。第四に、産業構造審議会特許制度小委員会が AI 利用発明の発明者性・記載要件について新たな方針を示した場合である。

4. 一次ソース確認上の注記

本稿の記述のうち、以下は確認状況に留保がある。OpenAI のナビエーストックス「解決」主張は、強制項付きの設定に限られ、クレイ数学研究所による認定はなく、査読前であり、かつデータ利用をめぐる論争が継続中である^{18,19,20}。Anthropic が公表したタンパク質設計の成功率は自己申告値である¹⁶。AlphaEvolve は査読前のテクニカルレポートである¹⁰。AI 予測データの記載要件充足性については、各国とも拘束的な判例・指針が存在せず、実務家および学界のコメントが中心である^{46,47}。中国「知財第 15 次 5 年規画（国発〔2026〕30 号）」は一次ソースを確認できなかった。

参考文献

- [1] Stevens, R., Taylor, V., Nichols, J., Maccabe, A. B., Yelick, K., Brown, D., AI for Science, Argonne National Laboratory / U.S. Department of Energy, 2020. DOI:10.2172/1604756. <https://www.osti.gov/biblio/1604756>
- [2] JST 研究開発戦略センター (CRDS) 「ショートレポート AI for Science Spotlights Vol.1『AI for Science の 現 在 地 ― 研 究 プ ロ セ ス を 変 え る AI の 広 が り ― 』」 2026 年 。 <https://www.jst.go.jp/crds/column/kaisetsu/shortreport1.html>
- [3] JST 研究開発戦略センター (CRDS) 「ショートレポート AI for Science Spotlights Vol.2『AI エージェント の 科 学 研 究 へ の 展 開 』」 2026 年 8 月 7 日 。 <https://www.jst.go.jp/crds/column/kaisetsu/shortreport2.html>
- [4] Merchant, A. et al., "Scaling deep learning for materials discovery", Nature, Vol.624, 2023 年 11 月 29 日 (GNoME) 。 Google DeepMind 公式ブログ "Millions of new materials discovered with deep learning", 2023 年 11 月 29 日。
- [5] Szymanski, N. J. et al., "An autonomous laboratory for the accelerated synthesis of novel materials", Nature, Vol.624, 2023 年 11 月 29 日. DOI:10.1038/s41586-023-06734-w
- [6] Chemical & Engineering News (C&EN) による A-Lab 論文の訂正および GNoME 論文への指摘に関する報道、2026 年 1 月。Palgrave, R., Schoop, L. M. ほかによる反論プレプリントを含む。
- [7] Abramson, J. et al., "Accurate structure prediction of biomolecular interactions with AlphaFold 3", Nature, 2024 年 5 月。
- [8] Google DeepMind, "AI achieves silver-medal standard solving International Mathematical Olympiad problems", 2024 年 7 月 25 日。 <https://deepmind.google/discover/blog/ai-solves-imo-problems-at-silver-medal-level/>
- [9] Nature, AlphaProof に関する査読論文, 2025 年 11 月 12 日. DOI:10.1038/s41586-025-09833-y
- [10] Google DeepMind, AlphaEvolve white paper, 2025 年 5 月 14 日。 arXiv:2506.13131 (査読前テクニカルレポート) 。
- [11] "A generative AI-discovered TNIK inhibitor for idiopathic pulmonary fibrosis: a randomized phase 2a trial", Nature Medicine, 2025 年 6 月 3 日 。 Insilico Medicine プ レ ス リ リ ー ス 同 日 。 <https://insilico.com/news/tnrecuxsc1-insilico-announces-nature-medicine-publi>
- [12] The White House, Executive Order "Launching the Genesis Mission", 2025 年 11 月 24 日。
- [13] Yamada, Y. et al., "The AI Scientist-v2: Workshop-Level Automated Scientific Discovery via Agentic Tree Search", arXiv:2504.08066, 2025 年 4 月。 Sakana AI 公式発表を含む。
- [14] Google "AI co-scientist" に関する査読論文, Nature, 2026 年. DOI:10.1038/s41586-026-10644-y (JST-CRDS Vol.2 の引用[3]も参照) 。

- [15] King, S. et al., Science, 2026 年 8 月 6 日. DOI:10.1126/science.aec2657 (Evo 1/2 によるバクテリオファージ全ゲノム生成と実験検証)。
- [16] Anthropic, "How Claude is accelerating protein design and analytical chemistry", 2026 年 8 月 18 日。
<https://www.anthropic.com/>
- [17] Insilico Medicine, "Insilico Medicine Doses First Patient in GENESIS-IPF-3, the World's First Phase III Trial of a Generative AI-Driven Innovative Drug", 2026 年 9 月 9 日 (NCT07687459)。
<https://insilico.com/news/isn1009261-insilico-medicine-doses-first-patient-genesis-ipf-3>
- [18] Nature (ニュース記事), "OpenAI's Navier-Stokes claim", 2026 年 9 月。DOI:10.1038/d41586-026-02842-5。OpenAI 公式発表 (2026 年 9 月 6~8 日) を含む。
- [19] ナビエ-ストークス問題をめぐる OpenAI 発表の検証解説 (強制項付き設定であること、クレイ懸賞の対象範囲、クレイ数学研究所の扱いに関する整理)、2026 年 9 月。<https://miraflow.ai/blog/navier-stokes-ai-proof-controversy-openai-astra-explained-2026> ※二次情報。
- [20] TechCrunch, "OpenAI fought dirty on career-making math problem, says NYU mathematician", 2026 年 9 月 8 日。<https://techcrunch.com/2026/09/08/openai-fought-dirty-on-career-making-math-problem-says-nyu-mathematician/>
- [21] Kirgis, P., Kapoor, S., Narayanan, A.ほか (CRUX), "Can AI agents conduct open-ended AI research? Early evidence from two case studies", arXiv:2607.27191, 2026 年 8 月。<https://arxiv.org/abs/2607.27191>
- [22] JST 研究開発戦略センター (CRDS) 「ショートレポート AI for Science Spotlights Vol.4『米国ジェネシス・ミッシヨンが本格始動 (後編)』」2026 年。
<https://www.jst.go.jp/crds/column/kaisetsu/shortreport4.html>
- [23] JST 研究開発戦略センター (CRDS) 「ショートレポート AI for Science Spotlights Vol.3『米国ジェネシス・ミッシヨンが本格始動 (前編)』」2026 年。
<https://www.jst.go.jp/crds/column/kaisetsu/shortreport3.html>
- [24] 文部科学省「『AI for Science の推進に向けた基本的な戦略方針』の推進状況について」2026 年 5 月 21 日 (総合科学技術・イノベーション会議有識者議員懇談会資料 1)。
<https://www8.cao.go.jp/cstp/gaiyo/yusikisha/20260521/siry01.pdf>
- [25] 知的財産戦略本部「知的財産推進計画 2026~成長戦略を支える知財戦略の推進~」2026 年 6 月 12 日 (全 143 頁)。https://www.cas.go.jp/jp/seisakukaigi/titeki2/260612/keikaku_all.pdf
- [26] 知的財産戦略本部「知的財産推進計画 2026 (概要)」2026 年 6 月。
https://www.cas.go.jp/jp/seisakukaigi/titeki2/260612/keikaku_gaiyo_all.pdf
- [27] カレントアウェアネス・ポータル (国立国会図書館)「知的財産戦略本部、生成 AI の適切な利活用等に向けた知的財産の保護及び透明性に関するプリンシプル・コードを公表」2026 年 8 月 25 日付記事。
<https://current.ndl.go.jp/car/284024>
- [28] 内閣府知的財産戦略推進事務局「生成 AI の適切な利活用等に向けた知的財産の保護及び透明性に関する

プリンシプル・コード」 2026 年 8 月 25 日。
https://www.cas.go.jp/jp/seisakukaigi/titeki2/ai_kentoukai/kaisai/pdf/ai_principle_code.pdf

- [29] プリンシプル・コード（案）のパブリックコメント後修正に関する法律事務所解説、2026 年 8 月 18 日公表分。<https://www.lexology.com/library/detail.aspx?g=4c44ce66-1c60-400a-bba5-a5b678497bae> ※二次情報。
- [30] 内閣府「価値創造を加速する知財・無形資産投資・活用のガイダンス」2026 年 8 月 3 日。
https://www.cas.go.jp/jp/seisakukaigi/titeki2/tyousakai/tousi_kentokai/tousi_guidance.html
- [31] 中国国家知識産権局（CNIPA）「専利審査指南」改正（2023 年 12 月 21 日公布、2024 年 1 月 20 日施行）。
- [32] Beel, J. ほか, "Evaluating Sakana's AI Scientist: Bold Claims, Mixed Results, and a Promising Future?", arXiv:2502.14297, 2025 年 2 月。<https://arxiv.org/pdf/2502.14297>
- [33] Thaler v. Vidal, 43 F.4th 1207 (Fed. Cir. 2022)（連邦最高裁は 2023 年に上告を受理せず）。
- [34] USPTO, "Inventorship Guidance for AI-Assisted Inventions", 89 FR 10043, 2024 年 2 月 13 日。
- [35] USPTO, "Revised Inventorship Guidance for AI-Assisted Inventions", Federal Register, Docket No. PTO-P-2025-0014, 2025 年 11 月 28 日。<https://www.federalregister.gov/documents/2025/11/28/2025-21457/revised-inventorship-guidance-for-ai-assisted-inventions>
- [36] EPO 法律審判部 決定 J 8/20 および J 9/20、2021 年 12 月 21 日。
- [37] Thaler v Comptroller-General of Patents, Designs and Trade Marks [2023] UKSC 49, 2023 年 12 月 20 日。
- [38] ドイツ連邦通常裁判所（BGH）決定 X ZB 5/22、2024 年 6 月 11 日。
- [39] 東京地方裁判所判決 令和 6 年 5 月 16 日（令和 5 年（行ウ）第 5001 号）。
- [40] 知的財産高等裁判所判決 令和 7 年 1 月 30 日（令和 6 年（行コ）第 10006 号 出願却下処分取消請求控訴事件）。事案および判示の整理は、ユアサハラ法律特許事務所「人工知能ダバス（DABUS）に関する令和 7 年 1 月 30 日知財高裁判決と AI 発明に関する考察」<https://www.yuasa-hara.co.jp/lawinfo/5599/> も参照。
- [41] 最高裁判所第二小法廷決定 令和 8 年 3 月 4 日（上告棄却・不受理）。※判決文は未確認であり、報道に基づく。
- [42] Aboy, M., Liddell, K., "Revised USPTO guidance on inventorship for AI-assisted inventions: a pro-innovation pivot away from Pannu factors", Journal of Intellectual Property Law & Practice, 2026, jpag021. DOI:10.1093/jiplp/jpag021
- [43] Executive Order 14179, "Removing Barriers to American Leadership in Artificial Intelligence", 2025 年 1 月 23 日。
- [44] USPTO, "Updated Guidance for Making a Proper Determination of Obviousness", 89 FR 14449, 2024 年 2 月 27 日。<https://www.federalregister.gov/documents/2024/02/27/2024-03967/updated-guidance-for->

making-a-proper-determination-of-obviousness

- [45] USPTO, "Request for Comments Regarding the Impact of the Proliferation of Artificial Intelligence on Prior Art, the Knowledge of a Person Having Ordinary Skill in the Art, and Determinations of Patentability", 89 FR 25417, 2024 年 4 月 30 日。
- [46] IPWatchdog, "Is AI the Answer for Meeting Patent Enablement Requirements?", 2025 年 10 月 23 日。
<https://ipwatchdog.com/2025/10/23/ai-answer-meeting-patent-enablement-requirements/> ※二次情報。
Amgen Inc. v. Sanofi, 598 U.S. 594 (2023) を含む。
- [47] Carpmaels & Ransford, "Patenting AI: Artificial intelligence and drug discovery".
<https://www.carpmaels.com/patenting-ai-artificial-intelligence-and-drug-discovery/> ※二次情報。
- [48] 特許庁「AI 関連技術に関する特許審査の事例について」（2024 年 3 月 13 日に事例追加）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/ai_jirei.html
- [49] キリンホールディングス株式会社・株式会社 GenerativeX「キリンと GenerativeX、AI で研究開発の在り方を刷新 研究者と共創する新モデルを導入」2026 年 8 月 6 日。
https://www.kirinholdings.com/jp/newsroom/release/2026/0806_04.pdf
- [50] 株式会社 NTT データ「AI を活用した発明発掘を支援するサービスを先行提供開始」2026 年 8 月 19 日。
<https://www.nttdata.com/global/ja/news/topics/2026/081900/>