

AI 駆動科学革命：自動化された発見から新たな認識論へ

Gemini2.5

エグゼクティブサマリー

本レポートは、「AI 駆動科学」（AI-Driven Science）と称される、人工知能（AI）が触媒する科学研究のパラダイムシフトについて包括的な分析を提供する。AI が人間の研究者のための計算ツールから、仮説生成、実験計画、そして実行までを担う、ますます自律的なエージェントへと進化する過程を追跡する。本分析では、生命科学から数学に至るまでの分野を革命的に変えつつある AlphaFold や AlphaEvolve といった基盤的 AI モデルを詳述する。次に、科学的手法全体の自動化を 4 段階のフレームワークで描き出し、「自律走行型ラボ（self-driving laboratory）」の概念へと至る道筋を示す。この技術的変革は、科学的理解の本質、不透明な「ブラックボックス」システムによって生み出された知識の妥当性、そして科学的手法の定義そのものに関する深遠な哲学的問いとの対峙を強いる。本レポートは、二重用途（デュアルユース）のジレンマや、バイアスと再現性の課題を含む、重大な実践的、倫理的、安全保障上のリスクを批判的に検証する。最後に、この革命を世界的な地政学的文脈の中に位置づけ、日本、米国、中国の戦略的な政策およびインフラ投資を比較する。結論として、科学における中核的なパートナーとしての AI への軌道は明確であるものの、その成功裏かつ有益な統合には、科学者の役割の根本的な再考と、倫理、ガバナンス、教育のための新たな枠組みの緊急な開発が必要であると論じる。

第 1 章 発見のエンジン：科学における基盤的 AI モデル

本章では、AI 駆動科学の概念を導入し、AI が予測ツールから創造ツールへと進化したことを示す 2 つの画期的な事例、AlphaFold と AlphaEvolve を分析する。

1.1 AI 駆動科学（AI-Driven Science）の定義

AI 駆動科学とは、AI が単なるデータ分析ツールではなく、発見プロセスの能動的な推進力となる研究パラダイムである。これは、AI が得意とする膨大なデータ処理能力や演繹的予測能力と、人間が得意とする直感的な仮説創出を融合させることで、相乗効果的な研究モデルを生み出すものである¹。従来の計算科学が人間によって設計されたモデルやシミュレーションに依存していたのとは異なり、AI 駆動科学では AI 自身がパターンを発見し、仮説を立て、さらには次の研究ステップを設計することさえ可能になる。これにより、知識の再統合と新たな価値創造が導かれる可能性が示唆されている³。これは、人間主導の探求をコンピュータが補強する段階から、人間と AI の協働、あるいは AI 主導の発見プロセスへの移行を意味する。

1.2 ケーススタディ：AlphaFold と生命科学における革命

生命科学の分野では、タンパク質の立体構造を解明することが 50 年来の大きな課題とされてきた⁴。X 線結晶構造解析やクライオ電子顕微鏡といった従来の手法は、1 つのタンパク質あたり数年を要するだけでなく、Spring-8 のような巨大な加速器施設を必要とするなど、時間的にも費用的にも膨大なコストがかかるものだった⁴。

この状況を一変させたのが、Google DeepMind が開発した AlphaFold である。アミノ酸配列を入力するだけで、数十分から数時間でタンパク質の 3 次元構造を高精度に予測できるようになった⁴。これはまさに「ゲームチェンジャー」であり、生命科学の進展速度を劇的に加速させた⁴。

その技術的進化は目覚ましい。

- **AlphaFold 1 (2018 年):** 深層学習を用いてアミノ酸残基間の距離を推定し、配列データから距離マップを作成した⁸。
- **AlphaFold 2 (2020 年):** 改良された「Evoformer」という深層学習アーキテクチャを用い、単一タンパク質の構造予測において前例のない精度を達成し、中心的な課題を実質的に解決した⁹。その影響は 2 万件以上の論文で引用され、マラリアワクチンやがん治療の研究に貢献していることから明らかである⁹。
- **AlphaFold 3 (2024 年 5 月):** 単一タンパク質を超え、タンパク質と DNA、RNA、リガンド（薬剤を含む小分子）との相互作用を予測する能力を獲得した。これにより、相互作用予測の精度は既存手法に比べて少なくとも 50% 向上した⁸。これは静的な構成要素の予測から、創薬に不可欠な動的な生命機械のモデリングへの移行を意味する⁴。

さらに、DeepMind と欧州分子生物学研究所（EMBL-EBI）の協力による AlphaFold Protein Structure Database は、科学的に知られているほぼ全てのタンパク質を含む 2 億以上の構造予測データをオープンアクセスで提供している¹¹。これにより、高価で時間のかかる実験的決定というボトルネックが解消され、構造生物学の知見が民主化された。

1.3 ケーススタディ：AlphaEvolve と新たな論理の創造

AlphaFold が自然界に存在するものを「予測」するのに対し、AlphaEvolve は新しく改良された解決策を「創造」する。これは、アルゴリズムやコードベース自体を設計し、洗練させる AI エージェントである⁴。

そのメカニズムは、Gemini のような大規模言語モデル（LLM）の創造的可能性と進化的フレームワークを組み合わせたものである。LLM がコードの変更案を提案し、自動評価機がそれをスコアリングし、最も有望なアイデアが次の世代の改善にフィードバックされる。このサイクルを繰り返すことで、より優れた解決策へと「進化」させていく¹⁶。

AlphaEvolve が示した成果は、AI が抽象的かつ論理的な領域で推論し、創造する能力を持つことを証明している。

- **数学:** 50 以上の数学の未解決問題に取り組み、そのうちの 20% で既存の最良解を上回り、75% で最先端の解を自力で再発見した⁴。特に、56 年間破られていなかった 4x4 行列乗算に関するシュトラッセンのアルゴリズムを改善したことは画期的である¹⁹。
- **コンピュータ科学・工学:** Google のデータセンターのスケジューリングを最適化し、全世界の計算資源の 0.7% を回収した。また、自身が稼働する基盤である TPU（Tensor Processing Unit）のハードウェア回路設計を改良し、Gemini モデル自体の学習カーネルを最適化して学習時間を短縮した⁴。

AlphaFold と AlphaEvolve を比較すると、科学における AI の役割の明確な進化の軌跡が見えてくる。それは、自然現象の強力な「予測者」（このタンパク質の構造は何か？）から、新規で最適化された人工物の創造的な「生成者」（このタスクのためのより良いアルゴリズムは何か？）への移行である。これは、AI が分析ツールから創造的なパートナーへと、科学的ワークフローにおけるその地位を根本的に変えつつあることを示唆している。さらに、AlphaEvolve が自身の学習を支える Gemini や TPU を最適化できるという事実は、強力な再帰的自己改善ループの具体的な現れである。AI 開発のツールが AI 自身によって改良されるこのサイクルは、もはや理論上の概念ではなく、工学的に実証された現実であり、今後の AI 能力を人間主導の研究開発サイクルを凌駕するペースで加速させる主要な要因となる可能性がある。

第2章 科学的手法の自動化：4 段階の進化

本章では、安野貴博氏が提唱する4段階のフレームワークを基盤とし、科学的ワークフロー全体がどのように自動化されつつあるかを、多様な分野からの証拠を統合して詳細に分析する

⁴。

2.1 レベル1：研究者のアシスタントとしての AI

この段階では、AI は研究ワークフロー内の個別のタスクにおいて人間を補助する。これは今日、科学の世界で最も広く普及し、確立された AI の利用形態である⁴。

- **コード生成とデータ分析:** 研究者は AI を用いて統計分析用のコードを生成し、時間を節約するとともに、複雑な計算作業への参入障壁を下げている⁴。
- **文献レビューと統合:** AI ツールは、人間には不可能な規模で膨大な数の論文を処理し、パターンを特定して知識を統合する、自動化された文献レビューやメタアナリシスを実行できるようになった²⁰。Iris.ai や Semantic Scholar のようなプラットフォームがその代表例である。
- **論文執筆支援:** AI は科学論文の草稿作成、編集、推敲を支援する⁴。
- **問題解決:** MathGPT や AI Math Solver のようなツールは、AI チューターや計算機として機能し、学生や研究者が複雑な方程式を解き、概念を理解するのを助けている²¹。

2.2 レベル2：インシリコ科学者としての AI

この段階では、AI は物理的な実験を必要としない研究領域、すなわちプロセス全体がコンピュータ内で完結する分野で主導的な役割を果たす⁴。

- **コンピュータ科学と AI 研究:** レベル2 の最も肥沃な土壌である。新しい AI モデルの開発は計算タスクであり、AI 主導の研究に最も適している。AlphaEvolve が示した再帰的自己改善ループは、この究極の例である⁴。
- **計算生物学:** AlphaFold による予測はレベル2 の活動である。AI は物理的な実験なしに配列データを分析し、構造仮説を生成する⁴。
- **創薬（スクリーニング）:** AI プラットフォームは、合成可能な化合物の数兆個に及ぶ膨大な仮想ライブラリをスクリーニングし、有望な新薬候補を特定できる。これは物理的なハ

イスループットスクリーニングでは不可能なタスクである²³。

- **材料科学（予測）**：GNoME のような AI モデルは、数百万の新しい安定した材料の構造を予測するために使用され、実際に合成される前に潜在的な材料の既知の展望を大幅に拡大した²⁵。

2.3 レベル 3：クラウドラボを介した物理世界とのループにおける AI

AI は、完全に自動化され、遠隔制御される物理的な実験室で実験を指示することにより、純粋な計算タスクを超えていく。これにより、デジタルの仮説と物理的な検証の間のループが閉じられる⁴。この飛躍を可能にするのが「クラウドラボ」である。

- **コンセプト**：ユーザーがインターネット経由で遠隔からアクセスする、中央集権的で高度に自動化された実験室。科学者は実験プロトコルをコードとして記述・送信し、ロボットが 24 時間 365 日実験を実行する⁴。
- **主要な事例**：
 - **Emerald Cloud Lab (ECL)**：数百万ドル相当の機器を備えた最先端の生命科学ラボで、遠隔アクセスが可能。全ての生命科学実験の 100% 提供を目指しており、実験費用はわずか 25 ドルからという例もある²⁶。これにより、スタートアップや学術ラボの初期設備投資を数十万ドルからゼロにまで劇的に削減する²⁸。
 - **Strateos**：「サービスとしての自動化（Automation - as-a-Service）」を提供し、ラボを「スマートなデータ生成センター」に変える。創薬や合成生物学のワークフローへの遠隔アクセスを提供している²⁷。
- **用語に関する注意**：「クラウドラボ」という用語は、マーケティングやオンライン学習プラットフォームを指す文脈でも使用されることがあるが³⁰、本レポートにおける定義は、ECL に代表される完全自動化科学実験室である。

レベル 2（インシリコ）からレベル 3（物理実験）への移行は、単なる段階的なステップではなく、デジタル世界と物理世界の間の決定的な境界を越えることを意味する。クラウドラボは、この飛躍を可能にする不可欠なインフラ、いわば「物理世界への API」である。科学的真理は最終的に経験的検証を必要とするが、クラウドラボは計算 AI の絶大な力を実験科学の現実 に接続する物理的・デジタル的インターフェースを提供する。このインフラがなければ、次のレベル 4 の「自律走行型ラボ」は理論上の概念にとどまるだろう。

2.4 レベル 4：自律走行型ラボと自律的フィールド科学

これは AI 駆動科学の最終段階であり、AI エージェントがロボティクスと統合され、閉ループシステムを形成する。AI は自律的に仮説を立案し、ロボットを用いて実験を計画・実行し、結果を分析し、その新しい知識を用いて次の実験サイクルを人間の介入なしに設計する⁴。

- **実験室応用（自律走行型ラボ）：**

- **材料科学:** この分野はレベル 4 自動化の最前線にある。

- **A-Lab (バークレー研究所):** AI アルゴリズムが新しい化合物を提案し、ロボットがそれらを合成・試験するという緊密なループを回すことで、電池や電子機器用材料の発見を劇的に加速させている³⁴。

- **CREST (MIT):** 文献や画像からの情報を取り込み、ベイズ最適化を用いて実験を計画し、ロボットが材料を合成・試験するプラットフォーム。燃料電池用の新規触媒を発見し、ドルあたりの出力密度を 9.3 倍向上させた³⁵。

- **フィールド応用（自律的フィールドロボティクス）：**

- **コンセプト:** レベル 4 は実験室に限定されない。ロボティクスにより、AI は実験台を超えて、複雑で制御されていない現実世界の環境で研究を行うことが可能になる⁴。

- **生態学・環境科学:** 自律型ロボット（ドローン、AUV）が、危険な場所や遠隔地（深海、アマゾン熱帯雨林など）での種のモニタリングやデータ収集、生息地の回復に使用されている。AI/ML を活用してセンサーデータをリアルタイムで分析し、洞察を生み出す³⁶。

- **深海・宇宙探査:** ウッズホール海洋研究所の「CUREE」のような AI 駆動ロボットは、海洋生物を自律的に識別・追跡でき、他のロボットは海底火山を探査する³⁹。惑星探査ローバーは、通信の遅延やアクセス不能な環境のため、自律的に判断を下しながら他の惑星で地質学的分析を行う⁴⁰。

レベル 4 への進展は、単なる AI の物語ではない。それは、意思決定のための AI、物理的実行のためのロボティクス、そして大規模で標準化されたデータセットを生成するためのハイスループット実験という、3 つの異なる技術分野の強力な融合の結果である。AI が「脳」として実験を設計し、ロボットが「手」として合成や特性評価を行い、ハイスループットなプラットフォームが次の AI モデルの学習に適した速度と規模でデータを生成する。これにより、「より多くの実験がより多くのデータを生み、それが AI をより賢くし、より効果的な実験の設計を可能にする」というフライホイール効果が生まれる。この融合こそが、レベル 4 科学における加速のエンジンなのである。

2.5 表 1: 科学研究における AI の 4 段階フレームワーク

レベル	AI の主な役割	主要な実現技術	具体例（出典）
1	アシスタント	大規模言語モデル（LLM）、生成的 AI	文献統合 ²⁰ 、コード生成 ⁴ 、論文執筆支援 ⁴
2	インシリコ科学者	生成モデル、大規模シミュレーション	AlphaFold タンパク質予測 ⁴ 、AlphaEvolve アルゴリズム発見 ⁴ 、仮想創薬スクリーニング ²³
3	遠隔実験者	クラウドラボプラットフォーム、API 制御ロボティクス	Emerald Cloud Lab ²⁶ 、Strateos ²⁷
4	自律的発見者	閉ループ AI エージェント、フィールドロボティクス	CRESt 材料発見 ³⁵ 、A-Lab ³⁴ 、自律的深海探査 ³⁹

第 3 章 科学的「理解」におけるパラダイムシフト

本章では、AI 駆動科学が、知識、発見、そして科学的手法そのものといった従来の概念に対して突きつける、深遠な哲学的・方法論的課題に正面から向き合う。

3.1 人間の制約を超える：高次元空間の航行

人間の脳は、3 次元を超える空間を直感的に把握する能力に根本的な限界がある。これは科学理論における認知的な制約となってきた⁴。一方で、コンピュータは何百、何千という次元で動作し、人間が視覚化したり直感したりすることが不可能な複雑なデータセットの中からパタ

ーンや相関関係を見つけ出すことができる。これにより、AI は効果的ではあるが人間中心的な意味での「理解」が難しい解決策を発見することが可能になる⁴。科学は、我々の最も強力な理論や発見が、もはや人間の心が直感的に理解できるものに制約されない時代に突入しつつあるのかもしれない。

3.2 性能と人間的理解の分離

歴史的に、人間の「わかった感」は、世界を正しくモデル化していることの良い代理指標であり、より良い成果につながってきた。しかし、AI 駆動科学はこの結びつきを断ち切る。モデルは、人間に理解できなくても、非常に効果的な結果を生み出すことができる⁴。

その好例が、AlphaGo がイ・セドル氏に勝利した対局である。AlphaGo は、プロの囲碁棋士が「悪手」や「理解不能」と見なした手を打ったが、それらの手は長期的には戦略的に優れていることが証明された。人間の専門家はその手の背後にある論理を「理解」できなかったが、AI はその優れた性能を実証した⁴。

安野氏は、この「わかった感」を優先することの危険性を説明するために、「地球平面説」の類推を用いる。地球平面説は直感的に単純で強い「なるほど」感を与えるが、現実のモデルとしては不十分である。対照的に、AI は直感に反し不透明に感じられるかもしれないが、現実を完璧に予測するモデルを生み出す可能性がある⁴。

このことは、科学の目標としての人間的な「理解」の価値が低下する可能性を示唆している。新たな評価基準は、純粹に実用的なもの、すなわち「そのモデルは機能するか?」「有用で検証可能な予測や技術を生み出すか?」ということになるかもしれない。

3.3 「ブラックボックス」の難問と科学的手法

特に深層ニューラルネットワークのような多くの先進的な AI モデルは「ブラックボックス」である。我々は入力と出力を見ることはできるが、その内部の推論プロセス、すなわち何百万もの重み付けされた結合は不透明であり、人間が理解できる言葉で解釈することはできない⁴。

この不透明性は、科学的手法の核心的な信条に挑戦する。

- **説明:** 科学は「何が」起こるかだけでなく、「なぜ」そうなるのかを問う。ブラックボックスは正しい予測（例：このタンパク質はこのように折りたたまれる）を提供できても、

その因果的・機械論的な説明（それがそのように折りたたまれる物理的原理）を提供することはできない⁴⁴。

- **検証と信頼:** AI がどのように結論に至ったのかが分からなければ、それを信頼したり、それが正しい理由で正しいと検証したりすることは困難である。モデルはデータ内の見せかけの相関やアーティファクト（「賢いハンス効果」として知られる現象）を拾っている可能性があり、その結果は頑健性や一般化可能性に欠けるかもしれない⁴⁶。
- **再現性:** AI モデルとその学習環境の複雑さは、真の再現性を大きな課題としている。わずかな変動が異なる結果につながる可能性がある⁴⁷。

説明可能な AI (XAI) という研究分野は、ブラックボックスの内部を覗き込み、機械が「どの」データを重要と見なしたかを理解し、「なぜ」かを推測するための技術開発を目指している⁴⁴。しかし、これらは依然としてモデルの真の論理の近似に過ぎない。

3.4 認識論のフロンティア：AI 時代の「知識」とは何か？

ある研究論文は、AlphaFold がもたらす興味深い認識論的トリレンマを提示している。我々は以下の 3 つの主張のうち 1 つを否定しなければならない。(1) AlphaFold は科学的知識を生み出す、(2) 予測だけでは、確立された科学的原理から導出可能でない限り科学的知識ではない、(3) 科学的知識は強く不透明であってはならない⁵。

AI は、従来の認識論に挑戦する規模で仮説や発見を生み出している。これらのシステムはパターンを発見できるが、その科学的価値は厳密な検証に依存する⁴⁹。ブラックボックスからの信頼性の高い予測は「知識」と見なされるのか、それとも単に高度な計測器の出力に過ぎないのか。発見のプロセス自体が、人間中心の観察、仮説、検証から、人間と AI の協働プロセス、あるいは完全に自動化されたプロセスへと移行している。AI の哲学は、「機械は創造的でありうるか？」「機械は真に『理解』できるか？」といった問いに取り組まざるを得なくなっている⁵⁰。

強力だが不透明な AI モデルの台頭は、科学の哲学における分裂の可能性を強めている。伝統的に、予測と説明は絡み合った目標であった。AI 駆動科学は、この 2 つを根本的に分離することを可能にする。我々は、実用的で予測的な力を優先する科学（そしてブラックボックスを許容する）と、因果的で機械論的な理解を引き続き要求する科学（そして不透明な手法を拒絶するかもしれない）という、2 つの並行した科学のトラックが存在する時代に入りつつあるのかもしれない。これは単なる技術的な問題ではなく、我々が科学の「目的」と考えるものにおける根本的な分岐である。

この認識論的ジレンマに対する一つの解決策は、AI に対する我々の概念を再構築することであ

る。AI を人工的な科学者や「知る者」と見なす代わりに、それを信じられないほど強力な新しい種類の科学的「測定器」として捉える方がより正確かもしれない。我々は顕微鏡や粒子加速器に、それらが生成するデータを「理解」することを求めない。それらは人間の感覚を拡張し、人間が解釈して知識を構築するためのデータを生成するツールと見なされている。同様に、ブラックボックス AI も、入力（生物学的配列、データセット）を受け取り、出力（予測された構造、新しい仮説）を生成する測定器と見なすことができる。この視点は、AlphaFold が人間の意味でタンパク質の構造を「知っている」のではなく、非常に信頼性の高い予測を生成する測定器であると位置づけることで、前述のトリレンマを解決する助けとなる。この場合、「知識」は、人間の科学者がその出力を受け取り、実験的に検証し、既存の生物学的理論の網の中に統合する時に創造される。この視点は、AI「測定器」の超人的な能力を最大限に活用しつつ、知識の最終的な裁定者としての人間科学者の役割を維持するものである。

第4章 リスクと倫理的フロンティアの航行

本章では、AI 駆動科学の責任ある発展のために取り組まなければならない、実践的、倫理的、そして安全保障上の課題について批判的な評価を行う。

4.1 実践的な障壁：自動化ラボの現実

- **統合と相互運用性:** 新しい自動化システムを既存のレガシー機器と統合することは大きな課題である。互換性の欠如は、データのサイロ化やワークフローの混乱を招く可能性がある⁵¹。
- **コストとスケーラビリティ:** 自動化は長期的なコスト削減をもたらす可能性があるが、ハードウェアとソフトウェアへの初期投資は法外なものになり得る。段階的な導入がしばしば必要となる⁵²。
- **メンテナンスと陳腐化:** 自動化システムは、エラーの急速な伝播を防ぐために、慎重なメンテナンスと校正を必要とする。さらに、急速な技術進歩により、高価な機器がすぐに陳腐化する可能性もある⁵³。
- **データ管理:** 自動化システムは膨大な量のデータを生成し、保管、セキュリティ、完全性の面で課題を生じさせる。不適切なデータ管理は、不正確な結果につながる可能性がある⁵¹。
- **人的要因:** 手作業に慣れた人員からの変化への抵抗や、自動化システムを操作・管理するための新しいスキルセットの必要性は、重大な組織的障壁となる⁵¹。

4.2 AI 科学における再現性とバイアスの危機

- **バイアスの増幅:** AI モデルは学習データに依存する。データに歴史的なバイアス（例：医療データにおける特定の人種や性別への偏り）が含まれている場合、AI はこれらのバイアスを学習し、増幅させる可能性がある。これは、不公平または不正確な科学的結論につながる⁴⁷。
- **再現性の課題:** AI の学習プロセスにおける非決定的な性質は、ソフトウェアのバージョンやハードウェアの違いと相まって、AI による発見の再現を困難にし、科学的手法の礎を揺るがす⁴⁷。
- **ハルシネーション（幻覚）と偽陽性:** 生成 AI や LLM は、自信に満ちているが虚偽の情報（例：文献レビューにおける捏造された引用）を「幻覚」として生成することがある。これは、人間の専門家による慎重な検証がなければ、科学的記録を汚染する可能性がある⁵⁷。

従来のデータサイエンスの格言「Garbage In, Garbage Out（ゴミを入れればゴミが出る）」は、AI 駆動科学においてより危険な側面を帯びる。AI モデルの「ブラックボックス」性とその權威の増大により、このリスクは「偏ったデータが入り、疑われることのない科学的『真実』が出る（Biased Data In, Unquestioned Scientific 'Truth' Out）」という形に変化する。AI モデルが学習データからバイアスを受け継ぐことは知られており、その不透明性ゆえに、なぜ特定の結論に至ったのかを特定することは困難である。AlphaFold のようなツールが成功を収め、広く統合されるにつれて、人間が機械の出力を過度に信頼する「自動化バイアス」の危険性が生じる。これらの要因が組み合わさることで、偏ったモデルが科学的に厳密に見える発見を生み出すという新たな増幅されたリスクが生まれる。そのプロセスが不透明かつ自動化されているため、根底にあるバイアスが見過ごされ、その発見が客観的な真実として受け入れられ、科学文献やその後の応用（例：臨床実践）にバイアスが定着してしまう可能性がある。

4.3 二重用途（デュアルユース）のジレンマ：発見と破壊のための AI

- **核心的な脅威:** 新薬の発見のような有益な目的のために設計された AI 技術が、悪意のある目的のために容易に転用されうる。これが「二重用途（デュアルユース）」のジレンマである⁵⁹。
- **戦慄の概念実証:** ある国際安全保障会議でこのリスクが探求された際、創薬 AI に毒性分子を設計するタスクが与えられた。目的を「無毒」から「有毒」に変えるだけで、AI は 6 時間以内に VX ガスのような神経剤を含む 4 万もの潜在的な化学兵器を提案した⁵⁹。

- **参入障壁の低下:** これらの化合物を合成するには依然として専門知識が必要だが、AI は新しい脅威を「設計」するための障壁を劇的に下げる。これにより、兵器開発の「着想」段階が、はるかに広範な主体にとってアクセス可能になる⁵⁹。

この二重用途のジレンマは、偶発的な欠陥ではなく、生成 AI を科学にとって非常に強力なものにしているその能力に内在する、根本的かつ不可避の結果である。有益な分子を見つけるために広大な化学的または生物学的な「解空間」を探索する能力は、有害なものを見つける能力と全く同じである。創薬 AI の事例が示すように、変更されたのは最適化目標（低毒性から高毒性へ）だけであり、中核となるアルゴリズムとその能力は同じままであった。これは、科学における生成 AI が、ある領域の根本的なルールを学習し、そのルールを用いて特定の目的に合致する新規の出力を生成することによって機能するためである。したがって、ある科学領域のための強力な生成モデルは、定義上、潜在的な二重用途技術であると言える。「ガードレール」は技術自体にあるのではなく、人間のユーザーが提供する意図（目的関数）にある。このリスクを軽減するためには、純粋に技術的な解決策だけでは不十分であり、これらの強力な生成ツールの使用方法に関する堅牢なガバナンス、アクセス制御、および監視が必要となる。

第5章 科学 AI をめぐる世界的な競争：政策、インフラ、戦略

本章では、AI 駆動科学革命の地政学的文脈を分析し、この新時代のリーダーを決定づけるであろう国家戦略とインフラ投資を比較する。

5.1 国家戦略の比較分析

- **日本:**
 - **哲学:** 厳格な規制よりも企業による自主的なガイドラインを重視する「ソフトロー」アプローチを通じて、イノベーションと産業競争力を促進することに重点を置いている⁶²。目標は「世界で最も AI フレンドリーな国」になることである⁶³。
 - **法整備:** 「人工知能関連技術の研究開発及び活用の推進に関する法律」（2025 年 5 月）は、リスクを管理しつつイノベーションを促進するための枠組みを確立し、内閣総理大臣を議長とする中央集権的な AI 戦略本部を設置した⁶³。
 - **研究機関:** 理化学研究所（RIKEN）の革新知能統合研究センター（AIP）が主要な役割を担い、AI 基盤技術の開発、戦略的分野（医療、材料）における科学研究の加速、

ELSI（倫理的・法的・社会的課題）の研究に取り組んでいる⁶⁵。

- **米国:**
 - **哲学:** 研究開発におけるリーダーシップを維持し、AI 人材を育成するために、AI リソースへのアクセスを民主化することに焦点を当てている。戦略は学術・研究コミュニティの強化に大きく傾いている⁶⁷。
 - **イニシアティブ:** 2024 年に開始された国家 AI 研究リソース（NAIRR）パイロットプログラム。これは、政府機関（NSF、DOE、NIH）や民間パートナーからの計算資源、データ、モデル、トレーニングリソースに研究者を接続する共有国家インフラである⁶⁷。
 - **ガバナンス:** 2020 年の国家 AI イニシアティブ法と、NAIRR の立ち上げを命じた 2023 年の「安全、安心、信頼できる AI」に関する大統領令によって推進されている⁶⁸。
- **中国:**
 - **哲学:** 2030 年までに AI における世界のリーダーになるという明確な目標を掲げた、国家主導のトップダウン戦略。「新世代 AI 発展計画」や「中国製造 2025」にその概要が示されている⁷⁰。
 - **戦略:** AI の自給自足、国産モデル（例：DeepSeek）の開発、全セクターへの AI 統合に焦点を当てている。デジタルシルクロード構想を通じて、自国の AI 技術とガバナンスモデルを輸出し、世界的な影響力を拡大している⁷⁰。

日本と米国の戦略を比較すると、AI エコシステム育成に対するアプローチに根本的な違いが見られる。日本の「ソフトロー」フレームワークは、経済成長と競争力を優先し、産業界がイノベーションを起こしやすい低摩擦の環境を作り出すことを目的としている。対照的に、米国の NAIRR は、基礎科学の進歩と人材育成を優先し、学術・非営利セクターが AI の莫大なリソース要件に取り残されないようにするための、政府主導の構造化された取り組みである。日本は産業界が先導することを期待する「パーミッションレス・イノベーション」モデルを追求し、米国は公的インフラを用いてアクセスを民主化する「構造化アクセス」モデルを追求している。国家の AI 能力を育成するためのこれら二つの異なる賭けは、長期的な結果をもたらすだろう。

5.2 基盤層：計算インフラの必要性

最先端の科学研究を行う上で、特に GPU クラスタのような大規模な計算資源へのアクセスが前提条件となりつつある⁴。これは AI 駆動科学が構築される基盤インフラであり、「GPU は新しい顕微鏡」と表現できる。

- **日本の投資:**
 - **ABCI (AI Bridging Cloud Infrastructure):** 産業技術総合研究所（AIST）が運用す

る、日本最大級の AI 特化型スーパーコンピュータ。産業界や学界向けのオープンプラットフォームとして設計され、2025 年には ABCI3.0 が稼働開始するなど、継続的にアップグレードされている⁷¹。

- **ABCI-Q:** 2025 年に稼働予定の新しいシステムで、2,000 基以上の NVIDIA H100 GPU を搭載し、特に量子コンピューティング研究を推進するために設計されている⁷⁴。
- **東京大学:** Wisteria/BDEC-01 システムを運用。これは、従来のシミュレーション用 CPU (A64FX) と、データ・学習タスク用の大規模な NVIDIA A100 GPU パーティションを組み合わせたユニークな構成を特徴とする⁷⁵。

GPU を多用するスーパーコンピュータへの大規模な国家投資は、単に科学を可能にするためだけではない。それは国家安全保障と経済主権に関わる問題である。21 世紀において、国家が独立して最先端の研究を行い、独自の基盤モデルを開発する能力は、エリート級の計算資源を自国で所有し、制御できるかどうかに関係している。ABCI の目的が日本の AI 研究開発を加速し、外国のクラウドプロバイダーへの依存を減らすことにありと明記されているように⁷²、「主権 AI (Sovereign AI)」の概念はこれを地政学的戦略に直接結びつけている⁷⁶。これらのスーパーコンピュータは、科学的装置以上の存在であり、かつての時代の航空母艦や粒子加速器に匹敵する戦略的国家資産である。高性能計算資源のグローバルな分布とアクセスは、今後数十年の地政学的状況を決定づける特徴となるだろう。

5.3 表 2 : 科学のための国家 AI 戦略の比較分析

国	指導理念	主要な法律/イニシアティブ	主導機関	インフラの焦点
日本	産業主導、「ソフトロー」によるイノベーション	AI 推進法 (2025 年)	AI 戦略本部、産総研、理研	主権スーパーコンピュータの開発 (ABCI)
米国	研究・教育のためのアクセス民主化	国家 AI イニシアティブ法 (2020 年) / NAIRR	NSF、DOE、ホワイトハウス OSTP	官民計算資源への連合型アクセス (NAIRR)

中国	国家主導による世界的リーダーシップの追求	新世代 AI 発展計画 (2017 年)	国務院、科学技術部	AI ハブへの大規模な国家主導投資
----	----------------------	----------------------	-----------	-------------------

第 6 章 結論：科学と科学者の未来

本章では、レポートの調査結果を統合し、変容した科学の展望と、進化し続ける人間研究者の役割について、未来を見据えた視点を提供する。

6.1 統合：人間と AI の協働という不可避のパラダイム

AI がアシスタント（レベル 1）から自律的な発見のパートナー（レベル 4）へと進化する道のりを振り返ると、これが未来の可能性ではなく、材料科学や創薬といった先駆的な分野における現在の現実であることがわかる。AI の統合は、科学的探求の速度、規模、そして本質そのものを根本的に再構築している。

6.2 進化する人間科学者の役割

- **データ生成者から問いを立てる者へ:** AI とロボティクスがデータ収集と分析を自動化するにつれて、人間の主要な役割は「上流」へとシフトする。最も価値のある人間の貢献は、洞察に満ちた問いを立て、重要な問題を定義し、研究全体の方向性を設定することになるだろう。
- **技術者から指揮者兼倫理学者へ:** 科学者は、AI ツールとロボットシステムのオーケストラを指揮する「指揮者」となる。これには、システム思考、プロンプトエンジニアリング、自動化プラットフォームのための実験計画といった新しいスキルが求められる。そして決定的に重要なのは、人間が倫理的な裁定者となり、これらの強力なツールの責任ある使用を確保し、バイアスや誤用を防ぐ役割を担うことである。
- **データ解釈者から不透明なモデルの解釈者へ:** 科学者の役割には、ブラックボックスモデ

ルの出力と格闘することが含まれるようになる。彼らは「AI フォレンジック」の専門家となり、XAI ツールや巧妙な実験設計を用いて、AI パートナーによる発見を検証し、それに異議を唱え、最終的に信頼を構築する必要がある。新たな発見を人間が持つ広大な科学的知識の網の中に統合するという最終的な行為は、依然として人間特有のタスクであり続けるだろう。

6.3 AI 時代に向けた戦略的提言

- **資金提供機関および政策立案者へ:**
 - **基盤インフラへの投資:** 主権的な計算資源（GPU クラスタ）やオープンデータプラットフォームへの持続的かつ大規模な投資を、重要な国家インフラとして優先する。
 - **新たなガバナンスフレームワークの開発:** 従来の倫理指針を超え、透明性、再現性、バイアス、二重用途のジレンマに対処する AI 駆動科学に特化したガイドラインを作成する。これらの基準に関する国際協力を支援する。
 - **学際的研究への資金提供:** AI、ロボティクス、特定の科学分野の交差点に位置する研究、並びに科学における AI の哲学・倫理に関する研究を奨励する。
- **学術機関および教育者へ:**
 - **STEM 教育の改革:** AI、データサイエンス、ロボティクスの基礎的なリテラシーを含むようにカリキュラムを更新する必要がある。焦点は、機械的な計算から、批判的思考、創造的な問題解決、技術倫理へと移行すべきである。
 - **人間と AI の協働スキル育成:** 次世代の科学者に、AI ツールを単に「使う」だけでなく、効果的かつ批判的に「協働」するための訓練を行う。
- **研究者へ:**
 - **学際的スキルの習得:** 科学者は、これらのツールを効果的に活用し、その落とし穴を避けるために、AI/ML の基礎に精通する必要がある。
 - **透明性とオープン性の擁護:** 再現性の危機に対処するため、コード、データ、モデルを共有するためのベストプラクティスを採用する。ブラックボックスモデルの出力を無批判に受け入れるのではなく、その課題に積極的に取り組む。

引用文献

1. smeai.org, 10 月 14, 2025 にアクセス、<https://smeai.org/index/ai-driven-science-human-role/#:~:text=%E4%BA%BA%E5%B7%A5%E7%9F%A5%E8%83%BD%E9%A7%86%E5%8B%95%E7%A7%91%E5%AD%A6%E3%81%A8%E3%81%AF%E3%80%81AI%E3%81%8C%E5%BE%97%E6%84%8F%E3%81%A8,%E3%81%99%E3%82%8B%E5%8F%AF%E8%83%BD%E6%80%A7%E3%81%8C%E3%81%82%E3%82%8A%E3%81%BE%E3%81%99%E3%80%82>

2. 人工知能駆動科学と人間の役割, 10 月 14, 2025 にアクセス、
<https://smeai.org/index/ai-driven-science-human-role/>
3. AI ロボット駆動科学 - 連続セミナー2025「AI が拓く次世代イノベーション」 - 情報処理学会, 10 月 14, 2025 にアクセス、
<https://www.ipsj.or.jp/event/seminar/2025/program07.html>
4. 科学者がいない時代が来る？AI が“理解を超える”世界へ | 安野貴博の未来予測 - YouTube.docx
5. The Epistemology of AI-driven Science: The Case of AlphaFold - PhilSci-Archive, 10 月 14, 2025 にアクセス、
https://philsci-archive.pitt.edu/26659/1/AlphaFold_preprint%20%282025%29.pdf
6. AlphaFold を理解したいけど生物学系の知識がないので勉強してみた 前編：事前知識、背景理解, 10 月 14, 2025 にアクセス、
<https://tech.fusic.co.jp/posts/2021-01-08-alphafold/>
7. AlphaFold ってなに？ - Zenn, 10 月 14, 2025 にアクセス、
https://zenn.dev/google_cloud_jp/articles/whatisalphafold
8. AlphaFold - Wikipedia, 10 月 14, 2025 にアクセス、
<https://en.wikipedia.org/wiki/AlphaFold>
9. AlphaFold 3 predicts the structure and interactions of all of life's molecules - Google Blog, 10 月 14, 2025 にアクセス、
<https://blog.google/technology/ai/google-deepmind-isomorphic-alphafold-3-ai-model/>
10. Google DeepMind Announces World-Leading AlphaFold 3 Model | TIME, 10 月 14, 2025 にアクセス、
<https://time.com/6975934/google-deepmind-alphafold-3-ai/>
11. AlphaFold Protein Structure Database, 10 月 14, 2025 にアクセス、
<https://alphafold.ebi.ac.uk/>
12. AlphaFold とは？AI によるタンパク質構造予測の革命と今後の展望を徹底解説 | AI 総合研究所, 10 月 14, 2025 にアクセス、
<https://www.ai-souken.com/article/what-is-alphafold>
13. AlphaEvolve - Wikipedia, 10 月 14, 2025 にアクセス、
<https://en.wikipedia.org/wiki/AlphaEvolve>
14. AlphaEvolve - AI Powered Coding Agent For Designing Algorithms, 10 月 14, 2025 にアクセス、
<https://alphaevolveai.dev/>
15. AlphaEvolve とは？Google DeepMind 開発の自ら進化する AI エージェントの仕組み・従来のコード生成 AI との違いを徹底解説！, 10 月 14, 2025 にアクセス、
<https://ai-market.jp/technology/alphaevolve/>
16. AlphaEvolve: A Gemini-powered coding agent for designing advanced algorithms, 10 月 14, 2025 にアクセス、
<https://deepmind.google/discover/blog/alphaevolve-a-gemini-powered-coding-agent-for-designing-advanced-algorithms/>
17. Google: AlphaEvolve - Gemini を活用した次世代アルゴリズム設計について - note, 10 月 14, 2025 にアクセス、
<https://note.com/repkuririn7/n/nc082aab192ff>
18. [技術紹介] AI がアルゴリズムを進化させる未来：AlphaEvolve が拓く、問題解

- 決の新境地, 10 月 14, 2025 にアクセス、<https://jobirun.com/google-alphaevolve-future-of-problem-solving/>
19. AlphaEvolve とは? Google DeepMind 発、アルゴリズム発見 AI の全貌を徹底解説 | AI 総合研究所, 10 月 14, 2025 にアクセス、<https://www.ai-souken.com/article/what-is-alphaevolve>
 20. AI, agentic models and lab automation for scientific discovery — the beginning of scAInce, 10 月 14, 2025 にアクセス、<https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2025.1649155/full>
 21. AIMath Solver & Calculator - Free Online, No Sign-up - NoteGPT, 10 月 14, 2025 にアクセス、<https://notegpt.io/ai-math-solver>
 22. MathGPT - AIMath Solver - Math Solver & Homework Helper, 10 月 14, 2025 にアクセス、<https://math-gpt.org/>
 23. 12 AI drug discovery companies you should know about in 2025 - Labiotech.eu, 10 月 14, 2025 にアクセス、<https://www.labiotech.eu/best-biotech/ai-drug-discovery-companies/>
 24. From Data to Drugs: The Role of Artificial Intelligence in Drug Discovery - Wyss Institute, 10 月 14, 2025 にアクセス、<https://wyss.harvard.edu/news/from-data-to-drugs-the-role-of-artificial-intelligence-in-drug-discovery/>
 25. AI can transform innovation in materials design – here's how | World Economic Forum, 10 月 14, 2025 にアクセス、<https://www.weforum.org/stories/2025/06/ai-materials-innovation-discovery-to-design/>
 26. Emerald Therapeutics: A Robotic Laboratory in the Cloud - Nanalyze, 10 月 14, 2025 にアクセス、<https://www.nanalyze.com/2015/10/emerald-therapeutics-a-robotic-laboratory-in-the-cloud/>
 27. Robotic Controlled Life Sciences Lab - Strateos Cloud Lab, 10 月 14, 2025 にアクセス、<https://strateos.com/>
 28. Publish Faster & Cheaper with a Cloud Lab, 10 月 14, 2025 にアクセス、<https://www.emeraldcloudlab.com/why-cloud-labs/efficiency/academia/>
 29. The Most Cost Effective Lab Space - Emerald Cloud Lab, 10 月 14, 2025 にアクセス、<https://www.emeraldcloudlab.com/why-cloud-labs/efficiency/startup/>
 30. 株式会社クラウドラボ, 10 月 14, 2025 にアクセス、<https://cloudlab.nigh.biz/>
 31. cloud-labo.jp, 10 月 14, 2025 にアクセス、<https://cloud-labo.jp/#:~:text=%E3%82%AF%E3%83%A9%E3%82%A6%E3%83%89%E3%83%A9%E3%83%9C%E3%81%AF%E3%80%81%E6%96%B0%E3%81%97%E3%81%84%E5%A0%B4%E6%89%80,%E4%BA%A4%E6%B5%81%E3%82%82%E5%90%AB%E3%81%BE%E3%82%8C%E3%81%BE%E3%81%99%E3%80%82>
 32. Automated Science | The Jensen Lab, 10 月 14, 2025 にアクセス、<http://jensenlab.net/automatedscience/>
 33. Autonomous laboratories in China: an embodied intelligence-driven platform to accelerate chemical discovery - RSC Publishing, 10 月 14, 2025 にアクセス、

- <https://pubs.rsc.org/en/content/articlehtml/2025/dd/d5dd00072f>
34. How AI and Automation are Speeding Up Science and Discovery, 10 月 14, 2025 にアクセス、<https://newscenter.lbl.gov/2025/09/04/how-berkeley-lab-is-using-ai-and-automation-to-speed-up-science-and-discovery/>
 35. AI system learns from many types of scientific information and runs experiments to discover new materials | MIT News, 10 月 14, 2025 にアクセス、<https://news.mit.edu/2025/ai-system-learns-many-types-scientific-information-and-runs-experiments-discovering-new-materials-0925>
 36. Environmental Monitoring Robots - Meegle, 10 月 14, 2025 にアクセス、https://www.meegle.com/en_us/topics/robotics/environmental-monitoring-robots
 37. Case Study: Autonomous Robotics - Ignitec, 10 月 14, 2025 にアクセス、<https://www.ignitec.com/insights/case-study-autonomous-robotics/>
 38. Emerald Cloud Lab: Remote Controlled Life Sciences Lab, 10 月 14, 2025 にアクセス、<https://www.emeraldcloudlab.com/>
 39. Did you know: Can AI help us explore the ocean? - Woods Hole Oceanographic Institution, 10 月 14, 2025 にアクセス、<https://www.whoi.edu/ocean-learning-hub/ocean-facts/can-ai-help-us-explore-the-ocean/>
 40. Planetary Robotic Exploration Driven by Science Hypotheses for Geologic Mapping, 10 月 14, 2025 にアクセス、https://www.ri.cmu.edu/app/uploads/2017/08/IROS17_1637_FLpdf
 41. How Robots Are Pioneering Space Exploration | Capitol Technology University, 10 月 14, 2025 にアクセス、<https://www.captechu.edu/blog/how-robots-are-pioneering-space-exploration>
 42. Why AI is a "Black Box" – And Why It Doesn't Work Like a Human Brain - Reddit, 10 月 14, 2025 にアクセス、https://www.reddit.com/r/ArtificialIntelligence/comments/1kph5tc/why_ai_is_a_black_box_and_why_it_doesnt_work_like/
 43. The Black Box Problem. By: Elizabeth Louie | by Humans For AI | Medium, 10 月 14, 2025 にアクセス、<https://medium.com/@humansforai/the-black-box-problem-c40d3c6f26fe>
 44. Explain the Black Box for the Sake of Science: the Scientific Method in the Era of Generative Artificial Intelligence - arXiv, 10 月 14, 2025 にアクセス、<https://arxiv.org/html/2406.10557v3>
 45. Explain the Black Box for the Sake of Science: the Scientific Method in the Era of Generative Artificial Intelligence - arXiv, 10 月 14, 2025 にアクセス、<https://arxiv.org/html/2406.10557v5>
 46. What Is Black Box AI and How Does It Work? - IBM, 10 月 14, 2025 にアクセス、<https://www.ibm.com/think/topics/black-box-ai>
 47. Trust, Bias, and Reproducibility in AI for Bioinformatics - KAMI Think Tank, 10 月 14, 2025 にアクセス、<https://www.kamithinktank.com/bioinformatician/blog-post-title-three-5spth>

48. Statistical Rigor and Reproducibility in the AI Era - PMC, 10 月 14, 2025 にアクセス、 <https://pmc.ncbi.nlm.nih.gov/articles/PMC12402953/>
49. The Need for Verification in AI-Driven Scientific Discovery - arXiv, 10 月 14, 2025 にアクセス、 <https://arxiv.org/html/2509.01398v1>
50. Philosophy of artificial intelligence - Wikipedia, 10 月 14, 2025 にアクセス、 https://en.wikipedia.org/wiki/Philosophy_of_artificial_intelligence
51. Avoiding Common Pitfalls in Laboratory Automation: Best Practices and Solutions, 10 月 14, 2025 にアクセス、 <https://www.wakoautomation.com/avoiding-common-pitfalls-in-laboratory-automation-best-practices-and-solutions>
52. Five challenges in lab automation and how to overcome them - Automata, 10 月 14, 2025 にアクセス、 <https://automata.tech/blog/five-challenges-in-lab-automation-and-how-to-overcome-them/>
53. Automation in the Life Science Research Laboratory - PMC, 10 月 14, 2025 にアクセス、 <https://pmc.ncbi.nlm.nih.gov/articles/PMC7691657/>
54. Challenges to the Lab of the Future - Labforward, 10 月 14, 2025 にアクセス、 <https://labforward.io/test-blog-labtwi-new-2021/smart-lab-of-the-future-through-digital-transformation-and-lab-automation-0>
55. What is AI Ethics? | IBM, 10 月 14, 2025 にアクセス、 <https://www.ibm.com/think/topics/ai-ethics>
56. The Ethical Considerations of Artificial Intelligence | Washington D.C. & Maryland Area, 10 月 14, 2025 にアクセス、 <https://www.captechu.edu/blog/ethical-considerations-of-artificial-intelligence>
57. (PDF) Artificial Intelligence in Systematic Reviews: Overcoming Reproducibility, Bias and Validation Challenges - ResearchGate, 10 月 14, 2025 にアクセス、 https://www.researchgate.net/publication/393046199_Artificial_Intelligence_in_Systematic_Reviews_Overcoming_Reproducibility_Bias_and_Validation_Challenges
58. Generative AI Ethics: Concerns and How to Manage Them? - Research AIMultiple, 10 月 14, 2025 にアクセス、 <https://research.aimultiple.com/generative-ai-ethics/>
59. The dilemma of dual-use AI - Project Ploughshares, 10 月 14, 2025 にアクセス、 <https://ploughshares.ca/the-dilemma-of-dual-use-ai/>
60. Dual Use of Artificial Intelligence-powered Drug Discovery - PMC, 10 月 14, 2025 にアクセス、 <https://pmc.ncbi.nlm.nih.gov/articles/PMC9544280/>
61. Dual Use of Artificial Intelligence-powered Drug Discovery | Request PDF - ResearchGate, 10 月 14, 2025 にアクセス、 https://www.researchgate.net/publication/359073288_Dual_use_of_artificial-intelligence-powered_drug_discovery
62. Norms in New Technological Domains: Japan's AI Governance Strategy - CSIS, 10 月 14, 2025 にアクセス、 <https://www.csis.org/analysis/norms-new-technological-domains-japans-ai-governance-strategy>

63. Understanding Japan's AI Promotion Act: An "Innovation -First" Blueprint for AI Regulation, 10 月 14, 2025 にアクセス、 <https://fpf.org/blog/understanding-japans-ai-promotion-act-an-innovation-first-blueprint-for-ai-regulation/>
64. Japan Artificial Intelligence - International Trade Administration, 10 月 14, 2025 にアクセス、 <https://www.trade.gov/market-intelligence/japan-artificial-intelligence>
65. About AIP | Center for Advanced Intelligence Project, 10 月 14, 2025 にアクセス、 <https://aip.riken.jp/about-aip/>
66. RIKEN Center for Advanced Intelligence Project - Science & Technology Office Tokyo, 10 月 14, 2025 にアクセス、 <https://www.stofficetokyo.ch/education-research/research-institutes/researchinstitute/riken-center-for-advanced-intelligence>
67. National Artificial Intelligence Research Resource Pilot | NSF - National Science Foundation, 10 月 14, 2025 にアクセス、 <https://www.nsf.gov/focus-areas/ai/nairr>
68. National Artificial Intelligence Research Resource (NAIRR), 10 月 14, 2025 にアクセス、 <https://nairratdoe.ornl.gov/>
69. NAIRR Program OpenMined, 10 月 14, 2025 にアクセス、 <https://openmined.org/programs/nairr/>
70. China's AI Strategy: A Case Study in Innovation and Global Ambition, 10 月 14, 2025 にアクセス、 <https://trendsresearch.org/insight/chinas-ai-strategy-a-case-study-in-innovation-and-global-ambition/>
71. About ABCI, 10 月 14, 2025 にアクセス、 https://abci.ai/en/about_abci/
72. AI Bridging Cloud Infrastructure - Wikipedia, 10 月 14, 2025 にアクセス、 https://en.wikipedia.org/wiki/AI_Bridging_Cloud_Infrastructure
73. Hewlett Packard Enterprise to deliver AISTs next-generation supercomputer, for generative AI, powered by NVIDIA | HPE, 10 月 14, 2025 にアクセス、 <https://www.hpe.com/us/en/newsroom/press-release/2024/07/hewlett-packard-enterprise-to-deliver-aists-next-generation-supercomputer-powered-by-nvidia-for-generative-ai.html>
74. NVIDIA Powers Japan's ABCI-Q Supercomputer for Quantum Research, 10 月 14, 2025 にアクセス、 <https://nvidianews.nvidia.com/news/nvidia-powers-japans-abci-q-supercomputer-for-quantum-research>
75. Overview | Supercomputing Research Division | Information Technology Center, The University of Tokyo, 10 月 14, 2025 にアクセス、 <https://www.itc.u-tokyo.ac.jp/supercomputing/en/overview/>
76. Governments are spending billions on their own 'sovereign' AI technologies – is it a big waste of money?, 10 月 14, 2025 にアクセス、 <https://www.theguardian.com/technology/2025/oct/09/governments-spending-billions-sovereign-ai-technology>