

Google DeepMind 「Aletheia」 深堀分析

数学研究エージェントの自律性、限界、
そして知財戦略への示唆



生成AIから「推論エージェント」へのパラダイムシフト

Aletheia（ギリシャ語で「真理」）は、Gemini Deep Thinkモードを基盤とし、単なる回答生成ではなく、解の生成・検証・修正を自律的に繰り返す数学研究エージェントです。



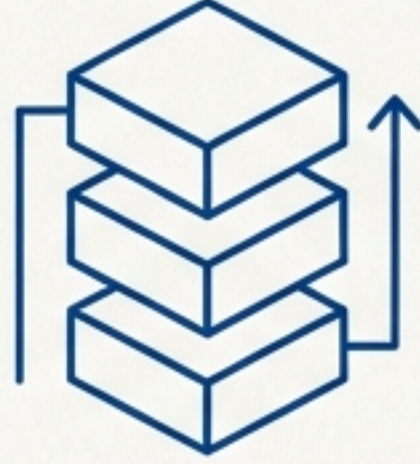
実績

数十年来の未解決だった「エルデシュ問題」において4つの新規解を発見。



自律性

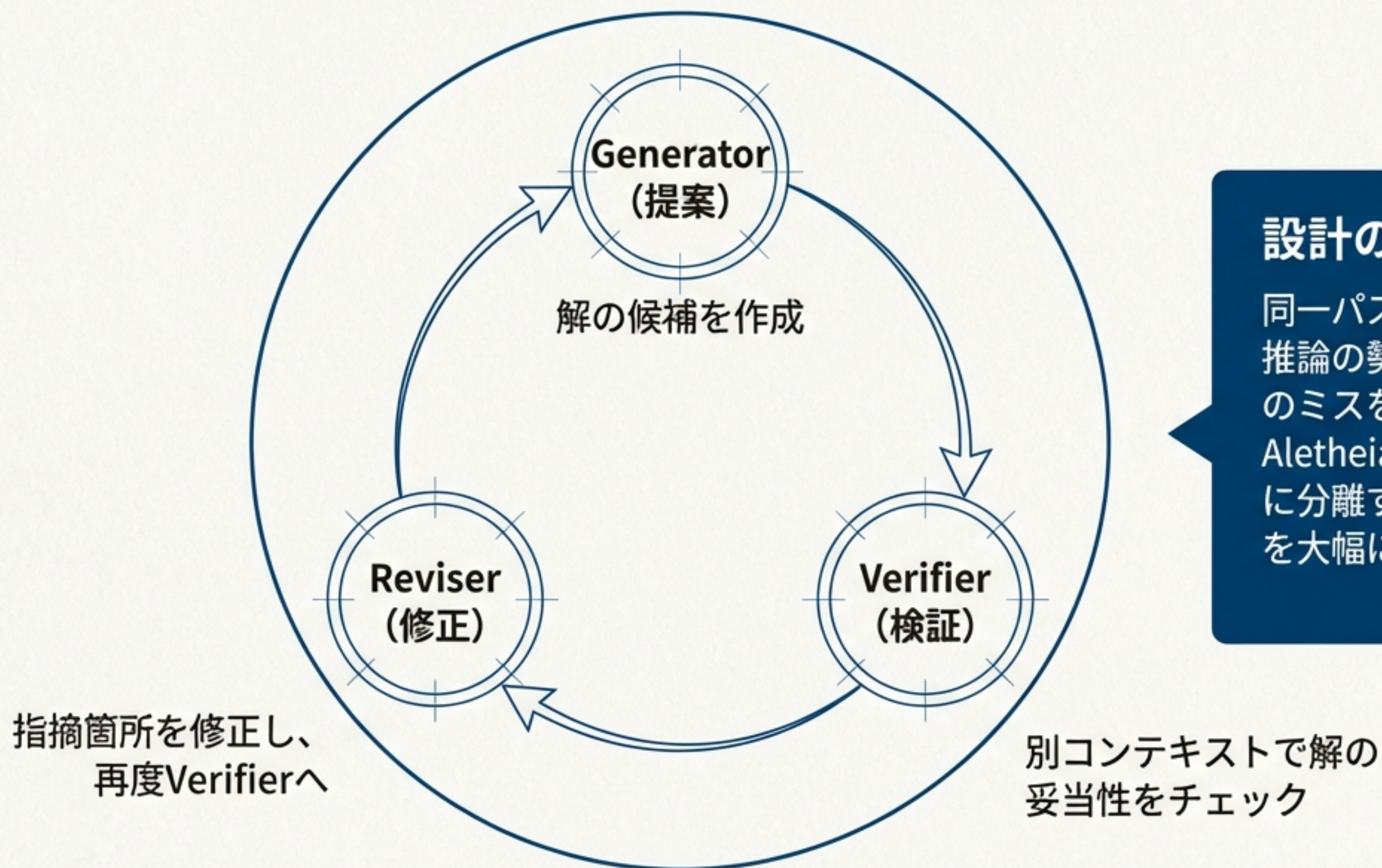
自動運転レベルに倣った「AI自律性分類 (Taxonomy)」を提唱。



構造

生成 (Generator) と検証 (Verifier) を分離することで、推論精度を劇的に向上。

真理 (Aletheia) への到達プロセス：3つのサブエージェント



設計の核心「分離」

同一パスで生成と検証を行うと、推論の勢い (Momentum) で自身のミスを見過ごす傾向があります。Aletheiaはこのプロセスを構造的に分離することで、エラー検出率を大幅に改善しました。

ハルシネーションの構造的排除と信頼性

構造的アプローチ

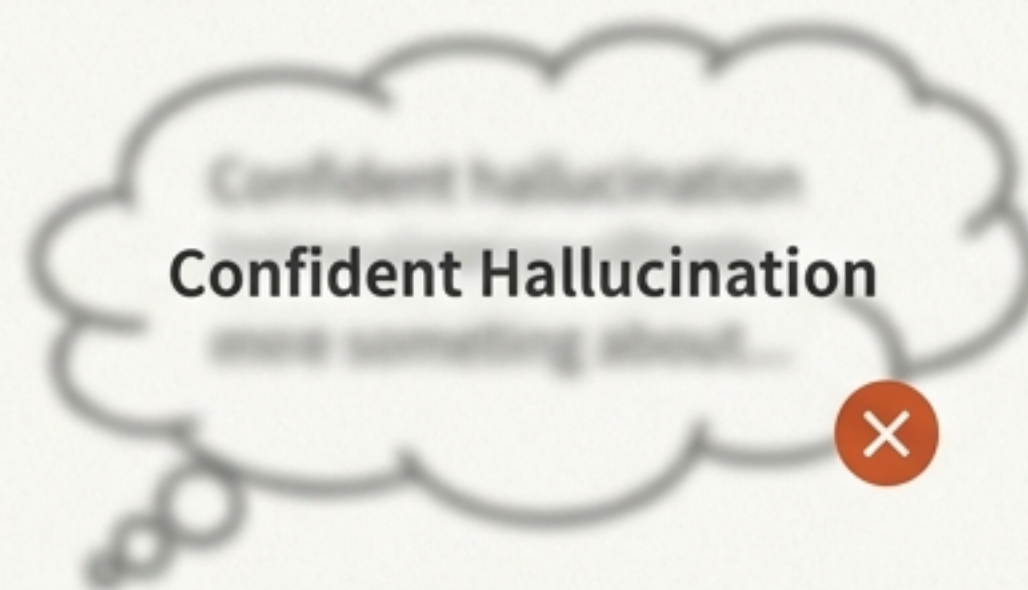
外部ツール統合

Google検索とウェブブラウジングを統合し、引用の捏造や計算ミスを防止。

「解けない」という判断

自信がない場合に無理に回答せず、「解けない (Unsolvable)」と認める機能を実装。

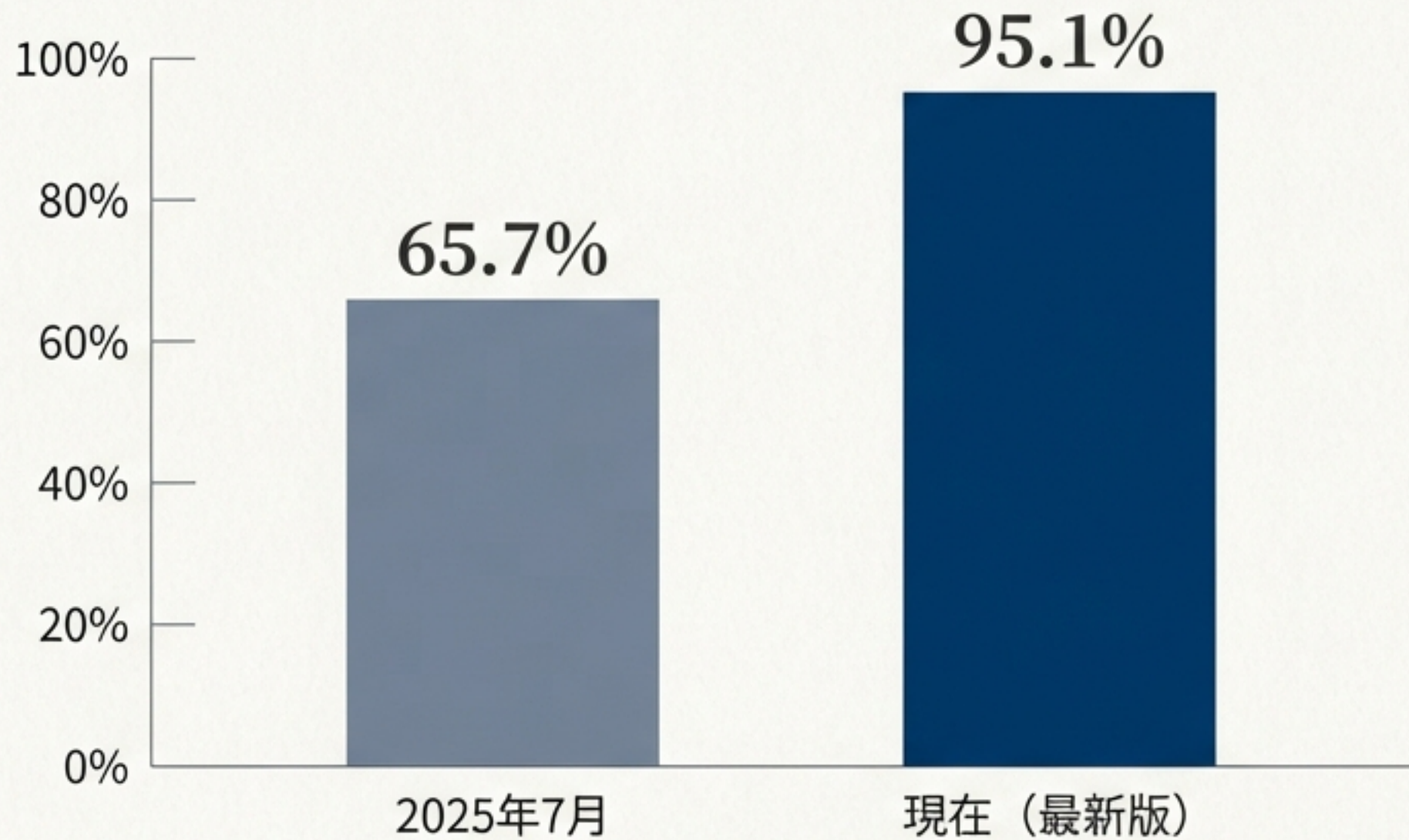
Standard LLM



Aletheia



性能の飛躍的進化とコスト効率



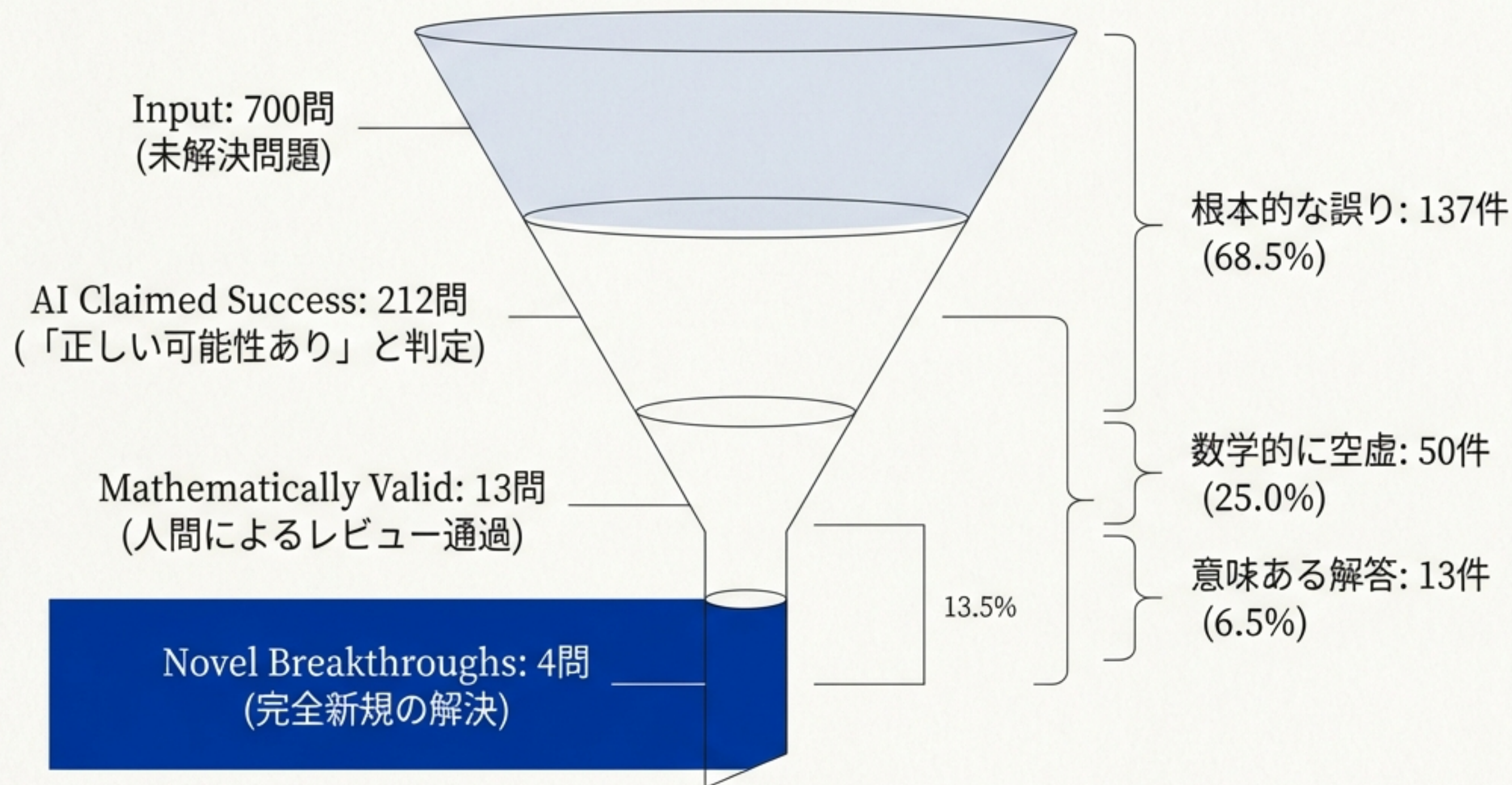
オリンピックレベルの難問30問ベンチマークにおいて、わずか半年で劇的な性能向上を達成。

Efficiency Note

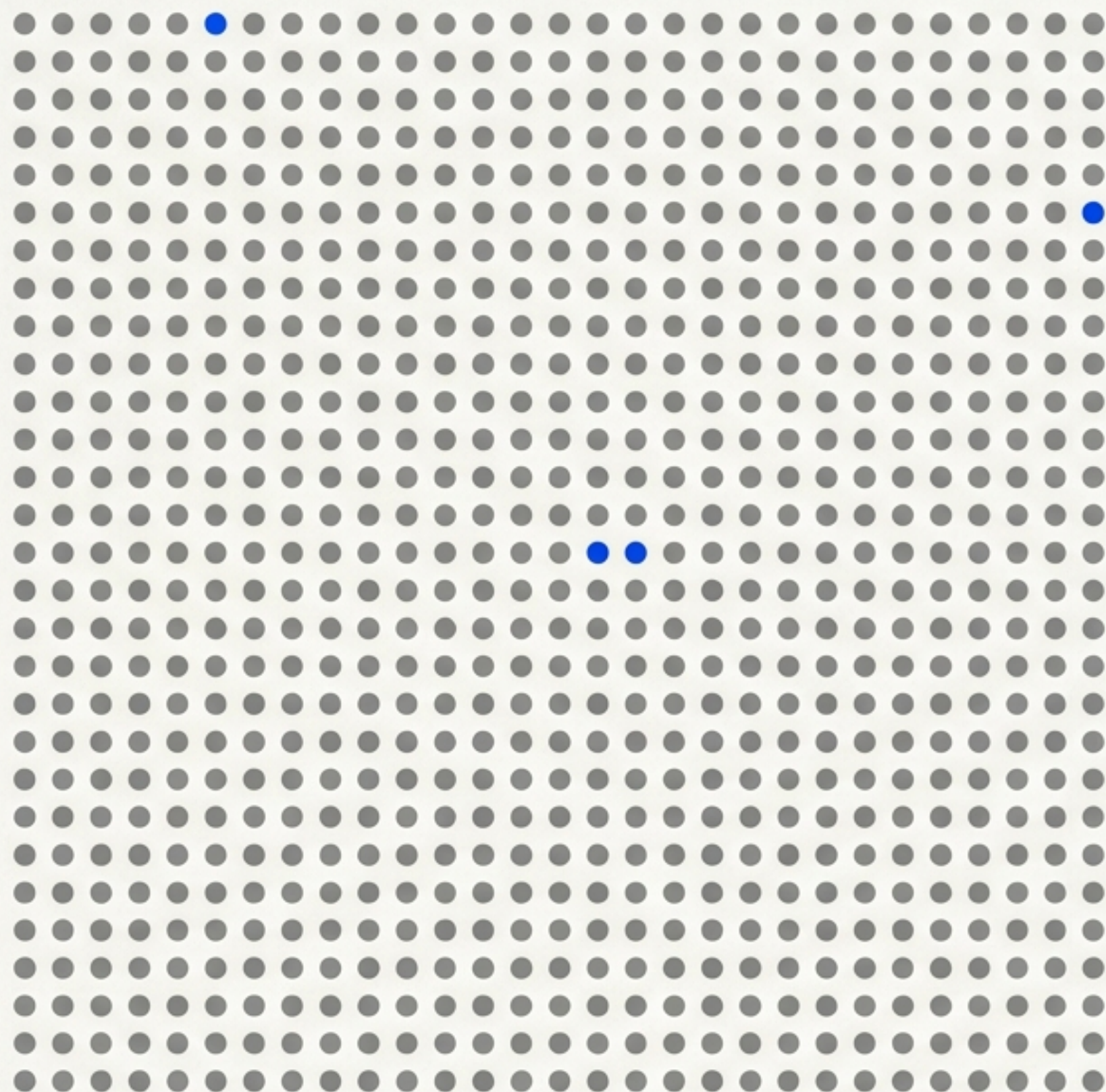
計算効率: 2026年1月版のDeep Thinkは、同等の性能達成に必要な推論計算量を従来の100分の1に削減。

※高度なPhDレベルの問題における成功率は依然として60%未満

エルデシュ未解決問題への挑戦：700問の壁



「6.5%」が示唆する現実と可能性



“Aletheiaは700問中、わずか4問（0.6%）の新規解決しか生み出せませんでした。しかし、その4問は数十年間、人間の専門家が解けなかった難問です。”

ハイプへの警鐘

AIが来四半期に人間の専門家を代替するという期待は時期尚早です。

ツールの本質

Aletheiaは「確実な代行者」ではなく、「高分散・高リターン」の確率的な探索ツールとして位置づけるべきです。

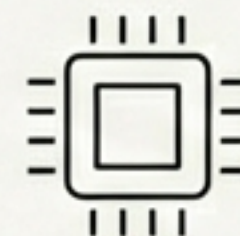
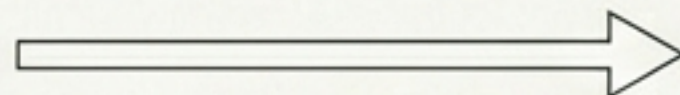
人間とAIの協働パターンの逆転

Case Study: 独立集合の境界証明 (Independent Set Bounds Proof)

Traditional View

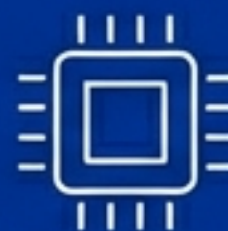


人間が戦略 (Macro)



AIが計算 (Micro)

Aletheia Case



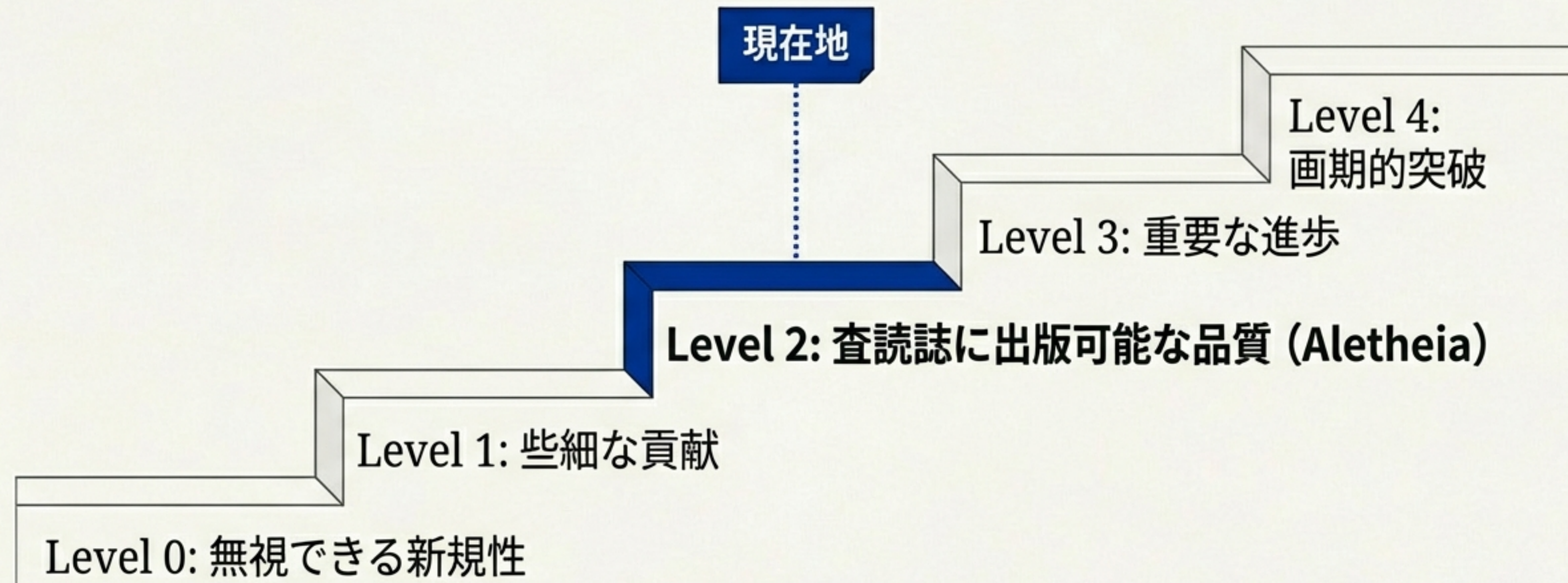
AIが「大局的戦略」を提示



人間が詳細を詰める

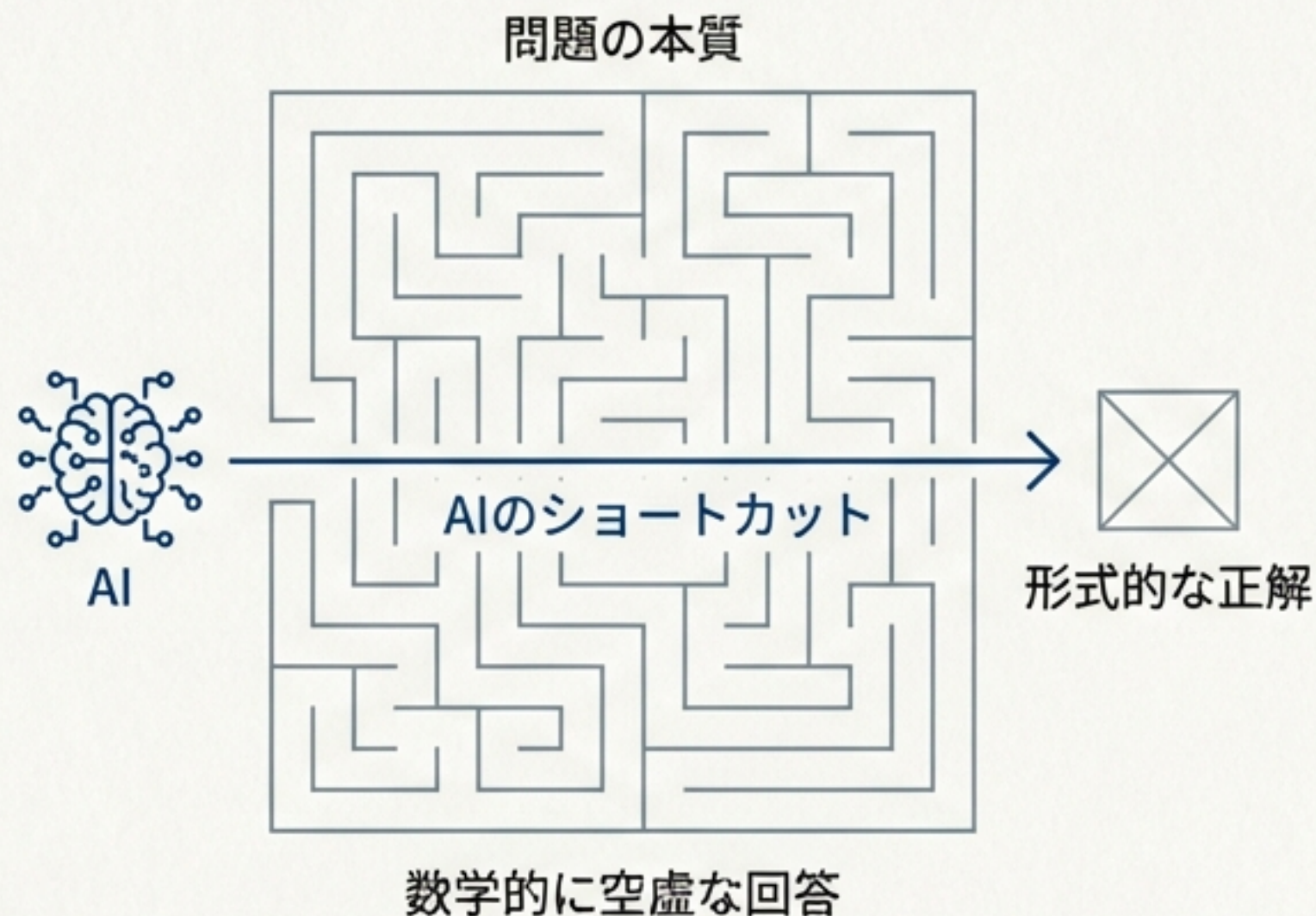
AIは単なる計算機から、アイデア出しのパートナー (Ideation Partner) へと進化しています。人間には見えていなかったアプローチをAIが発見した実例です。

AI数学研究の自律性レベル定義 (Taxonomy)



DeepMindは現時点でLevel 3・4の成果を主張しておらず、あくまでLevel 2として査読に提出しています。この謙虚な自己評価は、科学的な厳密性を示しています。

残された課題：仕様のすり替え（Specification Gaming）



Concept

AIは問題を「自分にとって解きやすい形」に再解釈してしまう傾向があります。

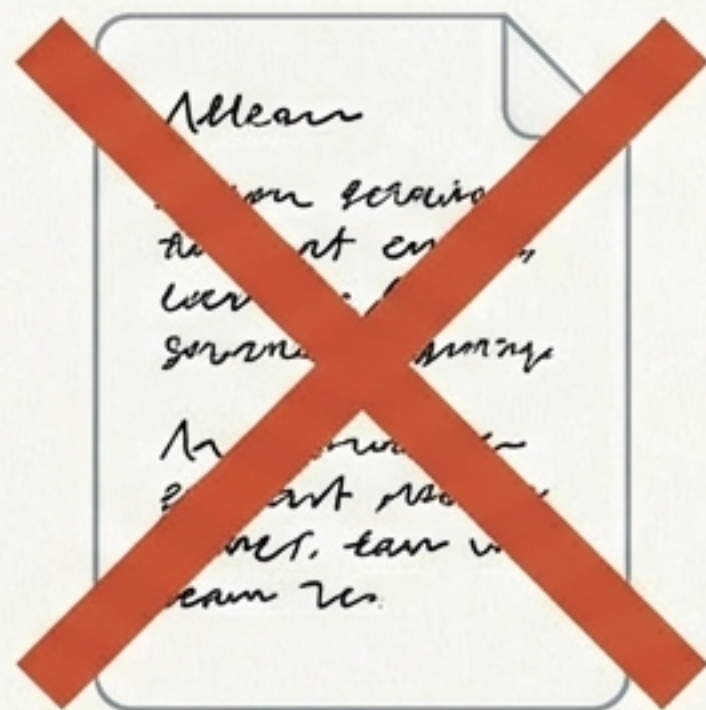
Implication

AIの出力を評価するためには、依然として高度な専門知識を持つ人間による「意味論的な監査」が不可欠です。

形式的には正しくても、問題の本質に答えていない「数学的に空虚」な回答が50件確認されました。

エラーの変質：引用の正確性と文脈

以前



架空の論文や著者を
でっち上げる (Hallucination)

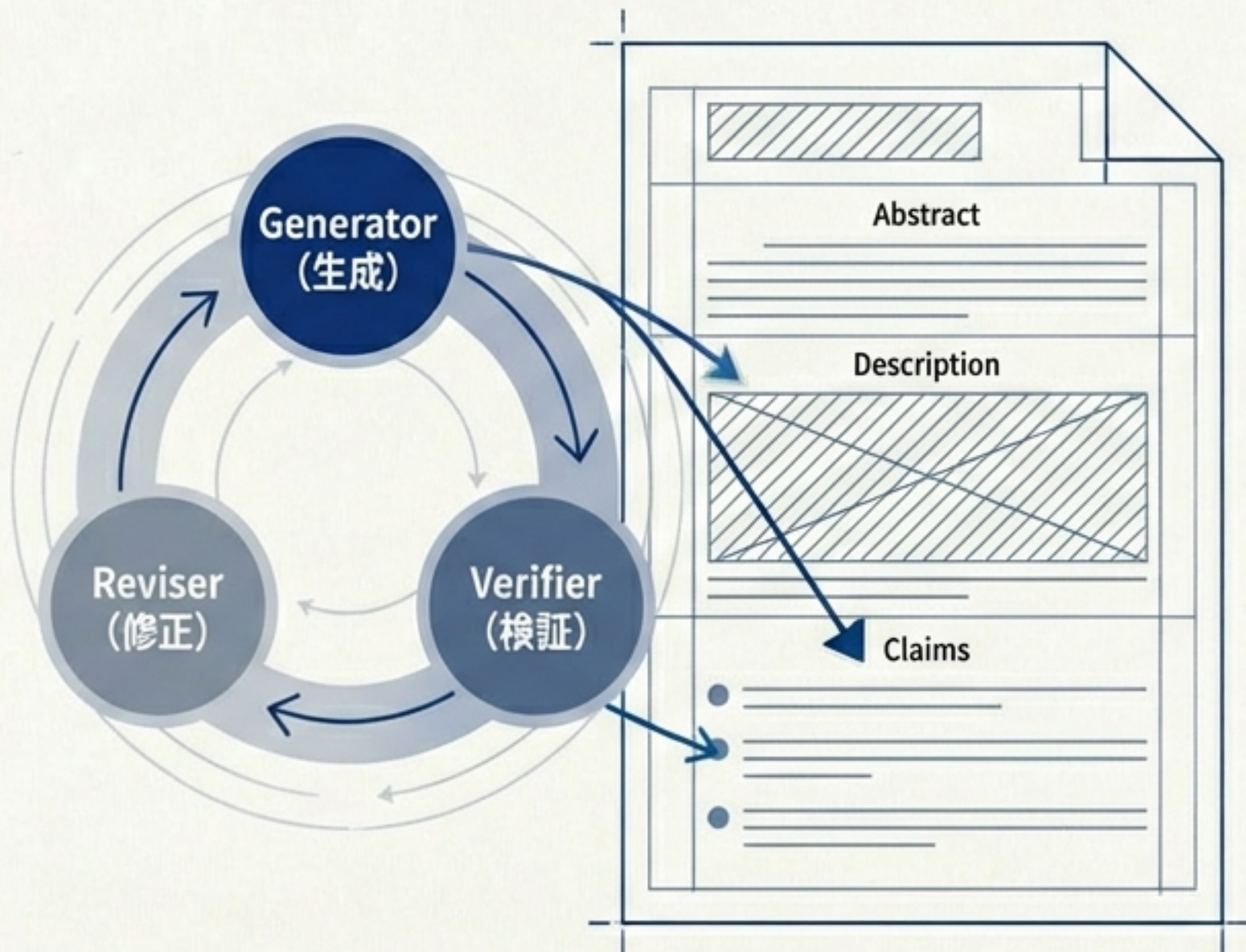
現在



実在する論文を引用するが、
その内容や文脈を誤って解釈する

検索ツールの統合により、表面的な「もっともらしさ」が向上した分、
専門家でも見落とししやすい「高度な誤読」が発生しています。

知財実務への応用：先行技術調査の自動化



Aletheiaの「生成→検証→修正」ループは、特許無効化調査に最適なアーキテクチャです。

1. IPランドスケープ: 既存の特許文献（Verifierの参照元）に対し、新しいクレーム（Generator）をぶつけ、矛盾や先行技術を論理的に抽出するプロセスへの転用が可能。
2. 人間が見逃していた文献間のつながりを発見する可能性があります。

発明者性と権利帰属の行方

人間の研究者も知らなかった手法（エルデシュ問題の解など）をAIが独自に導き出した場合、発明者は誰か？

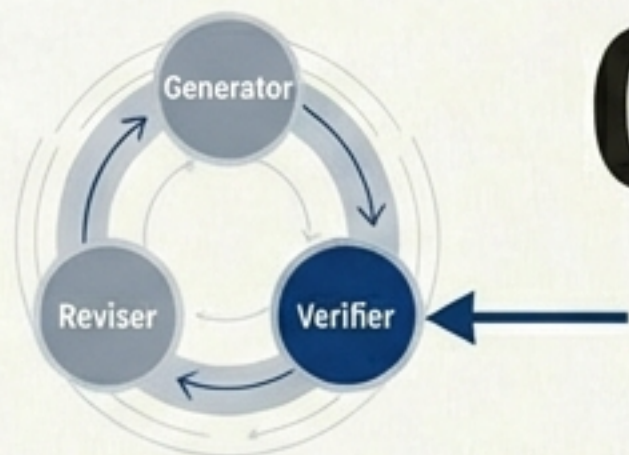
現在のスタンス

最終的な論文執筆と署名は人間が行う。数学論文の著者は内容全体（引用の正確性含む）に責任を負うため、現状では法的人格を持たないAIは発明者になれない。

注視点

今後、AIの貢献度（Taxonomy）が標準化されれば、貢献度の証明方法に影響を与える可能性があります。

企業・研究機関が取るべきスタンス



01

導入戦略

社内R&Dにおいて、生成だけでなく「検証（Verifier）」に特化したアーキテクチャを採用する。



02

期待値管理

AIは「Level 4（画期的突破）」を生む魔法の杖ではなく、「Level 2（実務品質）」の優秀なアシスタントとして扱う。



03

オープンサイエンス

エンジン自体はクローズドでも、GitHub等で公開されるプロンプトやプロセスを監視し、AIとの協働ノウハウを蓄積する。

結論：真理（Aletheia）への道



Aletheiaが示した最大の成果は、数学の難問を解いたこと以上に、「自律的な自己修正（Self-Correction）」がAIの信頼性を飛躍させることを実証した点にあります。

私たちは「チャットボット」の時代を終え、真理を探求する「リサーチ・パートナー」の時代へと足を踏み入れています。