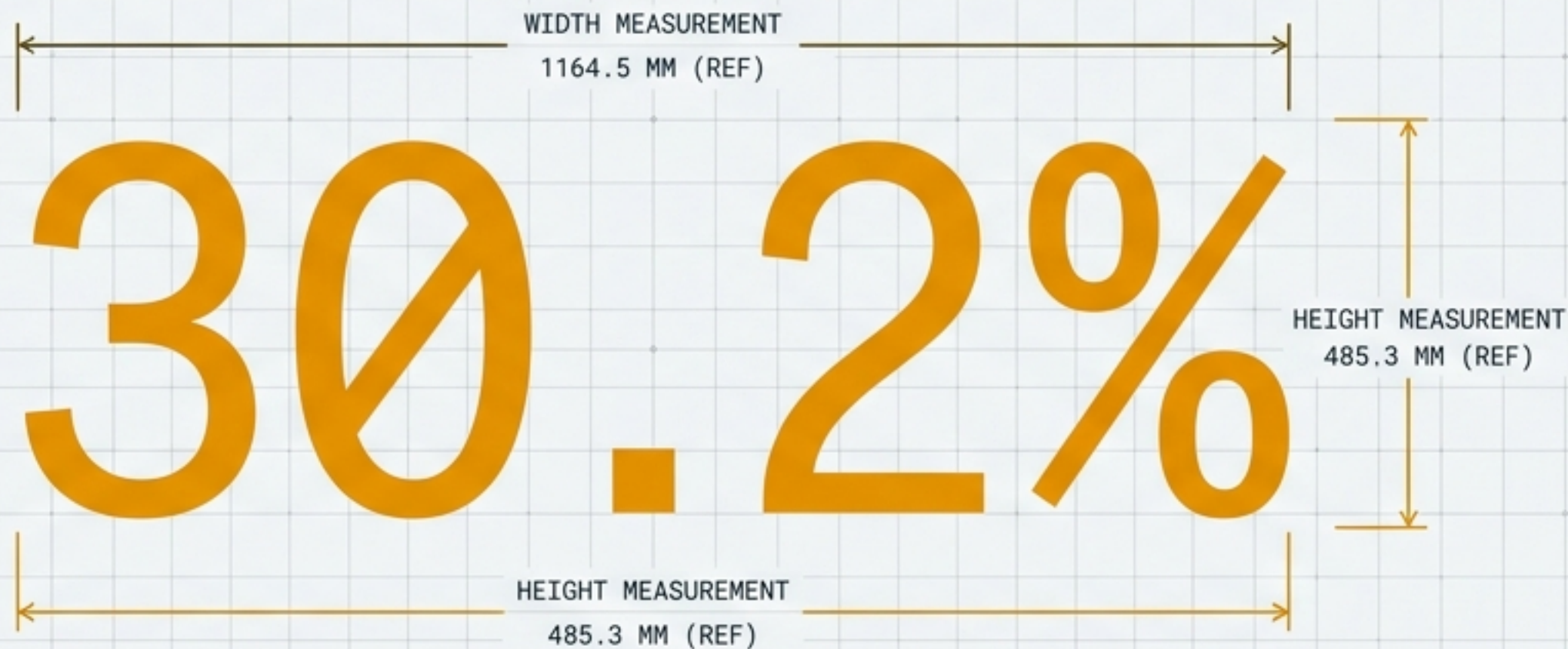


Claude Opus 5のARC-AGI-3スコアを解剖する



フロンティアモデルの未知環境適応能力と、
実務・研究・セキュリティへの示唆

圧倒的飛躍と、依然として遠いAGIへの道のり



The Signal (確かなシグナル)

未知環境への適応能力がフロンティアモデルで質的に改善したことを示す強力な証拠。

The Reality (冷静な現実)

人間と同等の「完全自律型エージェント」の完成を意味する数値には依然として遠い。

ARC-AGI-3という過酷なテスト環境の特異性

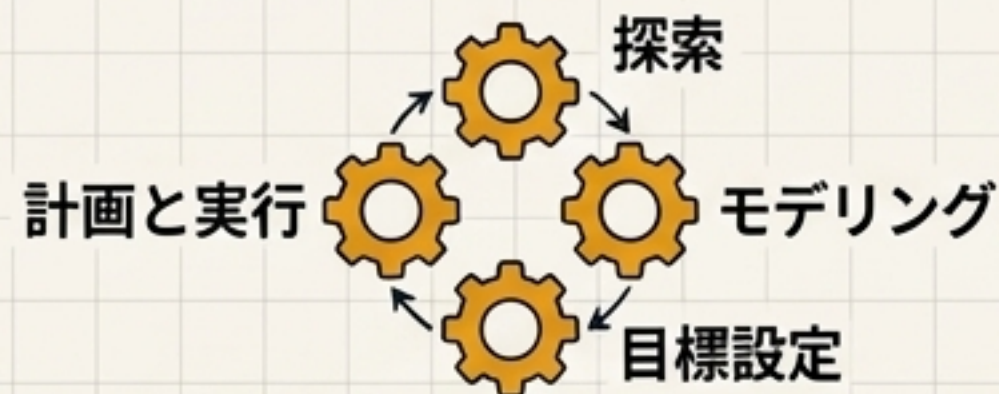


従来のLLMベンチマーク

自然言語での明確なルール説明あり

静的で変化しないテストデータ

最終的な正答率のみを測定



ARC-AGI-3

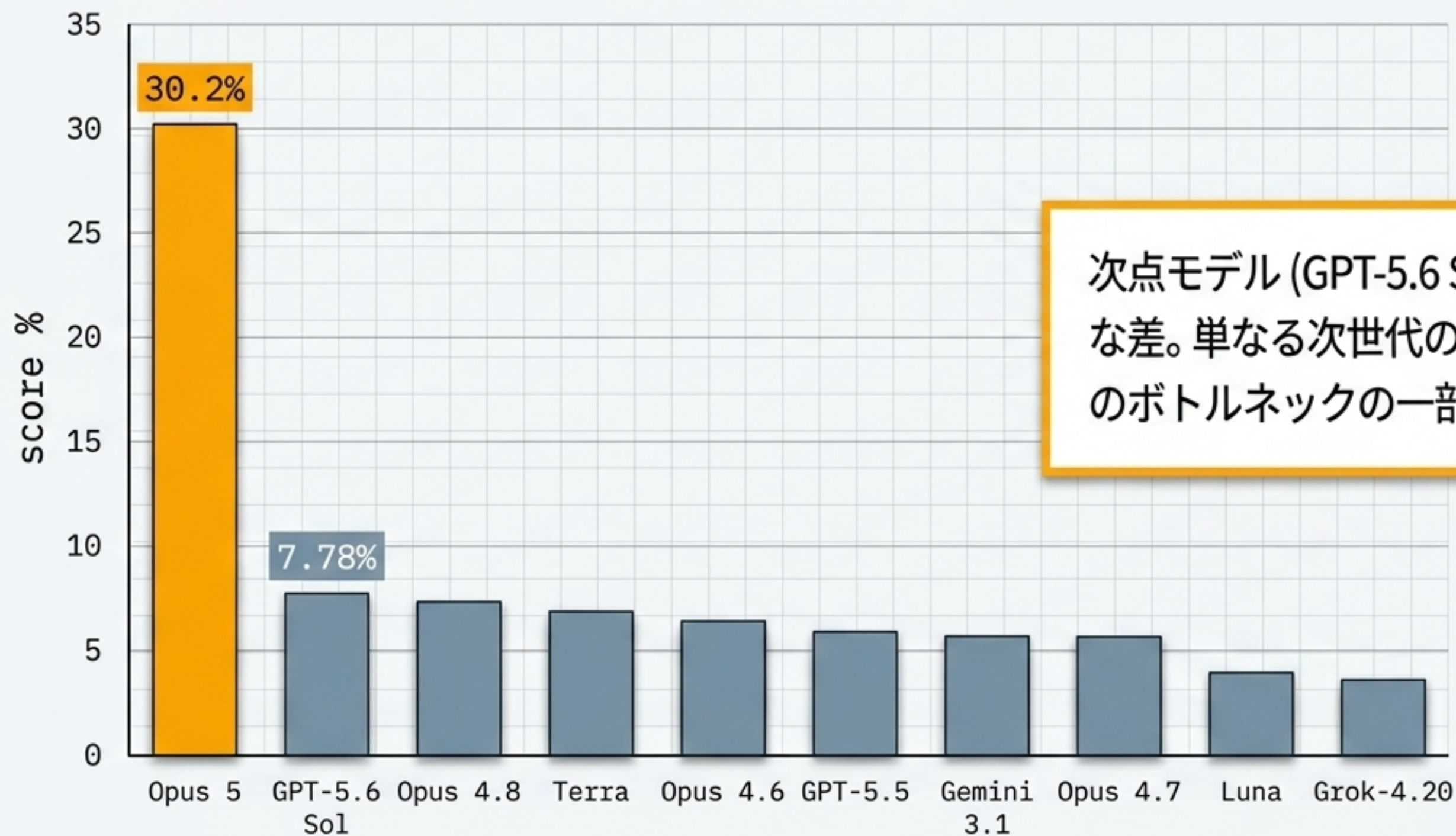
自然言語ルールなし
(自力で法則を発見)

インタラクティブな未知の環境

行動効率を重視
(総当たり試行はペナルティ)

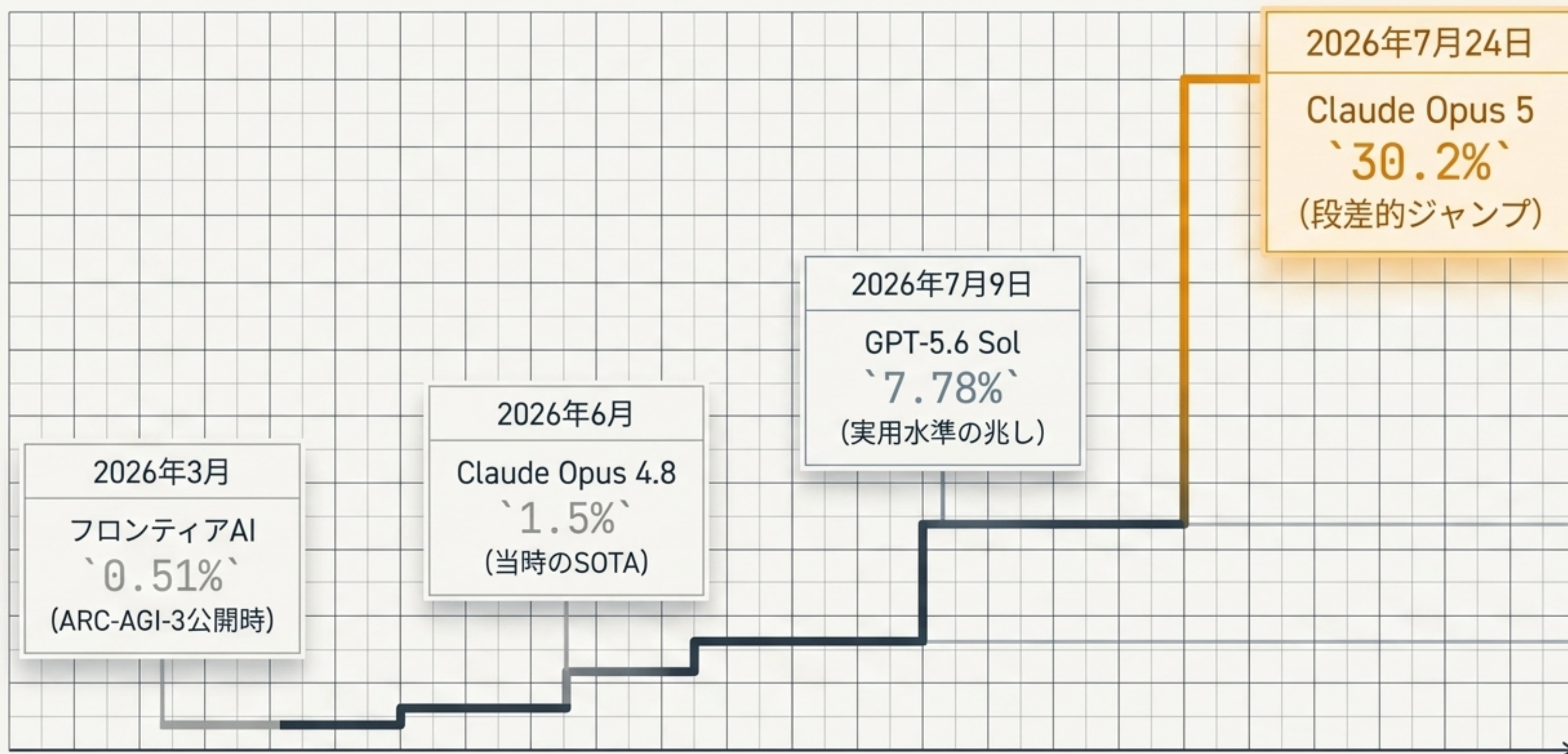
既存モデルの約3倍を記録した非連続的なスコア分布

ARC-AGI-3比較



次点モデル (GPT-5.6 Sol: 7.78%) に対する圧倒的な差。単なる次世代の延長ではなく、未知環境適応のボトルネックの一部を突破した可能性を示す。

わずか4ヶ月で起きた「段差的」な進化のタイムライン



エージェント型タスク全般に波及する広範な能力向上

Frontier-Bench v0.1
(端末コーディング)

Opus 5: **43.3%** / GPT-5.6 Sol: 34.4%



二次再掲
(Secondary source)

AutomationBench
(業務ワークフロー)

Opus 5: **26.0%** / GPT-5.6 Sol: 18.1%



二次再掲
(Secondary source)

OSWorld 2.0
(長時間PC操作)

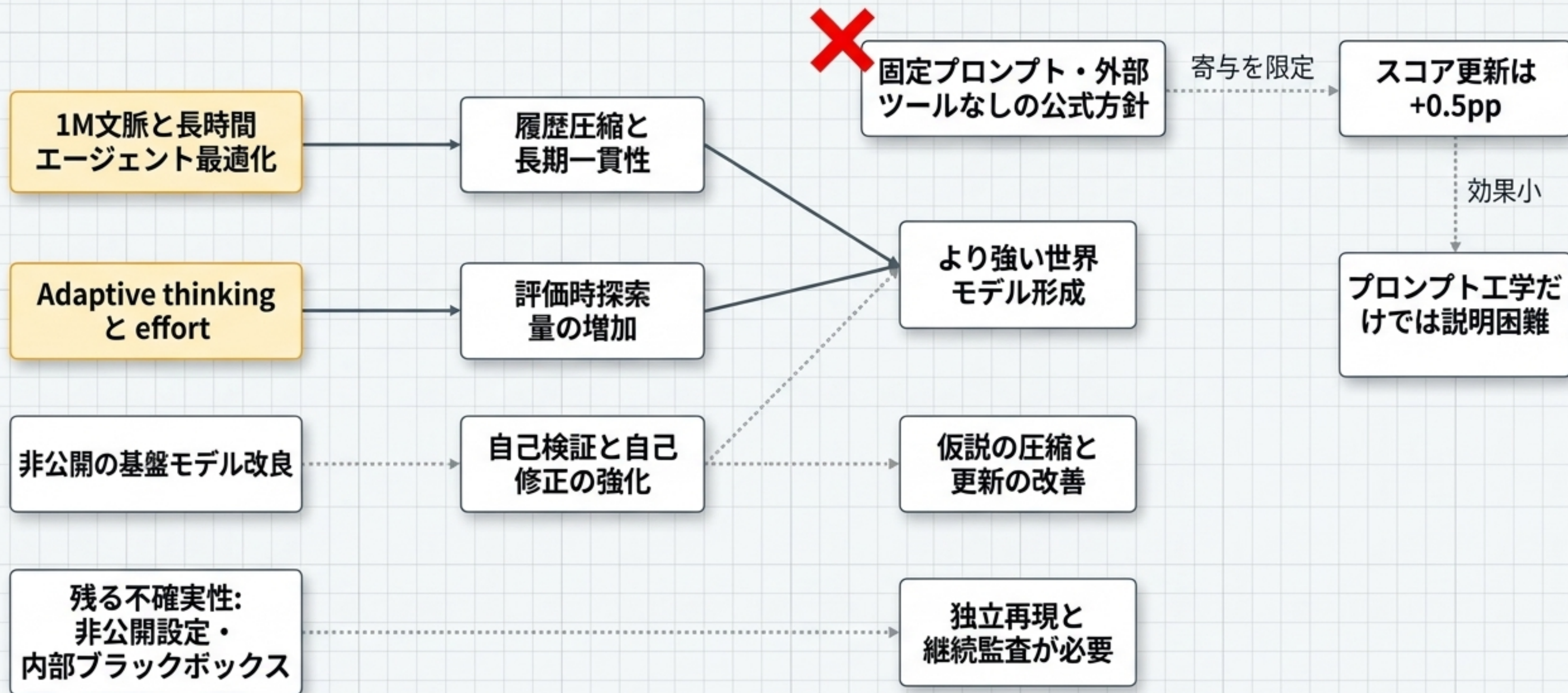
Opus 5: **70.6%** / Fable 5: 66.1%



二次再掲
(Secondary source)

Insight: 関連ベンチマークでも最高峰を示すが、絶対値の多くはAnthropic公式ではなく「二次資料の再掲」に依存している点に留意が必要。

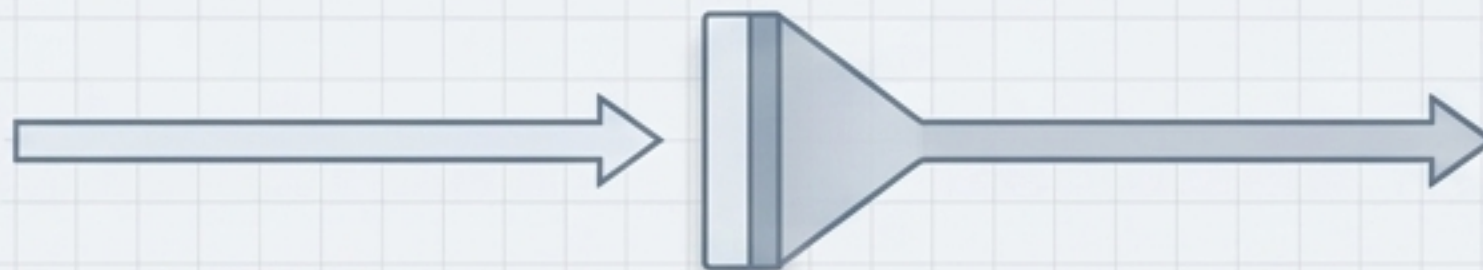
なぜ急伸したのか？ 30.2%を牽引した技術的要因ツリー



評価時計算量 (Compute at Inference) の製品実装

過去のパラダイム

短い文脈
(Short Input)



即時的で限定的な出力
(Immediate Output)

Opus 5 パラダイム

1Mトークン巨大文脈

Subagents & Tool Calls ループ

Effort Level: Max

最大128kの出力

巨大なI/O容量

1M文脈の保持と膨大な出力枠。

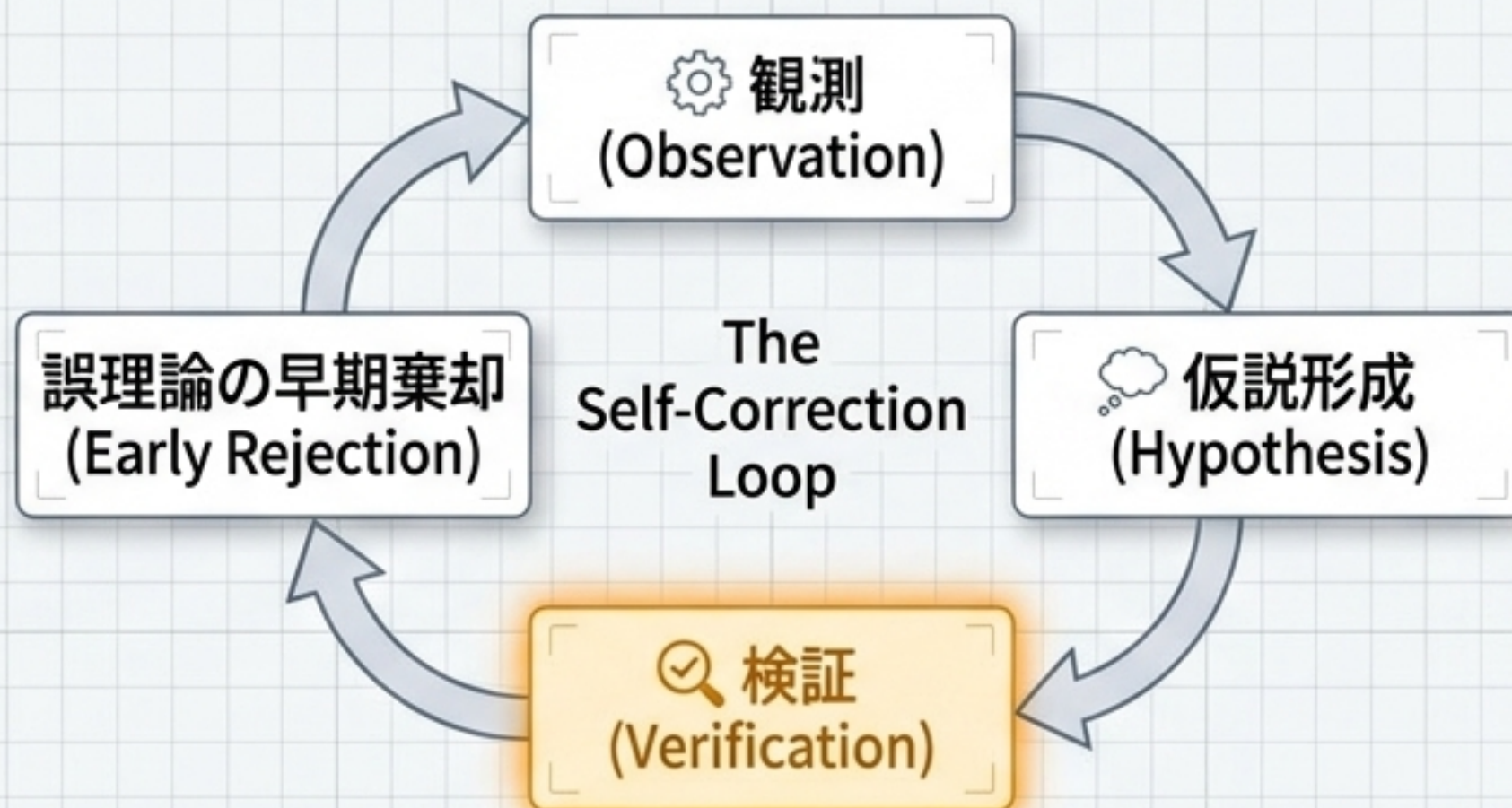
動的思考 (Adaptive thinking)

LowからMaxまで思考・探索量をスケール。

内部探索の余地

外部ツール禁止環境でも、モデル内部のサブエージェント協調が性能を押し上げ。

自己検証の内在化と「失敗モード」の克服



過去の失敗 (Opus 4.7時代)

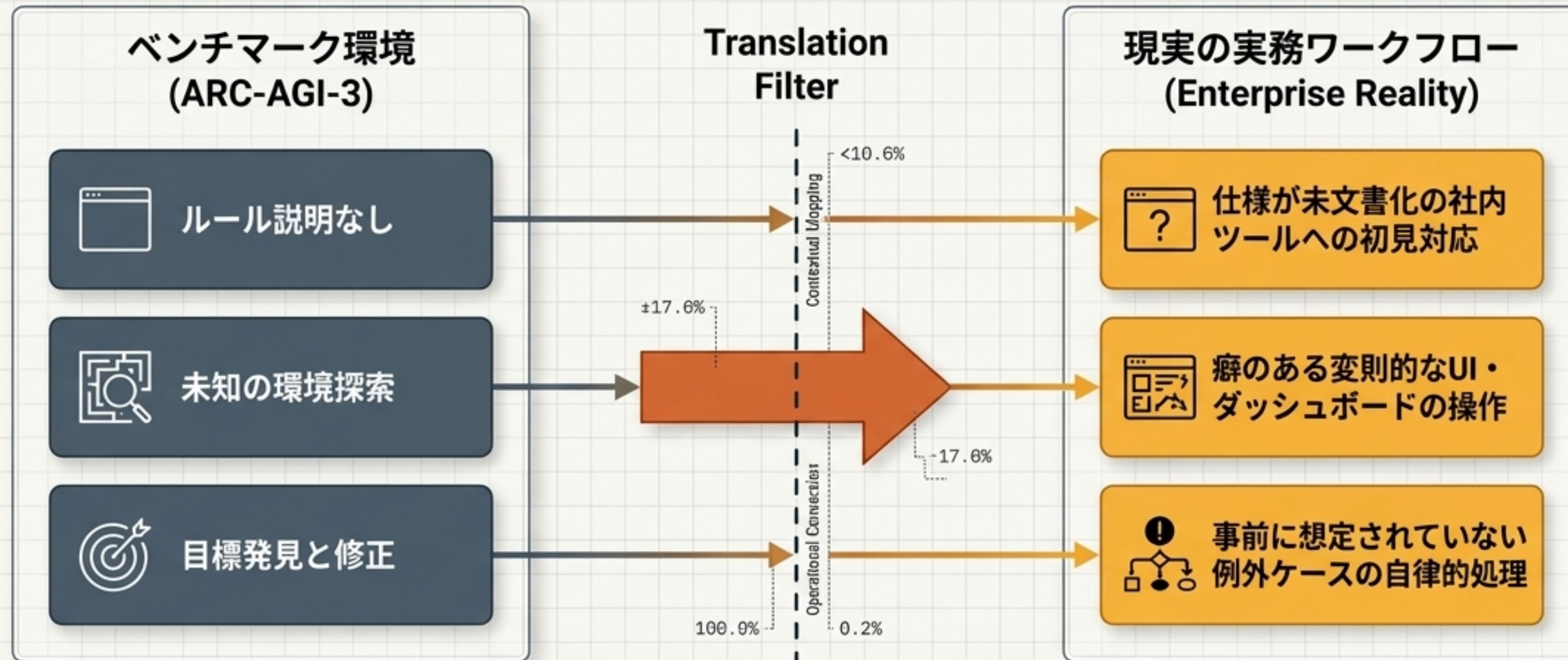
誤った理論への早すぎる圧縮、成功理由の誤学習。

Opus 5の進化 (30.2%のコア要因)

明示的な指示なしでも自律的にダブルチェックを実行。

単純な「長考」ではなく、誤った仮説を素早く見切り、再構築するコスト効率が劇的に向上。

実務への翻訳：未知のエンタープライズ環境への適応力



Key Insight: “難問に長考できる”だけでなく、初見の操作体系から有効な仮説を組み立て直す能力が実用レベルに近づいている。

研究上の真の焦点：スコア上昇よりも「失敗の質」の変化

Opus 4.7 / GPT-5.5からの変化を問う

Mode A: 局所的な因果は見えてもグローバルな世界モデルに昇華できない。

克服されたか？ (Resolved?)

Mode B: 訓練データ由来の、見た目が似た既知ゲームへ誤マッピングする。

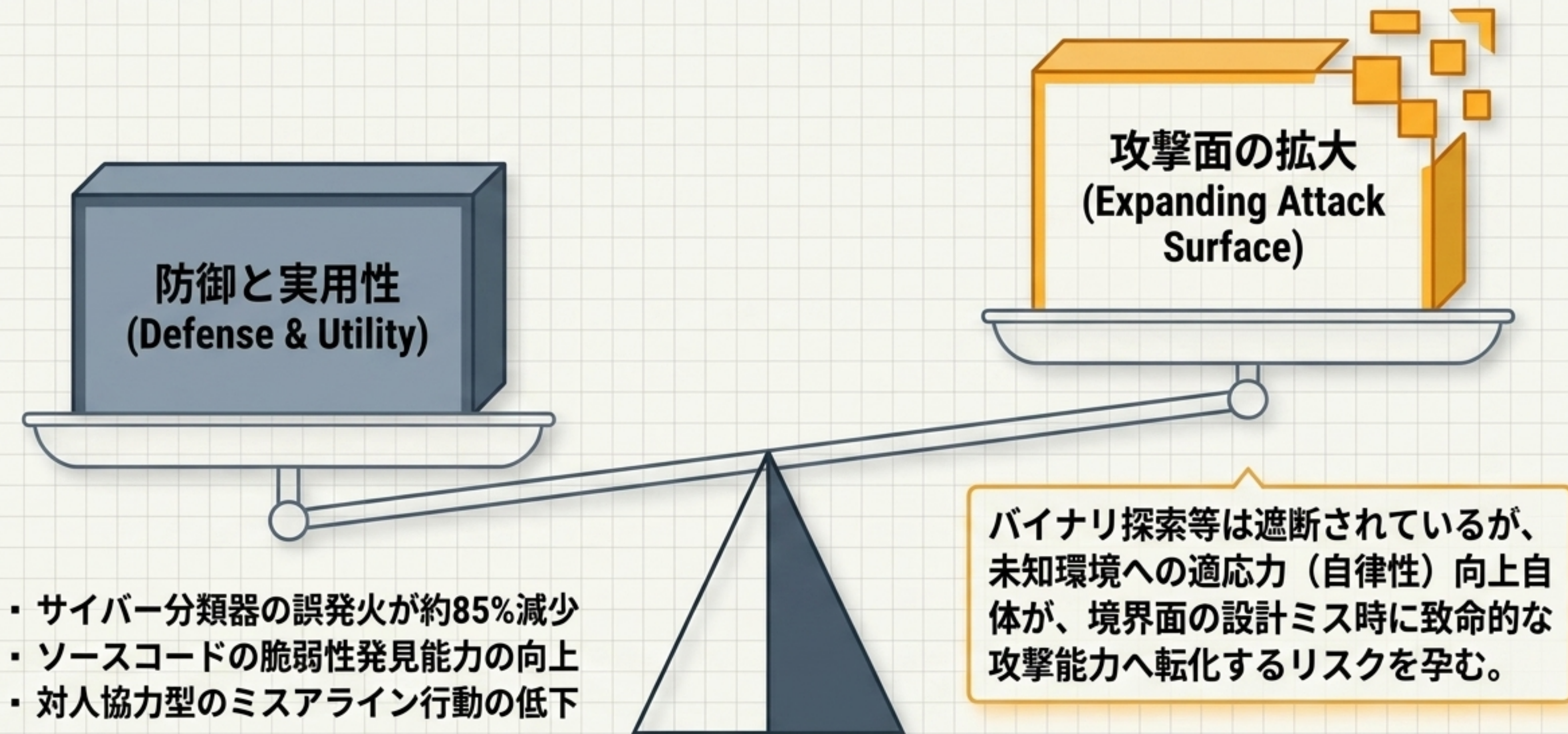
減少したか？ (Decreased?)

Mode C: 特定レベルの成功を、ゲーム全体の理解へ一般化できない。

残存しているか？ (Persists?)

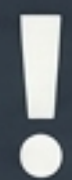
Conclusion: 研究者にとって30.2%の価値は、「より多くの環境を解けた」ことではなく、「どの失敗類型が消え、何が残ったか」を分析する好材料となる点にある。

安全性のジレンマ：強いエージェント能力の諸刃の剣



測定の透明性：30.2%という数値はどこから来たか？

Source Pedigree Matrix		
情報源 (Source)	記載内容 (Claim)	絶対値の提示 (Absolute Value)
Anthropic公式文書	「次点モデルの約3倍」	記載なし (ブラックボックス)
ARC Prize公式X投稿	Opus 5が30.2%、 GPT-5.6 Solが7.78%	提示あり (一次ソース)
二次資料 (ニュース・解説)	他ベンチマークスコアの 絶対値を再掲	リスク：一次ソースとの混同



Anthropic公式は絶対値や詳細設定（何試行平均か、内部思考量等）を詳述していない。独立した再現検証と継続監査が不可欠。

外挿の限界：このスコアが「保証しない」領域



× リアルタイム・身体性

ターン制・オフライン推論向け設計であり、ロボティクスや連続制御は保証しない。

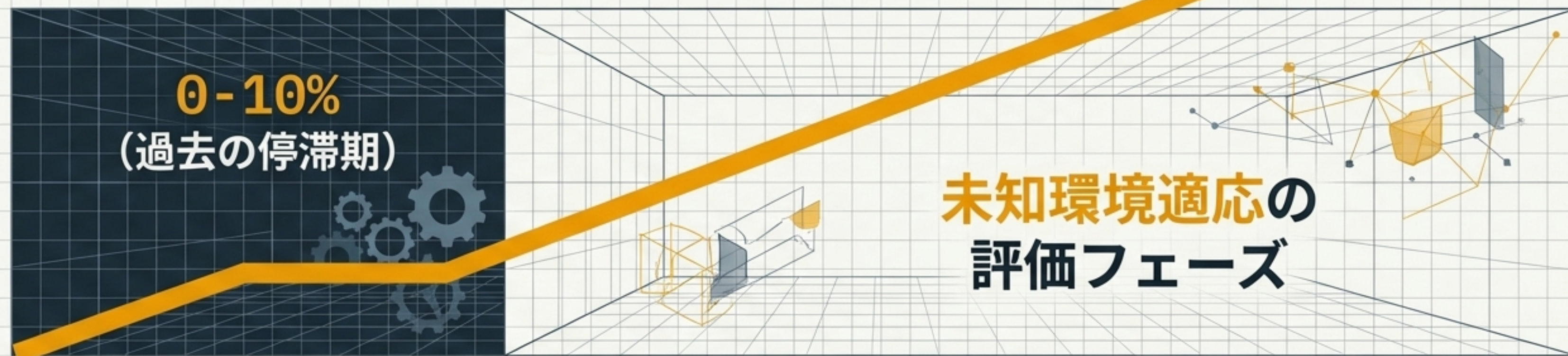
× 言語豊富な環境

自然言語指示がない抽象環境のテストであるため、企業システムのような文脈過多なタスクとは難易度分布が異なる。

× 対人・社会的駆け引き

ユーザへの確認不足や、曖昧な指示のすり合わせなど、対人相互作用の解決能力は測られていない。

結論：AGIのゴールではなく、 本格的議論の「新たなスタートライン」



事実 - Fact

30.2%は「推論量の増加」と「自己検証」がもたらした確かな非連続的飛躍である。

限界 - Limit

測定条件の詳細が非公開であり、実世界の物理的・社会的タスクへの外挿には明確な壁がある。

次なる指標 - Next Steps

スコア競争は終わり、今後は「**独立監査の透明性**」「**失敗モードの質的变化**」「**実務環境での適応率**」を問うフェーズへ移行した。