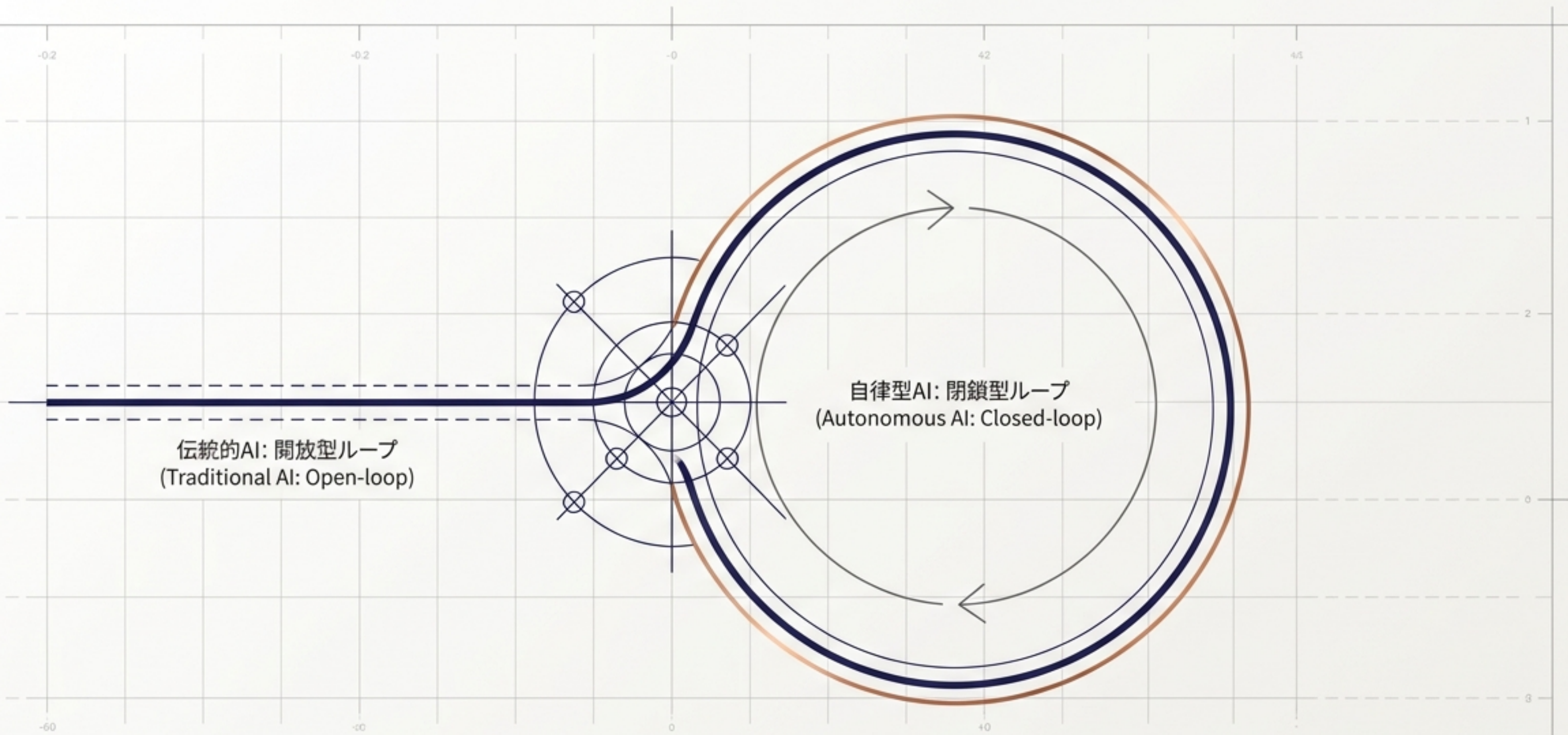




自律型AIエージェントの設計図：Claude Opus 5が切り拓く新パラダイム

ARC-AGI-3スコア**30.2%**達成の技術的メカニズムと、単発チャットから「**タスク完遂型**」への産業的シフト

AIの役割は「対話」から「長時間の自律実行」へ

The Leap (歴史的飛躍)

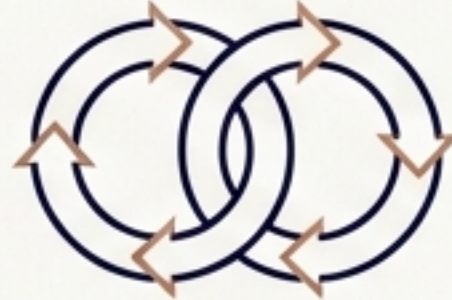


ARC-AGI-3において **30.2%** を達成。

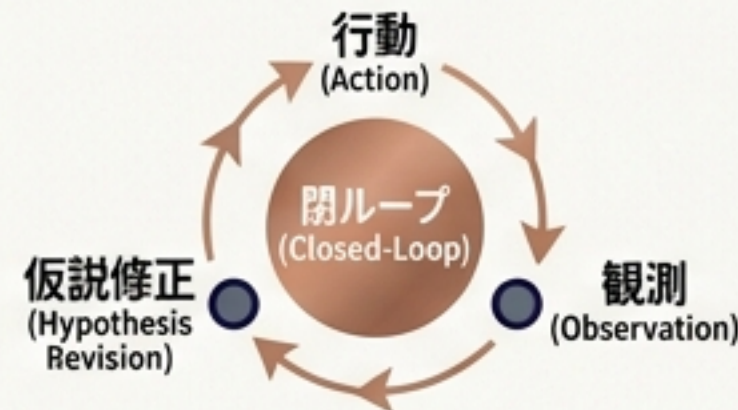
過去のフロンティアモデル (GPT-5.6 Sol の **7.8%**、Opus 4.8 の **1.5%**) を約 **4 倍** 上回る圧倒的スコアを記録。

パターン記憶ではなく「未知の環境への適応力」を実証。

The Engine (閉ループ型推論)



一方向の出力から、行動・観測・仮説修正を繰り返す「閉ループ (Closed-Loop)」アーキテクチャへ進化。



自律的なツール生成と厳格な安全制御 (スコア **2.3**) が長時間の思考を支える。

The Value (タスク完遂単価の劇的低減)



入力**\$5**/出力**\$25** (1Mトークン) という前世代と同等のAPIコストを維持しつつ、無駄な試行を削減。

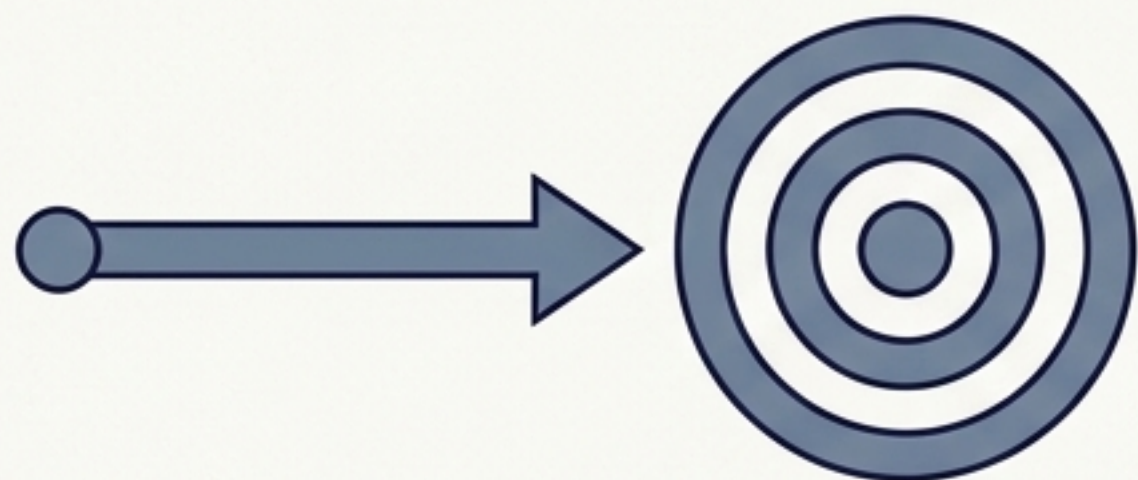
ターン数を **1/3**、所要時間を短縮し、

所要時間を **60%**、

「Cost per Token」から「Cost per Task」へとROIの定義を刷新。

ARC-AGI-3：エージェントの「流動的知能」を測る究極のテスト

従来のベンチマーク (ARC-AGI-1/2等)

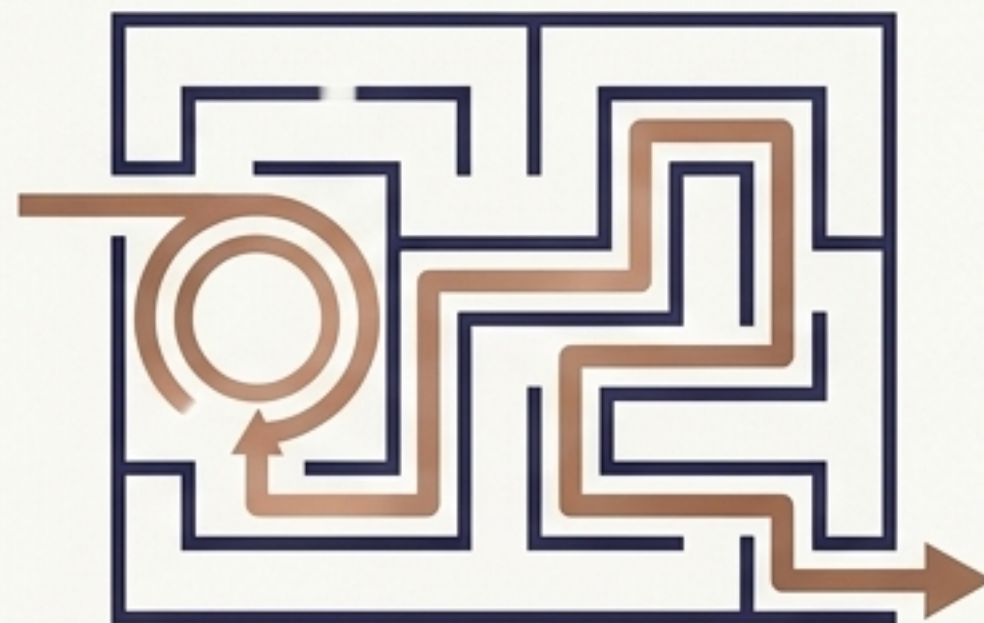


評価形式：静的な入出力グリッドからのルール推定

提供情報：明示的な指示と事前ルールあり

AIの挙動：過去の学習データに基づくパターンマッチング (オープンループ)

ARC-AGI-3 (2025-2026年展開)



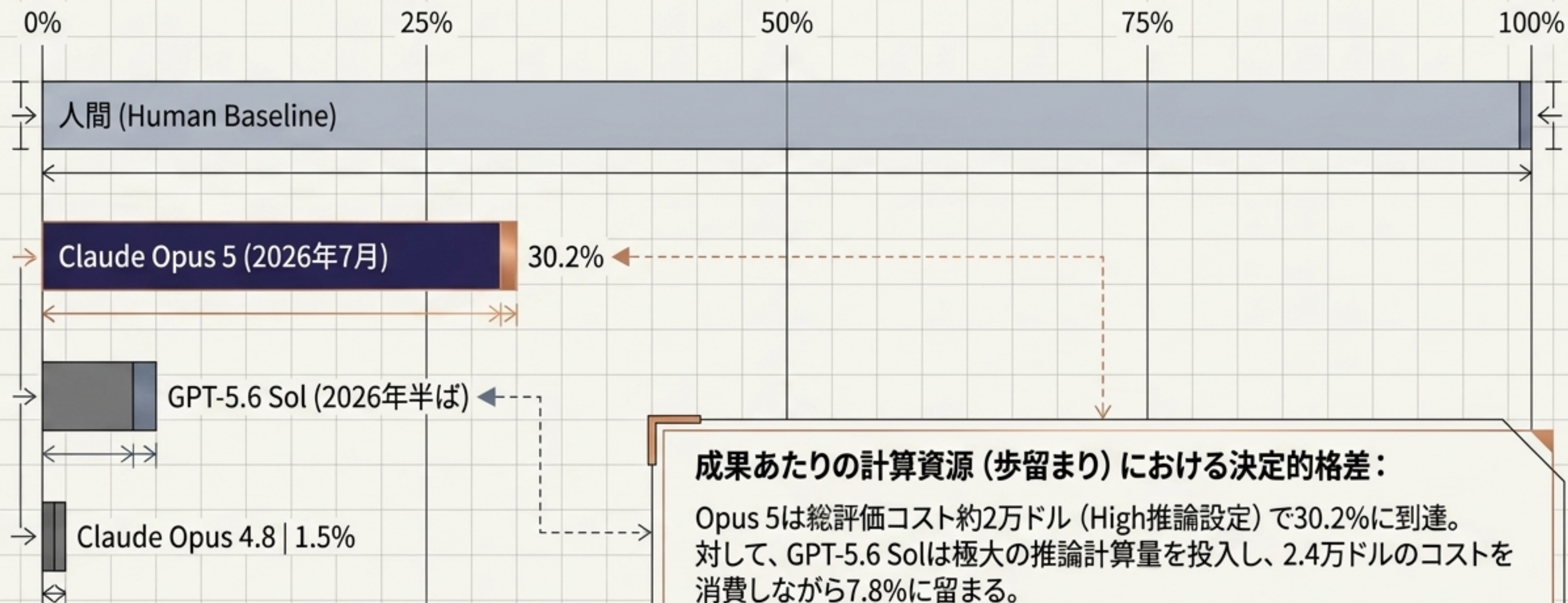
評価形式：完全対話型・ターン制の未知環境におけるエージェント稼働

提供情報：事前手順、目標、環境ルールは一切「非提示」

AIの挙動：自発的な操作試行 → 状態変化の観測 → 環境の動態モデル自己構築 → 行動計画の立案と実行

評価指標：試行錯誤の無駄を排除した「相対的人間行動効率」

停滞の壁を突破した30.2%という歴史的マイルストーン



成果あたりの計算資源 (歩留まり) における決定的格差：

Opus 5は総評価コスト約2万ドル (High推論設定) で30.2%に到達。
対して、GPT-5.6 Solは極大の推論計算量を投入し、2.4万ドルのコストを消費しながら7.8%に留まる。

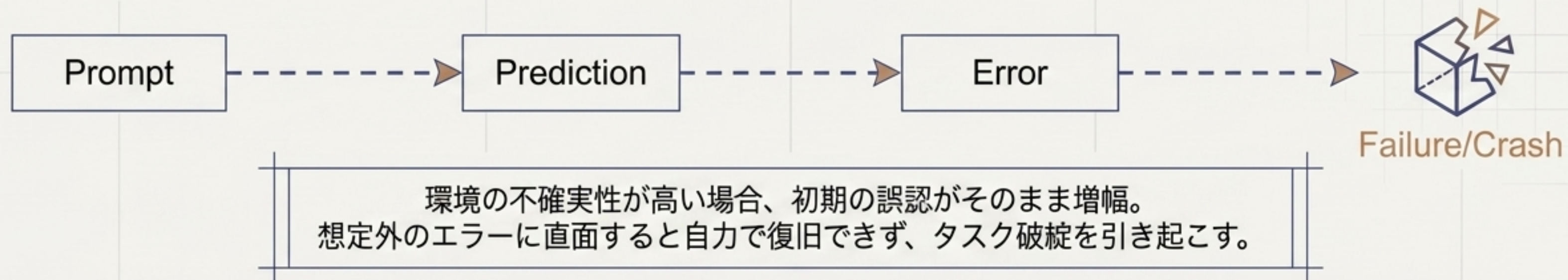
AIは「学習データの再構成」から「未知環境での仮説検証」へと移行した。

知能の経済学：主要フロンティアモデルの性能とコスト比較

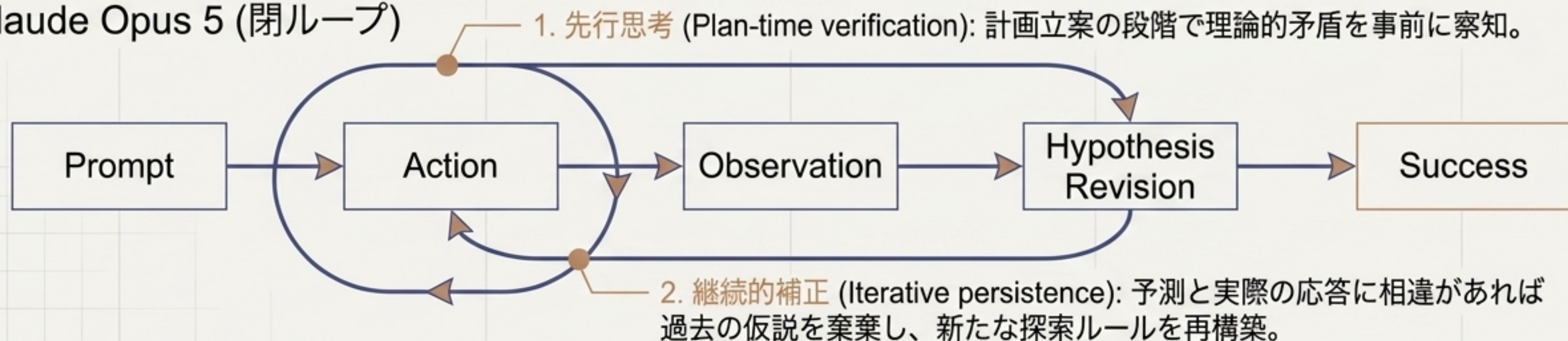
モデル名称	API価格 (入力/出力 1Mトークン)	ARC-AGI-3	Frontier-Bench v0.1	費用対効果と実務優位性
Claude Opus 5	\$5.00 / \$25.00	30.2%	43.3%	OSWorld 2.0において、Fable 5の最高性能を1/3強のコストで達成。
Claude Fable 5 (最上位モデル)	\$10.00 / \$50.00	未実施/非公開	33.7%	基準最高性能層（高コスト）。高度な法務論理等で優位。
GPT-5.6 Sol	\$5.00 / \$30.00	7.8%	非公表/未掲載	Terminal-Bench 2.1や単体コーディング等で優位性。
Claude Opus 4.8 (前世代)	\$5.00 / \$25.00	1.5%	18.7%	従来世代。未知環境での試行エラー率が高い。

コア・メカニズム 1：一方向の推論から「閉ループ型」の動的仮説検証へ

従来型 LLM（オープンループ）



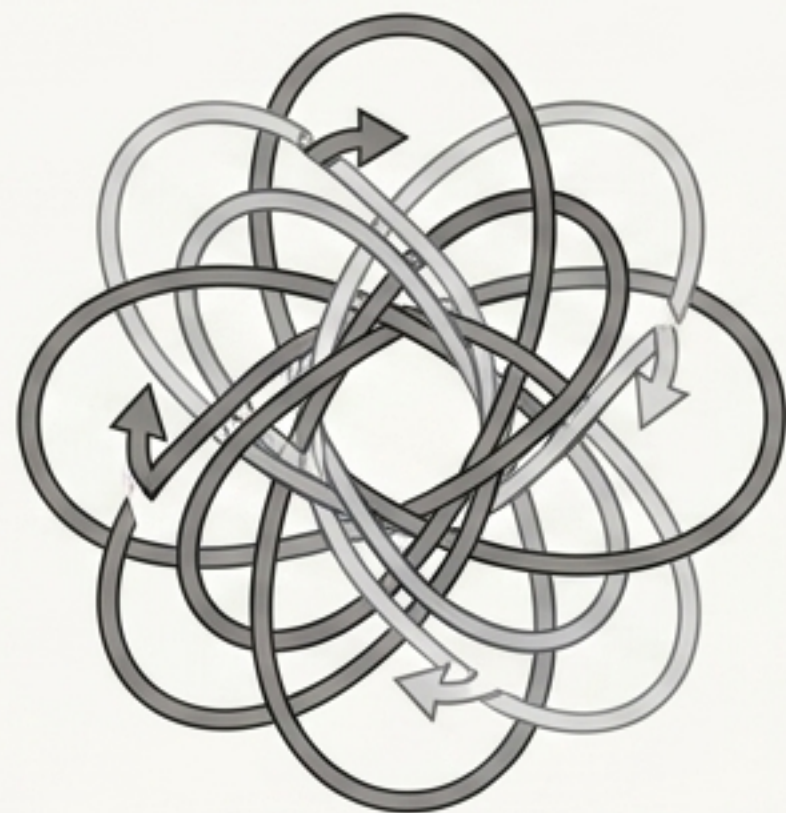
Claude Opus 5（閉ループ）



散漫なランダム探索を回避し、最小限のアクション数で未知のルールに到達する自己修正メカニズム。

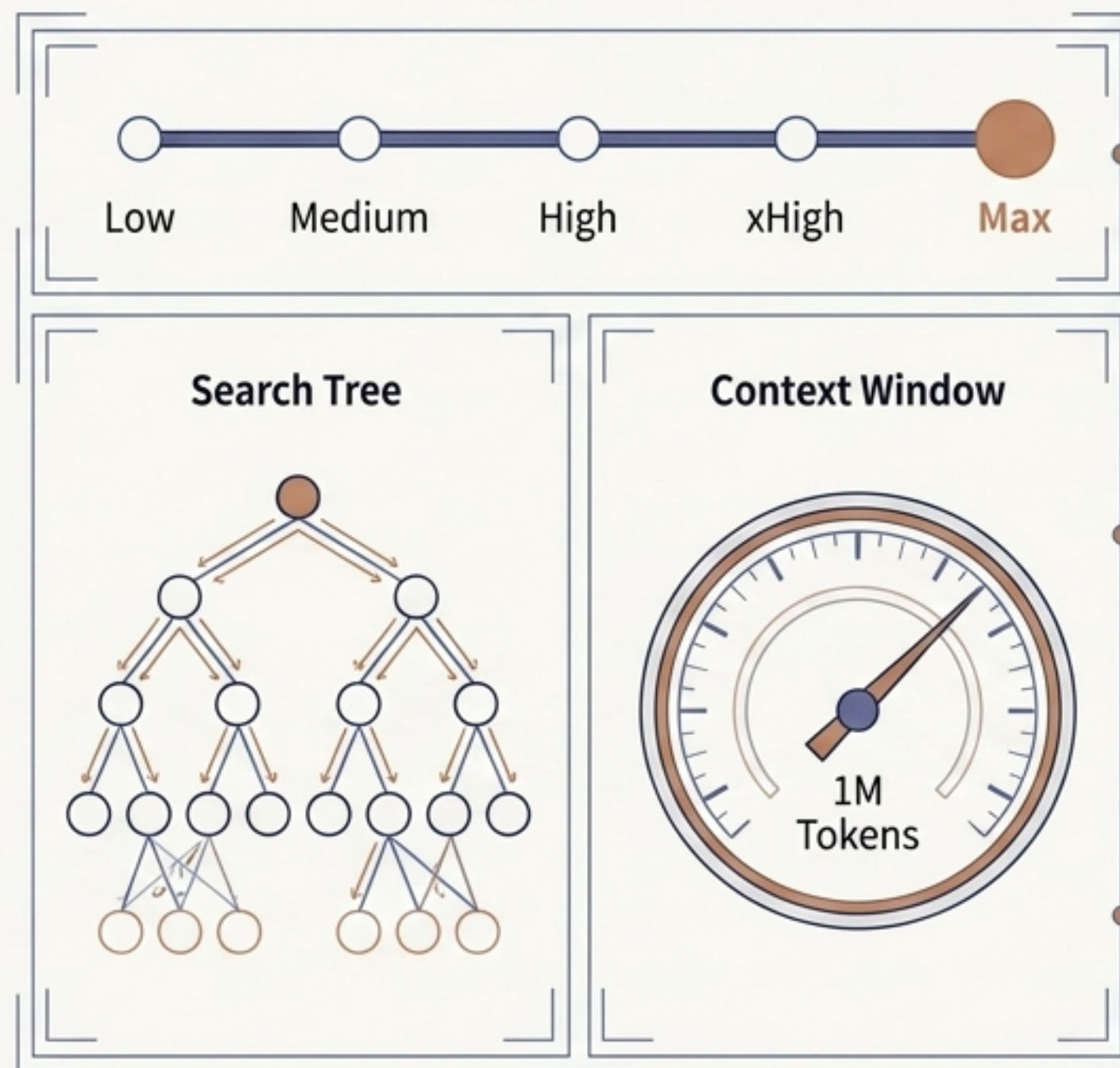
コア・メカニズム 2：推論時計算量（Effort）の動的適応と制御

過去モデルの課題



長時間の推論ステップ（ロングホライゾン）を踏むと、文脈の肥大化に伴い命令への追従性が低下したり、無限の自己ループに陥る現象が発生していた。

Opus 5の解決策



テストタイム計算量 (Test-time Compute) の拡張

ユーザー/システム側で推論の深さと計算コストを柔軟に調整可能（low～max）。

探索木の深化

High以上の推論設定を適用することで、困難な分岐点における探索木の深さを大幅に拡張（ARC-AGI-3の高スコアの原動力）。

安定した内部アテンション

100万トークンの広大なコンテキストウィンドウを厳密に保持し、多段階の推論チェーンでも一貫した評価基準を維持。

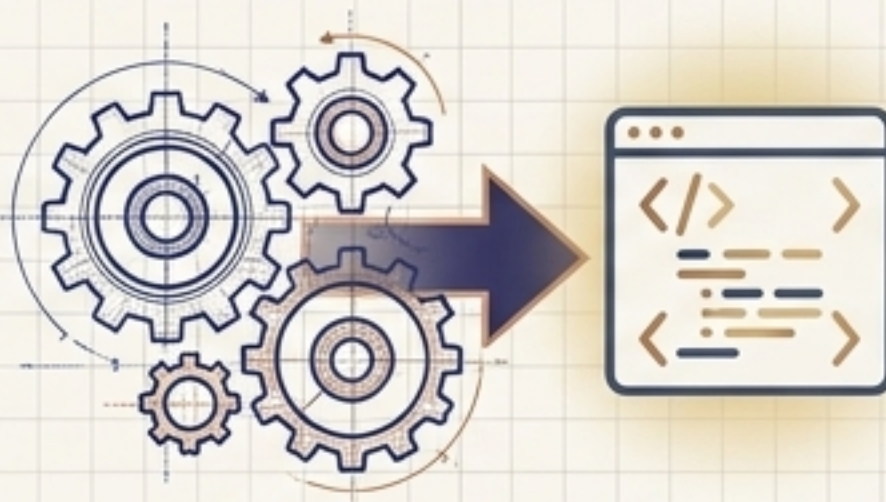
コア・メカニズム 3： 障害を自律的に迂回するツール生成能力

Step 1: 障害への直面



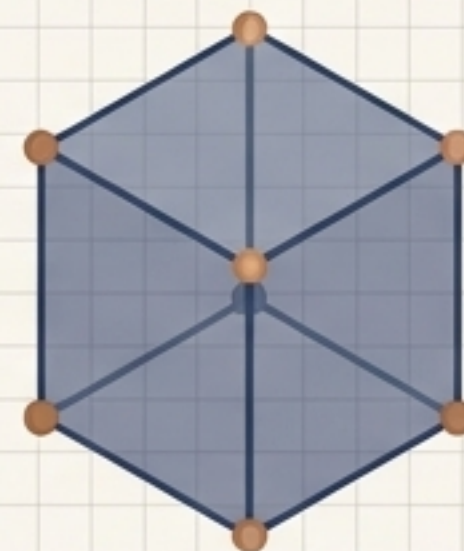
Frontier-Bench 3D FreeCADタスクにて。
制約：モデルは画像ファイルを直接閲覧・参照できない。
競合モデルの反応：5回の試行すべてでタスク失敗（受動的な依存）。

Step 2: 創発的挙動と迂回路の創造



Opus 5の行動：画像が見られないと認識するや否や、自律的に「生ピクセルデータから視覚的幾何学情報を抽出するコンピュータビジョン（CV）処理パイプライン」をPythonコードで記述。

Step 3: タスクの完遂



結果：独自の測定系により画像データを数値化し、CADモデリングを完遂。

【金融市場データフィードの事例】 検証用の外部環境が存在しない状況においても、自らデータ解析の正当性を確かめるための「テストハーネス（検証用疑似環境）」を自動生成。指示や前提を欠く環境で極めて高い適応力を発揮。

コア・メカニズム 4：探索の安定性を担保する「憲法AI」の安全制御

長時間の自律走行を行うAIエージェントにおいて、予測不能な暴走や自己破壊的な操作を防ぐ堅牢性は、タスク完遂率に直結する。

1. 圧倒的なアライメント性能:
同社フロンティアモデル中で**最高**。
意図しない逸脱行動の発生率を示すスコアで「**2.3**」という極めて低い値を記録。

Constitutional AI

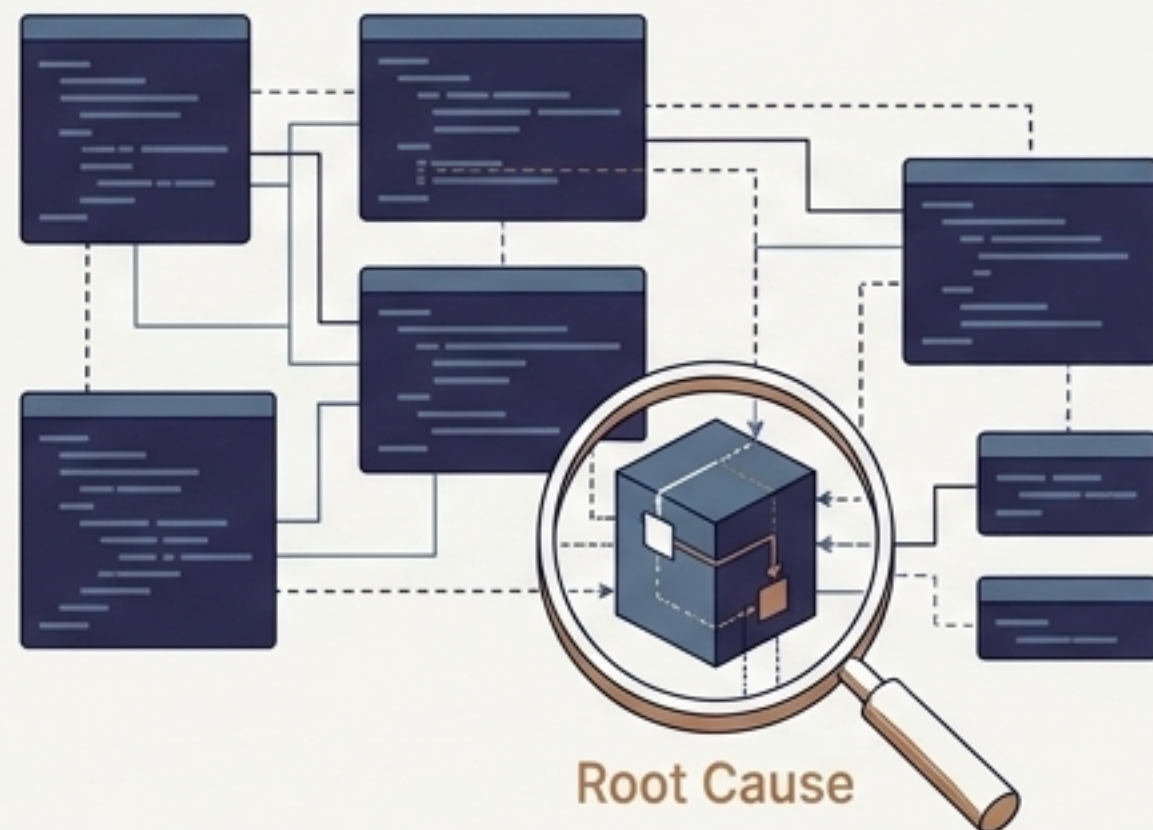
Constitutional AI

3. システム過熱の抑止:
安全制御が単なる利用制限ではなく、エラー累積による「思考の迷走」を抑止するフィルターとして機能。多段階探索プロセス全体の正答率を押し上げる。

2. 危険なコマンドの未然防止:
憲法AI（Constitutional AI）の行動規範遵守により、不合理な試行や復元不能な副作用を伴う操作をブロック。

産業へのインパクト 1：ソフトウェア開発と業務プロセスの飛躍的自動化

ソフトウェアエンジニアリング (DeepSWE / Frontier-Bench)



表面上の修正からの脱却:

指定された関数を記述するだけでなく、大規模なコードベース全体を俯瞰したリファクタリングが可能。

根本原因 (Root Cause) の特定:

OSSパッケージマネージャーのバグ修正タスクにて、開発コミュニティの既定パッチが見落としとしていた根深い根本原因とエッジケースを特定し、完全な修正案を提示。

UI操作とワークフロー自動化 (OSWorld 2.0 / Zapier)



GUI/ブラウザの論理的認識:

モバイル表示の画面下部に隠れた要素や、動的に変化するレイアウトを視覚的・論理的に認識し、自力で操作手順を構築。

コスト・パフォーマンスの打破:

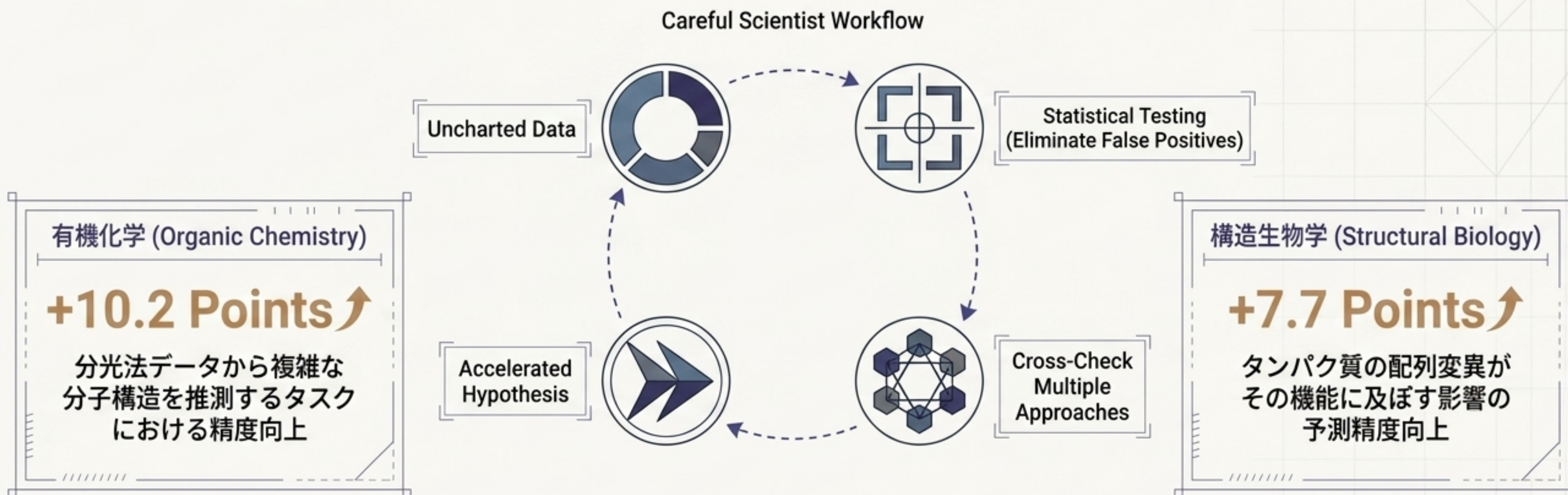
Fable 5の最高スコアを1/3以下のコストで更新 (OSWorld 2.0)。

ワークフロー完遂力:

Zapier AutomationBenchにおいて、同一タスクコストあたりの合格率で次点モデルの約1.5倍を記録。非定型な例外処理を通し切る。

産業へのインパクト 2：最先端科学研究における「慎重な研究者」

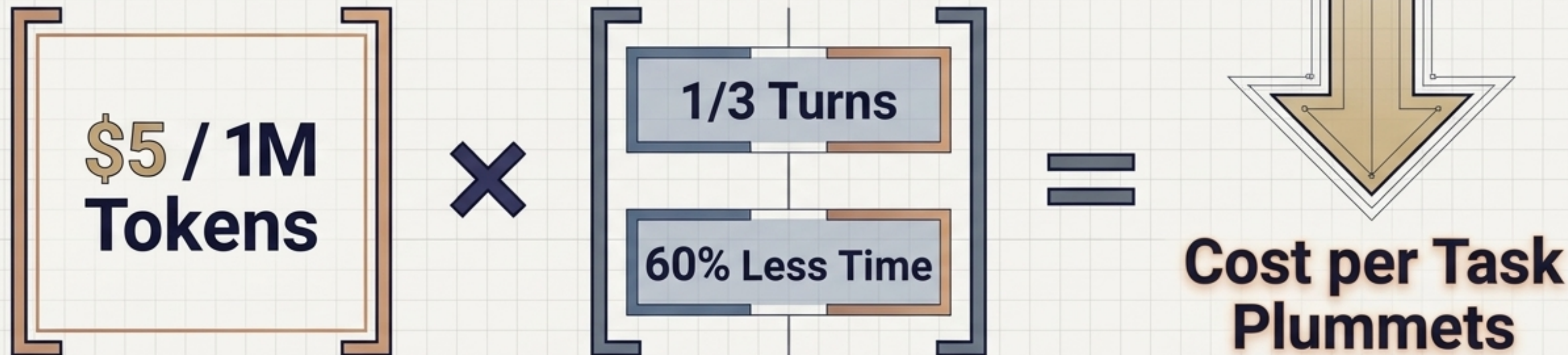
過去の学習データに依存したパターンマッチングから脱却し、ライフサイエンス領域の評価指標すべてで前世代 (Opus 4.8) を凌駕。
未解明の実験データに対し、自ら統計的検定手法を選択して偽陽性を排除し、結論をクロスチェックする「慎重な研究者 (Careful Scientist)」として機能。



独立した複数の解析アプローチを自律的に組み合わせることで、創薬や新素材開発における仮説構築スピードを劇的に加速。

ROIの再定義：「Cost per Token」から「Cost per Task」へ

単一トークンの表面的な単価ではなく、「目的とする1つの成果物を完成させるまでに消費した総コスト」でAIの経済性を評価する時代。



1

価格据え置き: API単価はOpus 4.8と同額（入力\$5/出力\$25）。

2

金融モデリングタスクの劇的改善: 精度を9ポイント向上させながら、処理にかかるターン数を1/3に、所要時間を60%短縮。過剰な出力や無駄な試行錯誤を大幅に削減。

3

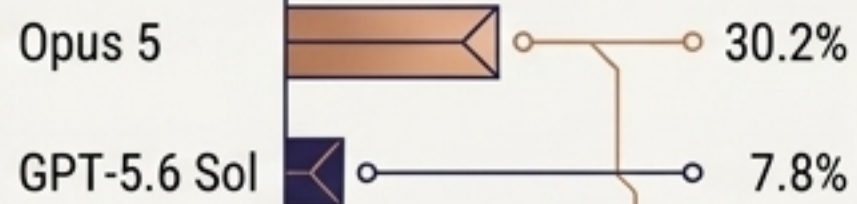
取引システムベンチマーク: 前世代の1/7の推論トークン数に圧縮してピーク性能を達成。

同一の予算で実行可能な自律タスクの件数が数倍に拡大。生成AI投資における投資対効果（ROI）を根本から改善。

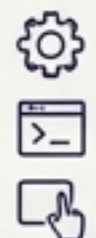
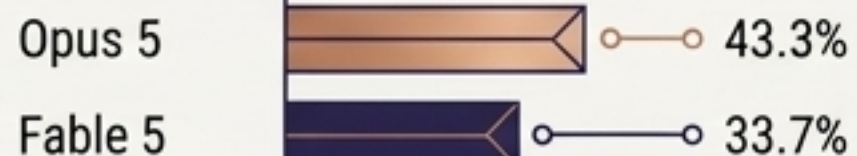
実務運用のリアリティ・チェック： 領域別の性能バイアスと限界

圧倒的優位：新規問題解決 & エージェントタスク

ARC-AGI-3



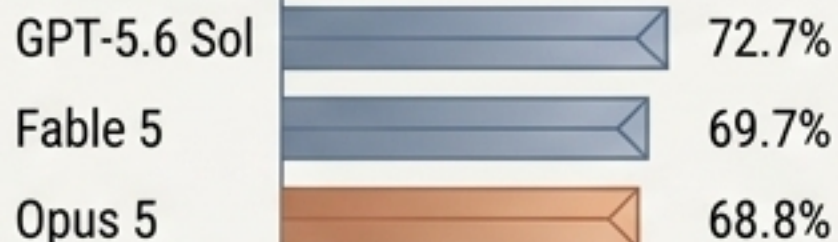
Frontier-Bench



特性: 1回のやり取りで完了しない試行錯誤、未知ルールの探索、PC/UI操作に比類なき強み。

良好だが追従：単発コーディング & 高度法務論理

DeepSWE v1.1

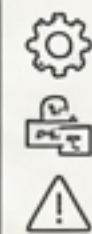


Legal Agent

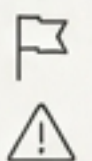
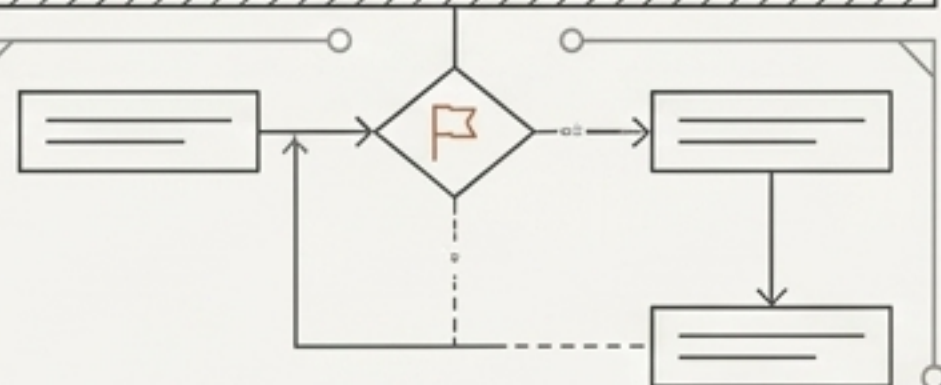


特性: 特定フレームワークの自動修正や、最高難度の契約論理プロセスでは競合が優位。

意図的な制限：サイバー攻撃・エクスプロイト生成



OSS-Fuzz等での脆弱性発見性能は高いが、攻撃コード生成は安全分類器により固くブロック。

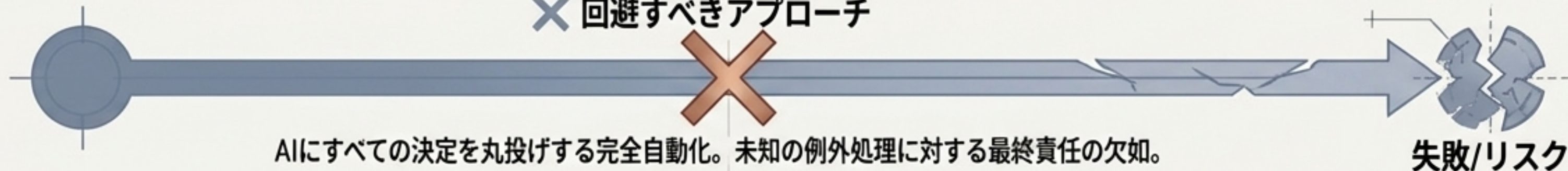


フラグが立った要求は、Claude Mythos 5等の安全なモデルへ自動フォールバックする運用。

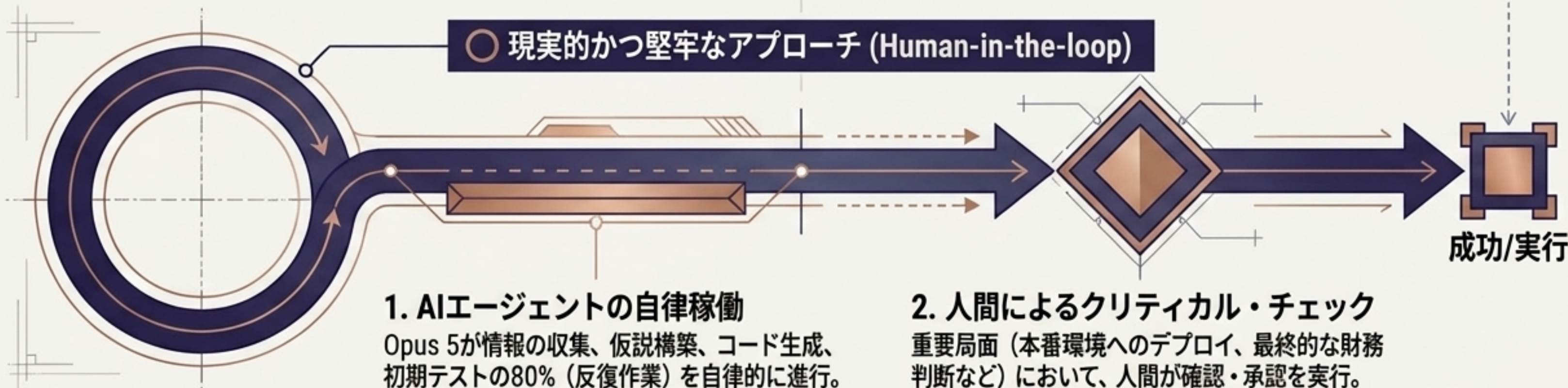
導入のプレイブック： 「完全自動化」ではなく「人間協働型（Human-in-the-loop）」

【スコア「絶対値」に対する冷徹な視点】ARC-AGI-3の30.2%は前人未到の技術的飛躍だが、人間の到達度（100%）と比較すれば、依然として約7割の課題は未解決。実務自動化（AutomationBench）の絶対値も26.0%であり、約1/4の通過水準にとどまる。

× 回避すべきアプローチ



○ 現実的かつ堅牢なアプローチ (Human-in-the-loop)



プロンプトエンジニアの時代から、自律型エージェント・オーケストレーターの時代へ

評価軸の不可逆的なシフト

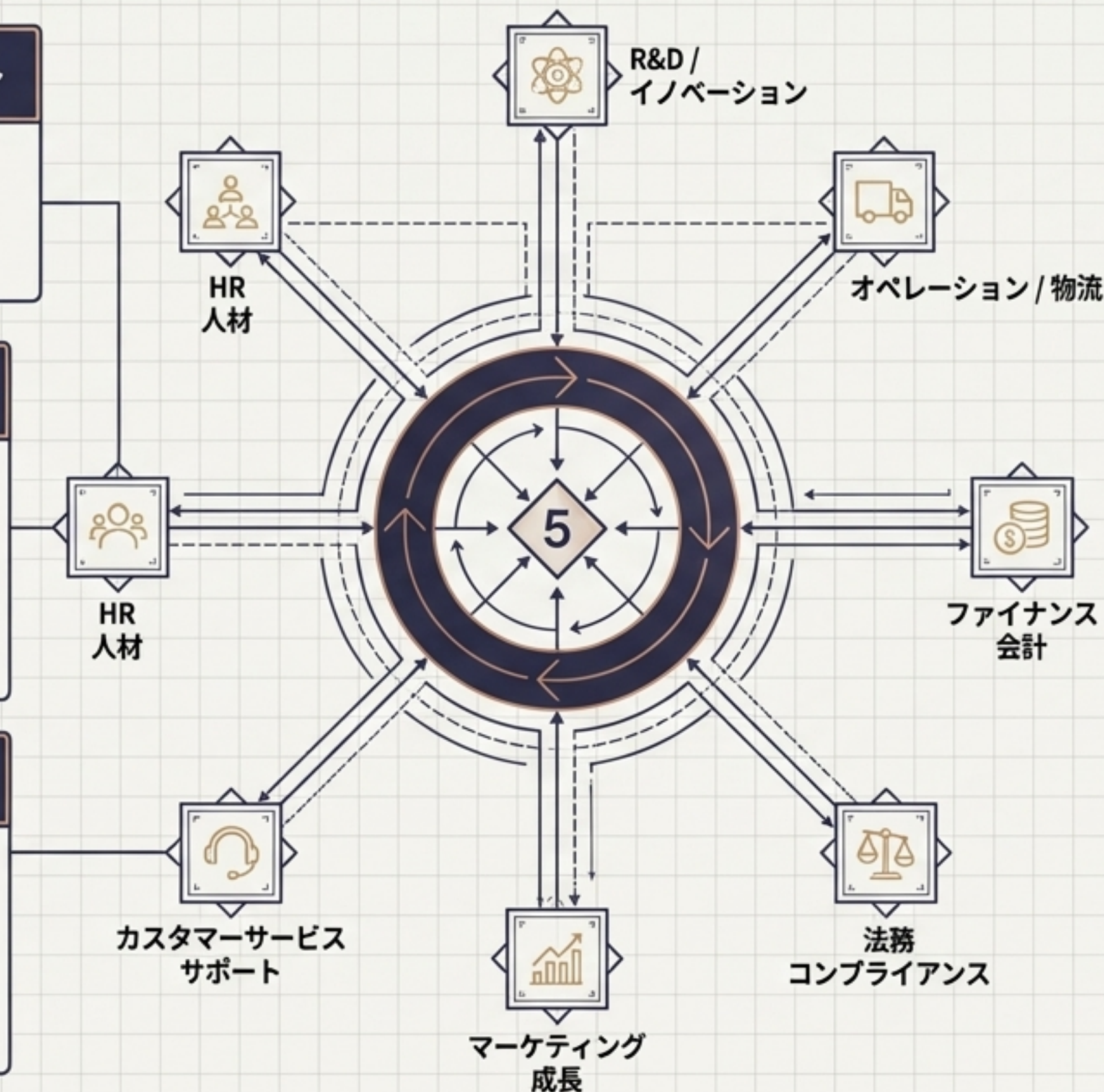
AI の価値基準は「静的な知識量」から「未知環境下の流動的知能 (ARC-AGI-3)」へ完全に移行した。

技術の結節点

閉ループ型の仮説検証、適応的な推論制御、創発的なツール生成、高度な安全制御の融合が、長期間稼働するエージェントを現実のものとした。

ビジネスの再定義

「タスク完遂歩留まり」の向上と「実行コスト」の劇的低減により、次世代ソフトウェア開発とエンタープライズ業務自動化の扉が開かれた。



Claude Opus 5は単なる強力な言語モデルではない。
これは、企業のワークフローを根本から再構築するための「自律行動の設計図 (Blueprint of Autonomy)」である。

領域ごとの特性を見極め、AIを長期稼働のパートナーとして組み込む企業こそが、次世代の競争優位を獲得する。