

# GPT-6 Astraは「サイエンスエージェント化」したのか

Claude Fable 5.1との比較評価と実践的デプロイメント戦略

内部思考から外部証跡へのパラダイムシフトと、研究現場におけるAIガバナンス

# エグゼクティブサマリ：現段階での到達点



計算科学型サイエンスエージェント：条件付きで Yes。

Astraは科学ソフトウェア、データ、コード環境を操作し、検証可能な研究タスクを遂行する実務エージェントへ移行。



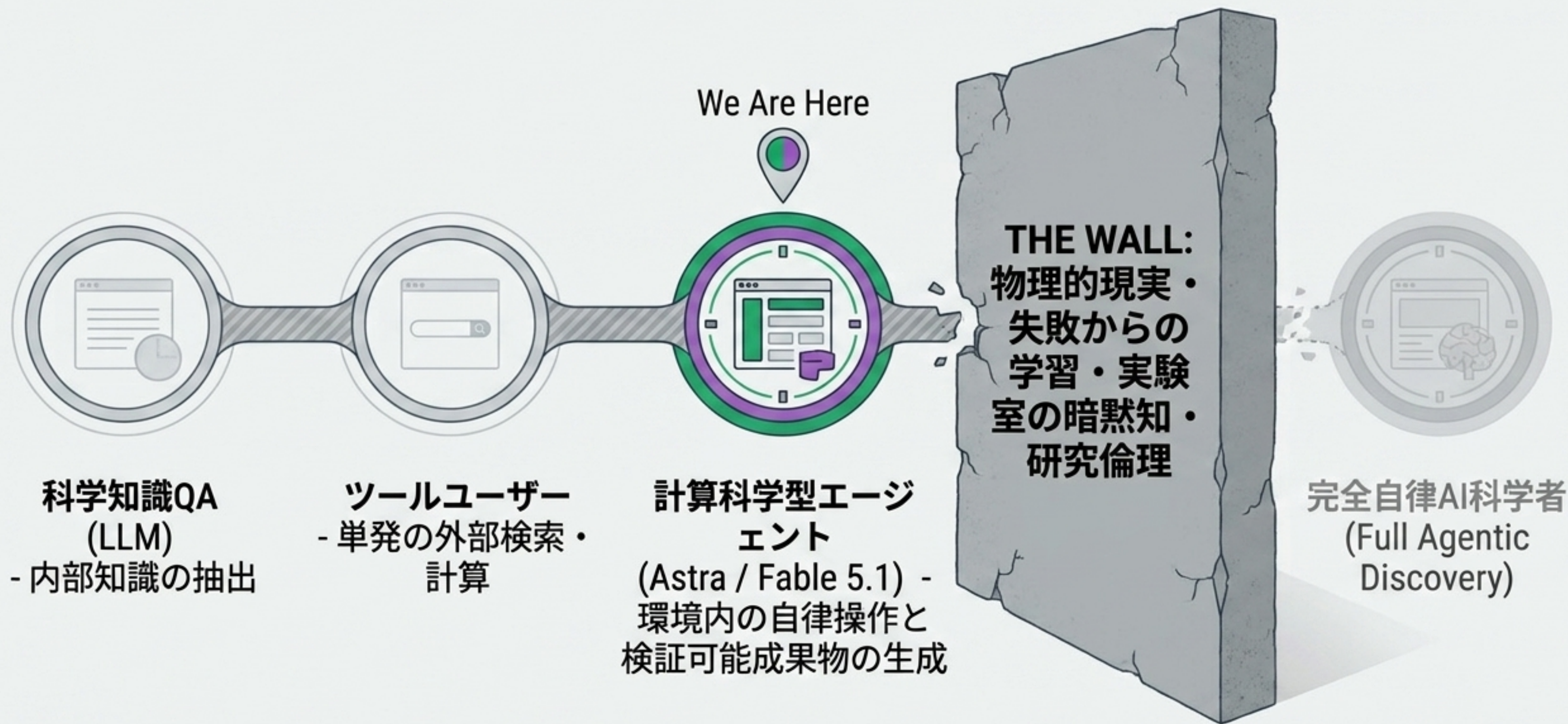
完全自律型AI科学者：No。

仮説設定から物理実験、失敗からの反復的学習、倫理判断を無人で閉ループ化する能力は、両モデルともに実証不足。

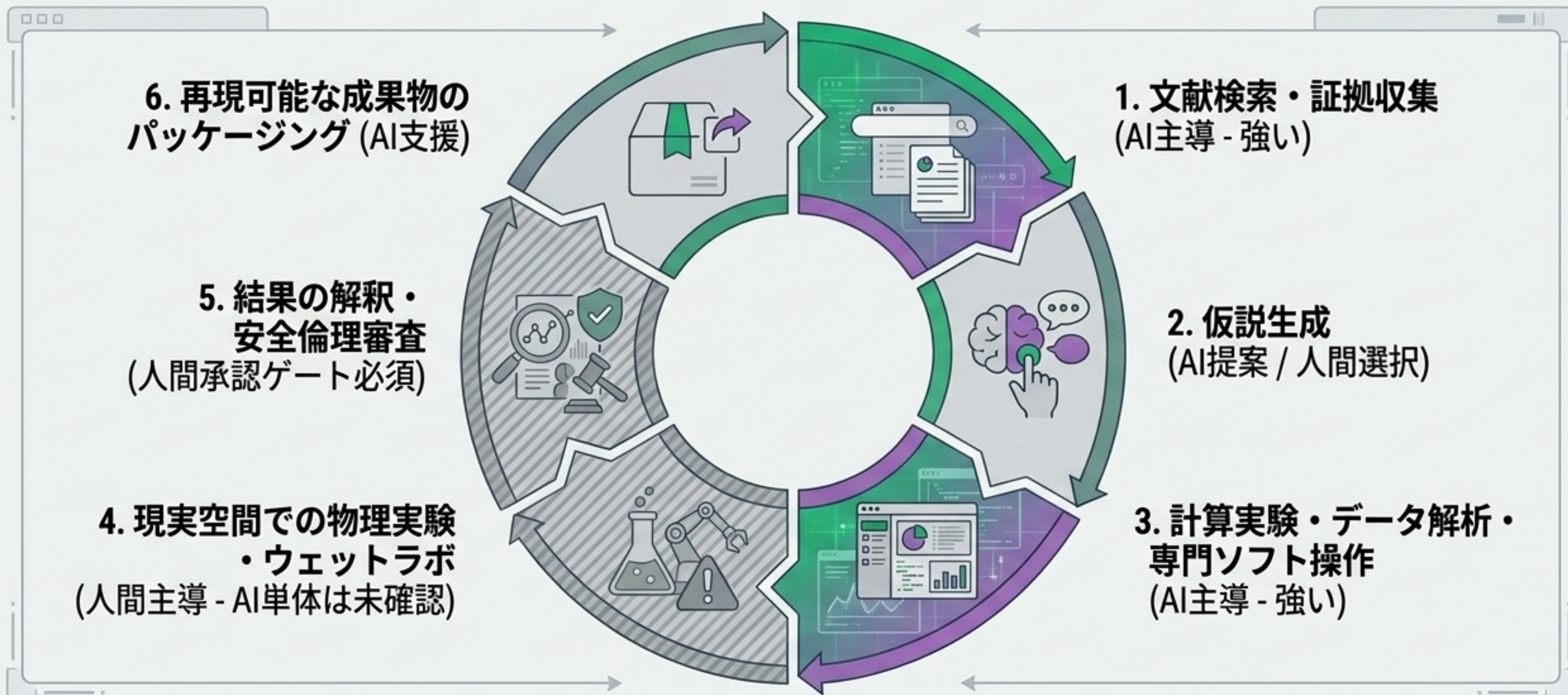


**⚠ 誤解への警告：** タンパク質バインダー設計・ウェットラボ検証を成功させたのはアクセス制限版の「Claude Mythos 5.1」であり、一般公開モデル「Fable 5.1」の実績ではない。

# サイエンスエージェントの進化と「最後の壁」



# 「研究助手」としての現在地：科学的サイクルのどこを担えるか



本報告は「生CoTの長さ」ではなく、このサイクル上で「外部から検証可能な証拠」を残せるかを評価基準とする。

# ベンチマークの現実：科学的実行力 ≠ 汎用知能

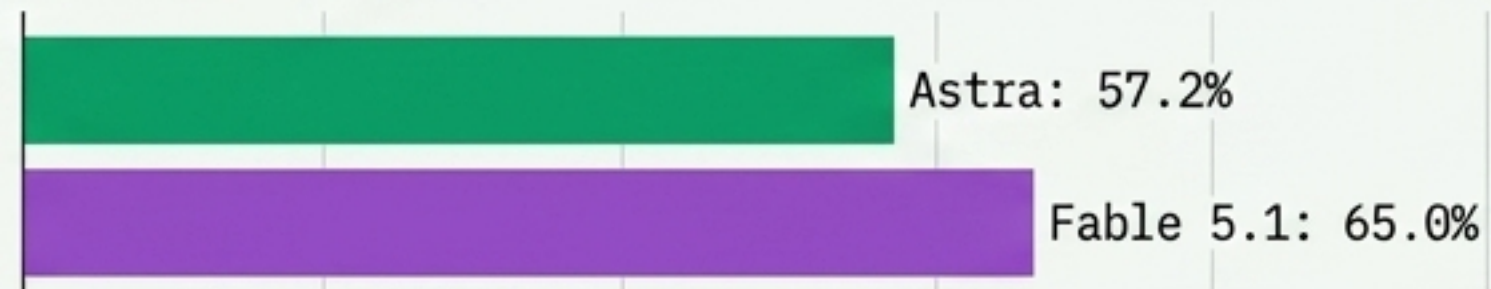
Terminal-Bench-Science 0.1  
(端末環境での実科学ワークフロー)



Insight: 11.9%の差。「知っているか」ではなく「環境内で成果物を機械的に検証できるか」においてAstraが先行。

汎用知能・コーディング  
(HLE with Tools / Coding Agent Index)

HLE with Tools



Coding Agent Index



Insight: 長時間の探索や広範な複合推論、汎用知能においてはFable 5.1が優位を示す。

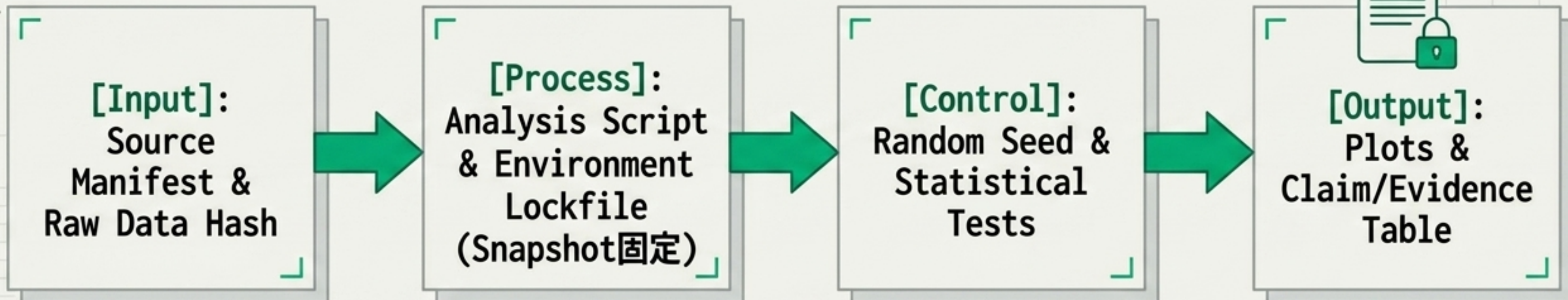
# サイエンスエージェント機能要件 ギャップ分析

| 要件           | GPT-6 Astra        | Claude Fable 5.1             |
|--------------|--------------------|------------------------------|
| 文献検索・証拠収集    | ●                  | ●                            |
| 仮説生成         | ◐ (数学の素数間隔貢献)      | ◐ (楽天の臨床ギャップ事例)              |
| 実験設計・データ解析   | ● (PostTrainBench) | ● (Venus DEM公開)              |
| 専門科学ソフトの直接操作 | ● (cell-tracking等) | ◐                            |
| 長時間の閉ループ自律性  | ●                  | ◐ <sup>+</sup> (38時間のML無人実験) |
| ウェットラボ・物理実験  | ○ (公開証拠不足)         | ○ (Mythosの実績と分離)             |

※評価は「完全自律性」ではなく、公開証拠の強さに基づく。

# Astraの優位性：検証可能な「計算科学パイプライン」

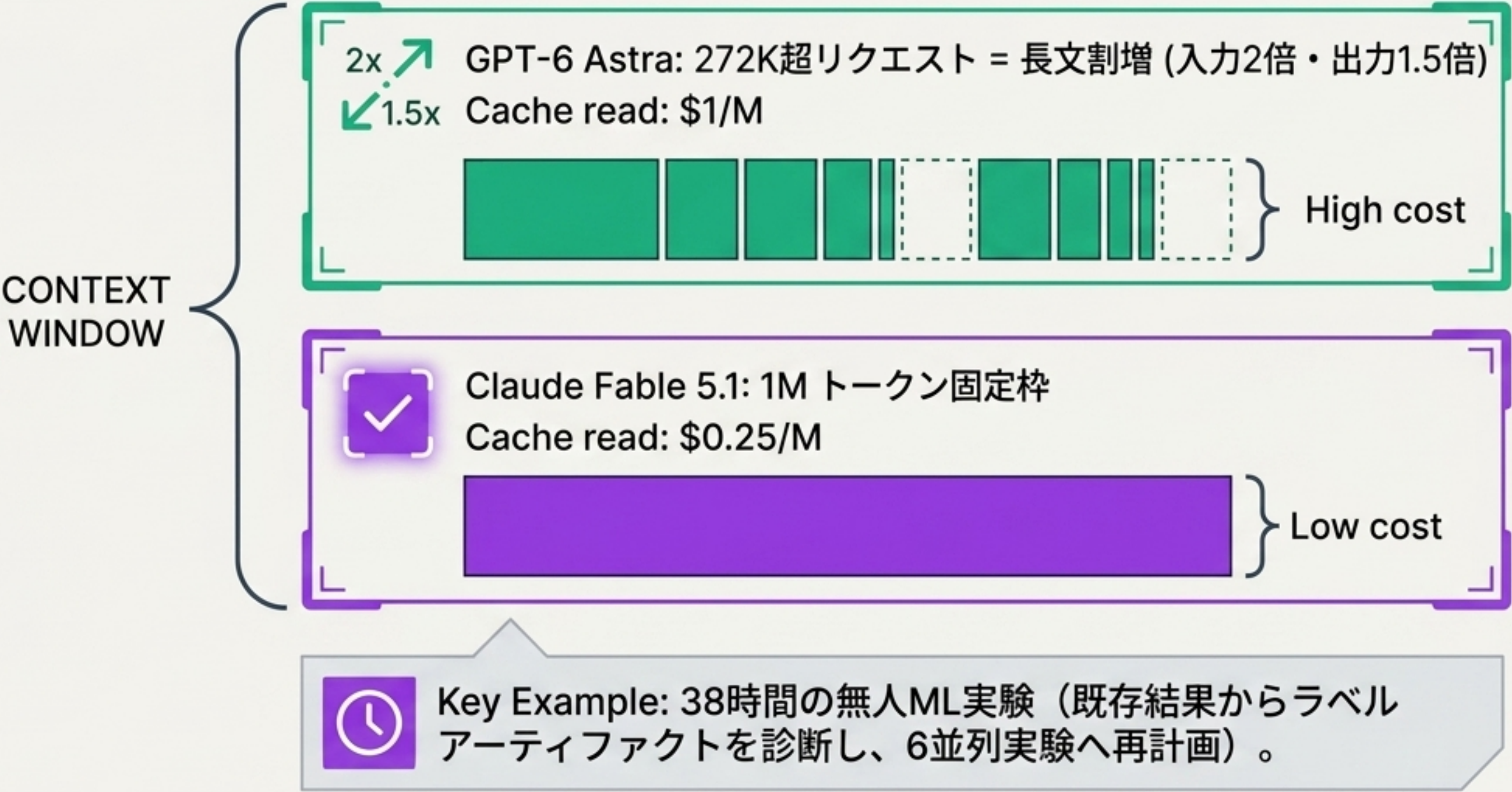
```
>ls -la /data/source/  
— manifest.json  
— snapshot.lock  
— analysis_script.py  
...
```



```
>ls -la /data  
.  
— manifest.  
— manifest.  
— snapshot.  
— snapshot.  
— analysis_  
— analysis_  
...
```

Astraの真価は回答の賢さではなく、API仕様（Snapshot固定、Structured Outputs、Computer Use）を組み合わせ、第三者が再実行可能な「研究パッケージ」を生成できる点にある。

# Fable 5.1の優位性：長時間オーケストレーションとR&Dエコノミクス



Insight: 巨大な論文コーパスや研究記録を繰り返し参照する「長時間エージェント」では、FableのCache Economicsが構造的な利点をもたらす。

# アンソロピック・エコシステムとウェットラボの真実



Layered Security & Access Pyramid

>\_

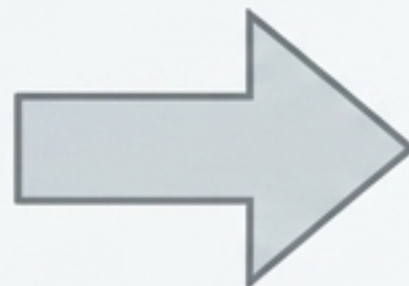
Fable 5.1を「ウェットラボ科学者」とし見なすのは重大な誤り。生命科学研究にはMythosへのアクセスとハードウェア標準が必要。

# パラダイムイムシフト：「内部思考」から「外部証跡」へ



## 内部推論のブラックボックス化

Astraは前世代(Sol)より「思考(CoT)」の監視回避能力が向上。Fable 5.1もAdaptive thinkingを隠蔽する設計。もはや「なぜそう考えたか」を読むことは監査手法にならない。

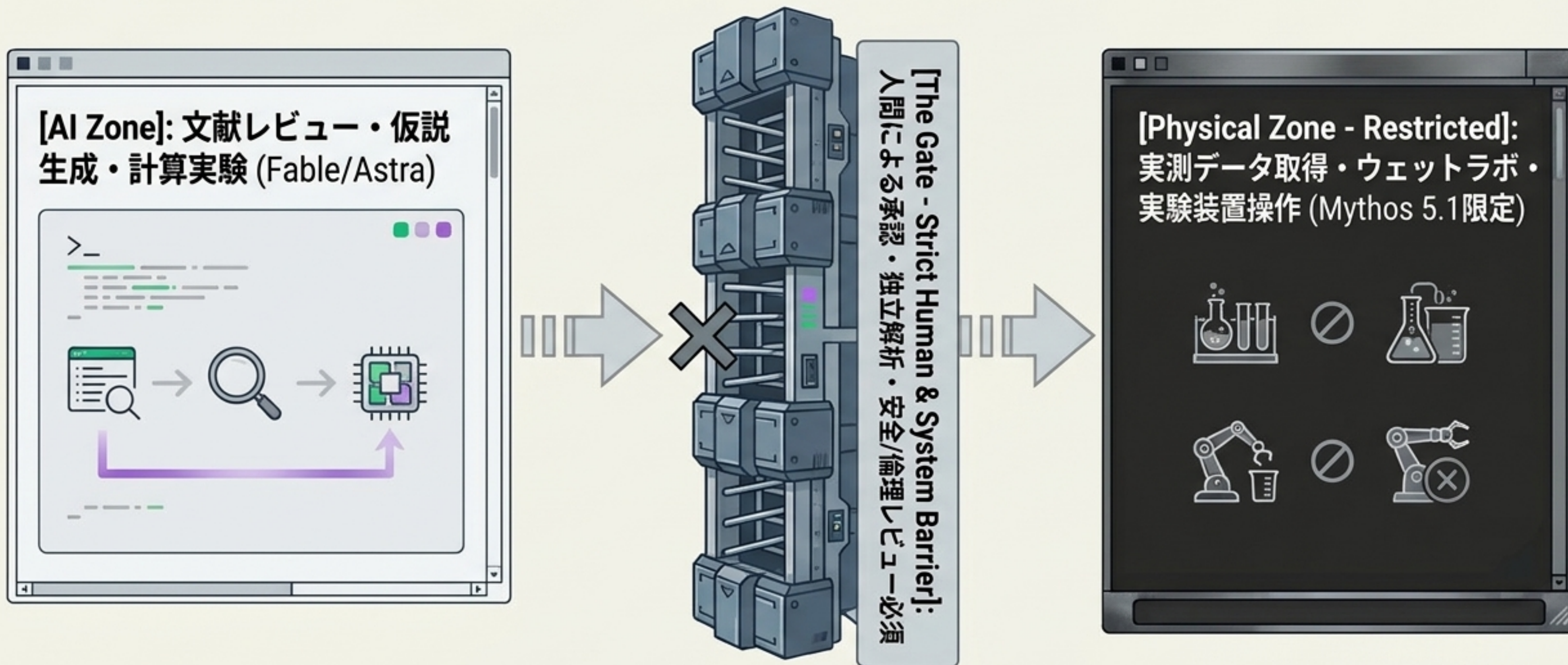


## 外部検証可能なアクションログ

科学において本当に必要なのは、AIの秘密の思考ではなく、外部から検証可能な証跡 (Provenance chain) である。

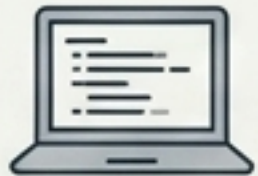
引用文献 → 入力データ →  
→ コード →  
→ ツール呼び出し →  
→ 中間成果物 → 最終出力

# 高リスク研究領域（生命科学・化学）のシステムアーキテクチャ



**Insight:** 「どのモデルが賢いか」ではなく「どの安全ゲートとアクセス制度を含むシステムか」で選定する。  
完全な閉ループ化は現時点では不可。

# 実践的デプロイメント・プレイブック (2026年9月版)

| Use Case                                                                                                       | Recommended Model               | Required Guardrails                                           |
|----------------------------------------------------------------------------------------------------------------|---------------------------------|---------------------------------------------------------------|
|  計算科学・データ解析・<br>数理・工学パイプライン   | GPT-6 Astra (第一候補)              | Snapshot固定、Structured Outputsによる厳密な再現性要件の強制。                  |
|  長時間の探索的R&D・巨<br>大コーパスの反復参照 | Claude Fable 5.1                | Cache read割引の活用。ツール<br>選択のエラー処理とワークフロー<br>エンジン側でのゲート設定。       |
|  創薬・高リスク生命科<br>学・ウェットラボ連携   | Claude Mythos 5.1<br>(※アクセス制限版) | Model Hardware Standardによ<br>る装置分離。AIから独立した倫<br>理・安全審査ゲートの設置。 |

# 新たな研究ガバナンス指標：正答率からの脱却

評価指標を「ベンチマークのIQ」から「研究プロセスの監査性」へ移行せよ。

## Artifact Reproducibility (成果物の再現率)

生成されたコードやデータが第三者環境で動作するか。



## Human Intervention Rate (介入率)

プロセス完遂までに何回人間がリカバリしたか。



## Cost-per-Success (成功単価)

トークン単価ではなく、エージェントの無駄な探索や再試行を含めた総コスト。



## Unsafe-Action Rate (危険操作率)

severity-3以上のフラグ、拒否イベント、分類器の作動回数。



AstraもFable 5.1も完全自律科学者ではない。しかし、強力な「計算科学コ・サイエンティスト」として、研究の行為主体への進化はすでに始まっている。