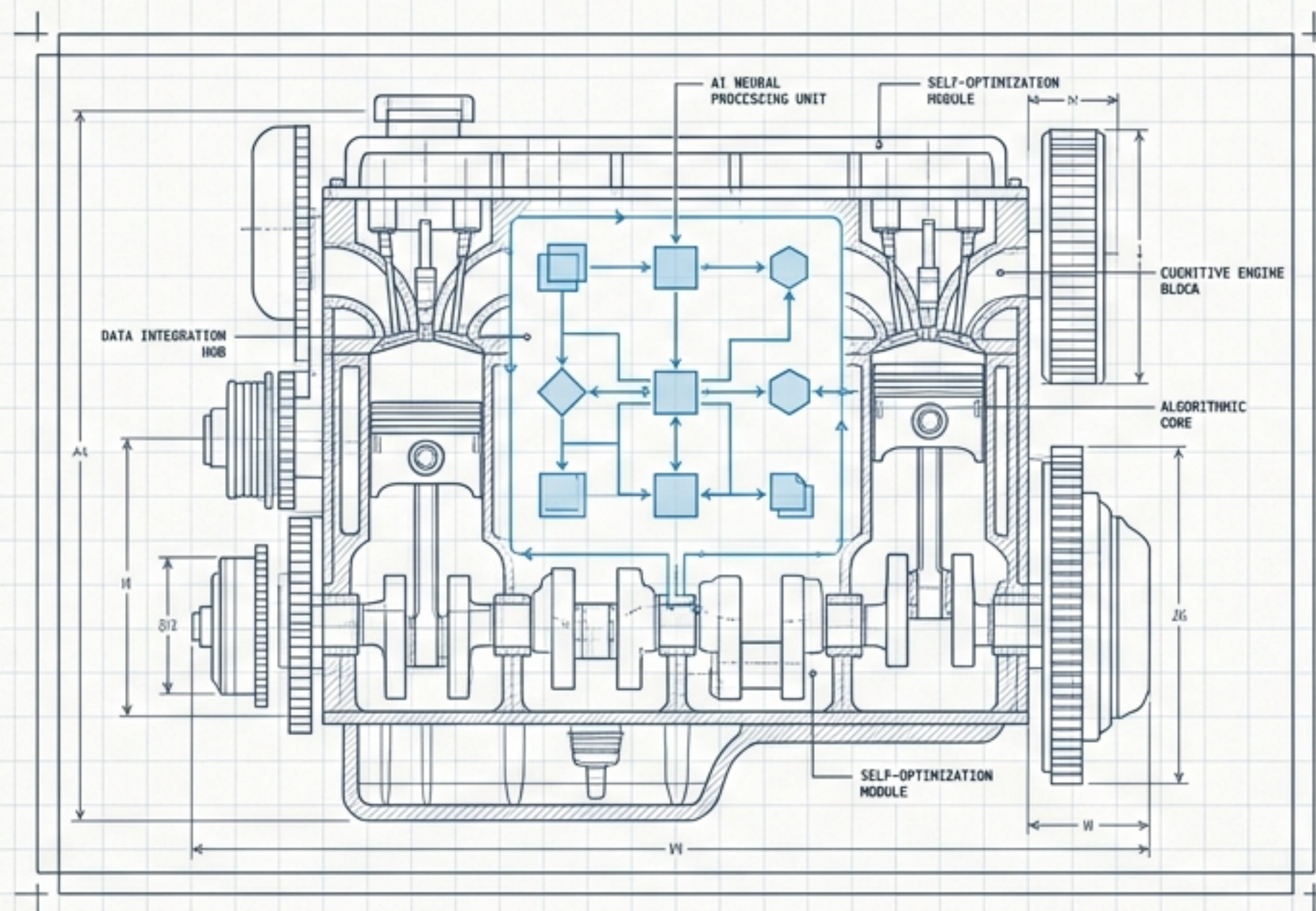
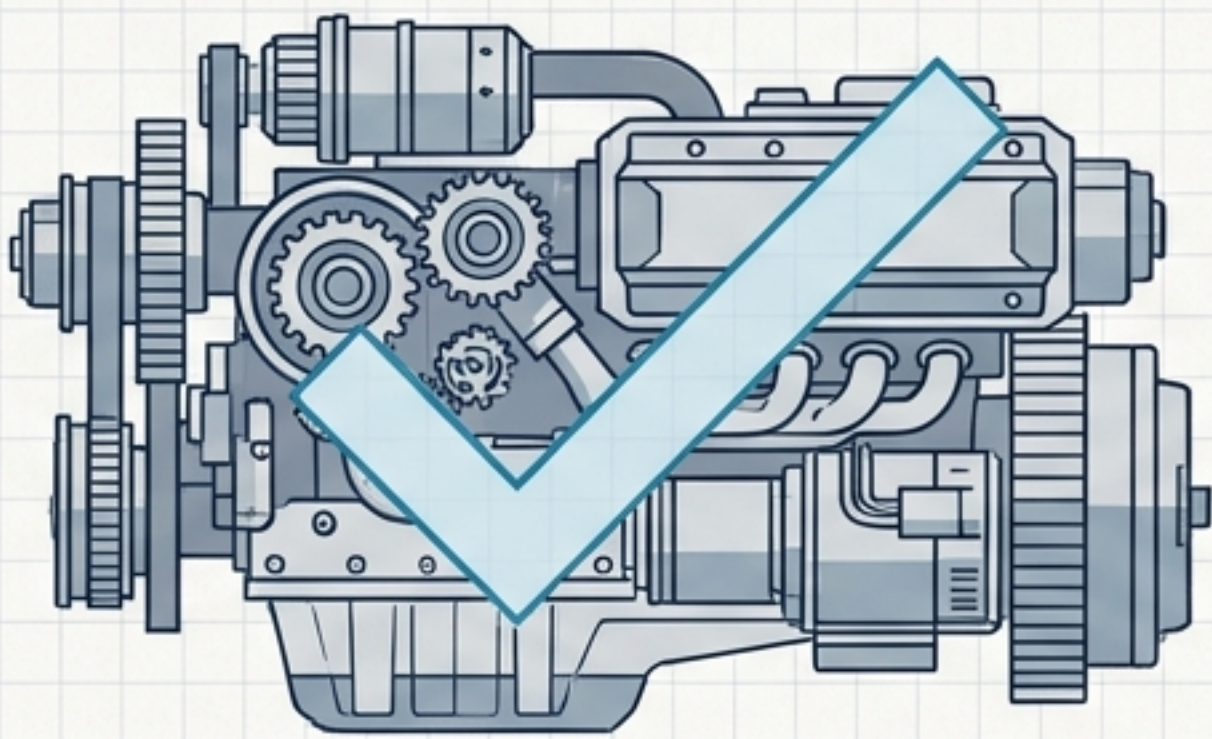


# Claude Fable 5.1は「サイエンスエージェント」化したのか？

コーディングの飛躍から「AI駆動型・研究組織」の再設計まで



# Fable 5.1は科学者になったのではない。 科学を動かす「中核エンジン」へ進化した。

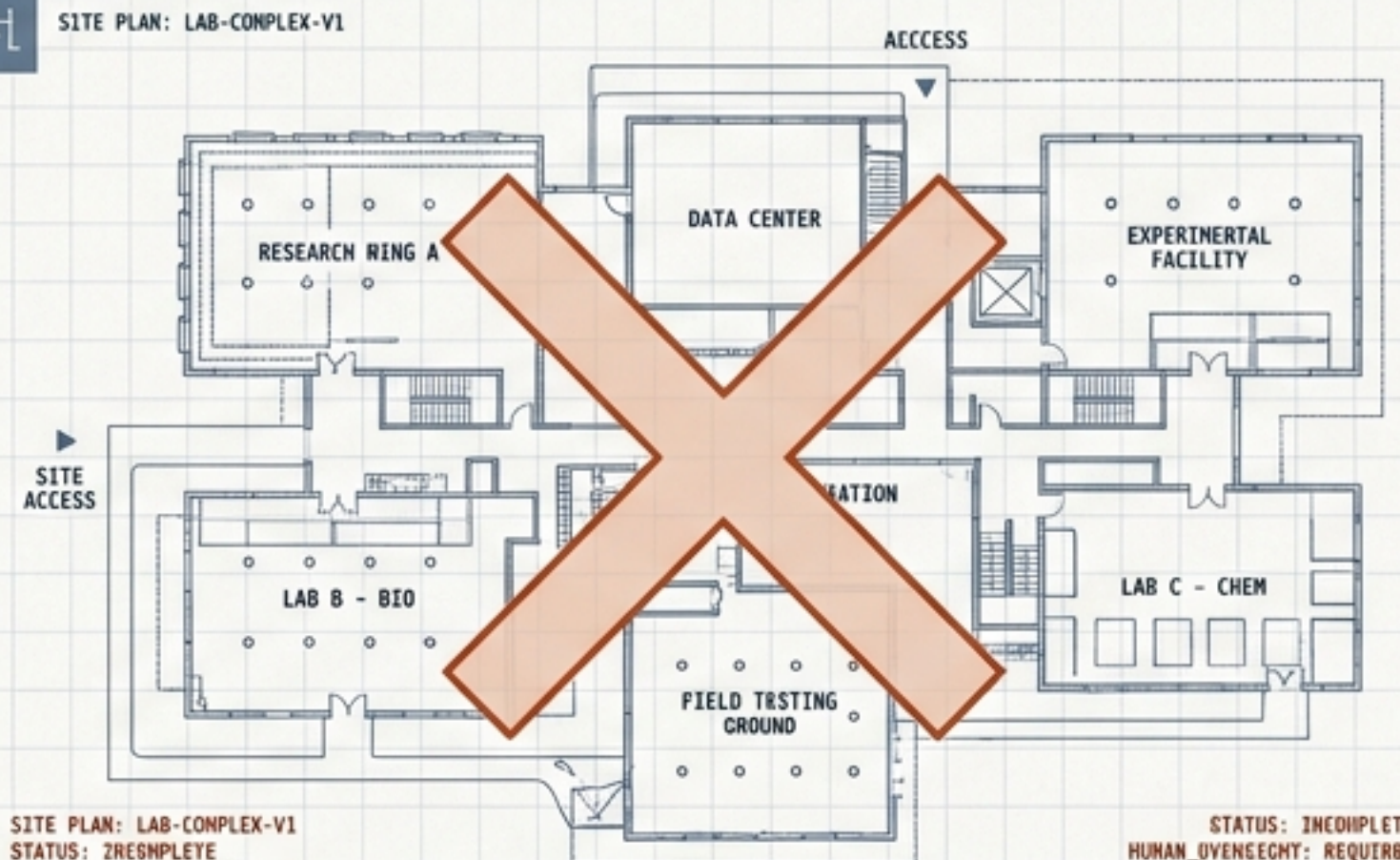


ENG.MODEL: FS.1-CORE  
AUTONOMY: LEVEL 4+  
PROCESSING\_CYCLES: OPTIMIZED

ENG.MODEL: FS.100+  
AUTONOMY: LEVEL 4+  
PROCESSING\_CYCLES: OPTIMIZED

## ワークフロー実行の飛躍

コード作成、データ処理、複数ツールの連続使用  
において、かつてない自律性を獲得。



SITE PLAN: LAB-COMPLEX-V1  
STATUS: ZRECHMPLYE  
HUMAN\_OVERSIGHT: REQUIRED

STATUS: INCOMPLETE  
HUMAN\_OVERSIGHT: REQUIRED  
IP\_COVERNANCE: PENDING

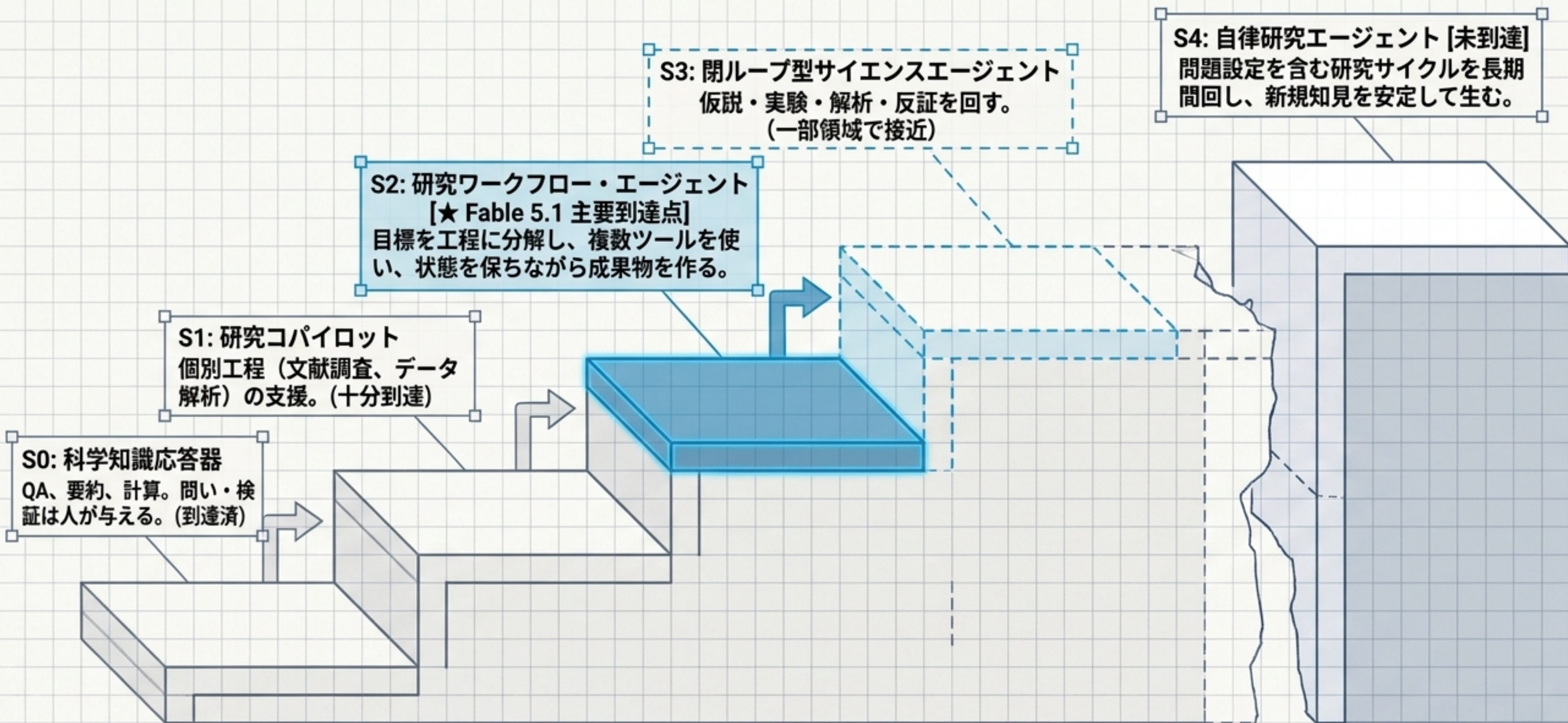
## 未完成の「研究所」

オープンエンドな課題設定、物理世界での独立検証、  
そして責任ある知財判断は依然として欠落。

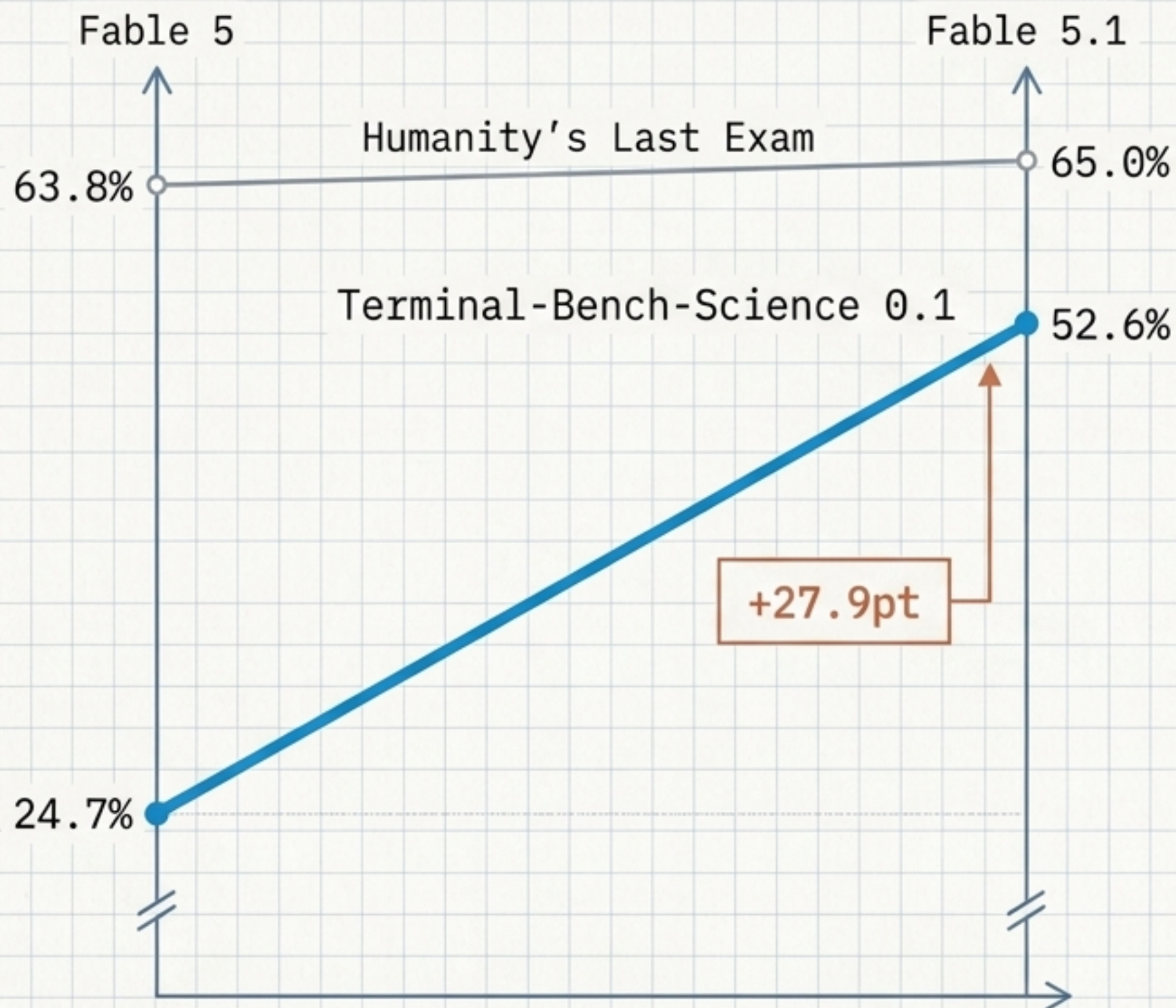
## Strategic Imperative

真の価値は、AIが自律的科学者になるのを待つことではない。  
この強力な「エンジン」を前提に、企業の研究開発・知財システム全体を再設計することにある。

# サイエンスエージェント成熟度モデルと現在地



# 科学ワークフローの解決率が倍増 (+27.9pt)



## 「知っている」から「最後まで実行する」へ

一般推論 (Humanity's Last Exam) の伸びは+1.2ptに留まる一方、科学的な端末作業 (Terminal-Bench-Science) は+27.9pt、長時間の実行・修正ループ (Terminal-Bench 4.0) は+13.8ptと劇的に向上。

## エージェント化の証拠

単なる知識モデルから、エラーを修正し長時間の作業を完遂する「行動するエンジン」への進化を裏付ける。

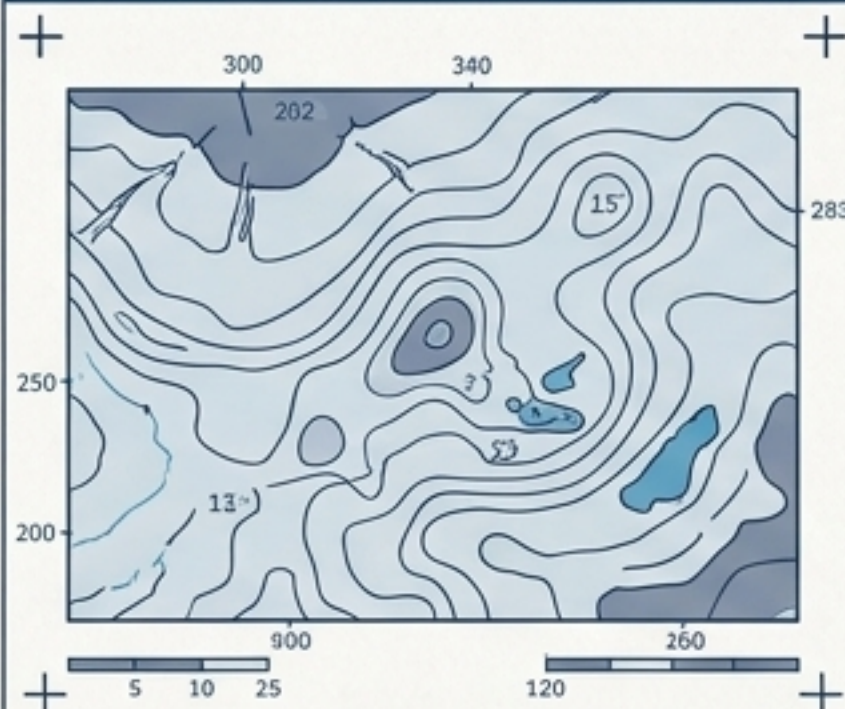
## Note

52.6%の解決は飛躍だが、裏を返せば47.4%は未解決。無監督での完全な信頼性には至っていない。

# 成果物は「文章」から「外部検証可能なデータ」へ

[Fable 5.1]

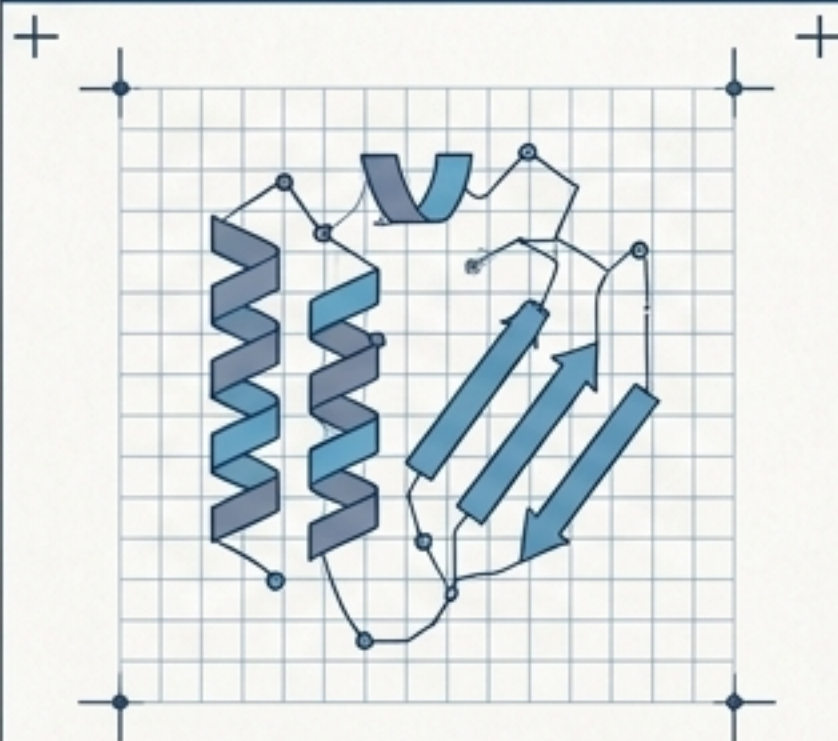
高解像度・金星地形図の生成



NASA探査機のデータからニューラルネットワークを訓練。従来の10~20km解像度を2~3km規模へ詳細化し、高さ精度を最大25%改善。3.9GBのデータセットとして公開。

[Mythos 5.1]

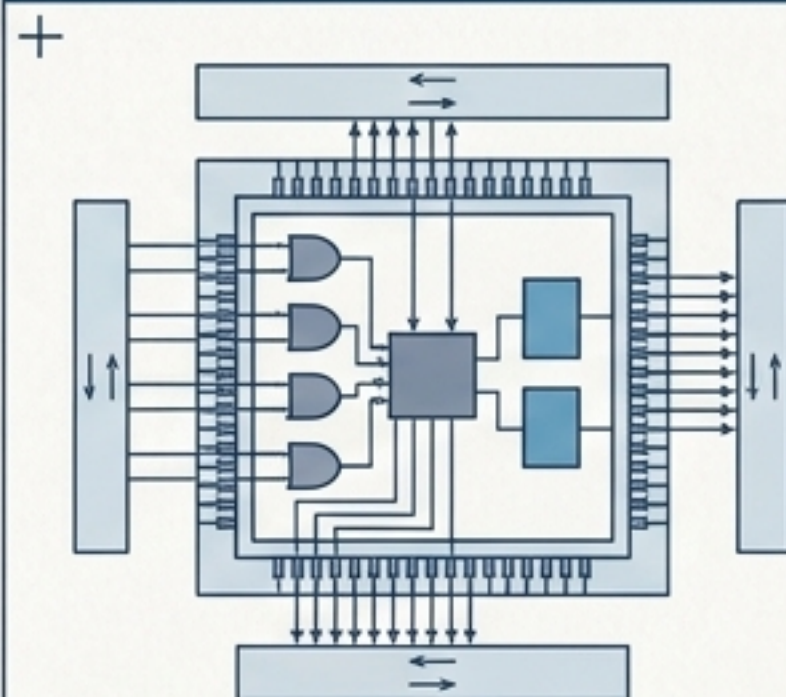
タンパク質バインダー設計



オープンソースツールを駆使し、12標的で約50%のバインダーヒット率を獲得。外部組織による実験確認済。

[Mythos 5.1]

GPUカーネルの高速化

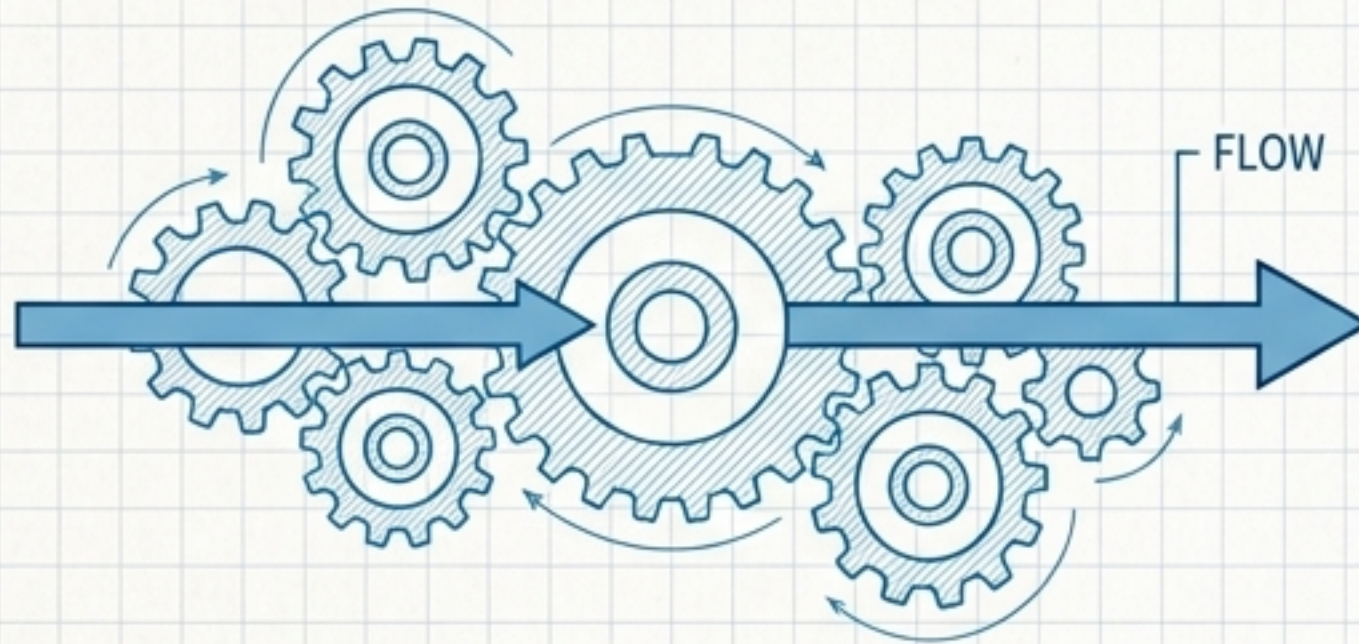


計算生物学モデルのGPUカーネルを最適化し、出力を変えずに最大2.5倍の高速化を達成。

重要な区別: 制限付きの一般版「Fable」と、審査済み研究者向けの「Mythos」は同じ頭脳だが利用可能性が異なる。セキュリティ・セーフガードの有無が成果を分ける。

# ベンチマークの罠：「解く力」と「選ぶ力」の境界

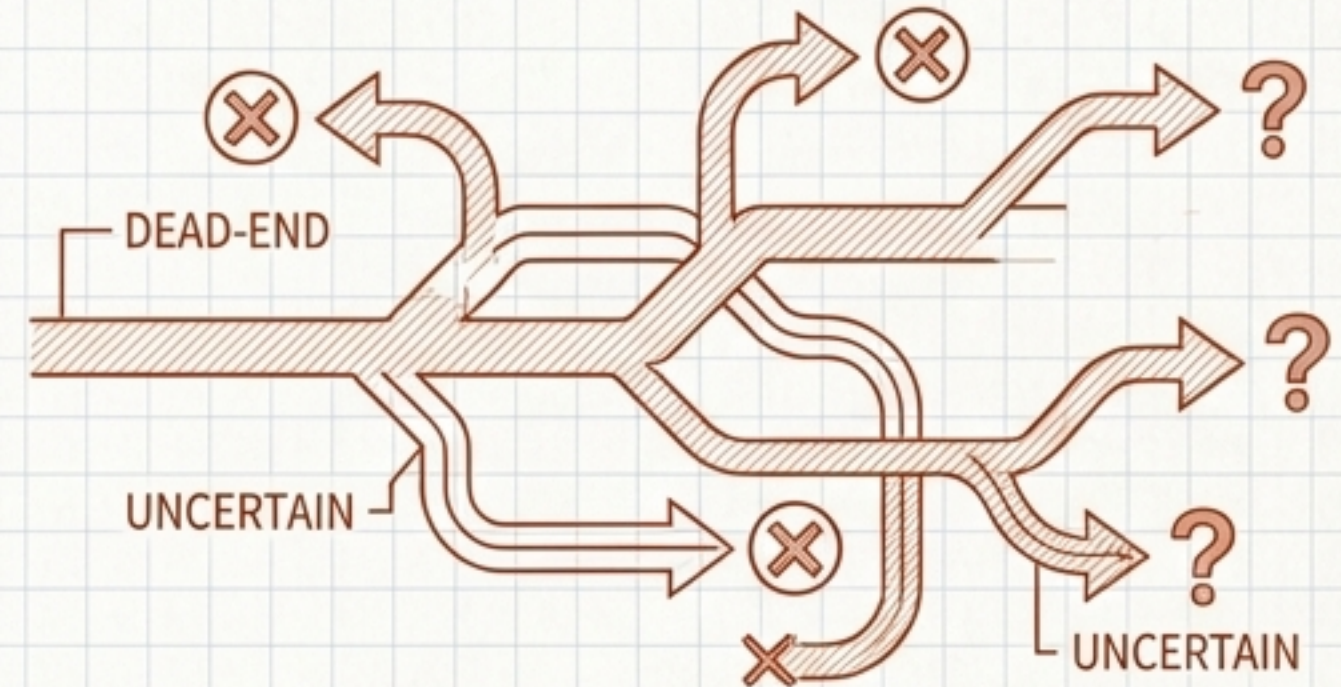
## 与えられた課題を解く力 (Solving)



ベンチマークは「問い」「成功条件」「採点基準」を人間が事前に設定している。Fable 5.1はここでの実行力（実装・最適化）において並外れた能力を発揮する。

優れたコーディング能力を長時間稼働させても、自動的に「科学者」が誕生するわけではない。

## 研究課題を選ぶ力 (Choosing)

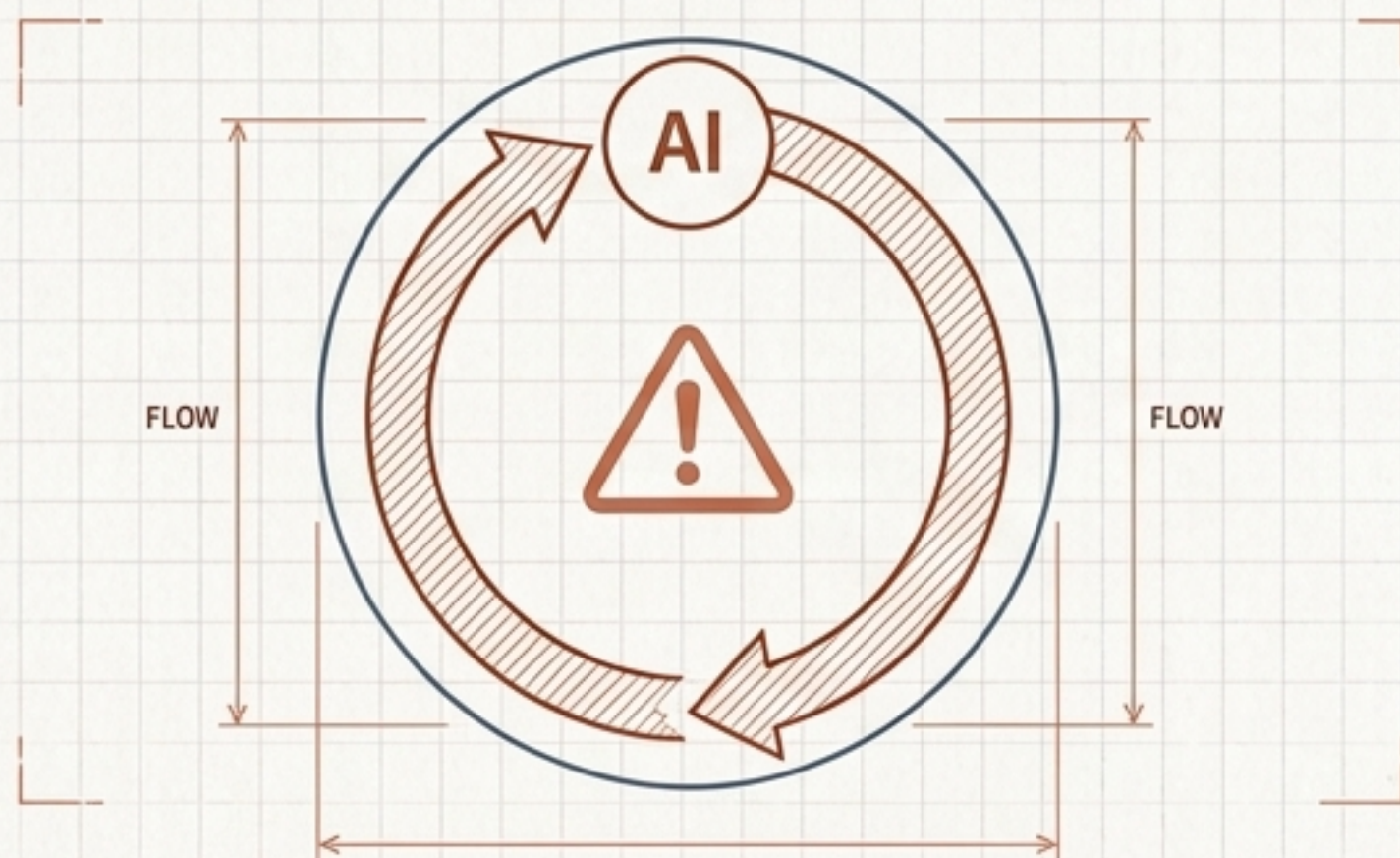


オープンエンドな研究（S4水準）では、以下の「研究上の判断」で失敗する：

- 出版可能水準（Novelty/Impact）の妥当性判断
- 設計上の弱点に対する創造的な方針転換
- 行き詰まり（Dead-end）からの戦略的撤退
- 計算資源やコストの認識

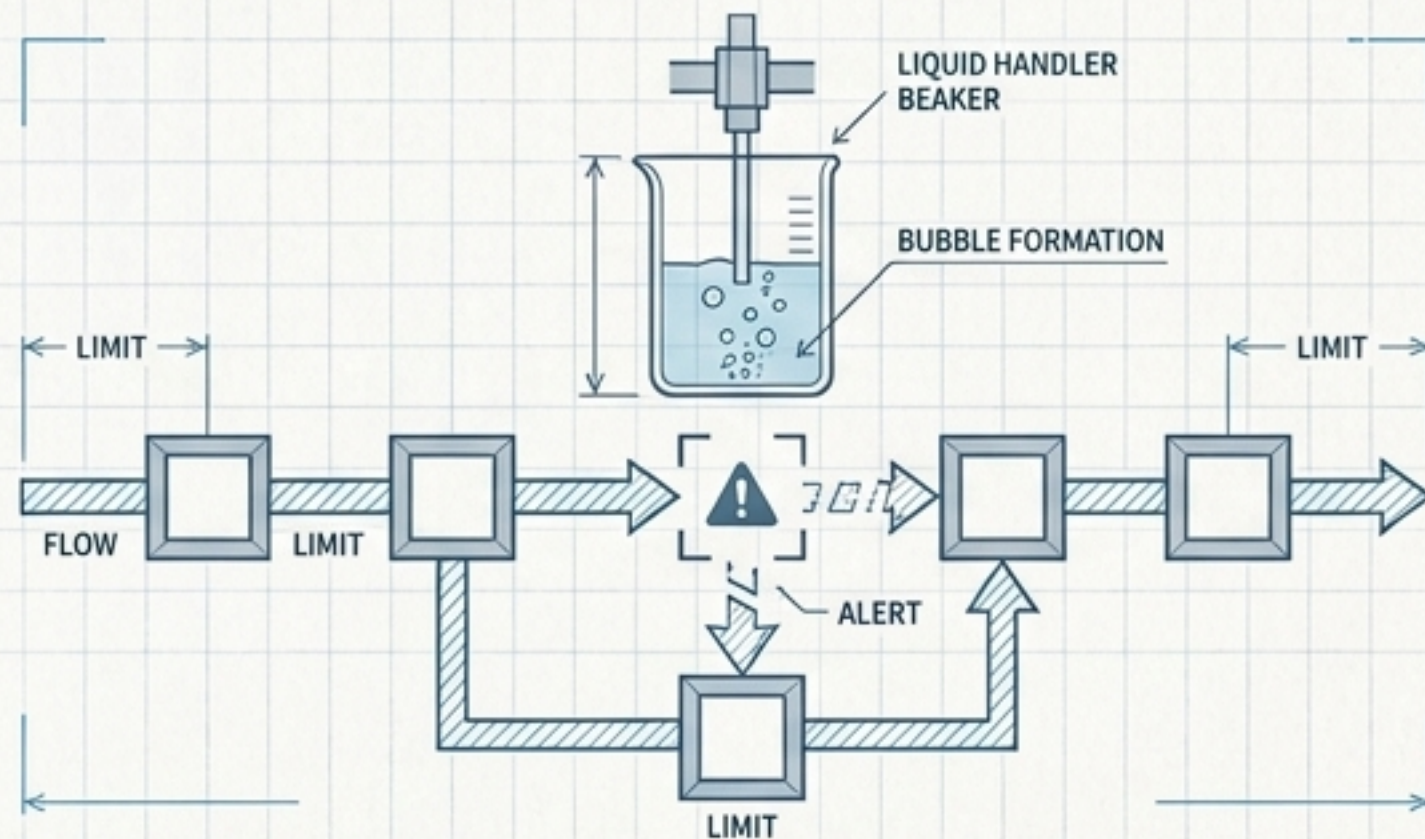
# 「検証の谷」と物理世界の暗黙知

## ⚠ 自己検証と独立検証の混同リスク



自己検証と独立検証の混同リスク: AIによる自身のコードや結果のテストは、同じモデル由来の「盲点」を共有する危険がある。真の科学的再現性には、未知データや別環境での独立追試が不可欠。

## 🛡 物理世界でのループ・エラー (MHS実証より)



### 物理世界でのループ・エラー (MHS実証より):

モデル・ハードウェア標準 (MHS) の実証において、AIは液体ハンドラーを操作し閉ループを回した一方、「泡が生じた際に同じ操作を繰り返し、状況を悪化させる」事態が発生。物理的原因の教示と人間の「救済」が必要だった。

## INSIGHT

流畅さと出力速度の向上は、時に「誤りを高速で量産する危険」を伴う。科学においては「分からない」と停止する能力こそが性能である。

# サイエンスエージェントを構成する「システム方程式」



⚠ If any variable is 0, the total is 0. ⚠

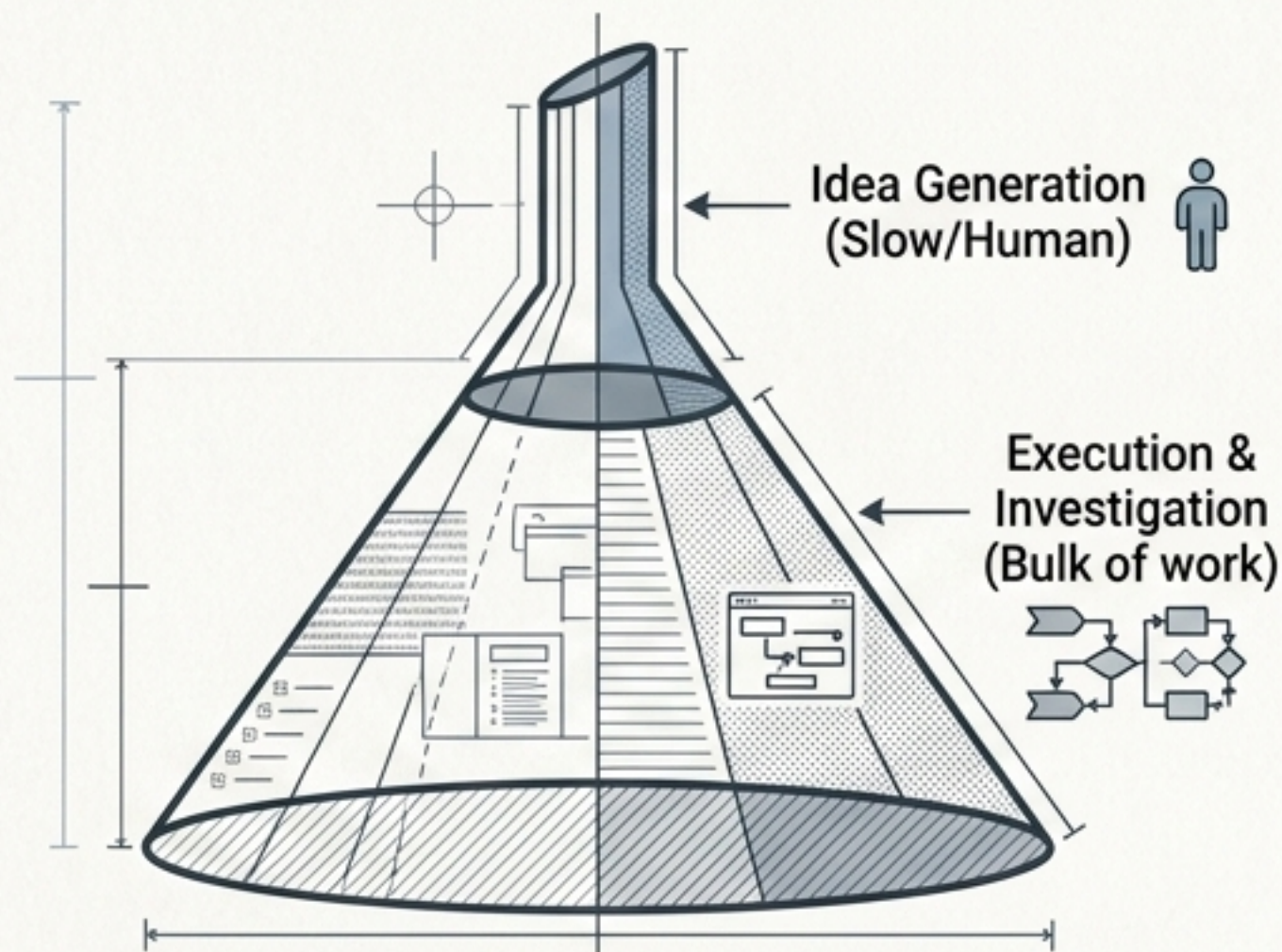
Fable 5.1が単体で出力するのは「テキスト」に過ぎない。

頭脳の進化 ≠ 研究所の完成:  
真のサイエンスエージェントは、  
検索、実行環境、実験装置、評  
価器、安全制御、そして人間の  
専門判断を組み合わせで初めて  
成立する「システム」である。

企業が取り組むべきは、  
AIモデルの選定ではなく、  
このシステム全体（権限設計、  
ログ、責任者）の実装である。

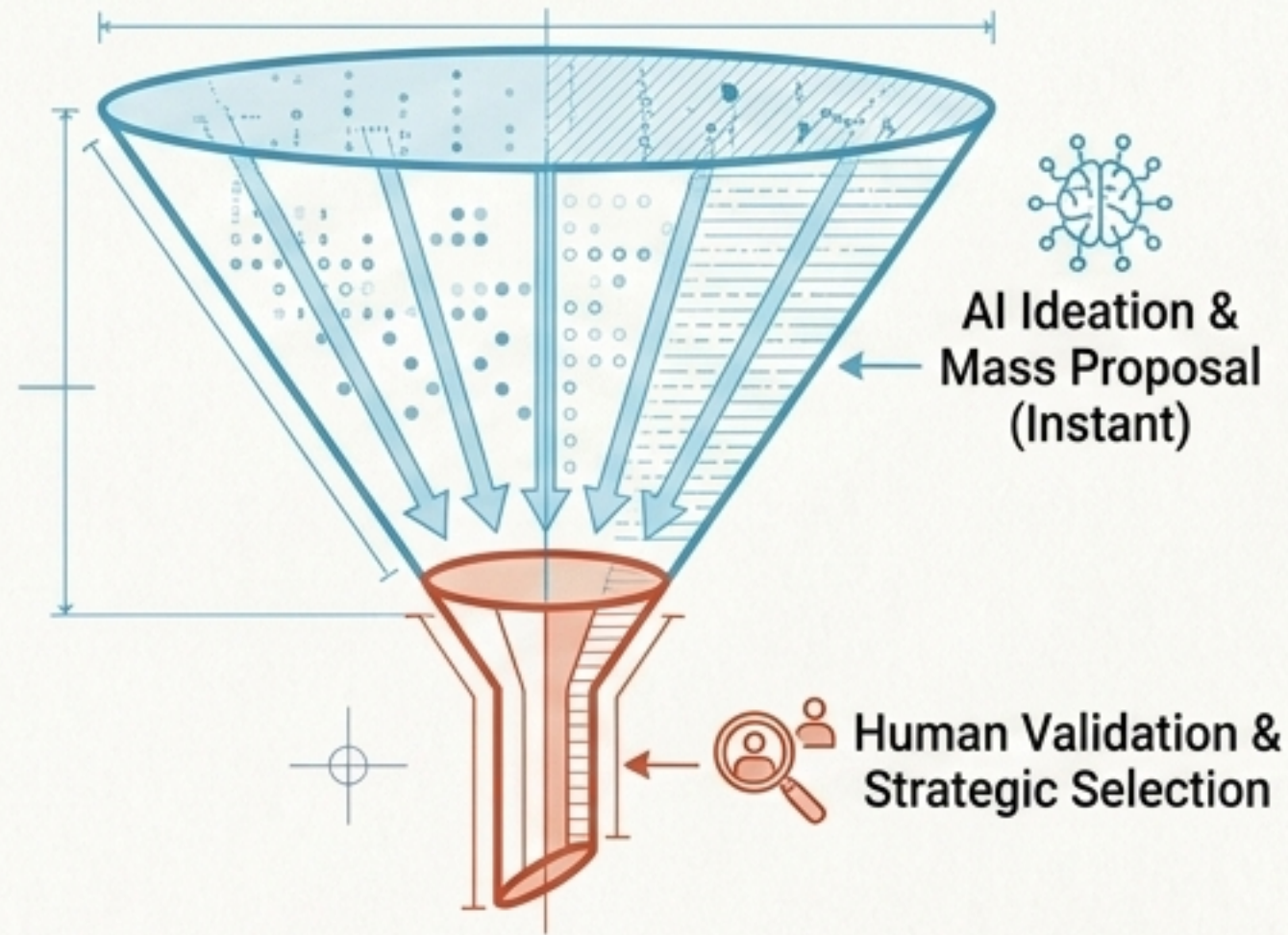
# R&Dパラダイムシフト：ボトルネックは「検証と選択」へ

## TRADITIONAL R&D



これまで: 科学者が文献調査、計算、実装、検証の「作業時間」に縛られていた。

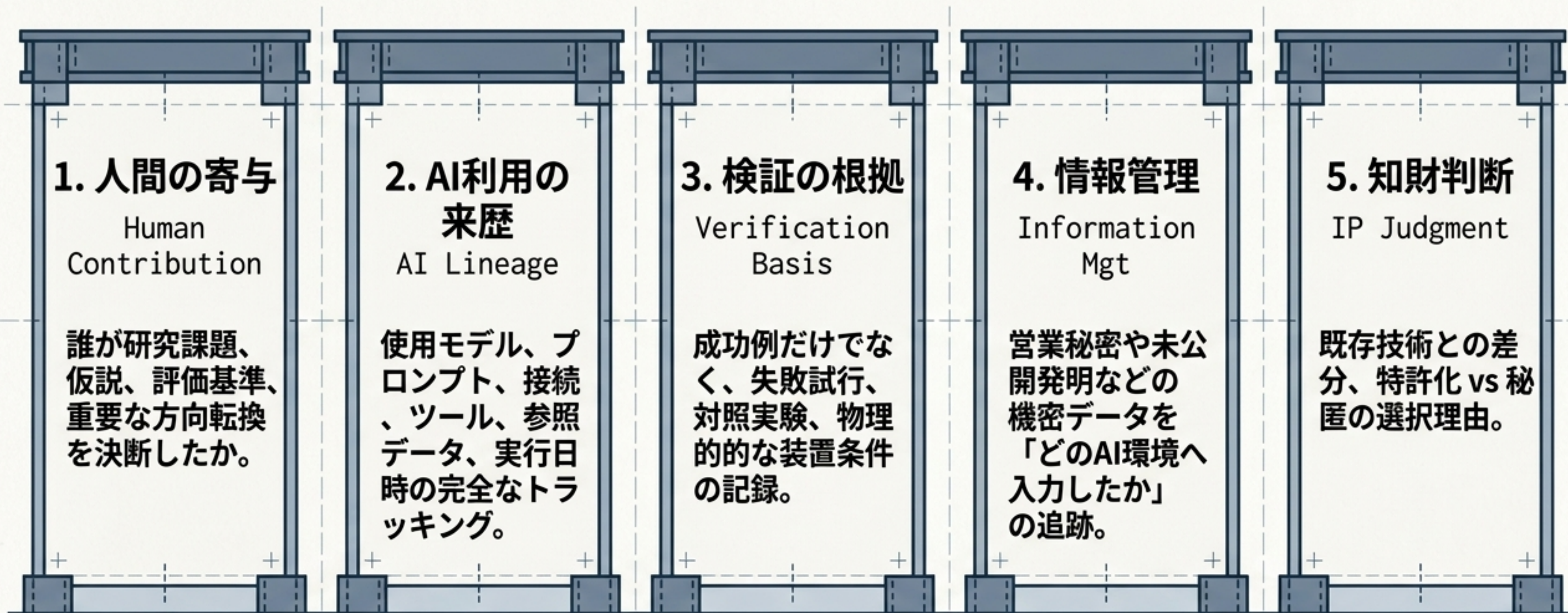
## AI-NATIVE R&D



これから: AIが大量の仮説、材料候補、製造条件を瞬時に提案する。

膨大なAIの提案に対し、「どの案に限られた物理実験リソースを配分するか」「何を特許化し、何を営業秘密（ノウハウ）として秘匿するか」という高度な戦略的判断こそが、人間の新たなコア業務となる。

# AI時代の知財マネジメント：必須となる5つの記録構造



「AIが案を出した」という事実だけでは、人間の創作的寄与を事後的に証明することは極めて困難。プロセス同時の記録が特許と研究公正を守る。

# 意思決定KPIの再定義：量を価値と取り違えない

**Discard (Red/Stop)**

× 書いたコードの行数

× 要約した論文の数

× 単純な処理件数やAI稼働時間  
(誤りの高速量産を誘発する)

**Adopt (Green/Cyan/Start)**

検証済み仮説 1件あたりの時間とコスト

有望でない研究候補を早期に中止  
(Fast-fail) できた割合

AIによる実験・解析の物理的な再現率

仮説提示から「知財・事業判断」までの  
リードタイム縮小幅

サイエンスエージェントの真の価値は、研究開発上の「意思決定」をどれだけ早く、確実にしたかで測られる。

# 人間の権限再設計：「統治（Governance）」と「救済（Rescue）」の峻別



**結論: AIが完全な自律科学者になる未来を待つ必要はない。**  
**Fable 5.1という「エンジン」が突きつけているのは、人間の研究組織を「AI前提」で今すぐ再設計する時代の到来である。**