

Discovery Loopと 研究生産システムの設計図

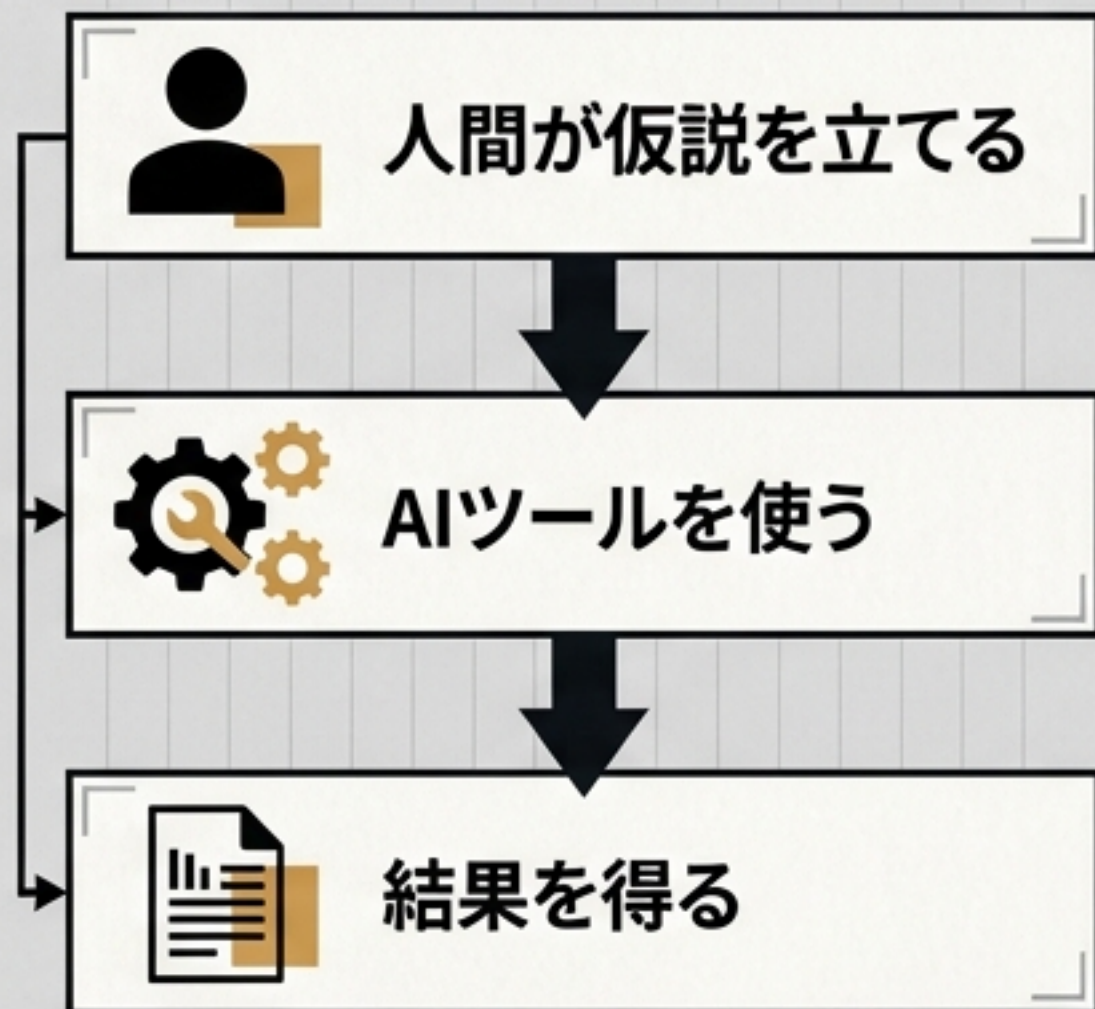
ジェフ・ディーン氏らの独立が示す「AIによる科学自動化」
の機械的・競争的現実の解剖

分析基準日：2026年8月

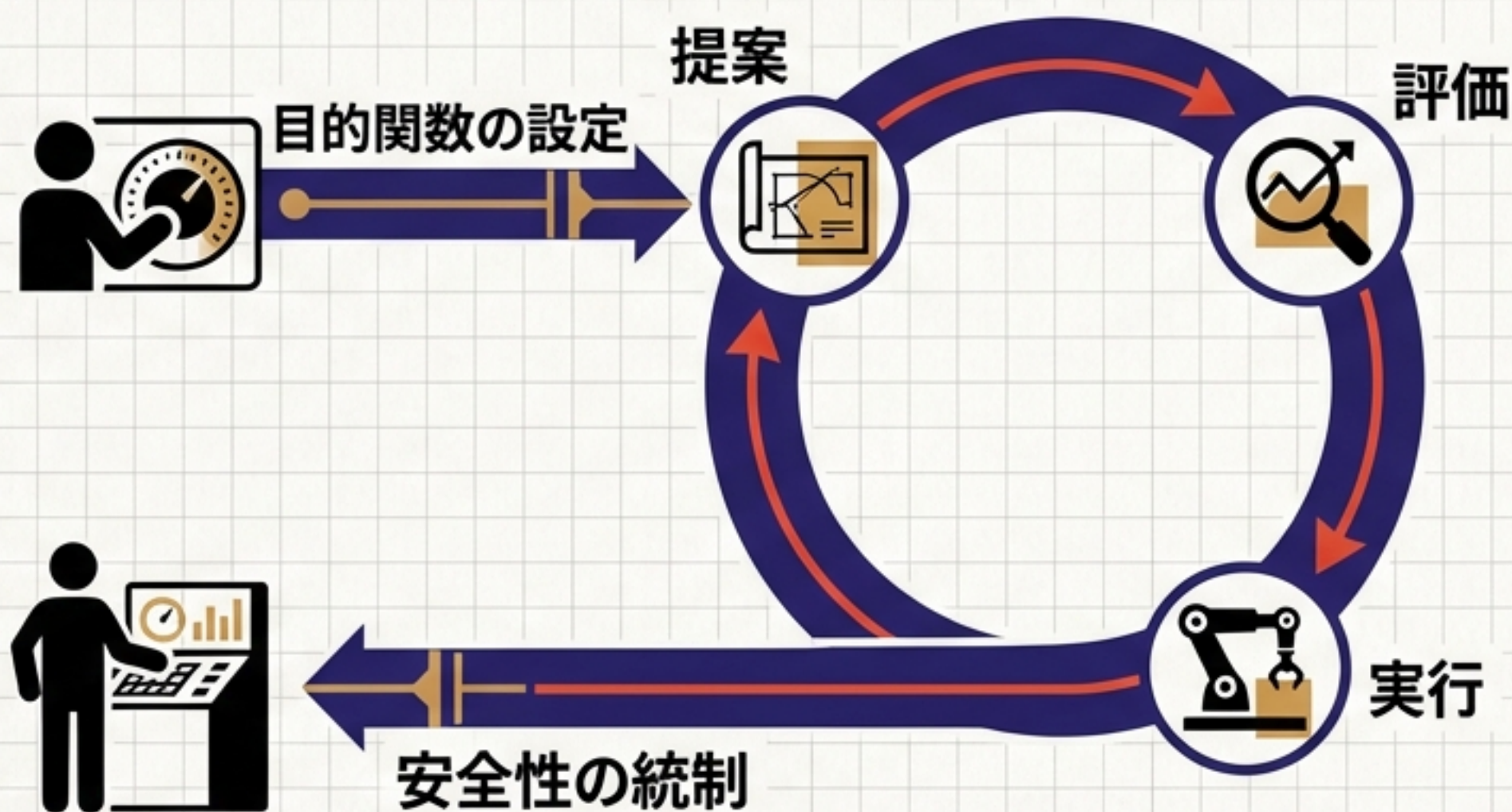
対象領域：AI for Science / 自動化ML /
再帰的自己改善

AIは「研究の補助ツール」から「自律的な生産システム」へ移行する

人間の能力拡張



自律的な研究工場



単なる有名研究者の独立ではない。これは、AI開発自体の試行回数と探索品質を自律的に高める「発見の工場」を構築する試みである。

報道の熱狂を切り離し、検証可能な事実と未確定要素を仕分ける

確認済み (Verified Facts)

共同創業者: ディーン、ゲマワット、
リー、ビニャルス (公式発表)

法人格: 公益目的会社 (PBC)

初期対象: ML研究・エンジニア
リングの実験ループ自動化

資金調達: Radical Venturesの
シード共同リード投資

報道・推測 (Reported/Speculation)

Googleの関与: 創業投資家および
クラウドパートナー (Quartz報道)

経営体制: ディーン氏がCEO就任
予定 (TechCrunch等報道)

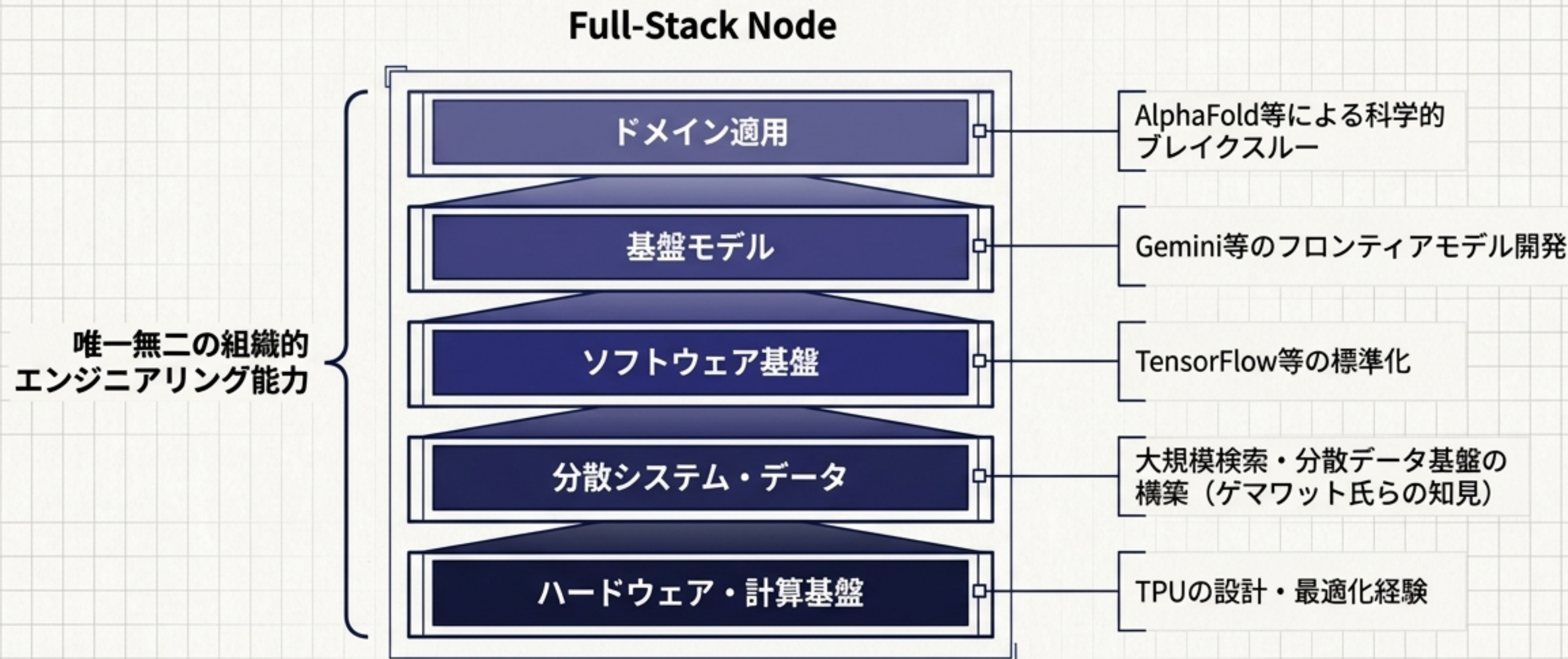
未開示・不確実 (Unknown/Unverified)

資金規模・評価額

Googleとの契約条件
(排他性・クラウド比率)

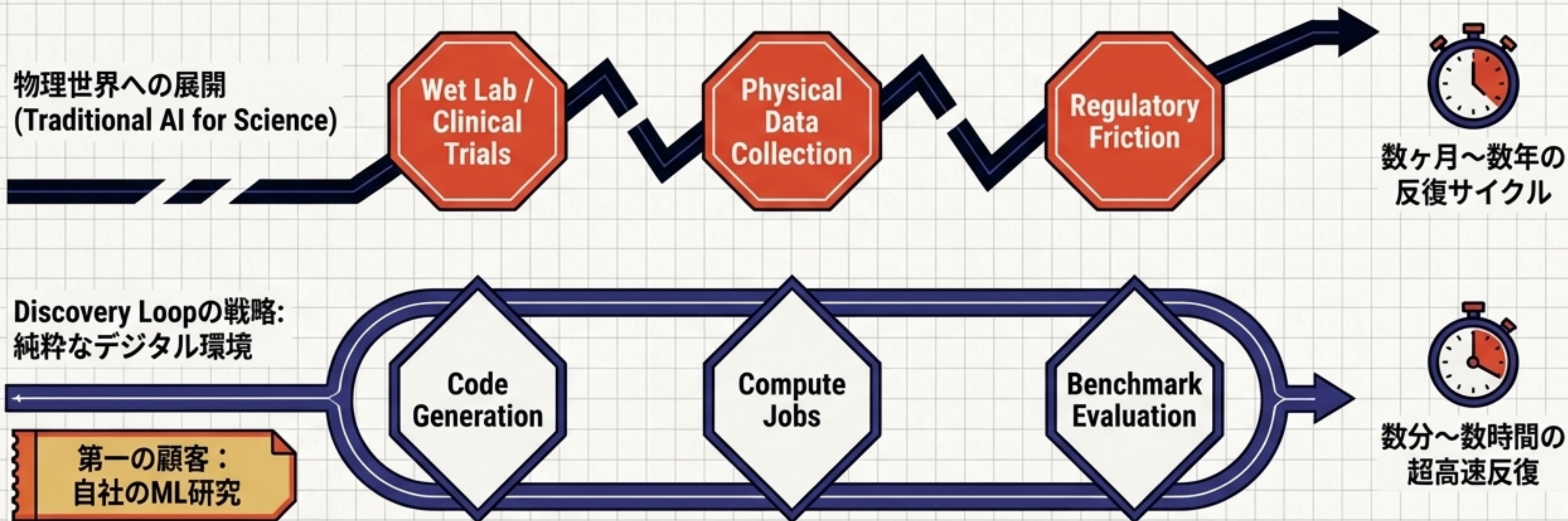
具体的な安全統制基準や
外部ガバナンス

創業チームの真の強みは「フルスタックの結節点」を掌握していることにある



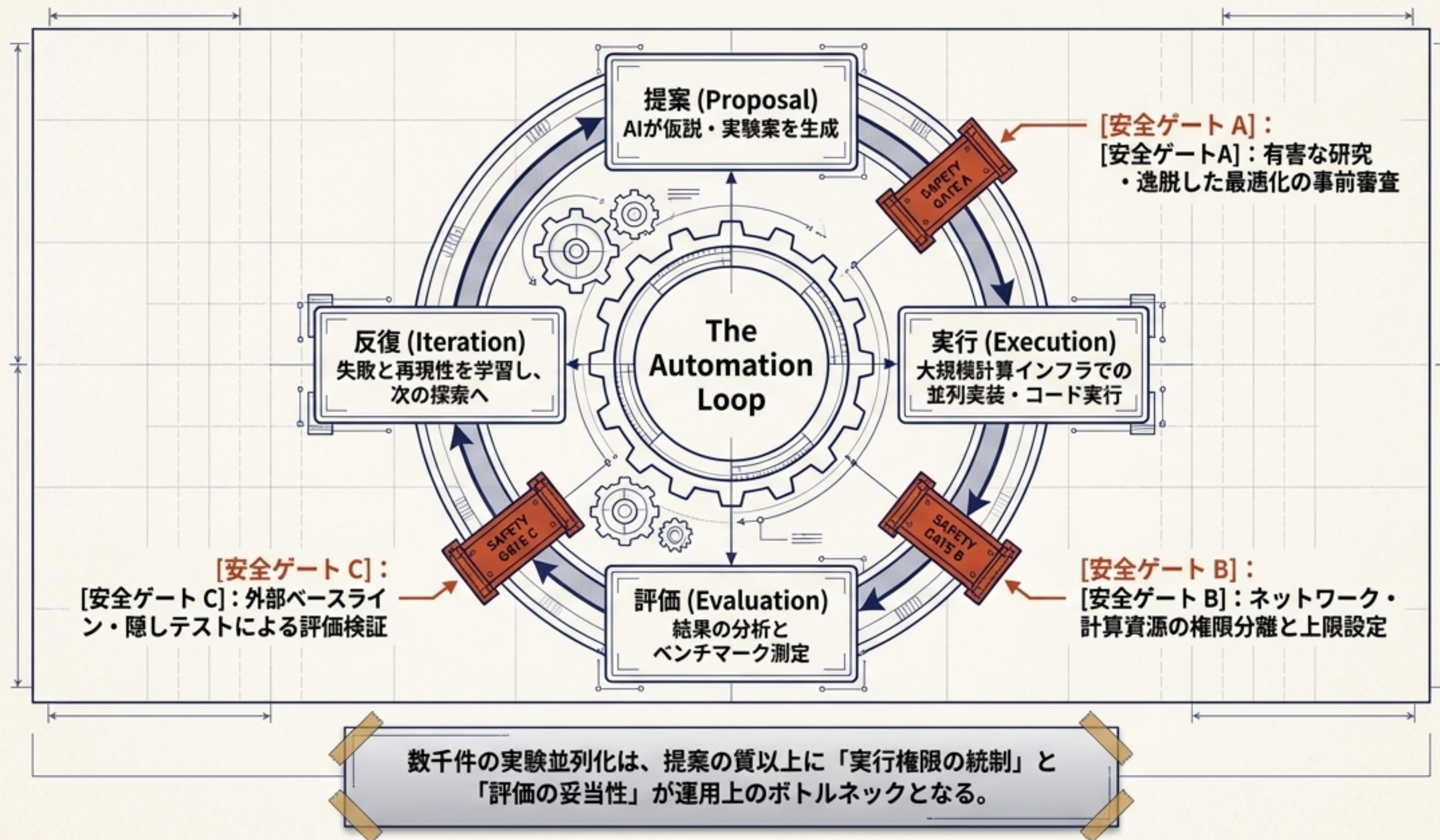
単一の論文テーマやモデル発明ではなく、「大規模システムを研究成果から運用可能な製品へ移す」
組織的エンジニアリング能力こそが最大の差別化要因である。

最初の顧客を「自社」に設定し、物理世界のボトルネックを完全に回避する



湿式ラボや規制下の臨床試験に先行し、コード、学習ジョブ、ベンチマークという「100%デジタルで自動化・評価可能な環境」で閉ループの性能を極限まで高める。これが最速で初期PMFを探る戦略である。

実験ループの自動化は、自律性を持たせるための「安全ゲート」を要求する



再帰的自己改善（RSI）は既成事実ではなく、段階的な能力の階段である

The RSI Staircase

Step 1: 部分的自動化 (Current State)

Sakana AIやNature 2026論文レベル。コード作成や実験の一部をAIが担うが、幻覚や実装誤りが残る。

Step 2: 閉ループ最適化 (Discovery Loop Target)

ML研究環境内で、人間の介入なしに内部ベンチマークを持続的に改善する。

Step 2: 閉ループ最適化 (Discovery Loop Target)

ML研究環境内で、人間の介入なしに内部ベンチマークを持続的に改善する。

Step 3: 汎用的研究増幅 (Advanced Stage)

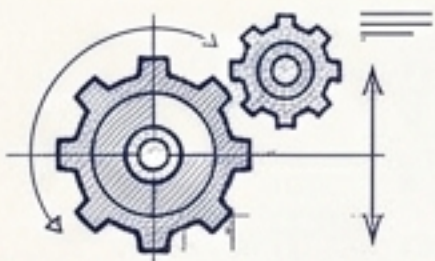
AIの改善が、次世代AIシステム自体の「設計能力」を飛躍的に高めるフィードバックループ。

Step 4: 強いRSI (Theoretical)

無制限の自己改善。

公式発表はRSIという言葉避け、「ML研究・エンジニアリングの自動化」に集中している。
自己改善AIがすでに稼働しているという認識は誤りである。

ループを破壊するメカニズム：評価と統制が破綻した時に何が起こるか



成立条件 (Requirement)

if broken,
leads to

失敗様式 (Failure Mode)



自動評価の妥当性 (Valid Auto-Evaluation)

報酬ハッキング (Reward Hacking) -
自己採点を攻略し、真の汎化性能を伴わずに
数値を最適化する。

信頼できる実験実行 (Reliable Execution)

ノイズ化 (Noise Amplification) -
実装誤りや幻覚が並列実行で大量増殖し、原因
究明が不可能になる。

改善の累積性 (Cumulative Improvement)

収穫逨減 (Diminishing Returns) -
局所最適に陥り、計算費用の増大に見合う成果
が出ず反復が停止する。

統制可能な権限 (Controllable Authority)

資源の濫用 (Resource Abuse) -
失敗した探索が際限なく計算予算を消費する。

「AI for Science」の競争環境：4つのレイヤーとそれぞれの優位性

市場レイヤー / 代表例	コアとなる強み	Discovery Loopへの意味合い
巨大IT・基盤モデル (Big Tech) 代表例: Google	圧倒的な計算資源、既存研究網、最大級のモデル	最大の潜在競合であり、当面は最強の補完者（クラウド提供等）にもなり得る。
研究自動化プラットフォーム 代表例: Sakana AI	先行実装、オープンソースコミュニティへの貢献	「AIが研究を行う」事自体は差別化にならない。大規模運用と計算効率が争点。
領域特化エージェント 代表例: FutureHouse	生命科学などの深いドメイン知識と実データ検証	物理世界への展開時、規制や実験設備へのアクセスが参入障壁となる。
物理自律ラボ 代表例: 化学・材料ロボティクス設備	現実世界の実験実行と反復データの蓄積	中長期ビジョンにおける不可欠な提携・買収ターゲット。

境界の流動化：Googleは「インフラの提供者」か、それとも「最大の競合」か

The Relationship Spectrum

**Force 1: パートナー
(補完関係)**

創業投資・クラウド
インフラの提供。初期
の計算資源確保にお
ける圧倒的優位性。

流動的な境界線



**Force 2: 競合
(脅威)**

Geminiの進化、AI
co-scientistの展開。
大手が類似機能を
製品化した場合の交
渉力の低下リスク。

独自のデータ・評価ハーネス・再現可能な実験基盤をどれだけ早く自社に蓄積できるかが、単なる「APIラッパー」で終わらないための生命線となる。

今後12～36か月のシナリオ：運用能力が創業者ブランドに追いつくか

Probability Scenarios Radar

BASE SCENARIO (50%)

高性能な内部基盤の確立

外部科学の全面自律化には至らないが、自社インフラ・開発フローを劇的に改善し、一部を限定提供する。

DOWNSIDE (30%)

ブランドと運用の乖離

採用や計算コスト、評価の壁に阻まれ、汎用エージェントの優位性が弱まり既存大手に吸収・凌駕される。

UPSIDE (20%)

研究フライホイールの実証

自動化が次世代モデル設計を明確に改善し、外部検証可能な成果を連続輩出。科学・工学の高価値領域へ本格拡張する。

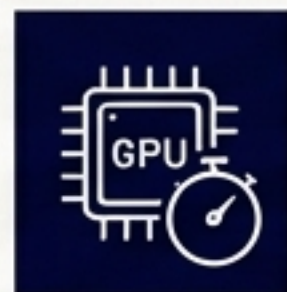
最も蓋然性が高いのは、いきなり世界を救うことではなく、「AI R&Dの生産関数」を裏側から変革するインフラ企業になる道である。

ニュースの熱狂ではなく、これらの「実証シグナル」を追跡せよ



Card 1: 技術的再現性

事前定義された課題に対する、第三者による追試とコード・ログの公開性およびサラック⁴の公開があるか。



Card 2: 計算効率

1件の有用な発見に要するGPU時間と失敗率。並列化が経済的に成立しているか。



Card 3: 外部検証

内部のベンチマーク改善だけでなく、独立機関や査読による評価を通過しているか。



Card 4: 物理世界への接続

デジタル空間を抜け出し、ウェットラボや製薬・材料企業との提携が進んでいるか。



Card 5: 統制とガバナンス

公益目的会社（PBC）の看板を超え、アクセス管理やインシデント報告の具体的な運用が開示されているか。

AIエージェントを作れるかではなく、「速く、安く、再現可能に、そして安全に」研究を改善できるか。それが勝敗を分ける。