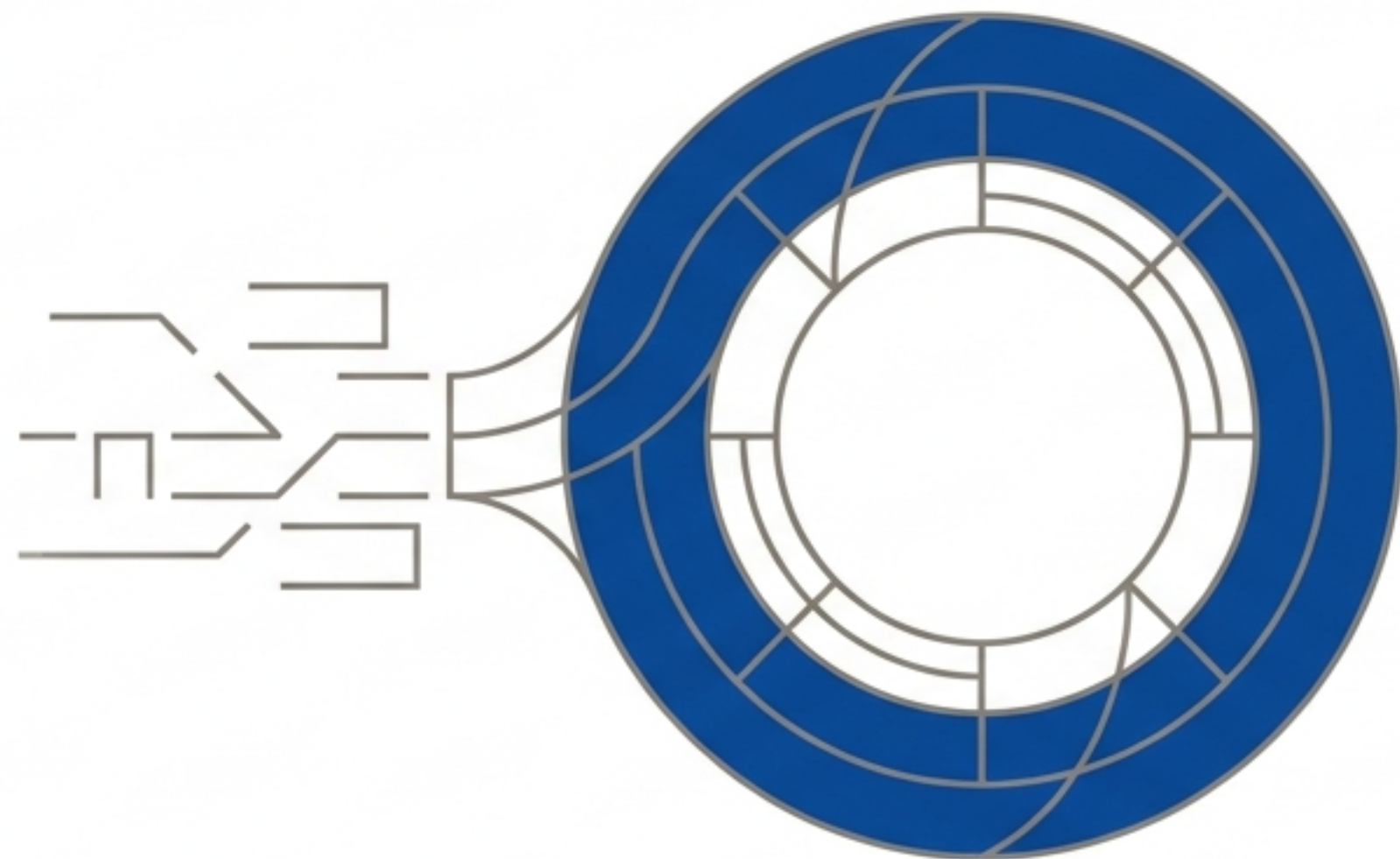


AI for Scienceと 知財戦略の パラダイムシフト

科学研究の「閉ループ化」がもたらす
ルールチェンジと次世代ガバナンス

2026年9月



エグゼクティブ・サマリー：AIは科学の方法そのものを自動化し始めた。競争優位は「発見結果」から「探索装置」へ移行する。



1. 新たな現実

支援ツールから、自律的な仮説生成・実験を行う『閉ループ』への進化



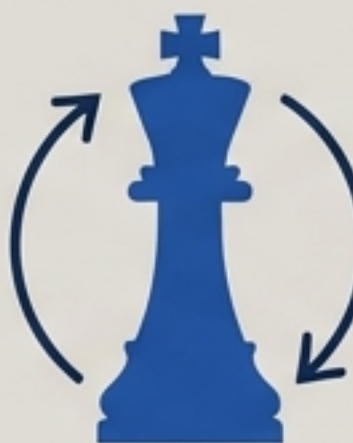
2. 制度の壁

再現性・第三者検証の脆弱性と、それに伴う特許要件（記載要件・進歩性）の地殻変動



3. 競争源泉の移行

公開されやすい『組成物・結果』から、秘匿すべき『データ・自律実験ノウハウ』への価値シフト



4. 勝者の戦略

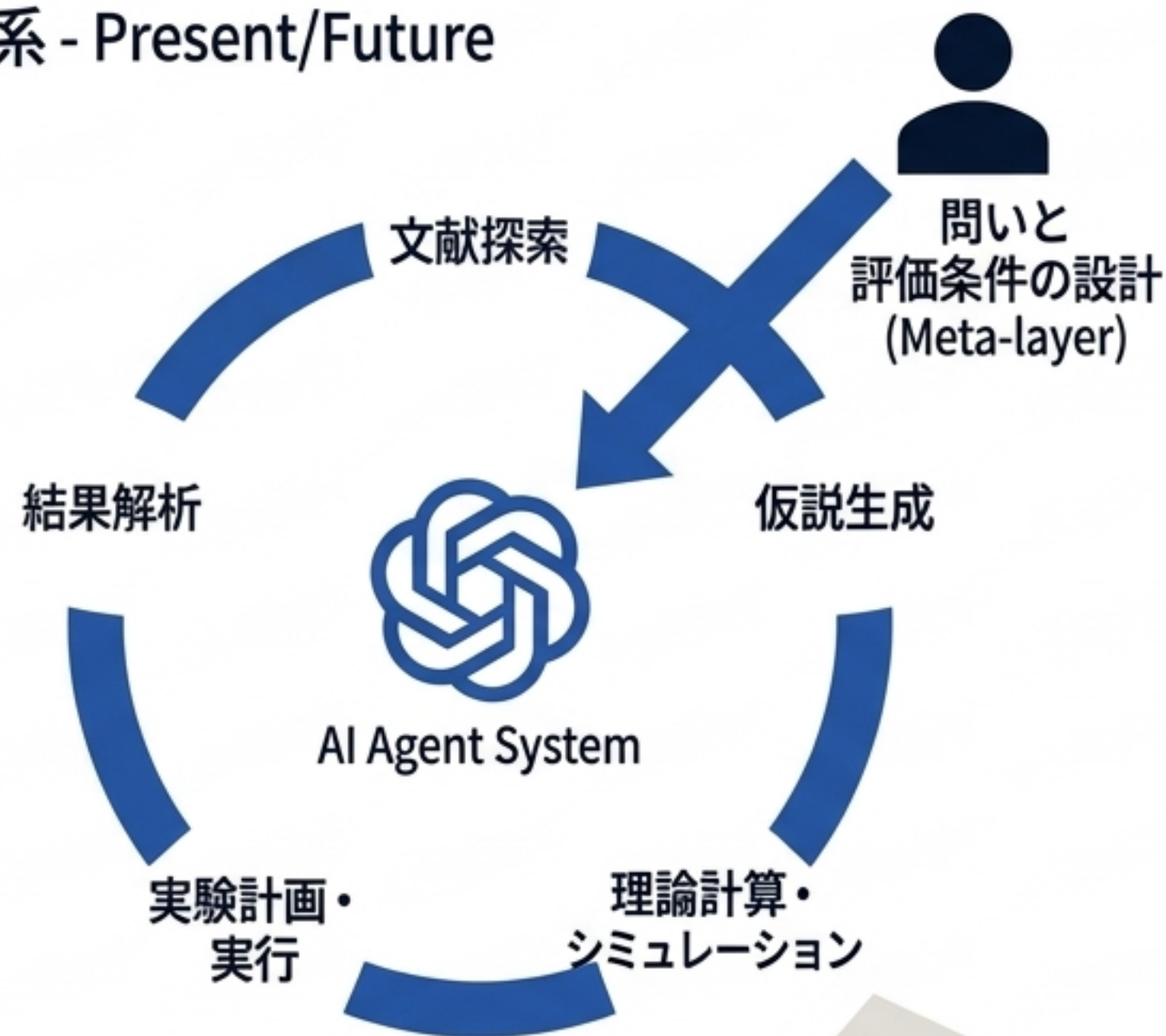
R&Dと知財を『Human over the loop』型へ再設計し、高速トリアージを実行する

評価指標は「計算速度」ではなく、「どこまで科学的ループを閉じられるか」へ。
評価指標は「計算速度」ではなく、「どこまで科学的ループを閉じられるか」へ。

支援系 - Past/Present

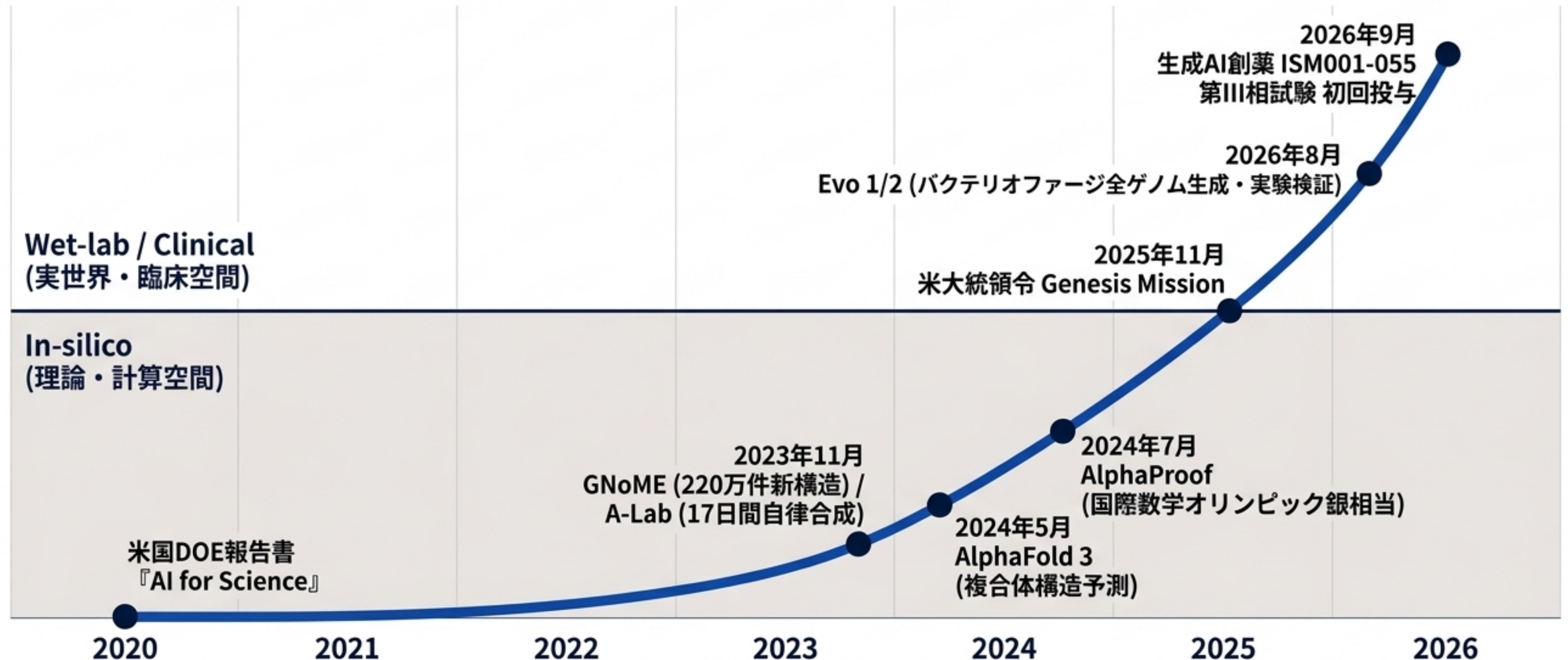


自律系 - Present/Future



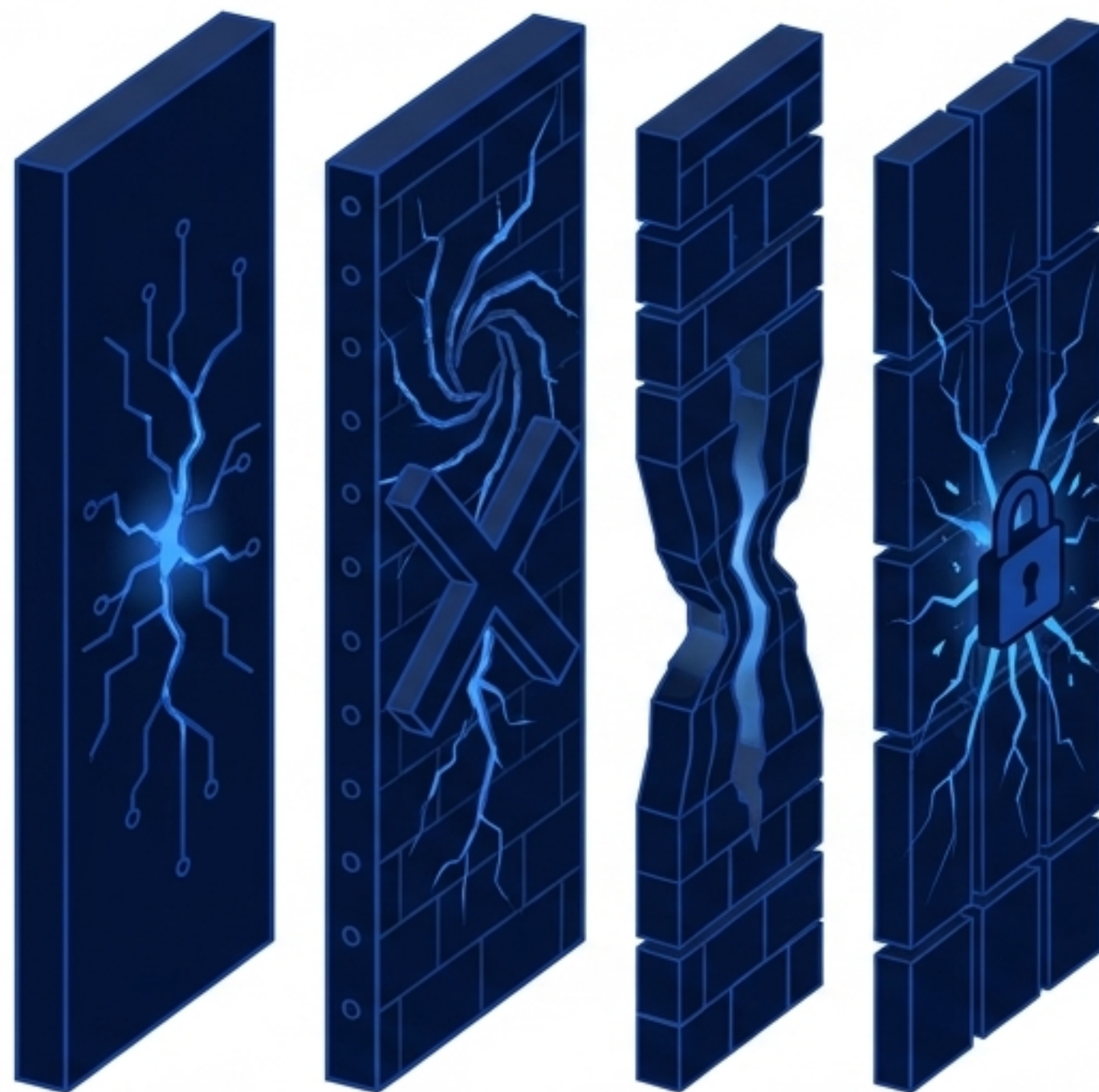
事例：Google Co-Scientist / Robin
(マルチエージェントによる仮説生成と自律ループの完遂)

理論空間の探索から、実社会での臨床・実験へと 「概念実証」のフェーズが移行。



探索能力と検証能力の「構造的なギャップ」。完全な閉ループ化を阻む4つの壁。

高速回転する
AIループ



再現性と第三者検証の脆弱性

(A-Lab論文の成功率下方修正、GNoMEの既知物質重複)

開かれた課題における判断力不足

(CRUX検証: エージェントは実験を完遂するが、創造性欠如で不合格)

自己評価バイアス

(AI査読機能における人間の判断との強い乖離)

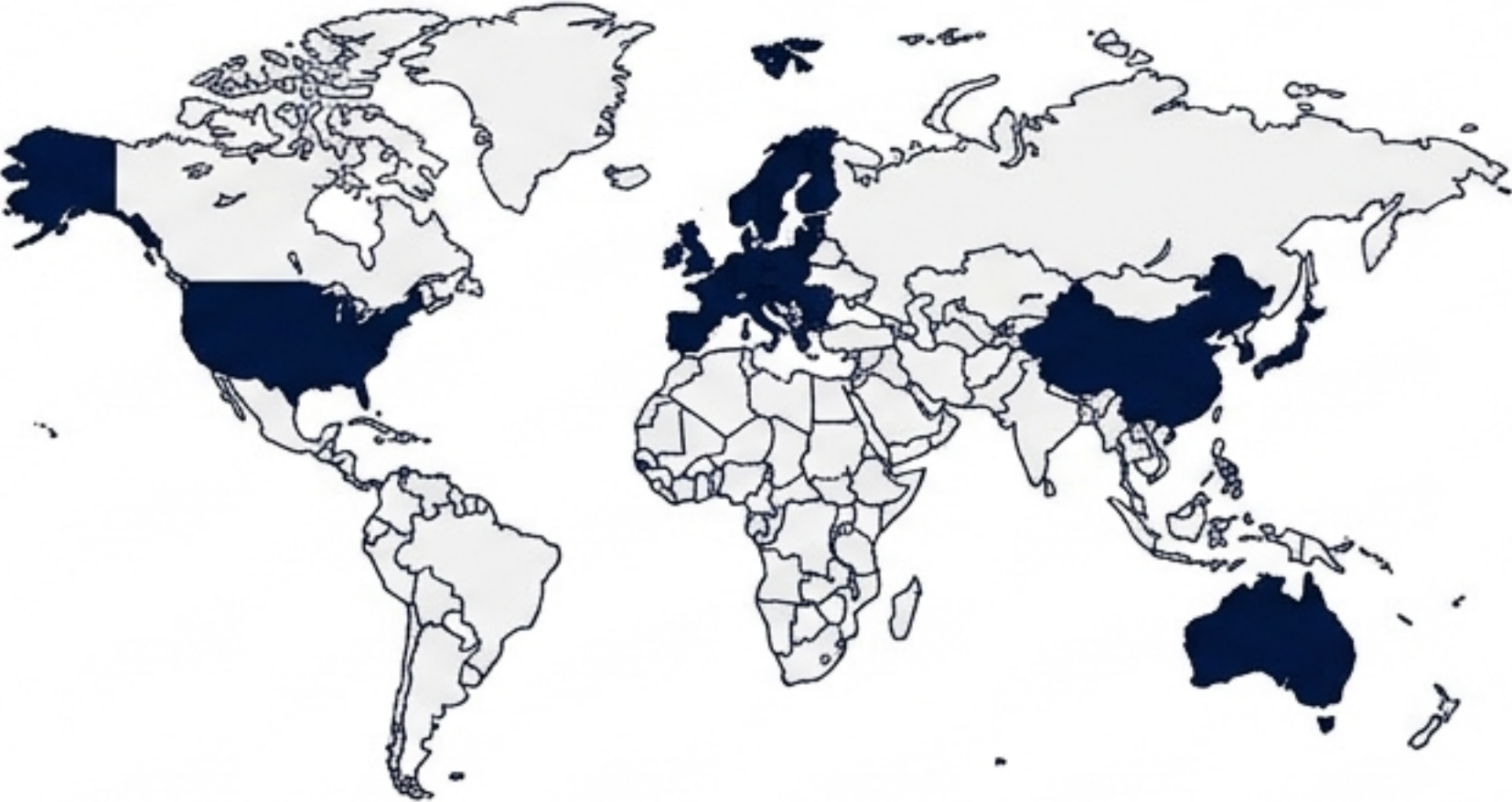
計算資源の寡占とデュアルユース

(フロンティア資源の集中、ゲノム生成の病原体転用リスク)

この「検証能力の欠如」が、特許制度における最大の壁（記載要件・進歩性）に直結する。

「AIは発明者になれない」で議論は終わらない。 争点は人間の貢献の立証へ。

発明者要件の帰結 (DABUS事件等)



米国:否定	欧州:否定	英国:否定
独・豪:否定	中国・韓国:否定	日本:否定(最高裁で確定)

日米ディープダイブ:残された実務課題

米国 (US) - 2025年11月 改訂ガイダンス	Pannu要件を排除し、伝統的な「着想 (Conception)」テストへ一律回帰。	実務影響: 着想の事実認定が訴訟リスク化
日本 (JP) - 最高裁令和8年決定	職務発明制度 (35条) は自然人発明者の存在を大前提とする。	実務影響: 発明者特定の失敗は、「相当の対価」対象の不明確化と権利帰属の不安定化リスクに直結

予測データのみでの広いクレームは無効化リスク大。実測による裏付けが標準へ。

AI予測データ (In-silico)

大量生成された予測構造、
GNoMEデータ等



湿式実験データ (Wet-lab)

物理的な合成、臨床データ、
in vitro / in vivo

Amgen v. Sanofi判決等の影響

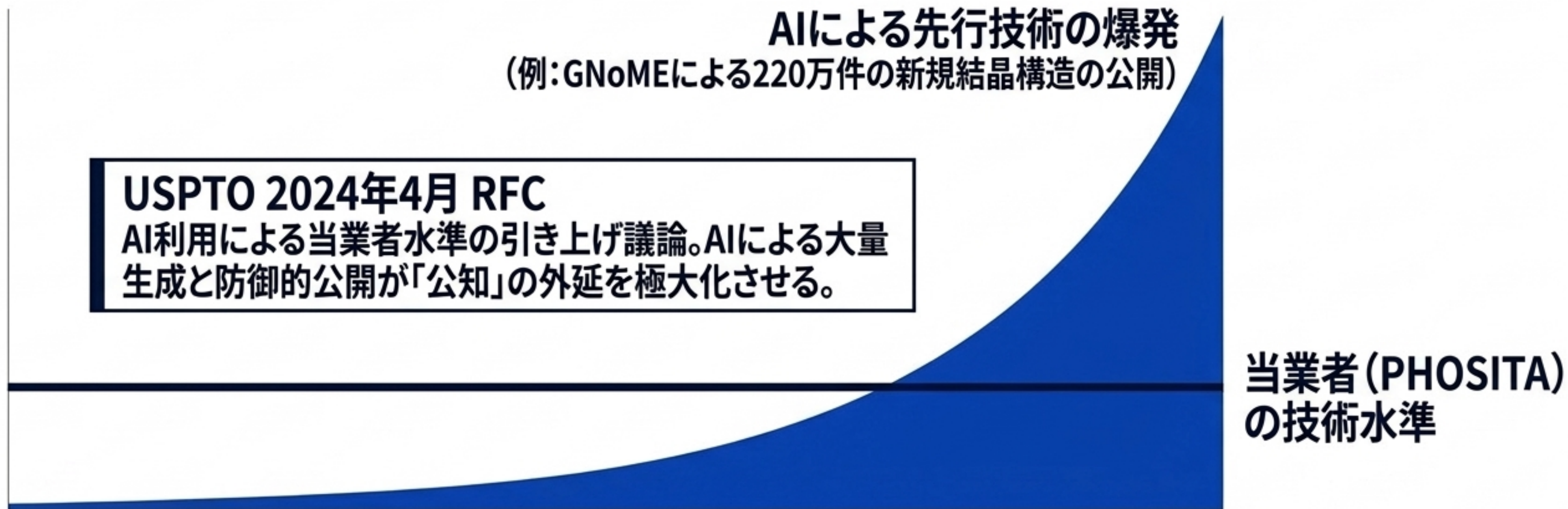
過度の実験を要しない機能的クレームの
サポート要件が厳格化。

A-Lab論争が示す現実

自律実験室の再現性リスクは、審査において
In-silico単独の信用性を低下させる。

当面の実務最適解: 「AI予測 + 実測データ」をセットで権利化する

「AIが大量に探索した中から選んだ」事実だけでは、
選択発明の技術的意義が否定される。



アクション: パラメータ発明や選択発明において、
これまで以上に厚い「予期せぬ効果」の実証データが要求される。

国家も企業も「横の基盤」制圧へ。発見結果から探索装置へ競争優位が移行。

過去・現在:Output
(発見結果)

発見された組成物・物質

公開されやすく、AIにより
コモディティ化が加速



現在・未来:Engine
(探索装置)

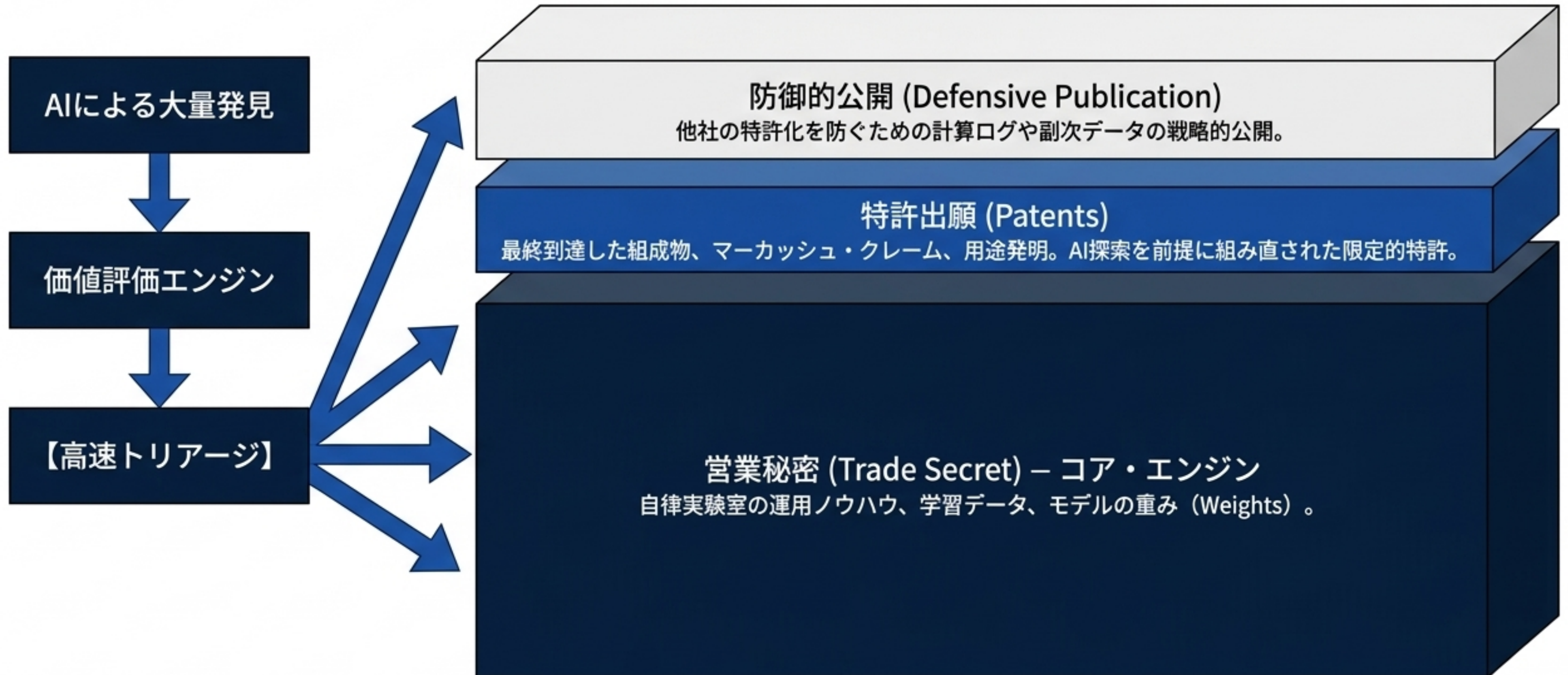
データセット、AIモデル、
自律実験室の運用ノウハウ

他社が模倣困難な
長期的な競争優位の源泉

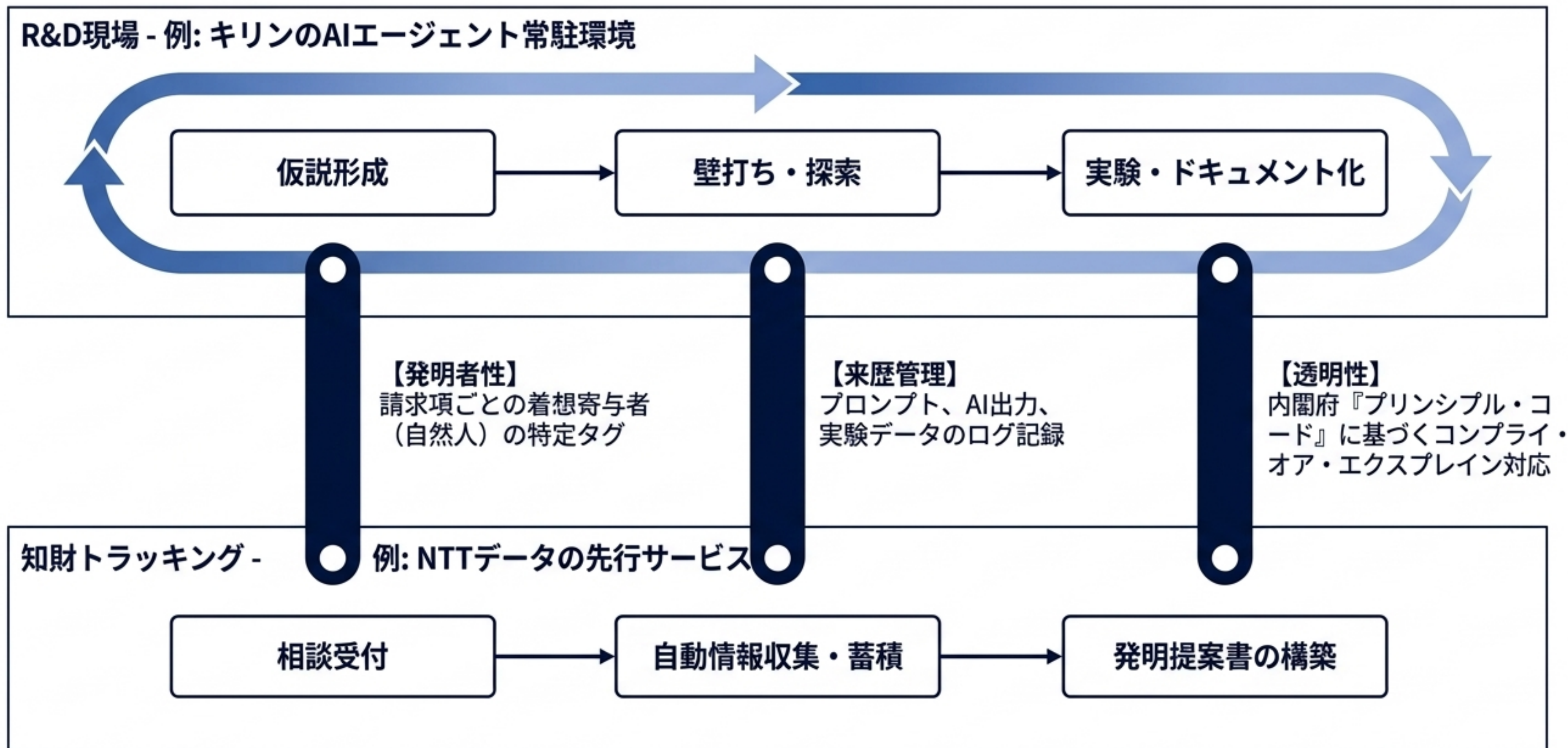
証拠:

米大統領令「Genesis Mission」(2025年11月)
個別領域の「縦の研究」と並行し、AIモデル・自律実験基
盤を共通化する「横の基盤(278プロジェクト)」へ
国家規模で集中投資。

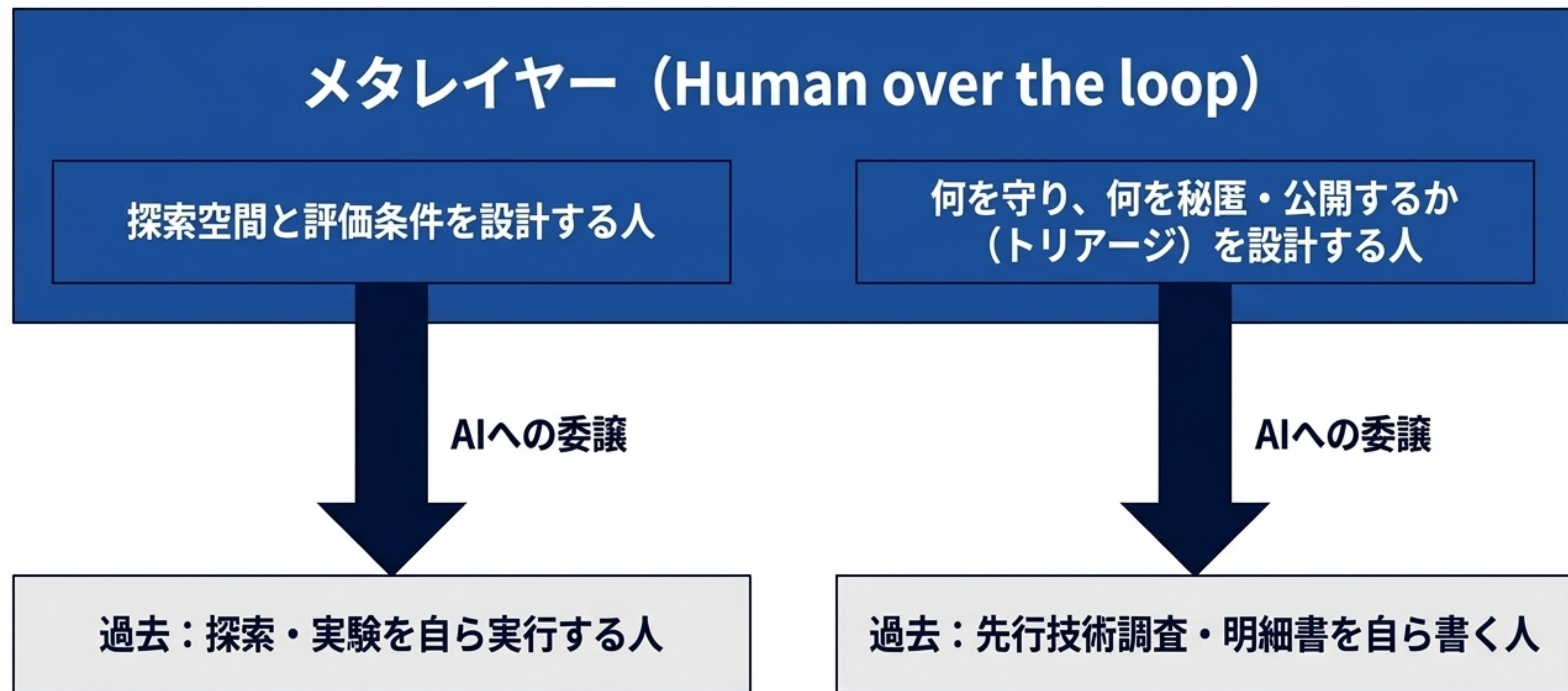
全件出願は崩壊する。出願・秘匿・公開を高速に振り分ける「トリアージ」が必須。



AIネイティブ環境における新たな知財ガバナンスと発明記録の同期。



知財業務も「Human over the loop」型へ。プロセスの実行からルール設計へ。



研究者と知財担当者の役割変化は、全く同じフラクタル構造をとる。

進化段階別の知財業務インパクトと対応マトリックス

		科学の実態（代表例）	知財業務への主な影響	主な対応
	支援系	文献探索・データ解析の支援 (AlphaFold等)	先行技術調査・ 明細書作成の効率化	AIツール品質管理
	仮説生成系	Co-Scientist, Robin, 麒麟のAIネイティブ環境	発明発掘の大量化、 着想寄与者の特定困難化	発明記録の制度化、 発明者性の立証設計
	自律実験系	A-Lab等の自律実験室	データ来歴と公知時期管理、 記載要件の厳格化	営業秘密と特許の二層構造 (トリアージ)
	自律研究系	AI Scientist、フロンティア モデルによる並列探索	進歩性論の再構築、 探索装置の流出リスク	ポートフォリオ・トリアージ、 制度動向の監視

ルールの不確実性を言い訳にせず、内部ガバナンスから即時着手する。

即時（～3ヶ月）

- ・ 発明記録テンプレート（着想寄与者・来歴）の整備
- ・ 共同研究契約・利用規約の「権利帰属条項」の点検
- ・ 生成AI利用ガイドラインとポリシー・コードの整合

短期（～1年）

- ・ 出願・公開・秘匿の「トライアージ基準」の策定
- ・ 材料・創薬における「AI予測＋実測」の標準化
- ・ 発明発掘へのAIエージェント試験導入

中期（1～3年）

- ・ 知財部門の「Human over the loop」型体制への移行
- ・ 自律実験室ノウハウの営業秘密管理体制（コア・エンジン）の確立

監視すべき4つのトリガー。引かれた瞬間、戦略のピボットが求められる。



Trigger 1: 米国進歩性基準の引き上げ

米国特許商標庁（USPTO）が、AI利用による「当業者水準の引き上げ」方針を正式採択した場合。



Trigger 2: In-silico要件の緩和

日本特許庁（JPO）または欧州（EPO）が、In-silicoデータ単独でのサポート要件充足について明示的な緩和指針を公表した場合。



Trigger 3: 未解決問題のAI解決

ナビエ-ストークス等、AI主導の「未解決問題の解決」主張が、正式な査読と第三者検証を完全に通過した場合。



Trigger 4: 国内ガイドライン改訂

産業構造審議会や、令和8年度中の改訂が予定される「知財・無形資産ガバナンスガイドライン」において、AI発明の新たな方針が示された場合。