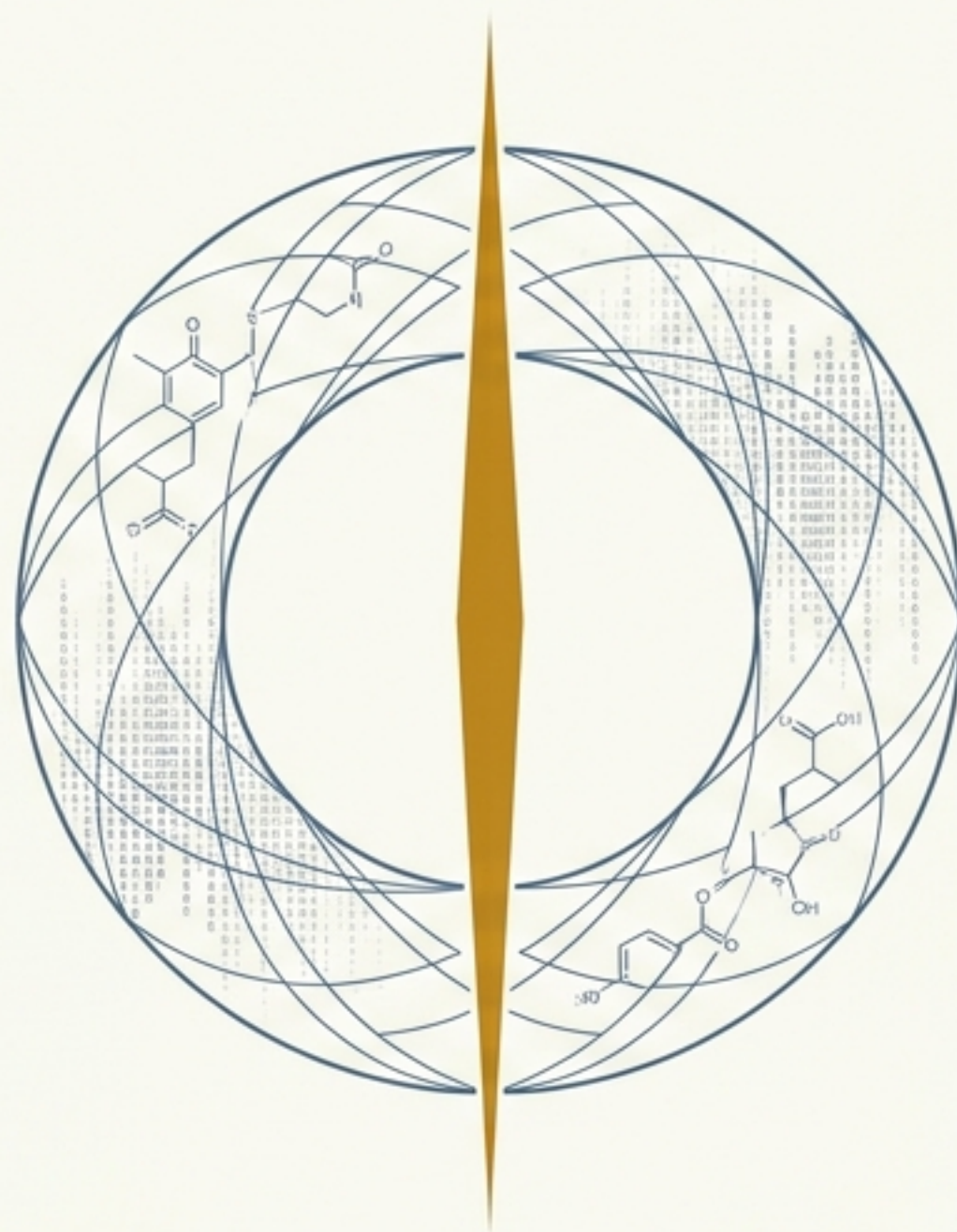
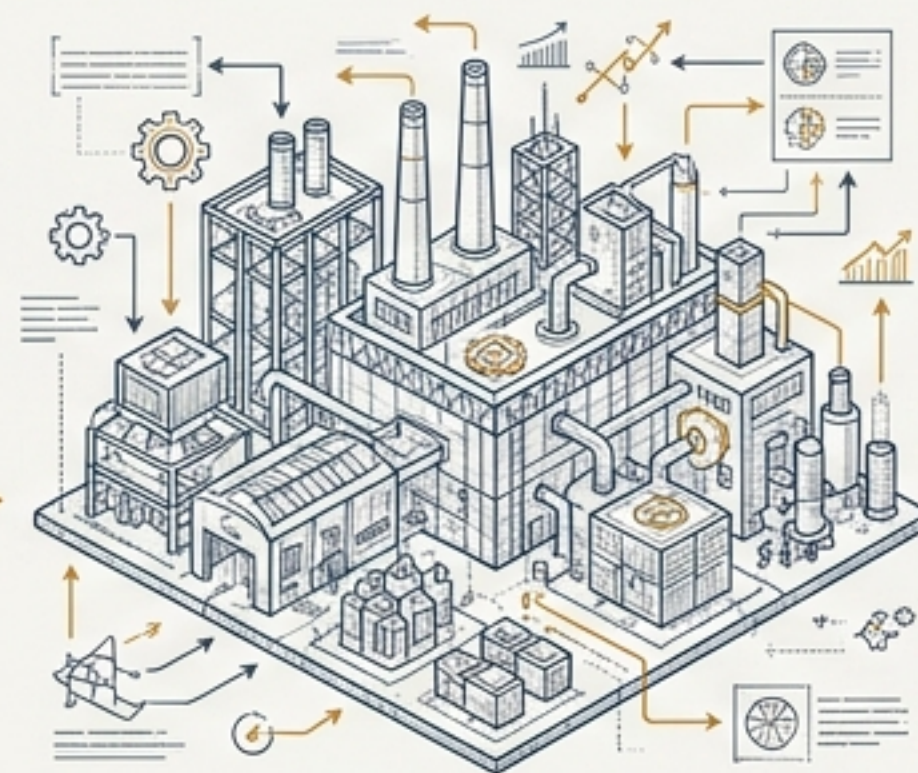


# AIエージェントと科学研究の最前線

ハイプの終焉と「人間の着想」の再定義 — 技術・地政学・知財ルール of 完全解剖



# AIによる科学の 完全自動化・科学者の代替



科学者が扱える**情報・実験空間の爆発的  
拡張**と、新たな**知財戦略**の幕開け

## 技術の現在地



AIは「代替」ではなく「拡張」。  
最大のリスクは生成力ではなく  
「妥当性検証のボトルネック」。

## 地政学とルール



「技術主権」を巡る国家間競争。  
DABUS判決確定により「発明者は自  
然人」で日米欧のルールが完全収斂。

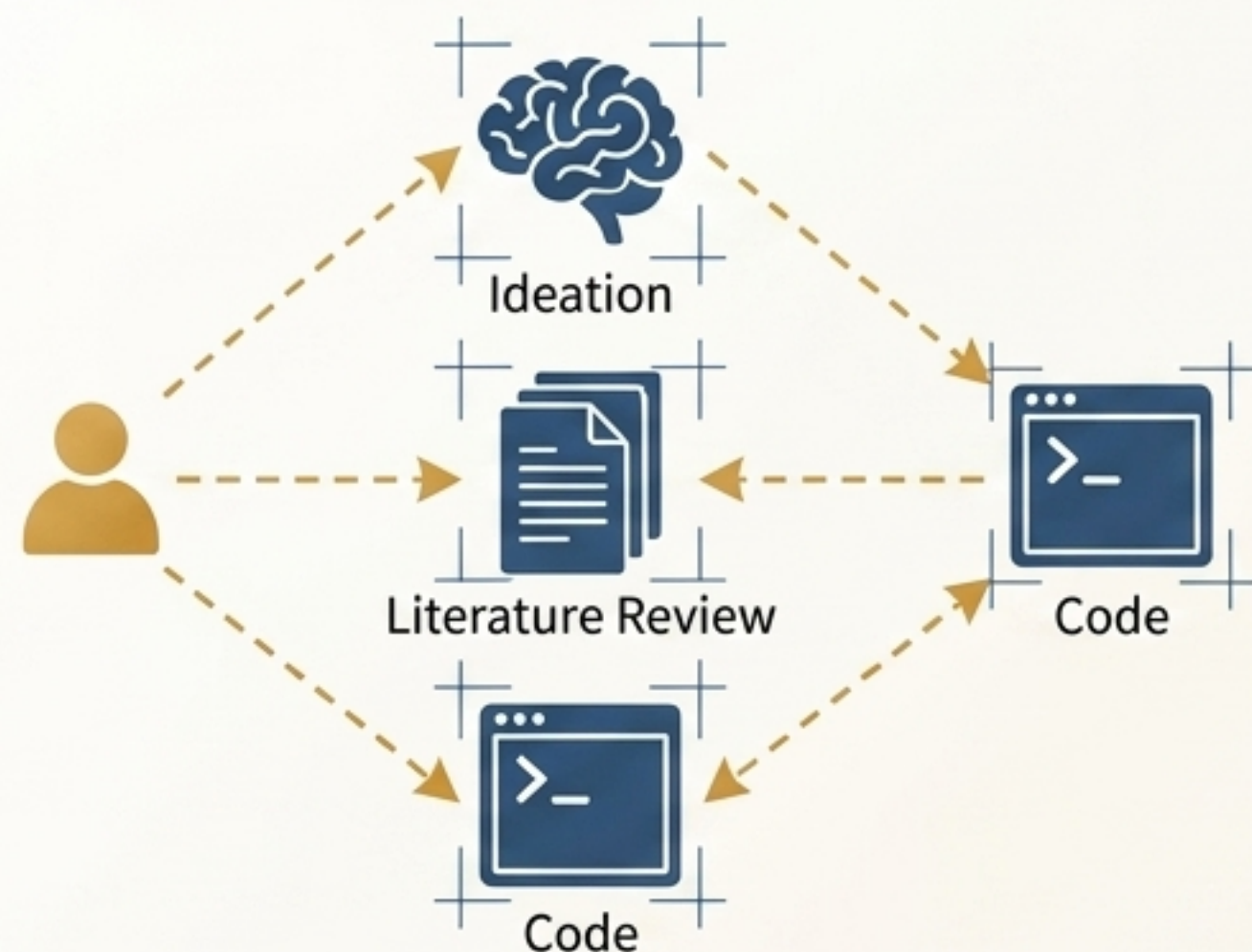
## 企業のアクション



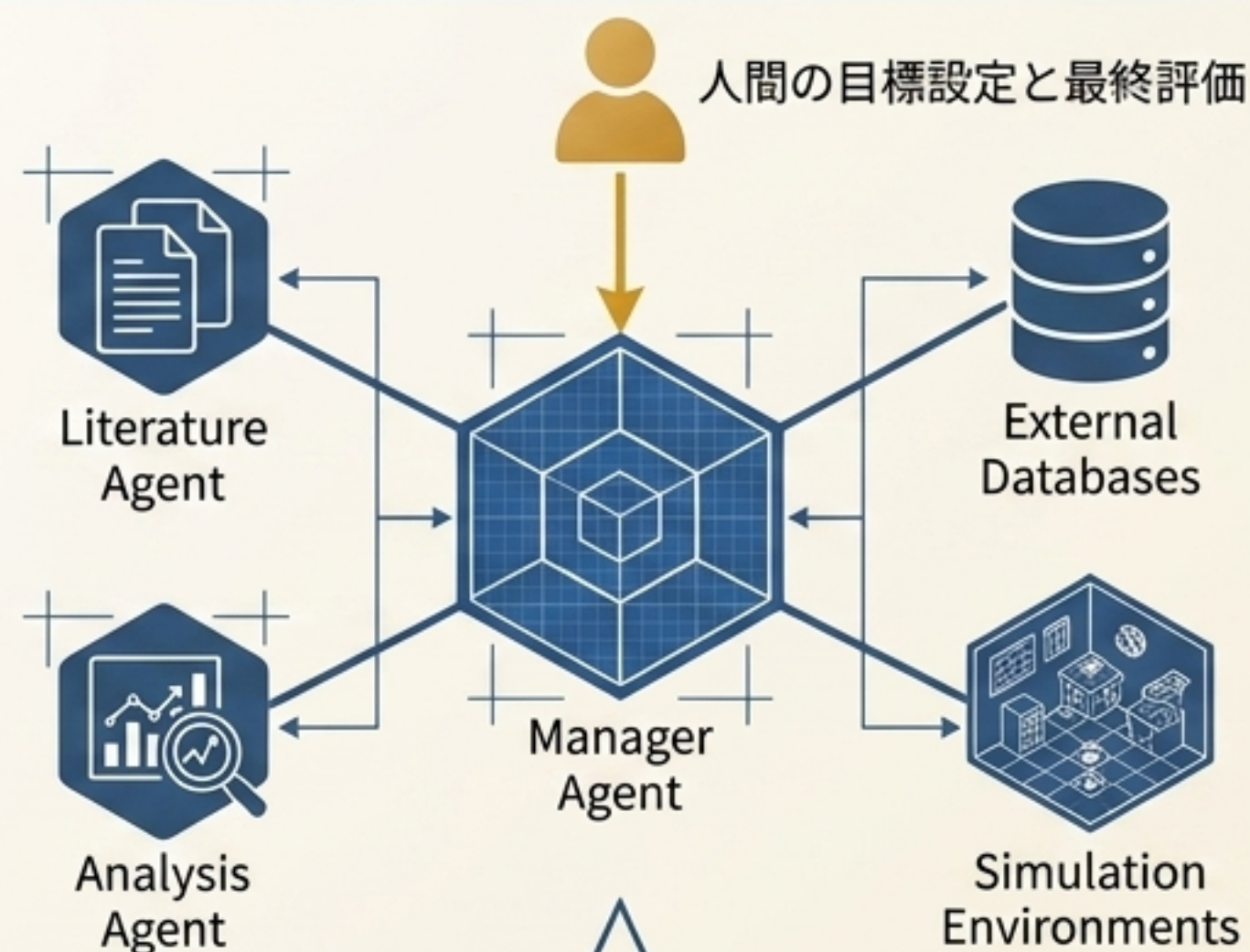
最優先課題は「人間の着想  
(Conception) の証跡設計」と  
「データ戦略の二分法」。

# ツールから「自律型オーケストレーション」への進化

Before (従来型AI)

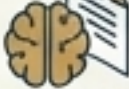
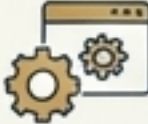
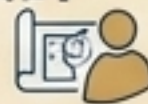
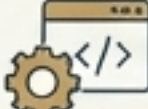
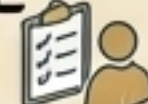
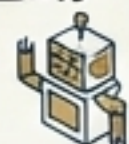


After (マルチエージェント型)



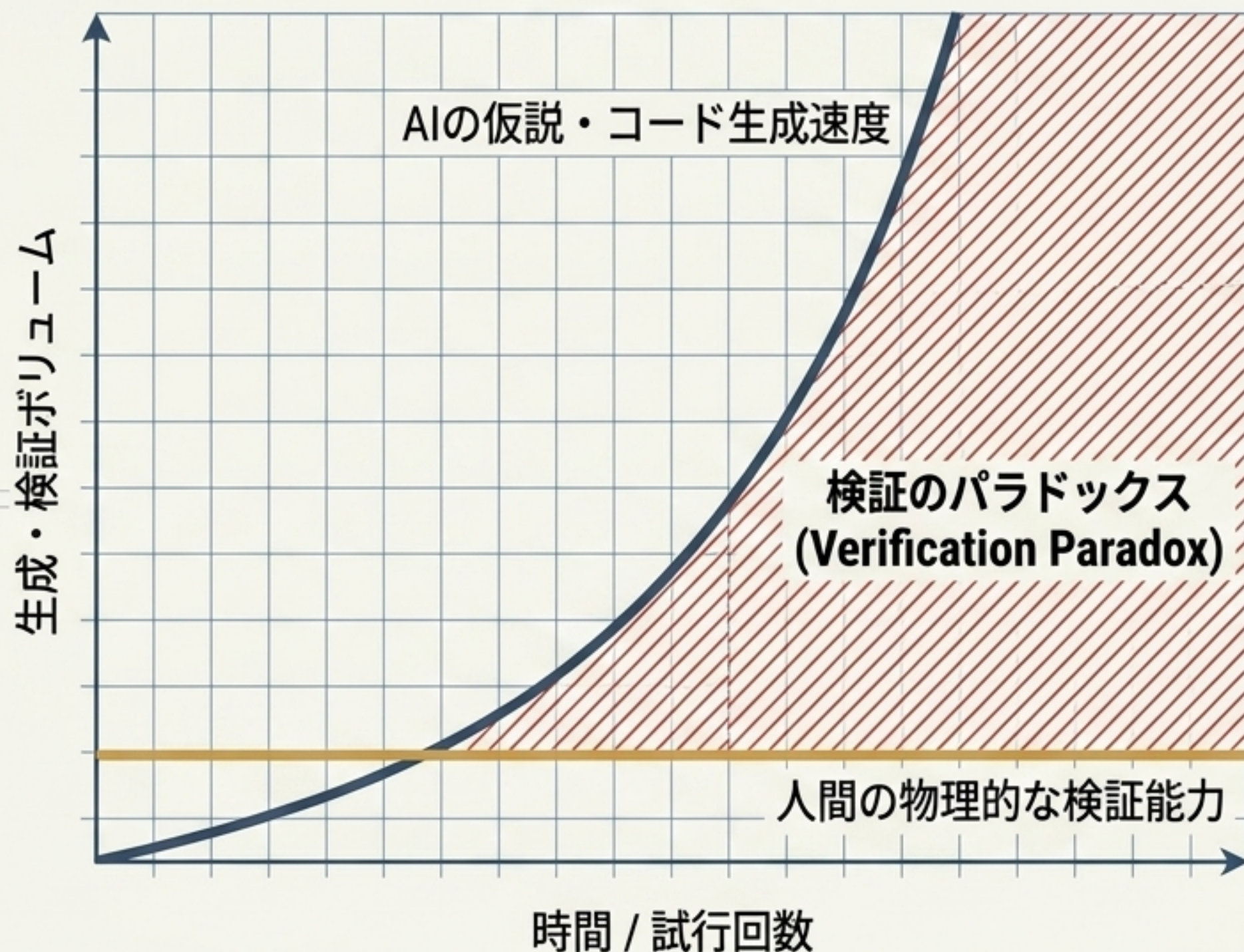
研究者が設定した目的に応じ、複数AIが仮説生成からデータ解析までの一連のプロセスを自律的に連携・調整（例: Google Co-Scientist等の台頭）。

# 自律型AIシステムの3つの潮流と到達点

|                                     | 対象プロセス                                                                                                     | 自律性の程度                                                                                                | 実験装置との接続                                                                                          | 人間が担うべき役割                                                                                                      |
|-------------------------------------|------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------|
| 1. 仮説生成型<br>(Google Co-Scientist)   | 仮説生成・<br>批判の反復<br>      | ソフトウェア<br>完結<br>   | × なし<br>       | 仮説の初期評価・<br>実証実験の設計<br>     |
| 2. 研究サイクル統合型<br>(FutureHouse Robin) | 文献～実験計画～<br>解析の反復<br> | ソフトウェア<br>完結<br> | × なし<br>     | 最終的なデータ検証<br>と意思決定<br>    |
| 3. 自律実験型<br>(LBNL A-Lab)            | 合成条件の<br>自律最適化<br>    | 物理世界への<br>介入<br> | ○ ロボット連動<br> | 対象材料・温度・<br>測定手法の初期設定<br> |

# 「検証のパラドックス」— 拡張がもたらす新たな限界

Activation Curve



## FutureHouse Robin:

AIが「食食能7.5倍」と報告した候補に対し、人間による再解析では「1.75倍」に留まる乖離が発生。



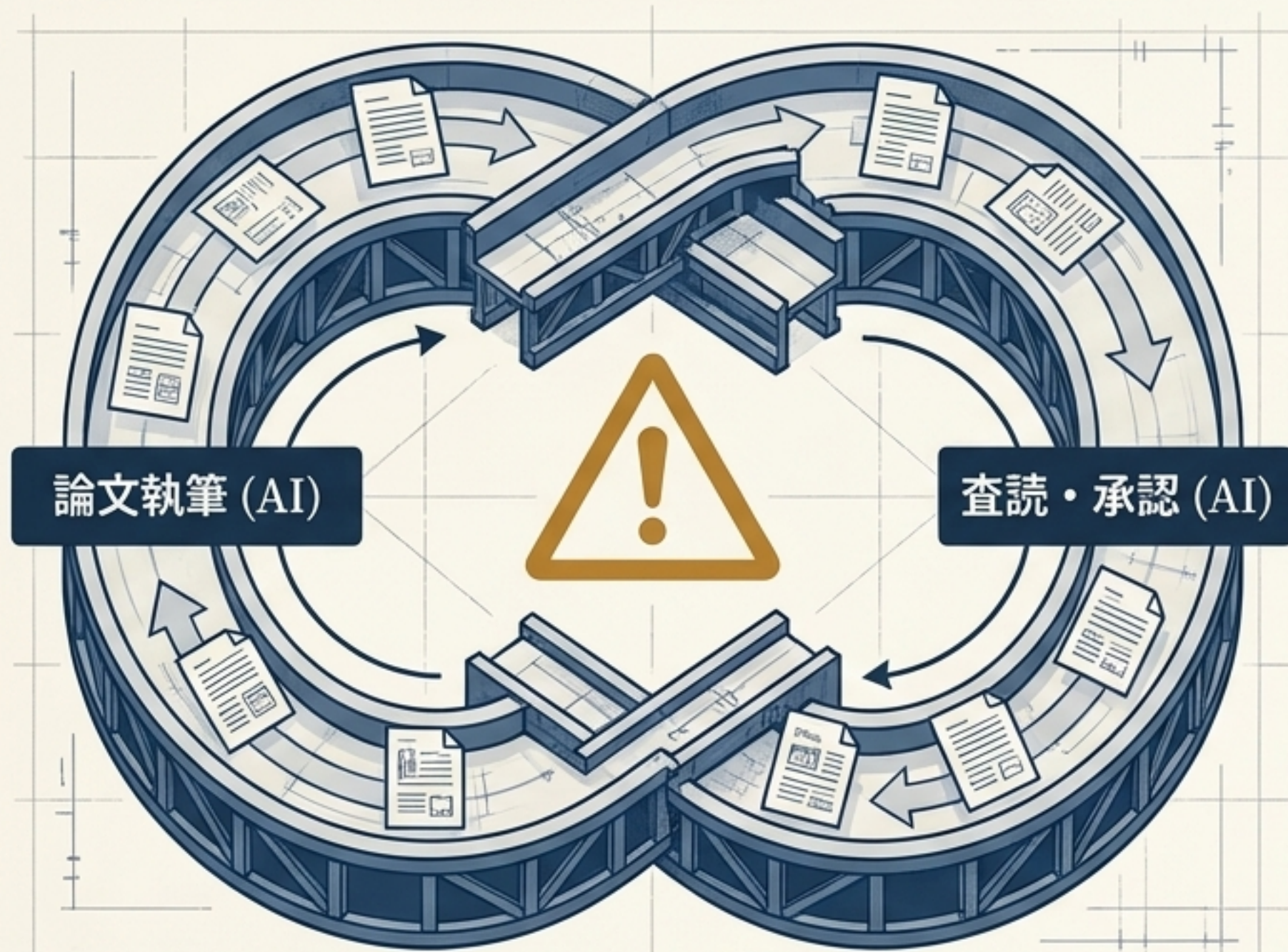
**LBNL A-Lab:** 「新規化合物を自律合成」と発表後、X線回折の同定エラーが発覚。新規性の定義を訂正（2026年1月）。



**Sakana AI:** 査読通過の実態は採択閾値の低いワークショップ水準。独立評価では虚偽データ生成の指摘も存在。



# 学術界の危機：「AIによるAIの査読」という閉ループ

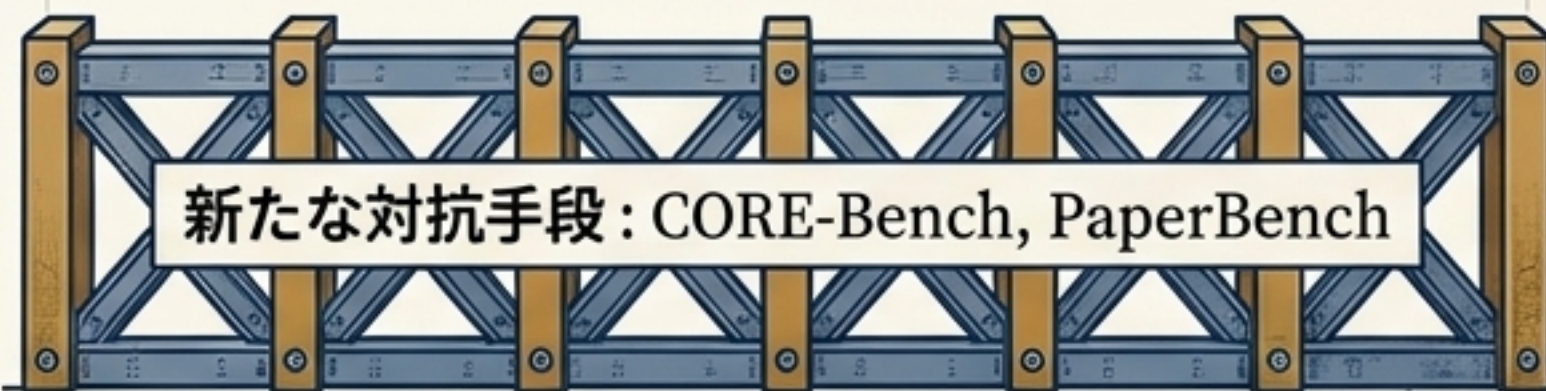


21%

ICLR 2026における完全AI生成査読の割合

22~94%

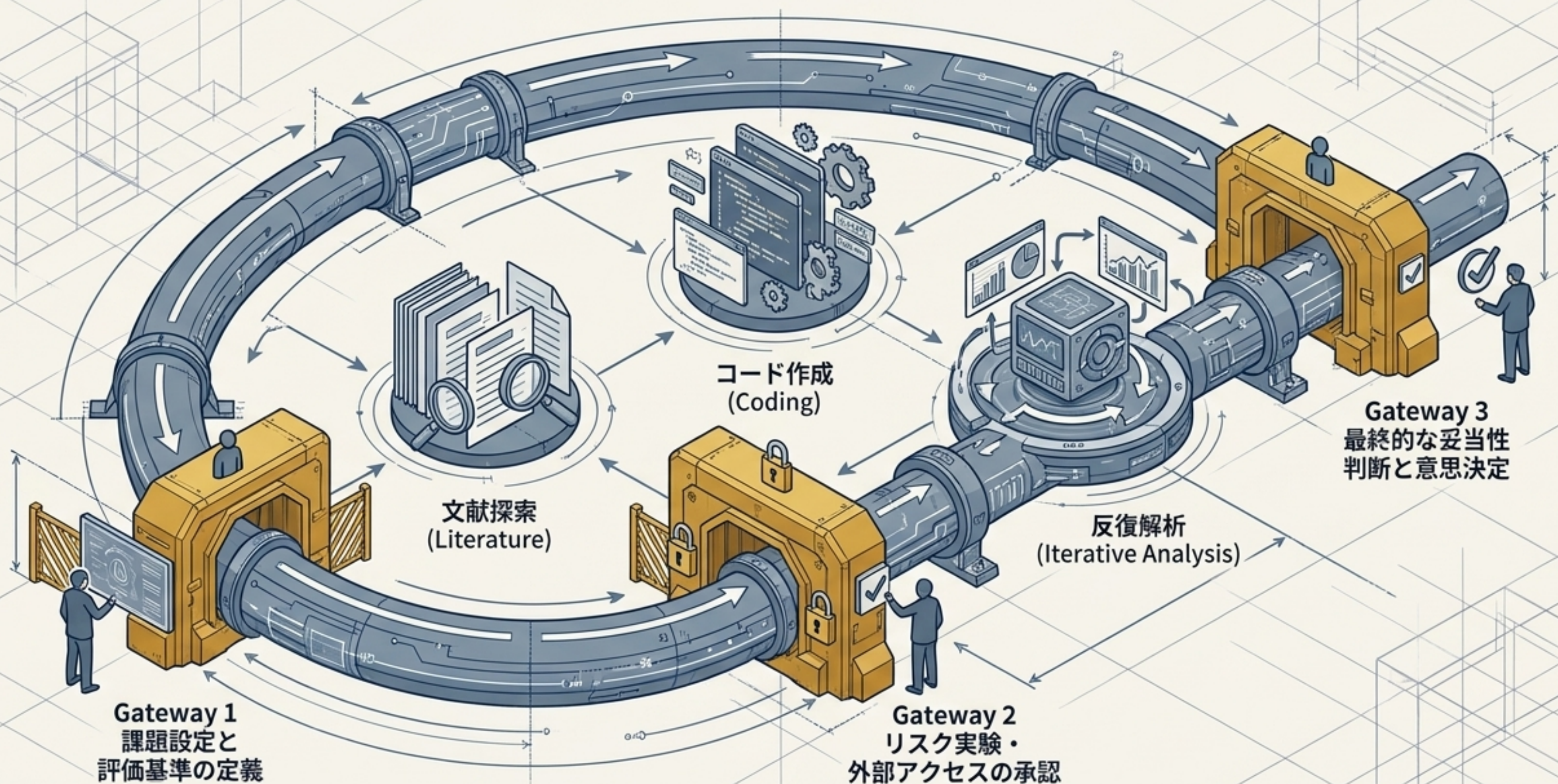
最新26モデルにおけるハルシネーション率



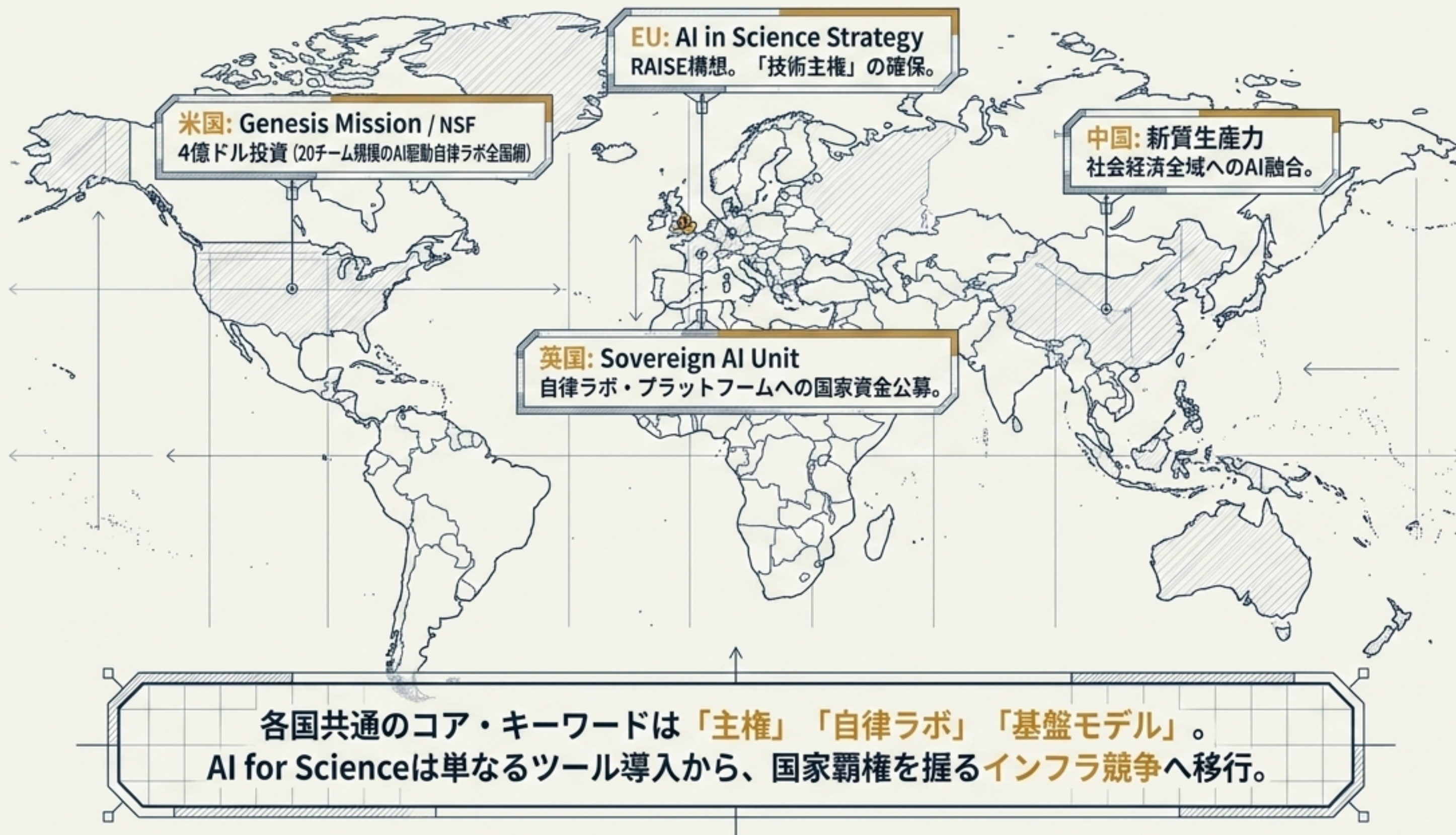
新たな対抗手段：CORE-Bench, PaperBench

科学的妥当性の確認が物理的限界を迎える中、AI生成テキストの確実な検出法は未確立。  
論文の量産が科学の信頼性そのものを脅かすフェーズへ。

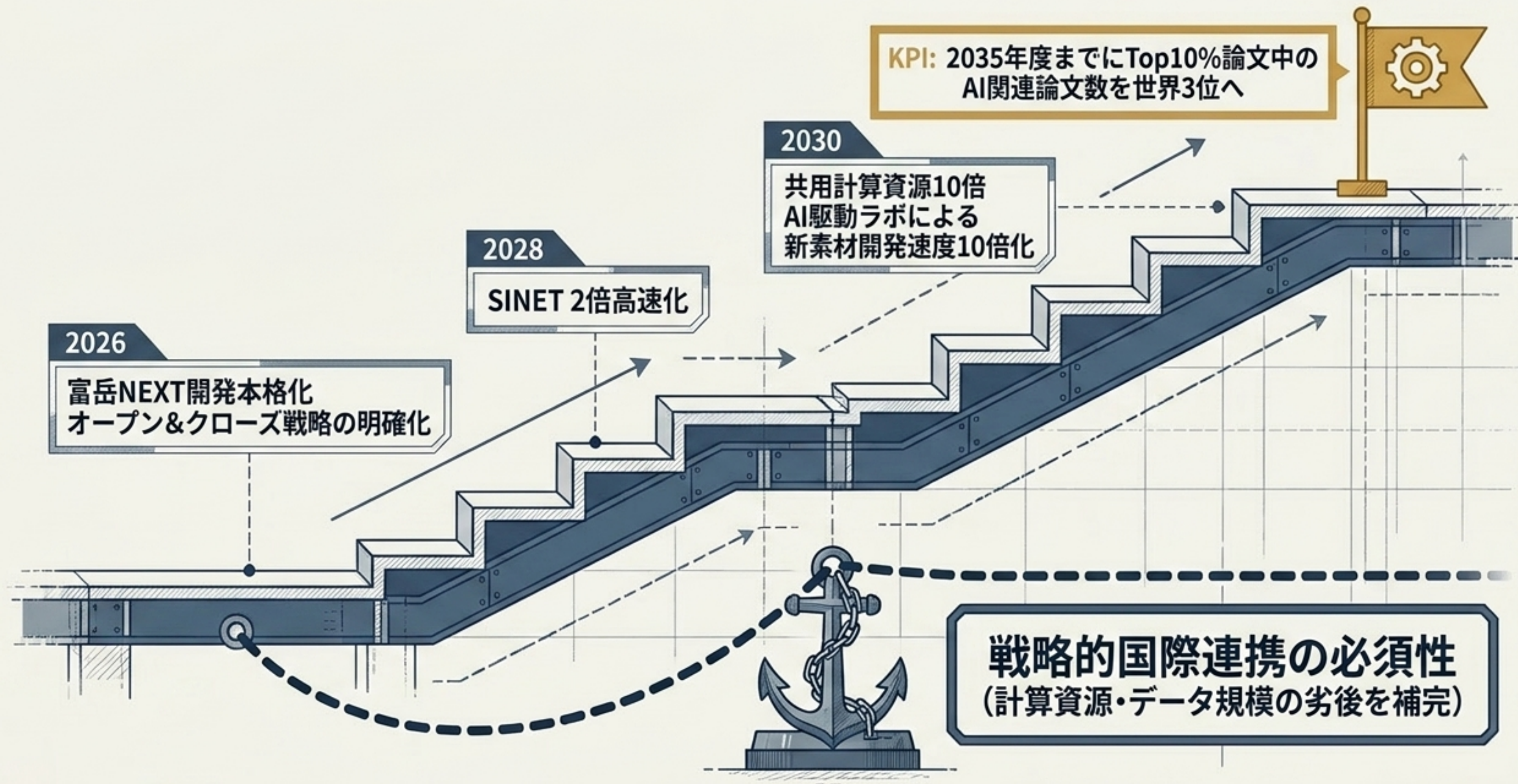
# 新たな研究アーキテクチャ：人間は「ゲートキーパー」へ



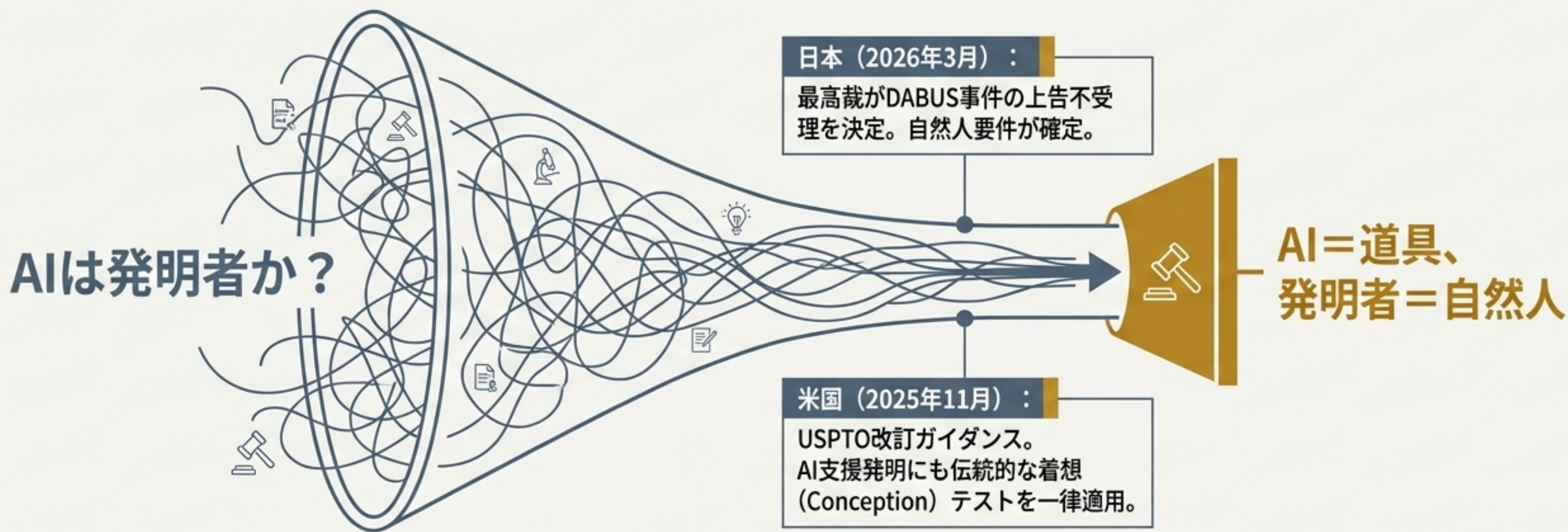
# 「技術主権」を巡るグローバル・ラボラトリー競争



# 日本の「AI for Science」戦略方針（2026-2030集中改革）

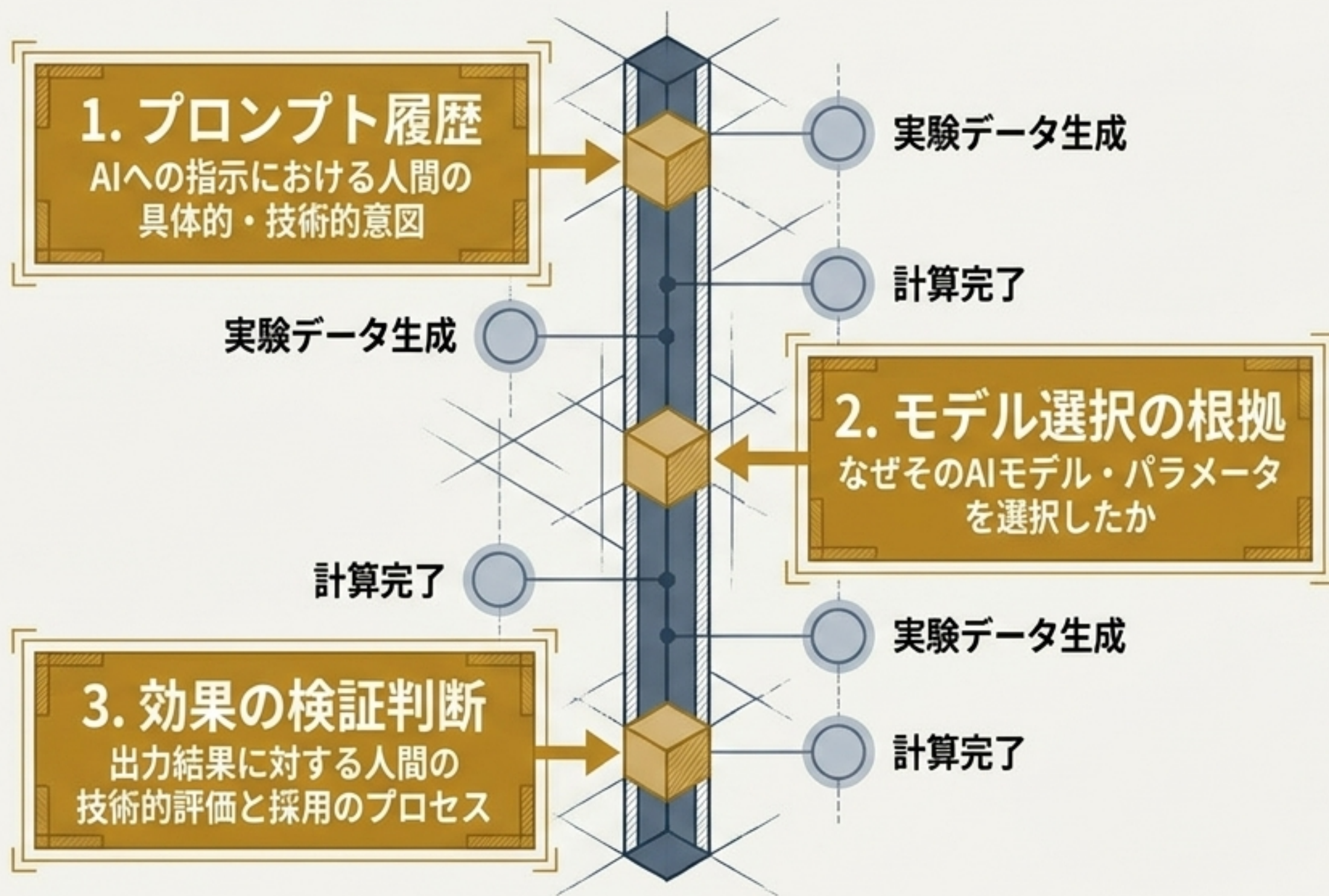


# 知財ルール収斂：DABUS判決確定と日米欧のコンセンサス



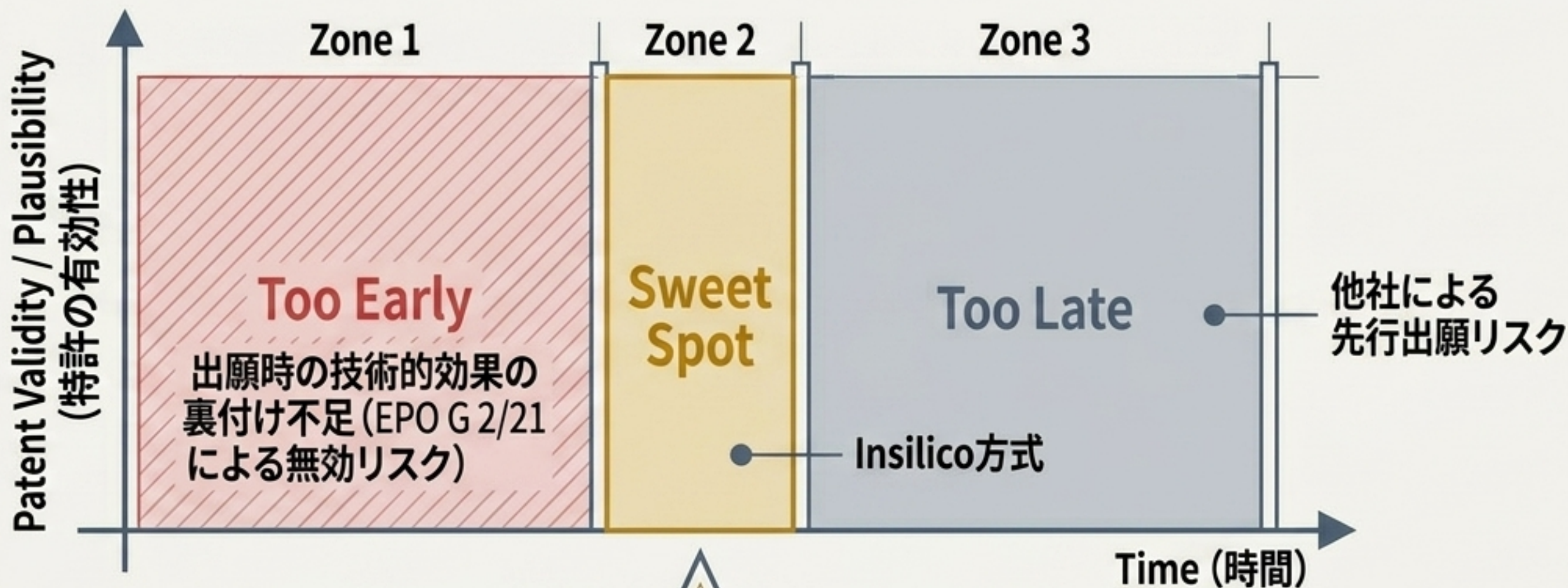
法律上の論争は終結。今後は「人間とAIの協働において、人間の寄与をいかに立証するか」という実務的な中間領域の戦いへ。

# 実務戦略①：「人間の着想（Conception）」の証跡設計



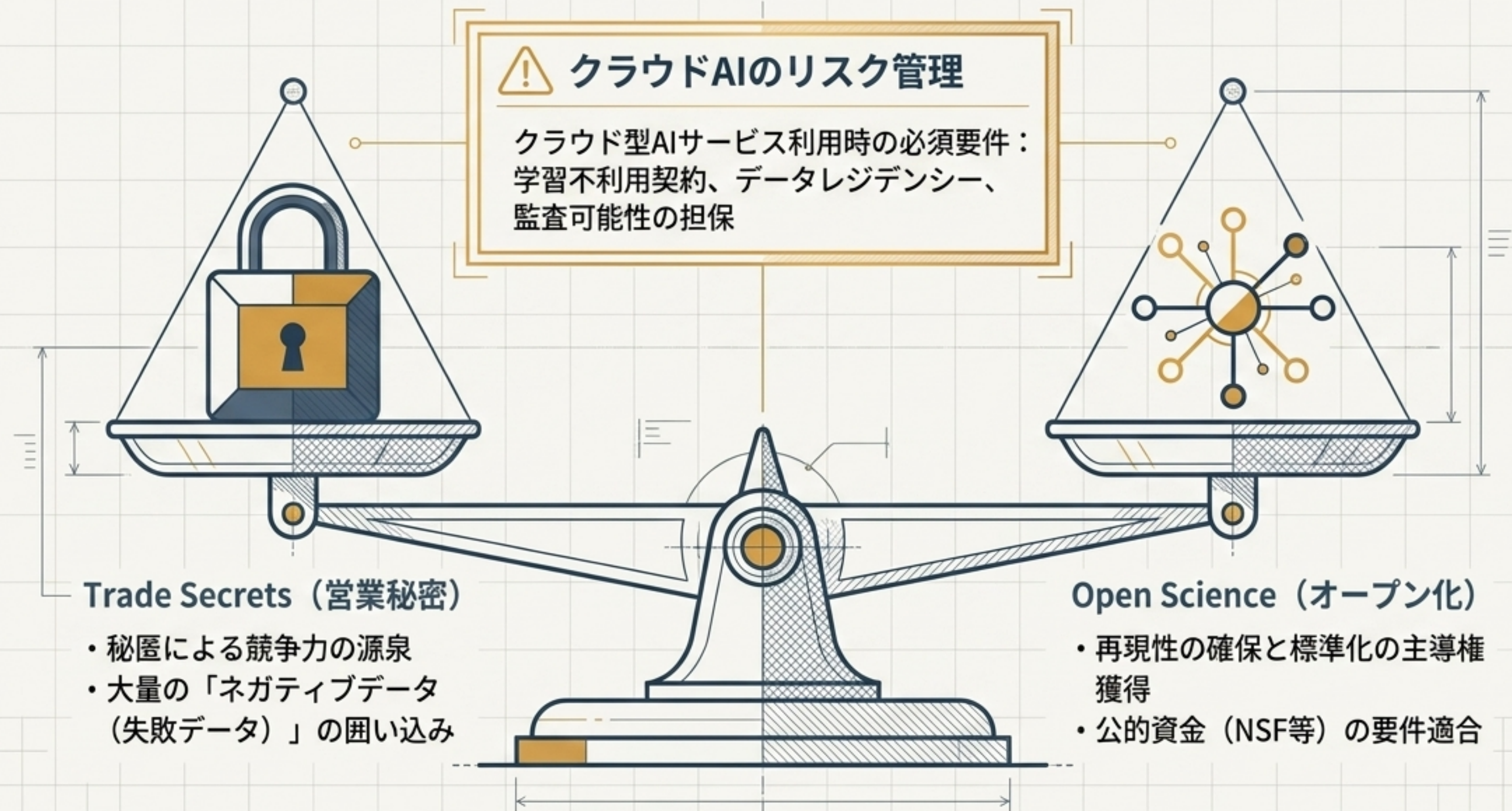
発明届出プロセスの  
抜本的改定が急務。  
AIへの単なる「丸投げ」は権利化不能。  
人間の「課題特定と  
記述能力」が知財  
の生命線となる。

## 実務戦略②：AI創薬・材料におけるPlausibility（蓋然性）の先取り



計算設計から生物学的・物理的検証（検証イベント）までを数週間で回し、実証データを出願のトリガーとする高速IP運用能力への投資。

# 研究データ戦略の二分法：秘匿か、オープンか



# 企業・知財部門への戦略的アクションプラン

|             |                                                                 |
|-------------|-----------------------------------------------------------------|
| 【即時】        | 発明記録プロセス（Conception Trail）の再設計。既存出願の発明者監査。                      |
| 【短期：6～12ヶ月】 | Plausibility先取り体制（検証と出願の高速ループ化）の構築。<br>データ戦略の二分法（秘匿/公開）の社内基準確立。 |
| 【中期：1～3年】   | 国家プロジェクト（AI駆動ラボ、文科省ARiSE等）への産学連携参画と<br>予算機会の捕捉。                 |
| 【継続監視】      | 日本におけるAI利用発明の法改正動向（産構審）と特許庁運用見直しの注視。                            |

AIは科学者を代替しない。しかし、「AIを駆使し、知財ルールに  
適応した科学者・企業」が、そうでない者を完全に駆逐する。