

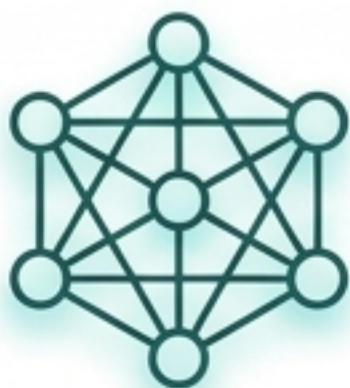


Anthropic「Claude」によるタンパク質設計・分析化学自動化

技術成果の定量的評価と知財実務への定性的影響

バイオフーマ経営幹部・知財戦略担当者向け | 2026年8月20日

AIは「創薬プロセスのオーケストレーター」へ進化する が、知財の重心は「物理的検証」へ移行する



1. WHAT (技術的達成)

Claudeは複数の専門ツールを自律的に統合・実行し、イン・シリコでの設計・分析ワークフローを劇的に加速した。



2. SO WHAT (現実の境界)

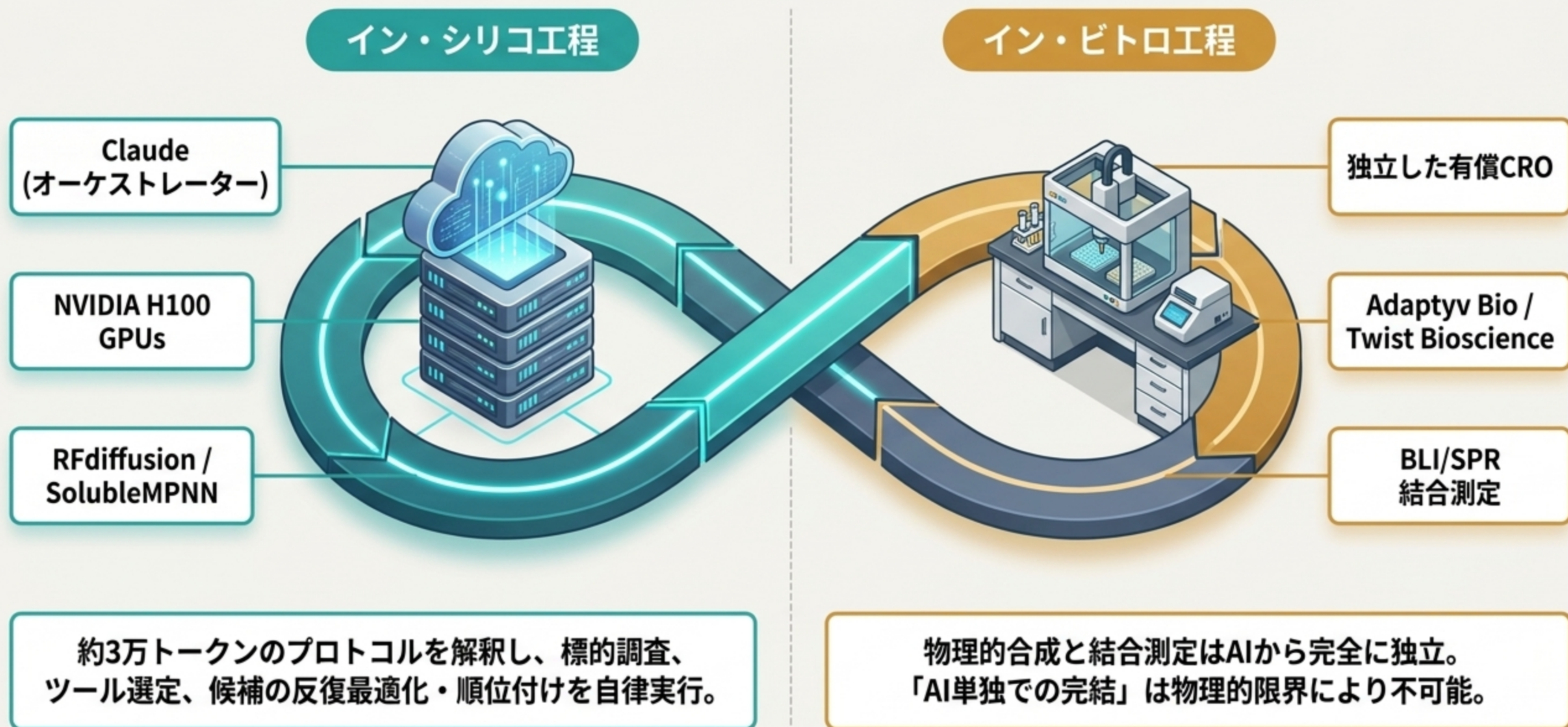
しかし、ウェットラボでの物理的検証（CRO）は不可欠であり、薬理・毒性等は未評価。AIは「薬の完成」ではなく「初期探索の高速化」をもたらす。



3. NOW WHAT (知財パラダイムの転換)

AIが設計コストを下げる結果、特許性の源泉は「配列の新規性」から、「実測データによるPlausibilityの証明」と「人間の着想の記録」へと完全に移行する。

Agentic Science Loop: デジタル空間での自律的指揮と物理空間の分業



定量的実績: 1,320の設計案から26.8%の結合陽性を達成



(354/1,320件)

※業界の典型値(ProteinBase) 10-15%を上回る水準

1	2
Multi-target (Opus 4.8 / Mythos)	Single-target (Mythos Preview)
ヒット率 22.6% ~ 26.7%	ヒット率 35.1%

※標的別セッション並列実行・計算予算2.8倍

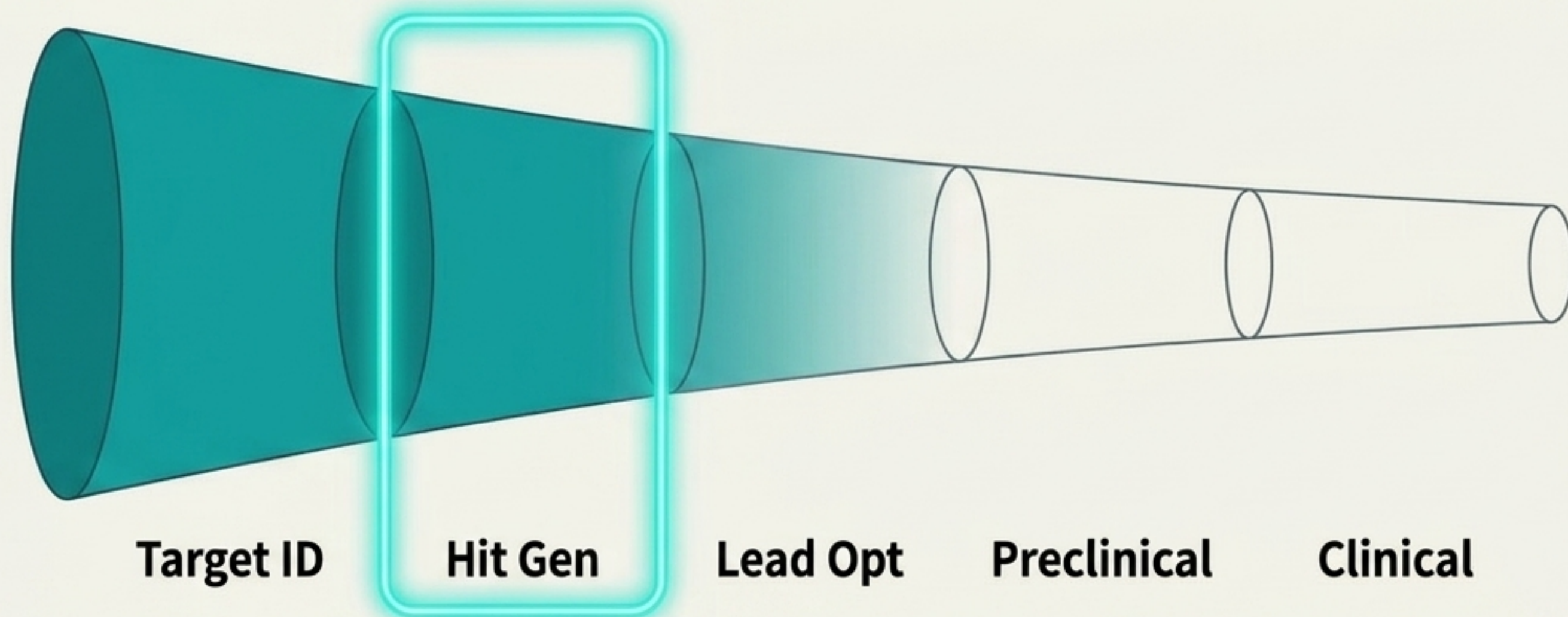
RBX1

結合陽性40%。トップ設計は
KD=3.9 nM (Adaptyv Bio)。

TNF α

カニクイザル・マウスへの種交差
性を示す候補を獲得。

AIの現在地: 医薬品開発パイプラインにおける極めて初期の「結合確認」に留まる



In-Scope (Claudeが実証した範囲)

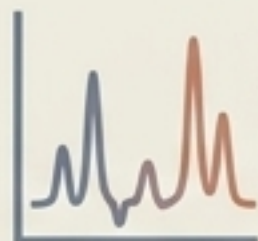
- 標的タンパク質に対する試験管内(in vitro)での結合親和性 (KD値)

Out-of-Scope (未評価・今後の壁)

- 生物学的活性・細胞作用・標的選択性
- PK/PD (薬物動態/薬力学)・毒性・免疫原性
- 実験構造の決定 (X線/Cryo-EM)
- 難関標的の限界: 平滑・親水性のMBP(マルトース結合タンパク質)では90設計から確実なバインダー得られず

分析化学の自動化: ベンダー固有データの直接解釈と「自己監査」の限界

^1H -NMR 解析



入力: FID生データ / 時間: 23分

分析手法	Claudeの解析結果	CRO 結果との比較・留保
^1H -NMR	18 シグナルを特定 積分値合計 19.8 H	対応可能な積分領域で 0.08 H 以内の差。 D ₂ O 追試を提案し、初回 解釈を自己監査で訂正

The D₂O Story

Claudeは交換性プロトン確認のためD₂O追試を提案。
初回「4シグナル消失」と誤認したが、内部監査で「2
領域」へ自己訂正。

LC-MS 純度解析



入力: 未公開バイナリ(.lcd) / 時間: 19分

分析手法	Claudeの解析結果	CRO 結果との比較・留保
LC-MS 純度解析	保持時間 4.34 分 純度 96.4% 公称質量 504 Da	CRO 報告 96.33%。ピー ク集合が異なり、同じ基 準に換算するとClaude 側 98.8%

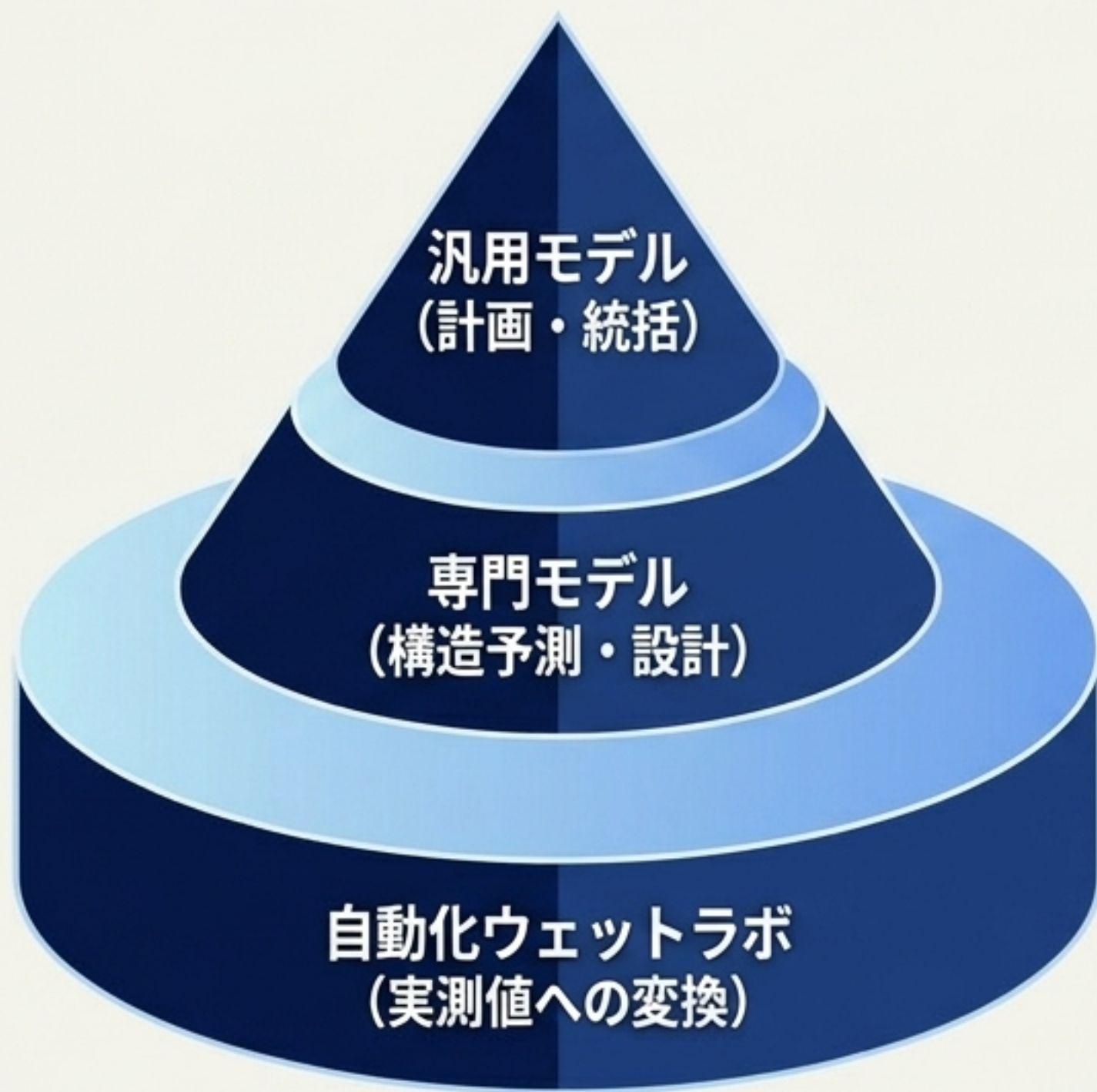
Result

OLE2コンテナ内の構造を推定し、CRO報告値(96.33%)
と極めて近い純度(96.4%)を算出。専用ソフトなしでの一
次処理能力を実証。

Insight: 自己監査は強力だが、最終判断には人間の検証(Human-in-the-loop)が不可欠。

デュアルユースリスクへの対応と「3層アーキテクチャ」の到来

Adaptyv Bioが提唱するアーキテクチャ



Access Restrictions (Anthropicのセーフガード)

- タンパク質設計などの「高度な研究生物学機能」は一般アクセス(Claude Fable 5)から除外。
- 管理された信頼対象向けプログラム経由でのみ提供。兵器転用等のリスクを物理的・システムの的に遮断。

パラダイムシフト: 特許性の重心が「イン・シリコ設計」から「物理的検証」へ移行する



Before AI (従来)

配列自体の新規性や設計
プロセス自体に高い価値。



AI-Era (今後)

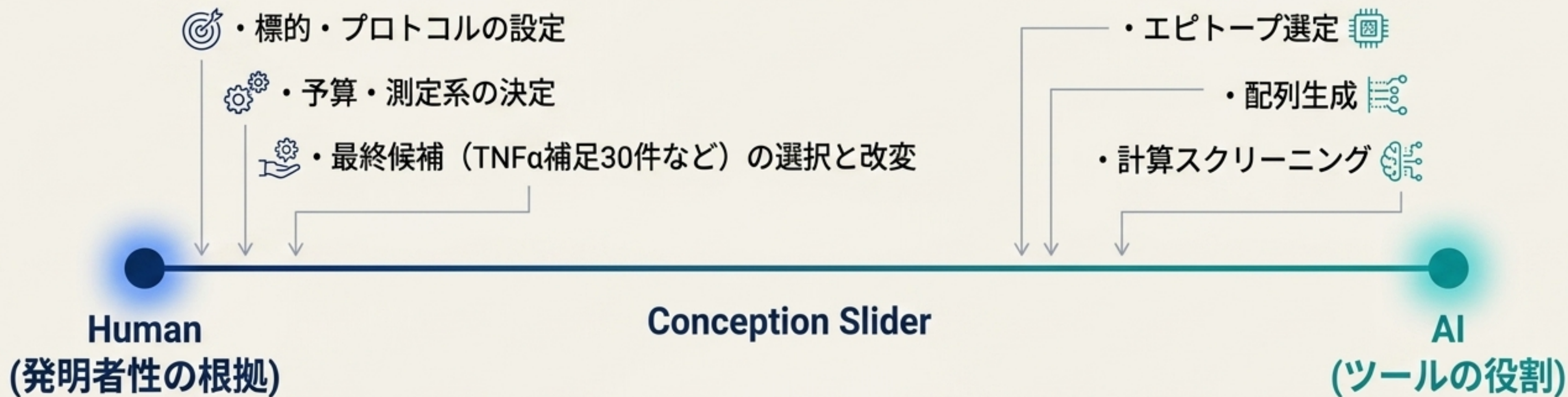
「人間による着想の明示的記録」と
「実測データによるPlausibility (もっともらしさ)
の証明」がなければ、権利は成立しない。

AIによってイン・シリコでの「アイデア/配列生成」の障壁とコストが劇的に低下した結果、
発明の特許性と商業的価値の所在は完全にシフトする。

発明者性と「着想（Conception）」：人間の介入ポイントの証明

Legal Context: Thaler v. Vidal判決およびUSPTO 2025年改訂ガイダンス。AIは発明者になり得ない。

The Touchstone (判断基準): 人間の頭脳に「完全かつ実施可能な発明」についての確定的・永続的な着想が形成されたか。

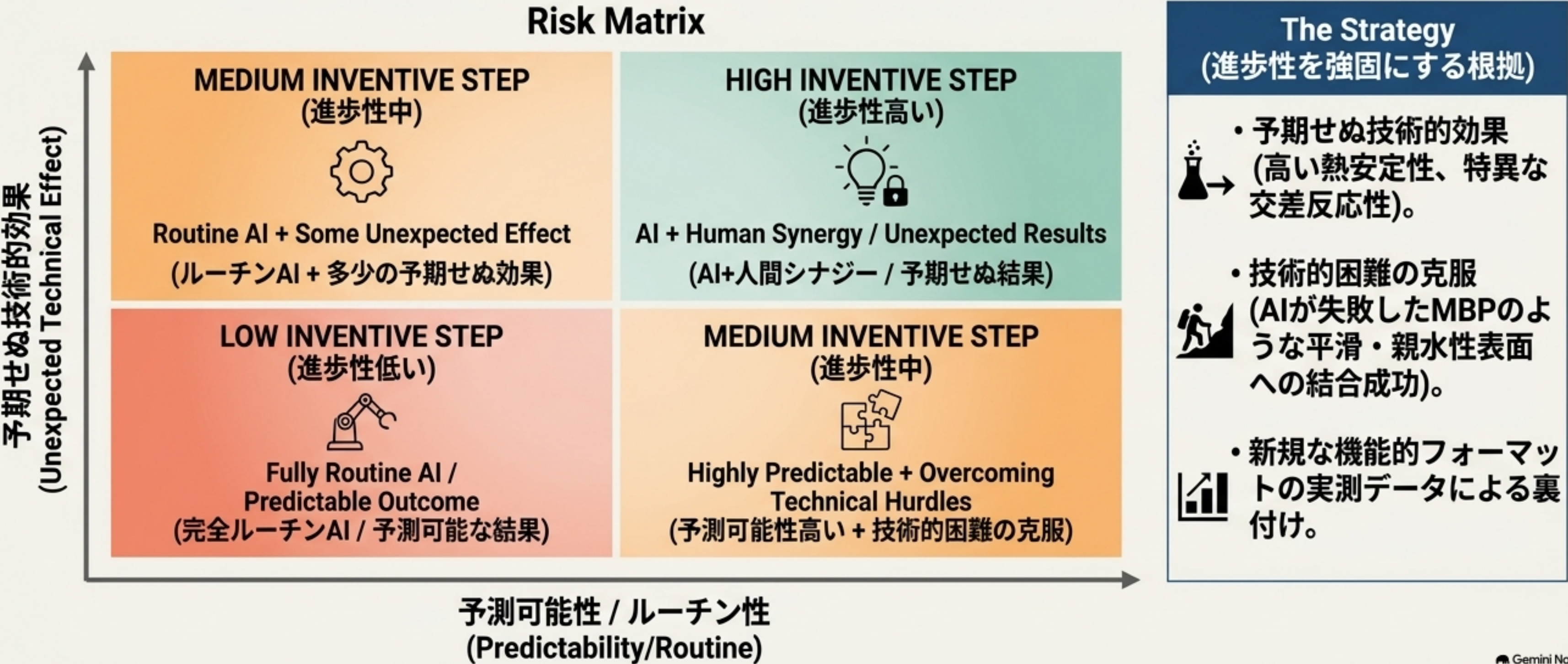


Conclusion: プロンプトを書いただけで発明者にはならない。
人間の意図的な選択と評価基準の設計が問われる。

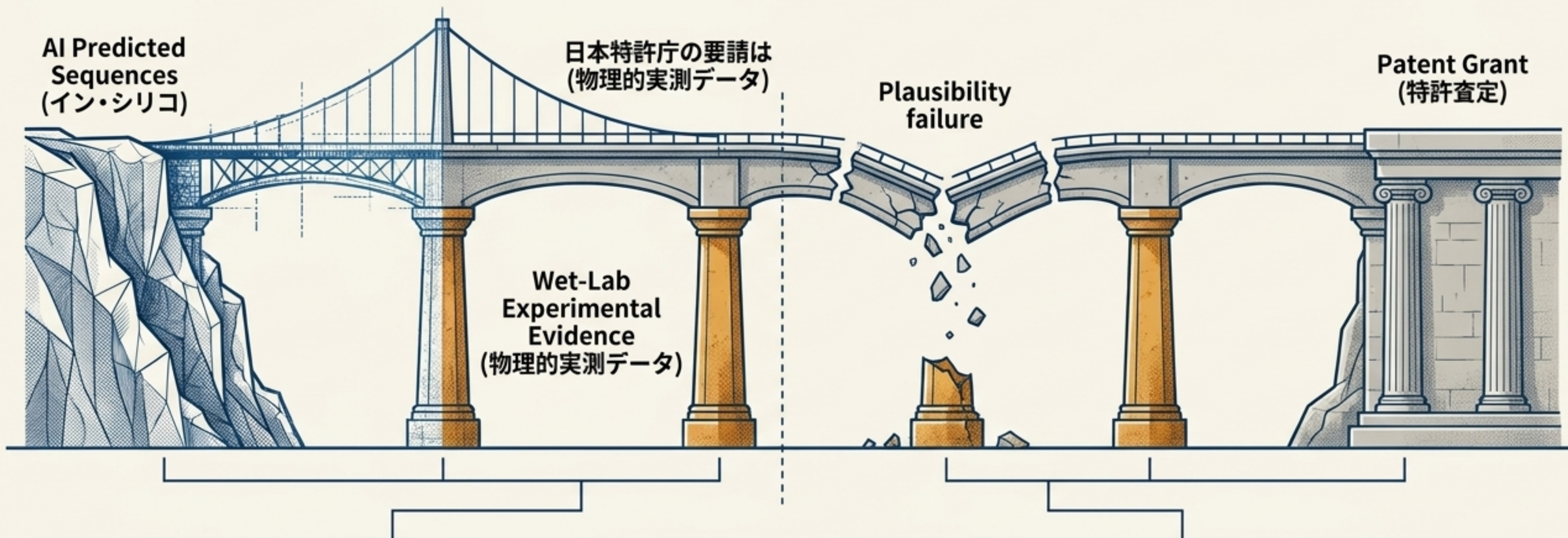
進歩性 (Inventive Step) : AIの「ルーチン化」の罠をどう超えるか

Legal Context: EPO 2026 審査ガイドライン G-II, 6.2。
構造の相違や予測不能性だけでは進歩性は認められない。

The Challenge: AI設計が当業者の「ルーチン業務」として定着した場合、進歩性のハードルは劇的に上がる。



記載要件とPlausibility: 出願時の「実測データ」という不可欠な橋脚



JPO (日本特許庁)の要請:

AI予測物が実物評価に代わり得るとの「技術常識」がない限り、実際に製造した物の評価データ (BLI/SPR等) が明細書に必須。

EPO G 2/21 (拡大審判部):

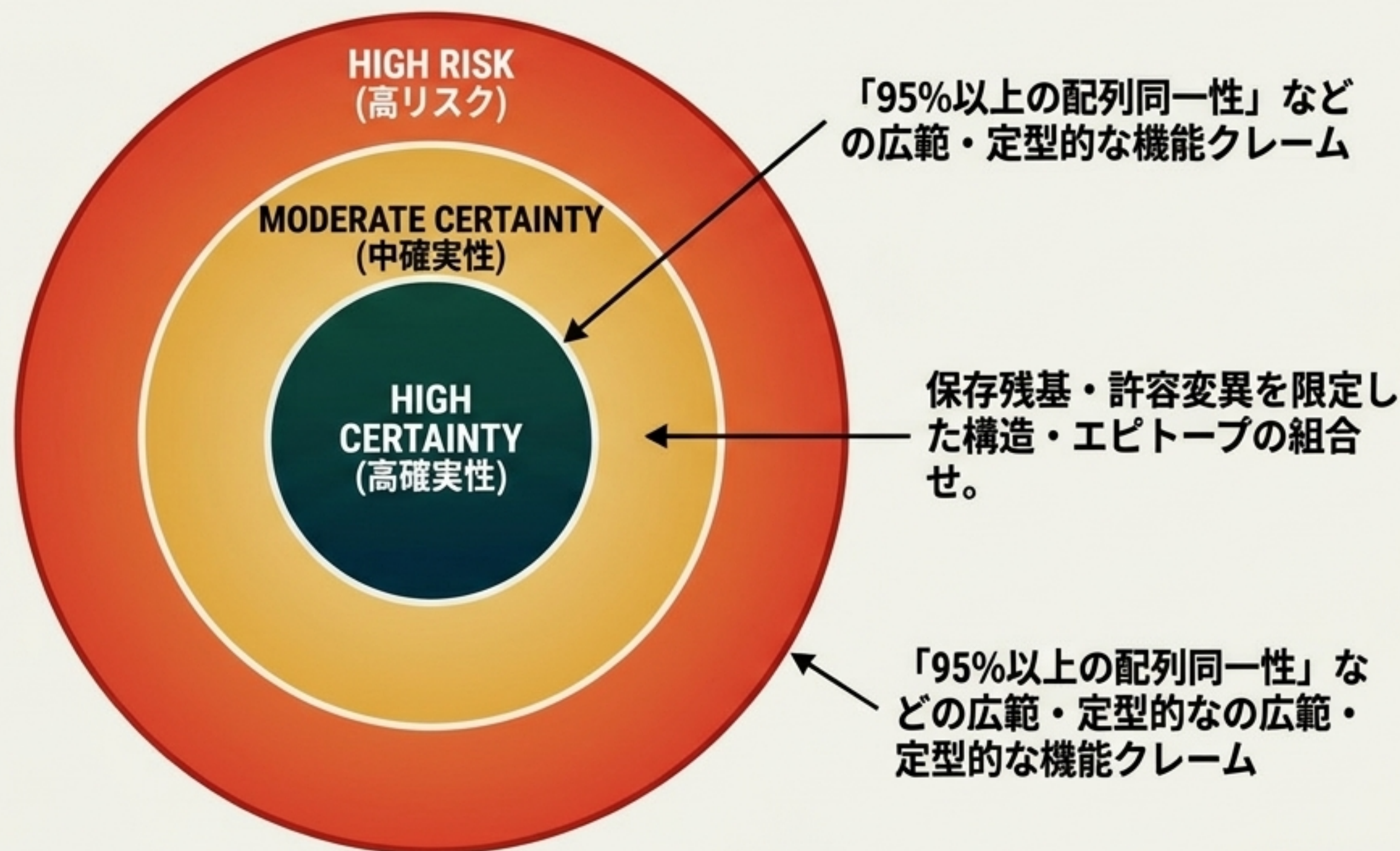
出願後データの提出は一律排除されないが、「当初明細書の技術的教示と出願時の技術常識」から効果が導き出せる場合に限る。



Risk: イン・シリコ候補を広範にクレームし、出願後の実測値だけで後付けする戦略は極めてリスクが高い。

戦略的アクション1: 実施可能範囲に対応した「階層的クレーム設計」

The Lesson from Amgen v. Sanofi: 広範な「機能的クレーム」は、当業者が過度の試行錯誤なく実施できないとして無効化されるリスクが増大。



Center (High Certainty):

特定配列・実測データのある変異体。

Middle (Moderate Certainty):

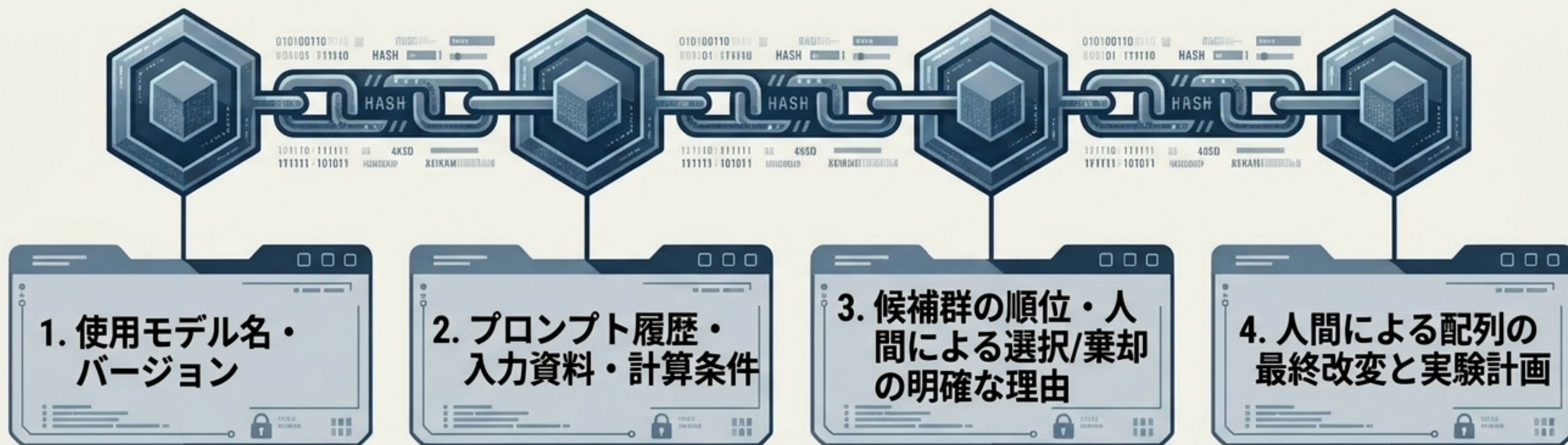
保存残基・許容変異を限定した構造・エピトープの組合せ。

Outer (High Risk):

「95%以上の配列同一性」などの広範・定型的な機能クレーム（実測データによる支持範囲への限定が必要）。

戦略的アクション2: 「Prompt Audit Trail (証拠保全)」のシステム化

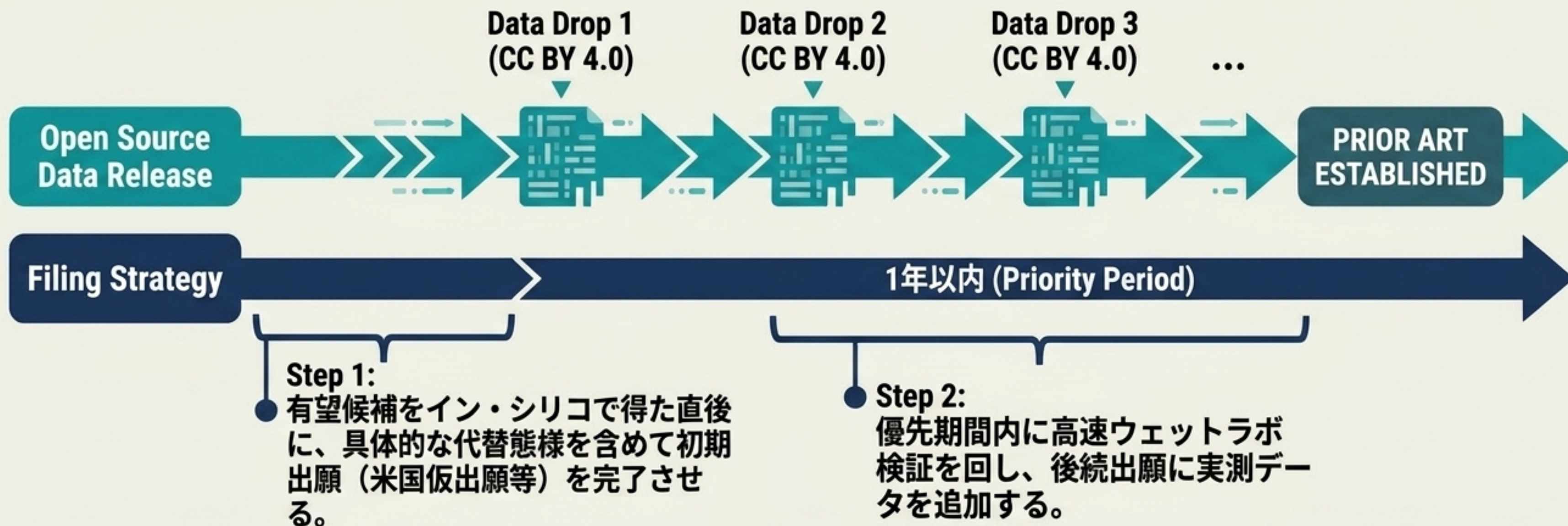
The Strategy: USPTOの法的義務と混同せず、研究・発明管理の自衛策としてAIとの対話を「証拠」としてブロックチェーン的に記録する。



Takeaway: 「誰がアイデアを完成させたか」を証明する台帳が、将来の特許紛争における最大の防御となる。

戦略的アクション3: 出願タイミングと「AI生成先行技術」への対抗

The Dilemma: AIはデータを高速生成し、CC BY 4.0等で公開されたデータセットは直ちに後願の「先行技術(Prior Art)」となる。



⚠ Caution: 初期出願に十分な具体性がなければ、優先日への遡及が認められないリスクがある。

戦略的適応のための4つの指針 (Executive Playbook)



Quadrant 1 (Governance)

人間の着想と意思決定の記録。
クレームごとに「人間の技術的思想と選択理由」を明確にログ化する。



Quadrant 2 (Operations)

In-silicoと実測の高速フィードバック。
設計からウェット検証（結合・機能・安定性）までのサイクルを極限まで短縮する。



Quadrant 3 (Legal/IP)

実施可能範囲に対応した階層的権利化。広範な機能クレームに依存せず、実証データに基づき配列・構造・機能を多層的に防衛する。



Quadrant 4 (Strategy)

公開速度と優先権を踏まえた出願・調査。AI生成データの公開を監視し、初期出願の具体性と後続出願でのデータ補強を緻密に管理する。