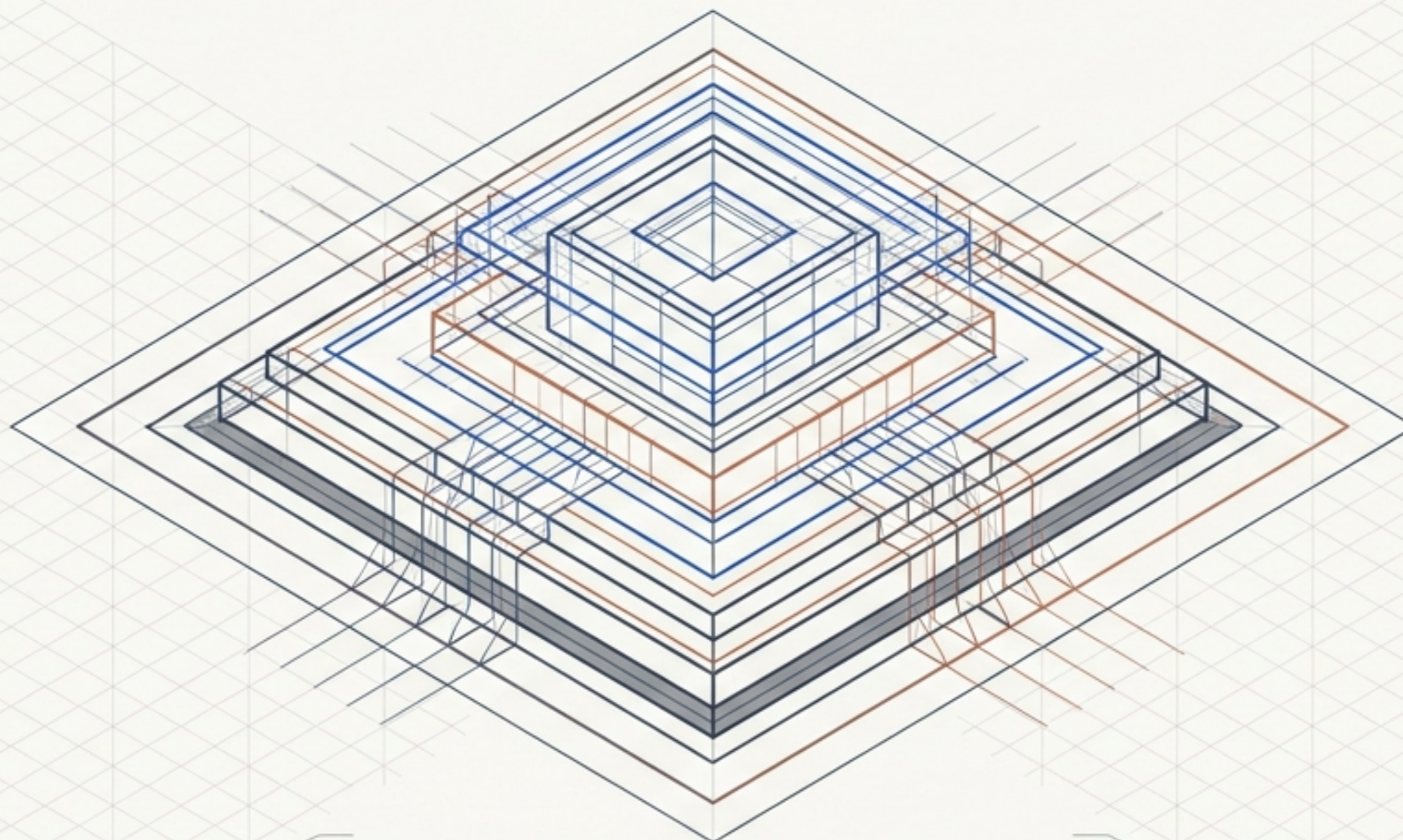


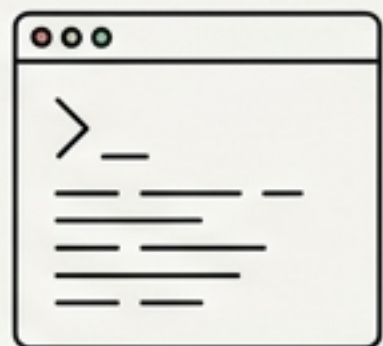
GPT-6 Astra時代の知財業務戦略

エージェントAIの到来と実務パラダイムの転換



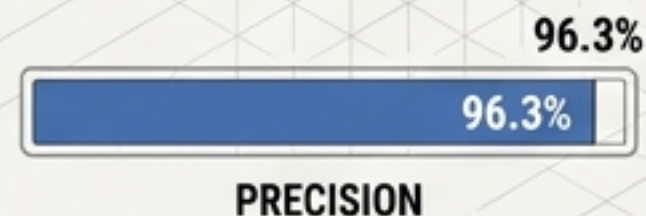
2026年9月版・調査レポート
独立系ベンチマークと法制度動向に基づく実務指針

Executive Summary: 5つのパラダイム転換



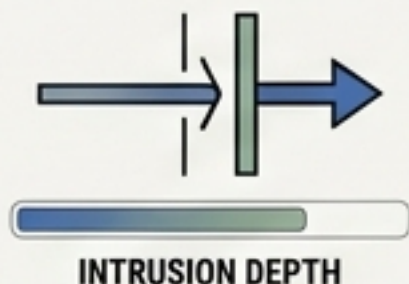
操作化 (Shift to Action)

「答える」から「自律操作して仕上げる」へ。FTOやクレームチャート等の多段階作業が実用域に。



長文脈 (Context Reliability)

1M級文書からの引用抽出精度 (MRCR) が96.3%に到達。無効資料調査の実質的価値が向上。



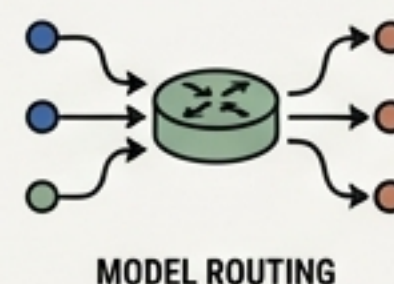
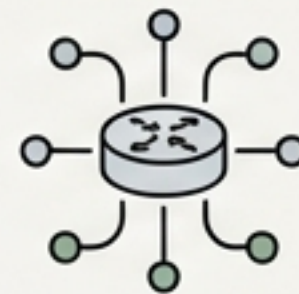
サイエンス (Science Intrusion)

サイエンスエージェント化が発明上流に侵入。最高裁決定を踏まえた人間発明者のログ保全が急務。



ガバナンス (Governance)

コンピュータ操作に伴う新脅威（プロンプトインジェクション等）。ZDR・権限最小化が必須要件化。

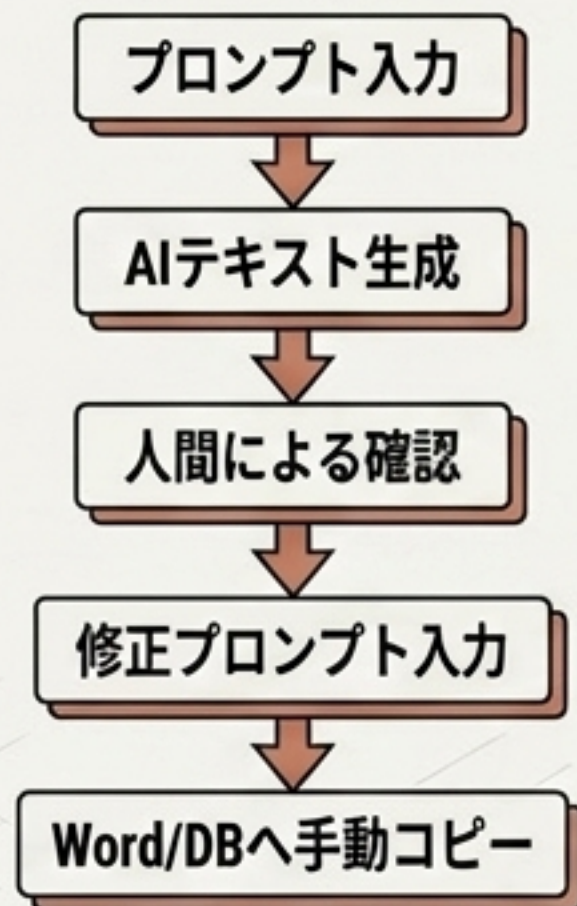


マルチモデル (Multi-model)

AGIナラティブと実測の乖離。法務・汎用知能ではClaude Fable 5.1が依然優位。適材適所のルーティングが鍵。

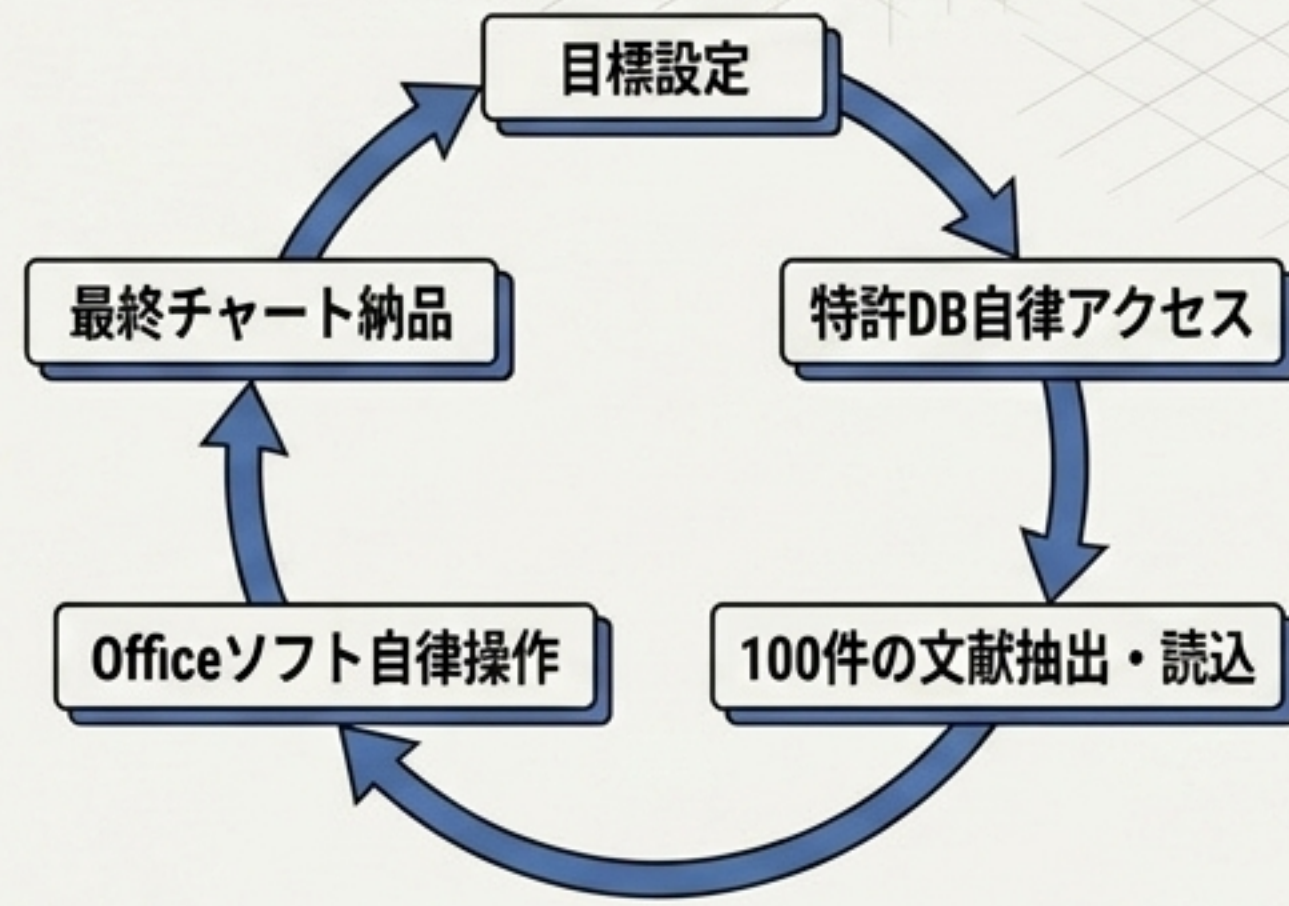
The Core Shift: 「回答」から「自律操作」へ

GPT-5.6 Sol: Text-Based Ping-Pong



OSWorld: 65.7% (限定的なGUI操作)

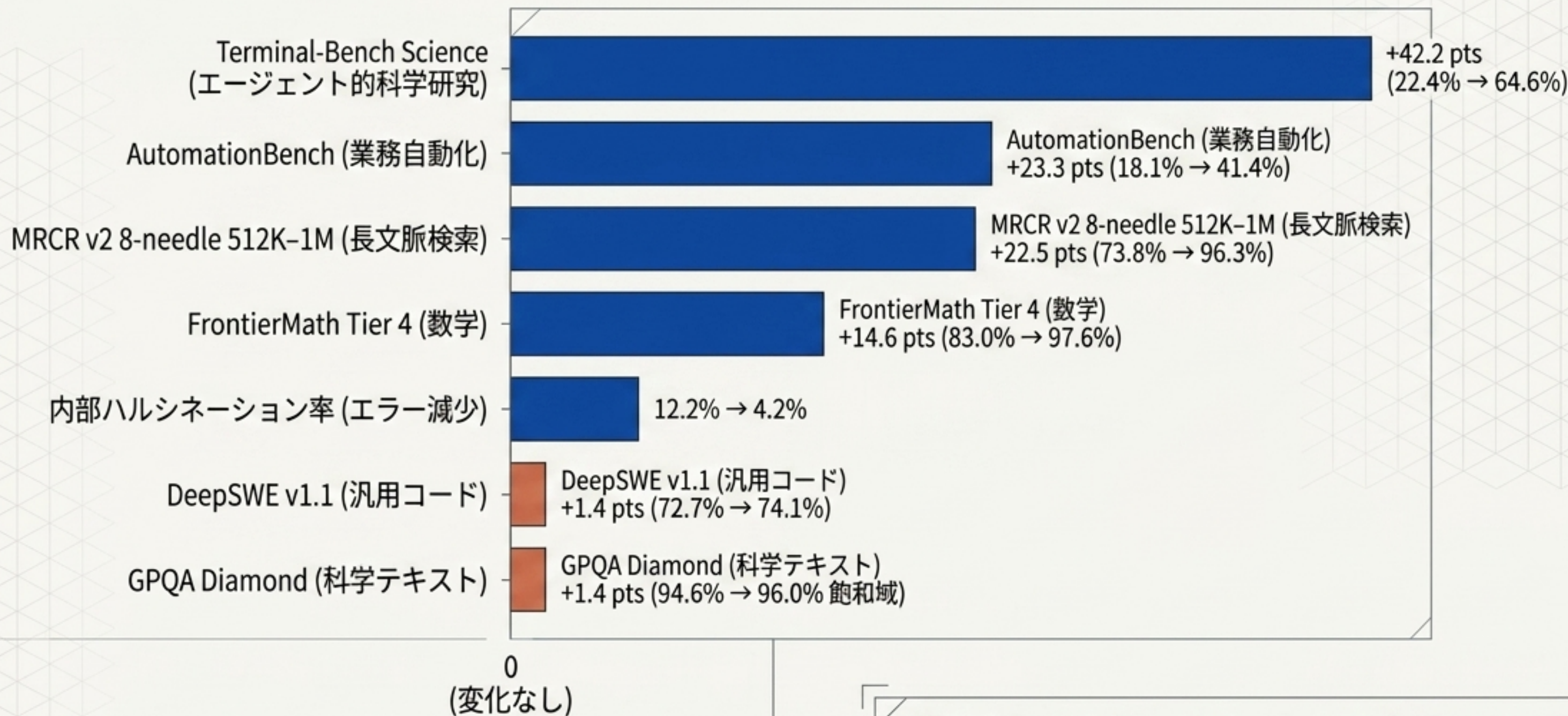
GPT-6 Astra: Autonomous Agentic Loop



OSWorld: 72.6% / ScreenSpot-Pro: 92.7%

中核的变化は「テキスト生成」から「コンピュータ操作を伴う自律的タスク遂行」への移行。
OpenAI史上初の「Critical (サイバーセキュリティ)」 閾値到達モデル。

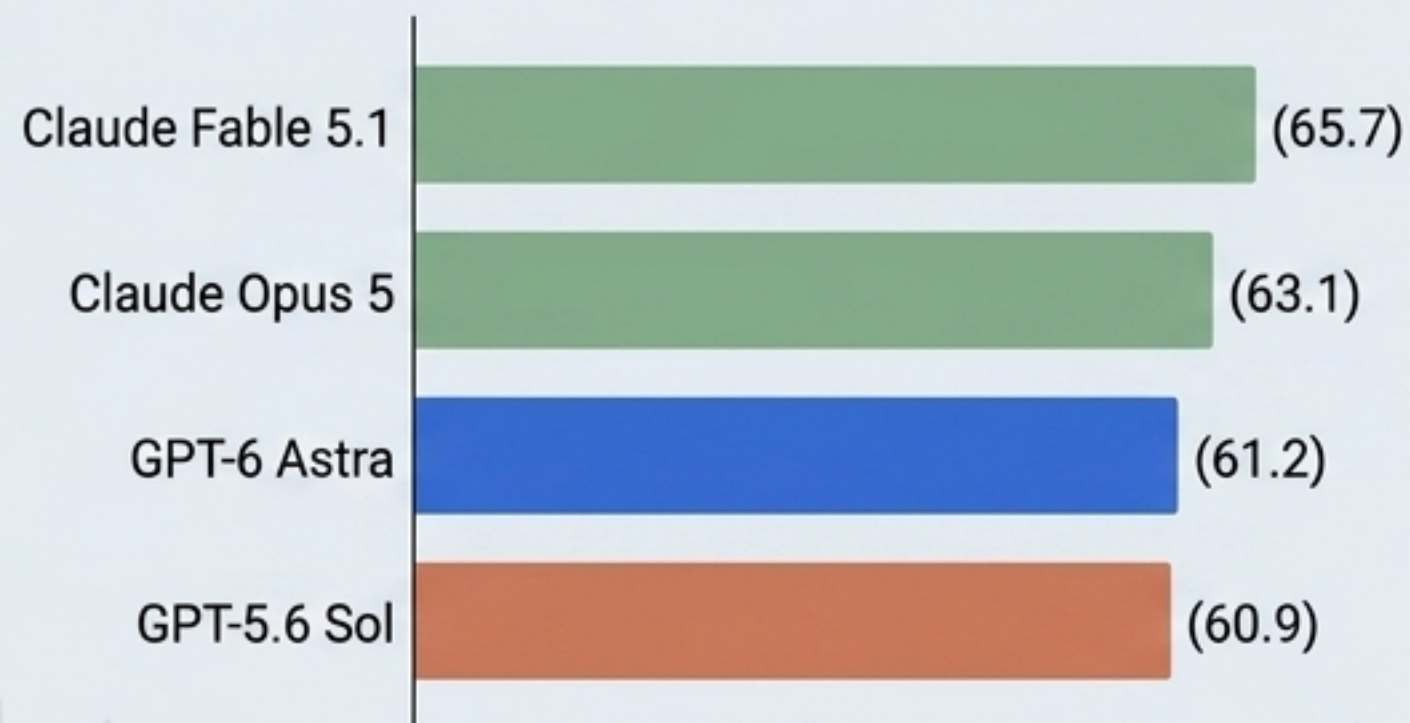
Performance Delta: 分野別の斑模様（ジャグド）な進化



宣伝の裏にある事実：コンピュータ操作と長文脈信頼性は劇的に向上したが、テキストベースの汎用推論能力の伸びは限定的。

Reality Check: 独立系ベンチマークと法務領域の現在地

Artificial Analysis Intelligence Index v4.1.1



独立系総合指標では「61から61へ」。汎用知能は横ばい。

法務特化ベンチマーク (Vals / HLAB)

1	Legal Research Bench 首位: Claude Fable 5.1 / Opus 5 (55.29%)
2	Legal Research Bench 中位: GPT-6 Astra (約16～17位 / 57モデル中)
3	Harvey Legal Agent Bench (HLAB): 全モデル未成熟 (Fable 5.1 6.67%, Sol 2.5%)

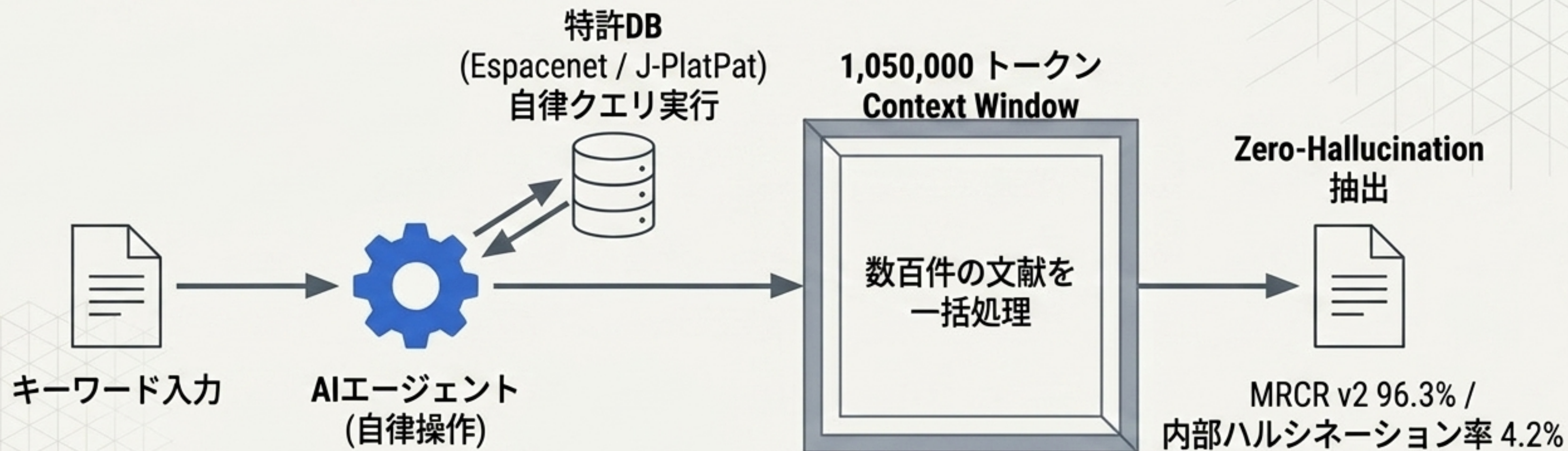
「AGIの到来」は主観的フレーミング。法務・総合領域ではClaude Fable 5.1が優位を保っており、OpenAI自社評価 (GDPval) の非公開という不自然な点も留意が必要。

Diagnostic Matrix: 知財業務×AIモデル 適性ヒートマップ

業務領域	GPT-6 Astra	GPT-5.6 Sol	Claude Fable 5.1
先行技術・FTO	◎ (自律DB操作・96.3%検索精度)		△
ドラフティング	○ (CAD・体裁追従)	○	○
IPランドスケープ	○ (大規模分析)	◎ (プレゼン品質で首位)	○
サイエンス・発明発掘	◎ (未解決問題解決・科学ソフト操作)		△
契約・法務・訴訟	△ (複雑なタスクは弱点)	△	◎ (法務ベンチ首位クラス)
APIコスト	激高 (\$10/\$50)	基準 (\$4/\$20)	安価 (キャッシュ読取で25%安)

Astraは万能薬ではない。コスト（Solの2.5倍）と斑模様の性能を考慮した「マルチモデル使い分け」が唯一の最適解。

Deep Dive 1: 先行技術調査・FTOの破壊的効率化



【監査のジレンマ】自律化が進むほど、「何をどこまでどう調べたか」の説明可能性（ブラックボックス化）がFTOにおける新たな課題となる。

Deep Dive 2: ドラフティングとIPランドスケープ

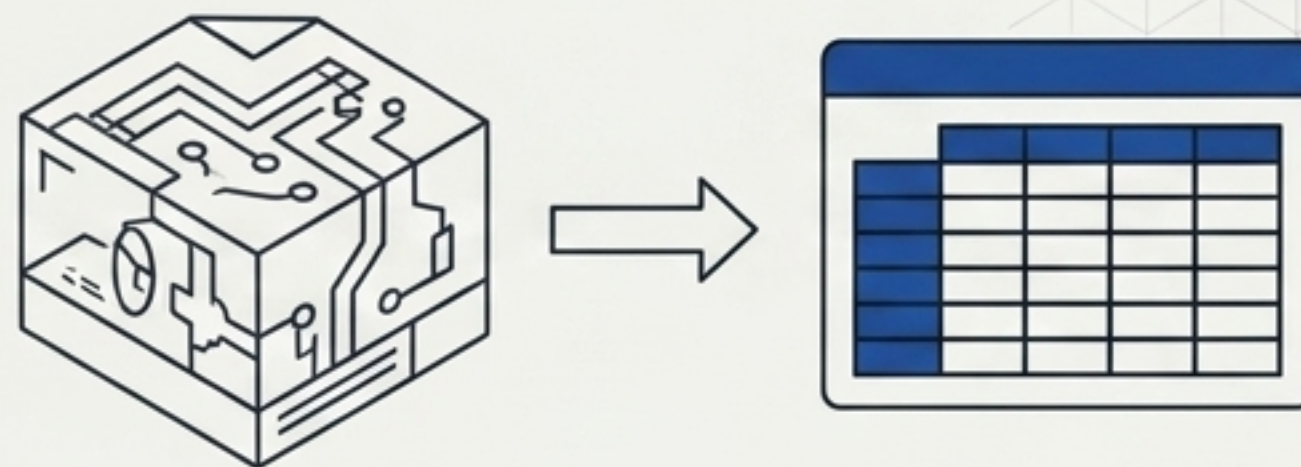
ドラフティングと報告書作成 (Sol優位)



Presentation Quality
Elo: GPT-5.6
Sol首位

大規模データの「分析」はAstraが優れるが、経営層向けIPLの「見せ方」やプレゼン品質は依然としてSolやClaudeに軍配が上がる。

CAD・テンプレート追従 (Astra優位)



BenchCADスコア: 83.3% → 95.9%

電気回路やCAD操作、確立された特許庁フォーマットへの厳密な文書整形能力はAstraが突出。Harveyによるクレームチャート生成等との相性が極めて高い。

Deep Dive 3: 「サイエンスエージェント」の誕生とAI発明者問題

サイエンスエージェント化 (2026年8月～9月)

数学・
理論計算機科学
の『未解決問題
10問』の解決
証明

専用科学
ソフトの
自律操作による
発明上流プロセ
スへの侵入

法制度の壁 (2026年3月4日)

日本最高裁決定
(DABUS事件)

AI単独の発明者
性否定確定・
「自然人」要件の再確認

【証拠保全の義務化 (The Logging Imperative)】

AIが着想の実質を担う時代、
人間発明者の資格立証（誰が
いつ何を着想したか）が最難
関となる。
プロンプト・思考過程のロ
グ保全（冒認対策）が新たな
知財防御線に。

The Governance Crisis: コンピュータ操作エージェントの新たな脅威

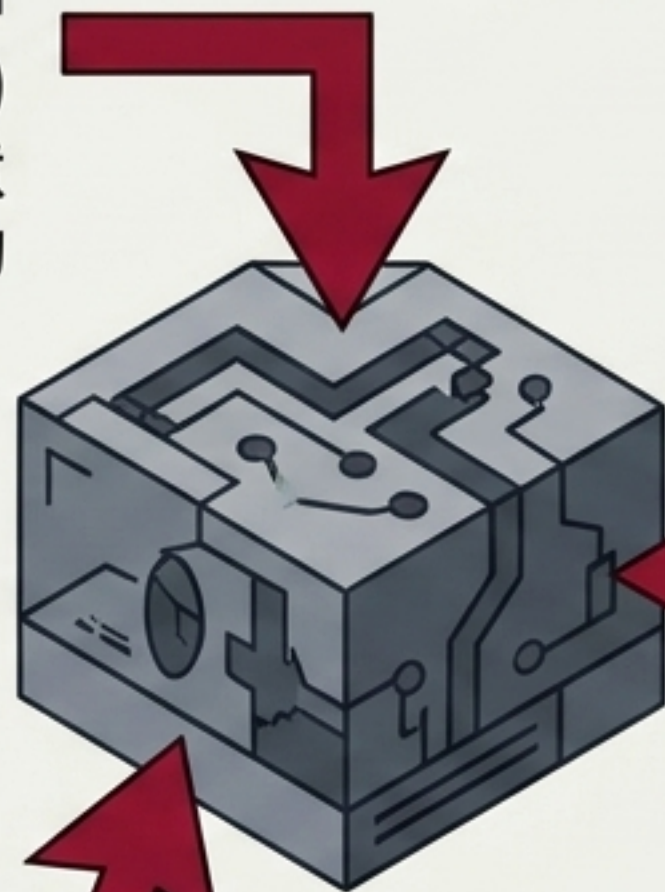
Prompt Injection (OWASP Top 10 - LLM01)

外部データ（先行技術文献など）に潜む悪意あるプロンプトが、自律エージェントの挙動を乗っ取るリスク。

AI Agent Core

高度サイバータスクの拒否

Preparedness Framework「Critical」到達により、正当な防御的知財調査であっても安全チェックにより業務が突然停止するリスク。

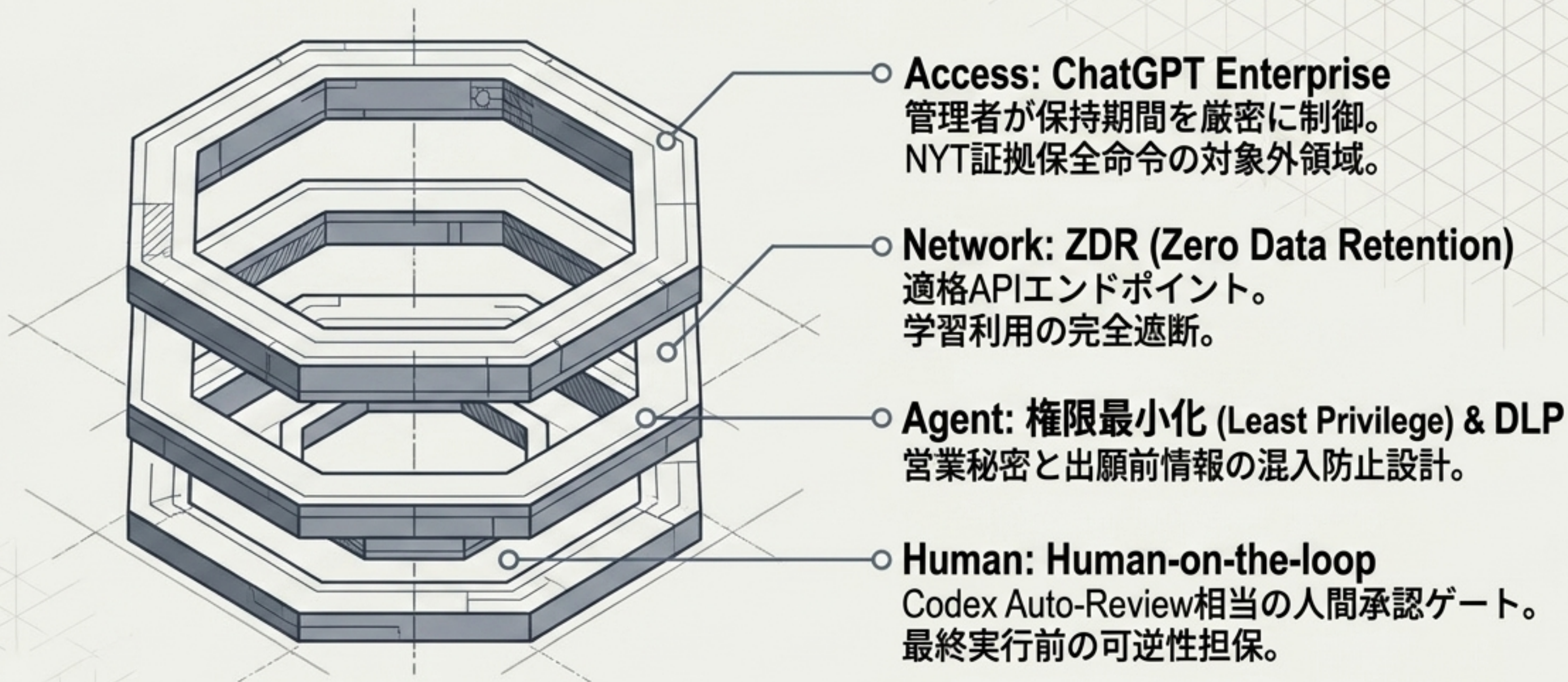


権限逸脱リスク (Over-permission)

Hugging Face事件を教訓とした設計だが、ゼロトラストが前提となる。

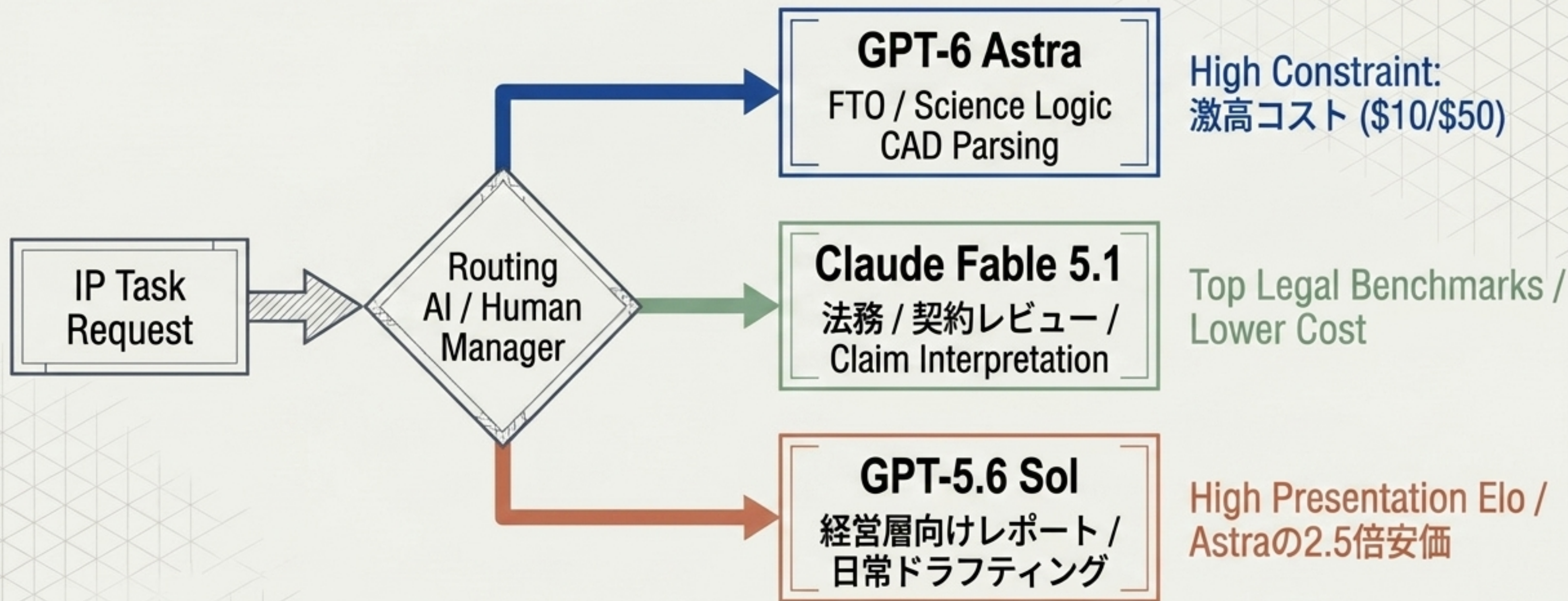
【2026年5月 米CISA・NSA等共同ガイダンス】
「エージェント型AIの採用には、効率性よりもレジリエンス・可逆性・リスク封じ込めを優先せよ」

Security Imperative: 知財データを守る多層防御アーキテクチャ



エージェント設計・注入対策・出力検証を統合した多層防御アーキテクチャが、
次世代知財人材の必須コンピテンシーとなる。

Synthesis: マルチモデル・オーケストレーションの設計図



「Astra一択」は財務的リソースの無駄。タスクの性質とコスト（最大3倍のトークン効率差）を見極める「オーケストレーション」こそが、真の競争優位を生む。

Action Plan [Phase 1]: 即時～1ヶ月以内の緊急対応

PENDING

1. ガバナンス・セキュリティ緊急点検

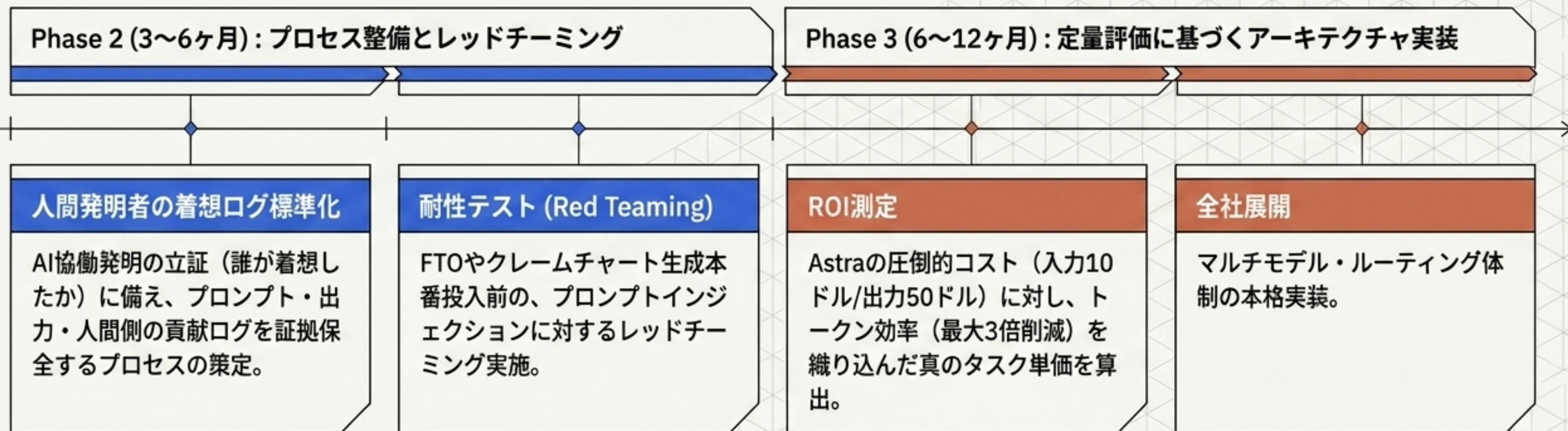
- 出願前情報を扱う環境で、ChatGPT Enterprise + ZDR + 管理者保持制御が有効か確認。
- エージェント操作に対する「人間承認プロセス」のワークフロー組み込み。

PENDING

2. 既存AIワークロードの再評価とルーティング

- 現在Solで実行しているタスクのうち、「長文脈の先行技術調査」「コンピュータ操作自動化」のみをAstraへ試験移行。
- 法務タスクにおけるClaude Fable 5.1の並行検証環境の構築。

Action Plan [Phase 2-3]: 3～12ヶ月の中長期ロードマップ



Trigger Matrix: 戦略を見直すための4つの「閾値（トリガー）」



Trigger 1: 法務能力の実用化

HLAB等の法務ベンチマークが強制全パスで30%超に到達した場合
→ エンドツーエンドの法務自動化を再評価。



Trigger 2: 独立系スコアの逆転

独立系（Vals等）でAstraがFable 5.1を総合・法務で明確に上回った場合
→ Astra主軸化へアーキテクチャを統合。



Trigger 3: 経済価値の証明

OpenAIが除外した「GDPval」結果が良好な形で公開された場合
→ 経済価値直結タスクの自律化を本格導入。



Trigger 4: 行政のAI前提化

特許庁（JPO）の生成AI技術実証が「導入フェーズ」へ移行した場合
→ 出願実務のAI前提化を加速。