

GPT-6時代の知財戦略設計図

SolとLunaがもたらすコスト破壊と安全へのシフト

単一モデル依存から、マルチモデル・パイプライン構築へのパラダイムシフト

Executive Summary

核心的インサイト

2026年9月公開の新モデルの真価は、知能の向上ではない。

Astra比で1/5～1/100となる圧倒的なコスト破壊と、不明な問いに対する安全側へのシフト（回答拒否）にある。

知財実務への示唆

簡潔さの代償として生じる網羅性の低下リスクを直視する。

Luna（大量スクリーニング）、Sol（要件抽出）、Astra（複雑な検証）を適材適所で組み合わせる検証のパイプラインの構築が急務。

GPT-6 エコシステムにおけるモデル・ポジショニング



主要スペック

公称コンテキスト長

両モデル 105万トークン

最大出力

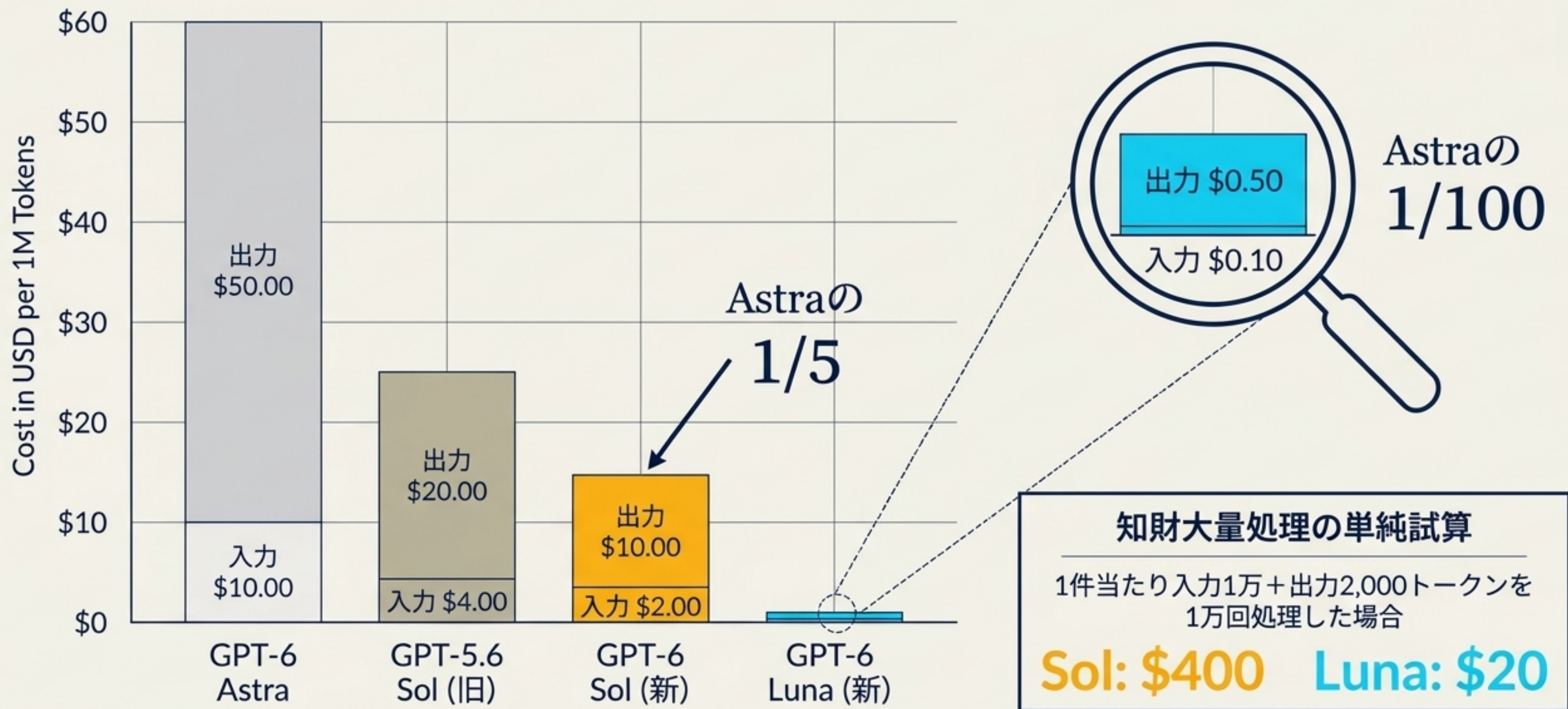
両モデル 12.8万トークン

知識の基準日

Sol: 2026年4月20日

Luna: 2026年5月18日

The Cost Destruction: 圧倒的なコスト破壊のスケール



知能と効率のフロンティア：性能微増と効率激変のトレードオフ

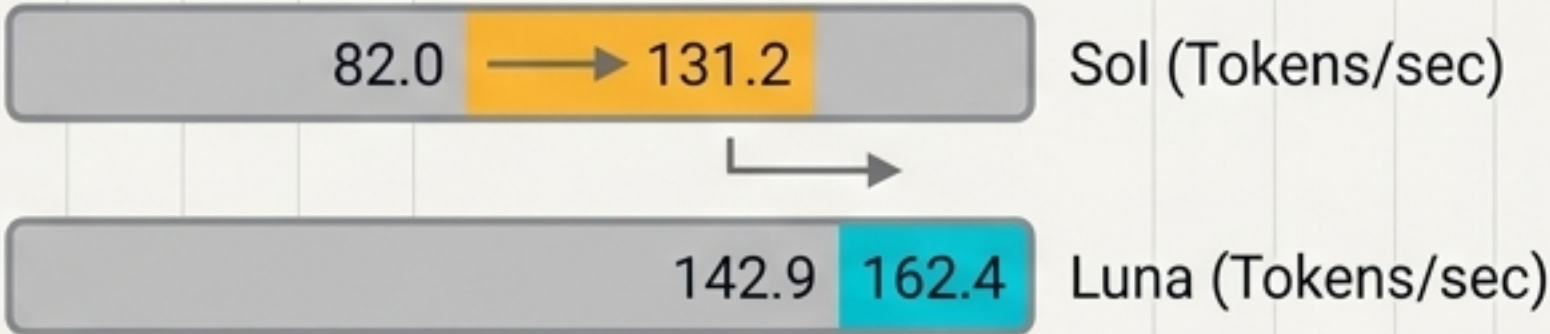
知能の停滞 (Intelligence Index)

Artificial Analysis Intelligence Index v4.3.2

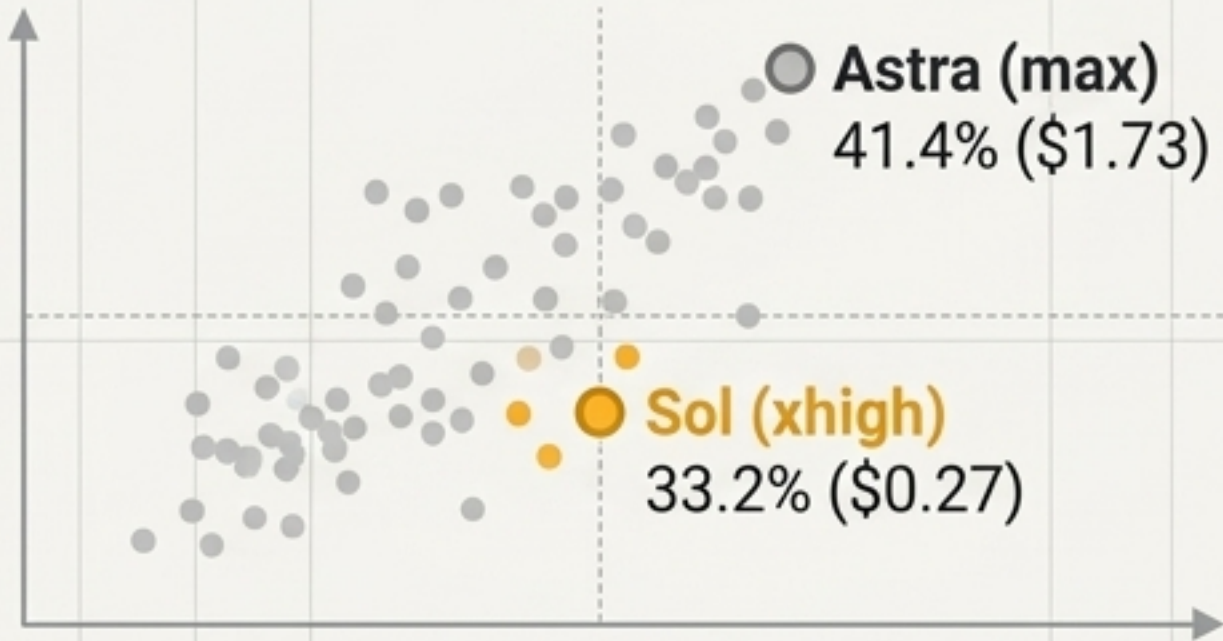
GPT-5.6 Sol	47
GPT-6 Sol	48 (微増)
GPT-5.6 Luna	37
GPT-6 Luna	37 (変化なし)

参考: Claude Opus 5.5 = 58

効率と速度の飛躍 (Efficiency & Speed)

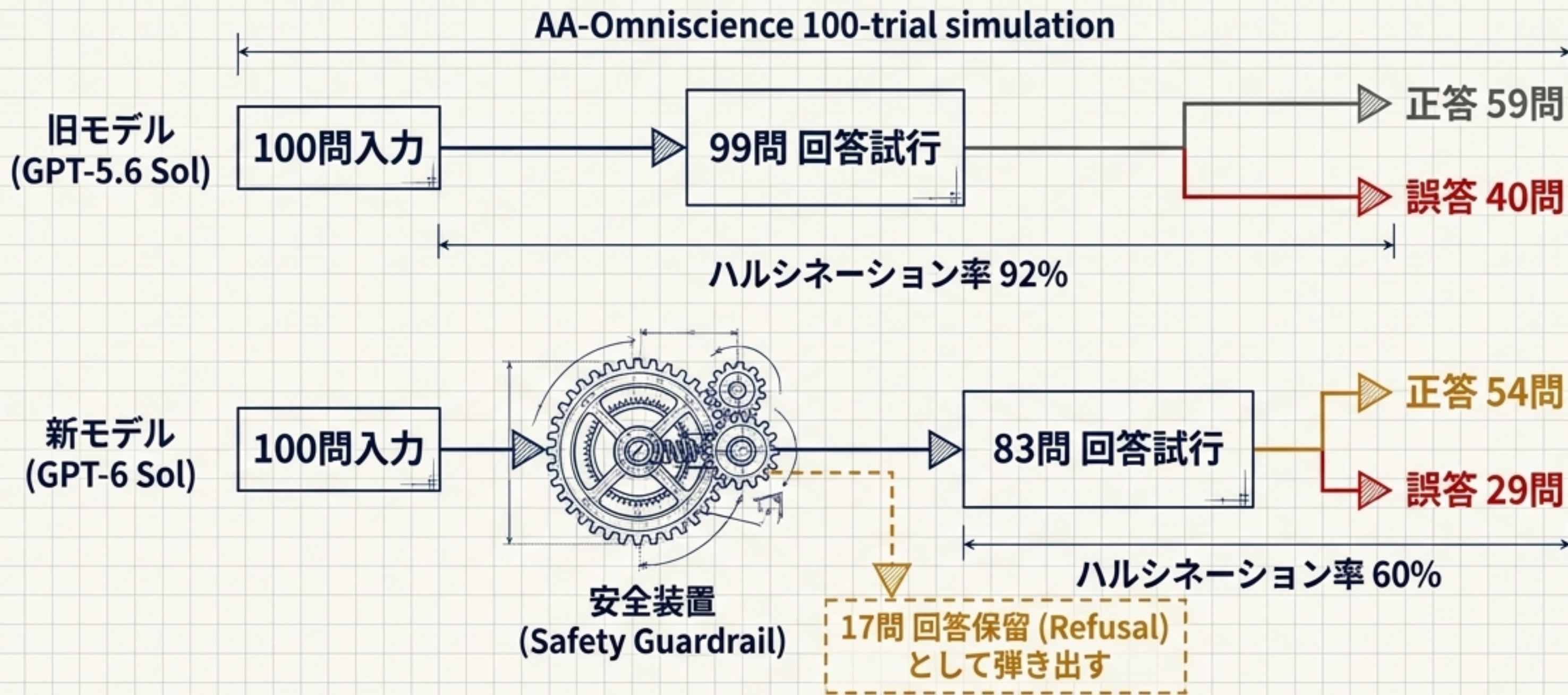


Zapier AutomationBench スコア vs コスト



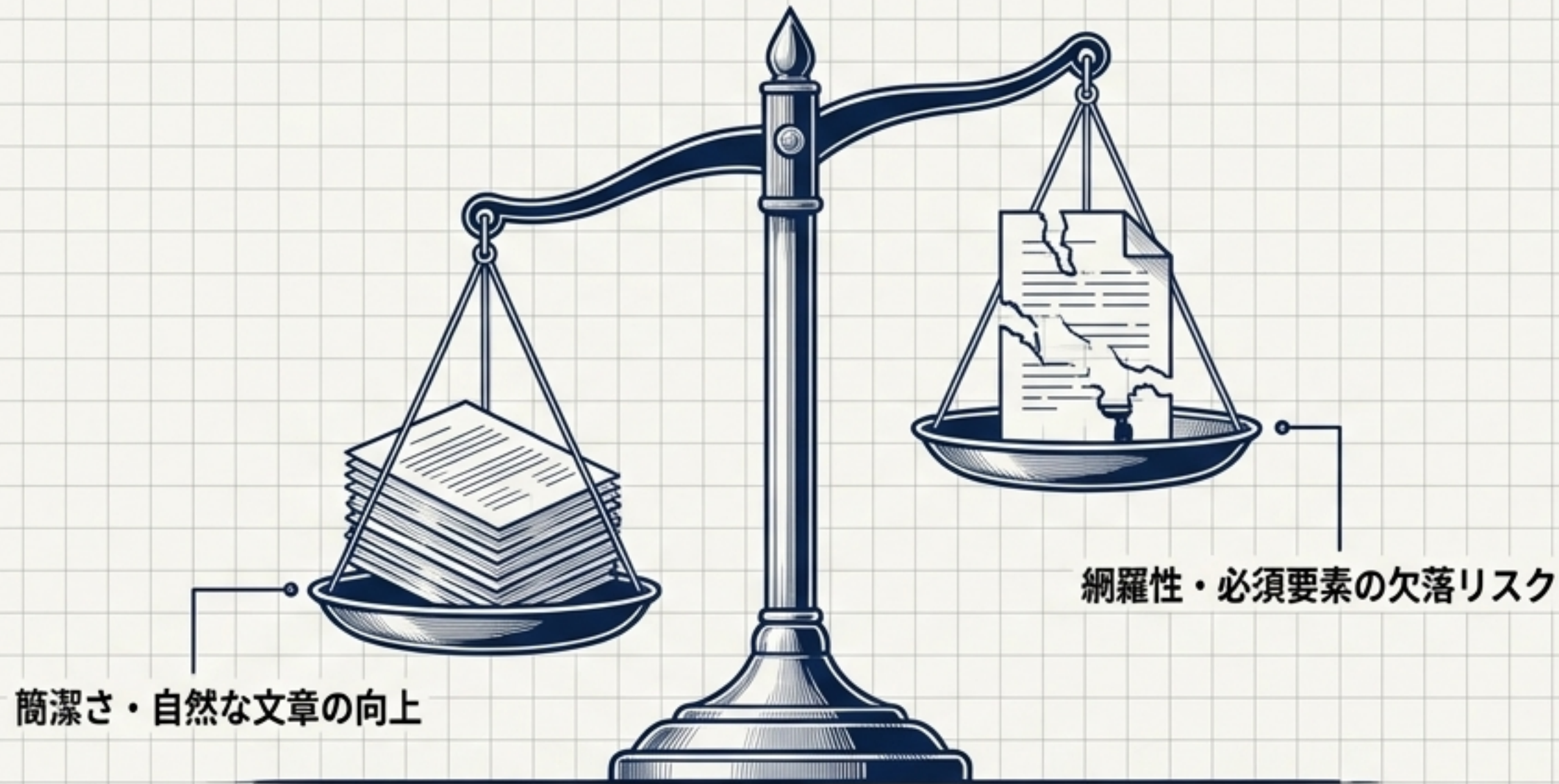
最高設定のmax(32.0%)より、xhigh設定の方が高得点かつ低コストを記録

ハルシネーション減少のパラドックス：回答拒否メカニズム



インサイト：わからない問題に対して、もっともらしい嘘（架空の公報や判例）をつくくらいなら、黙るようになった。
知財実務においては「関連文献なし」と「モデルの回答保留」を明確に区別し、人間へエスカレーションする設計が不可欠である。

出力品質のトレードオフ：「簡潔さ」と「網羅性」の天秤



知識労働の成果物評価
(GDPval-AA v2.1)

↓ Sol: 約 **100 Elo** 低下

↓ Luna: 約 **75 Elo** 低下

知財実務へのクリティカルな警告

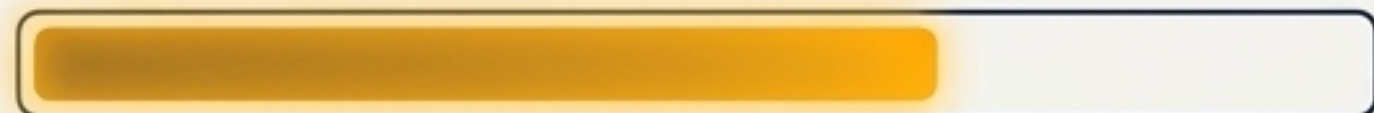
簡潔な出力スタイルは、特許請求項の対比や拒絶理由対応において、採点上必要な要素（構成要件、対比根拠、反対解釈など）の脱落を招く。

対策：対象請求項、相違点、留保事項などを事前に指定し、各項目の充足を厳格にチェックするプロンプト設計が必須となる。

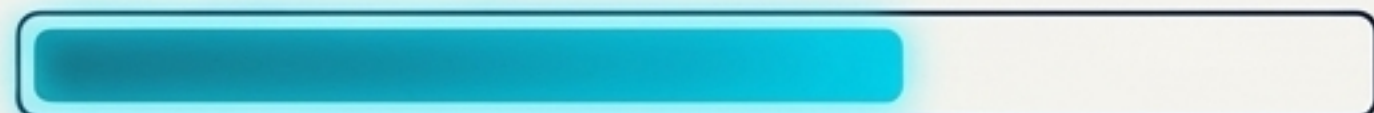
コーディングと自動化：現実的な限界と特化領域

基礎開発能力 (DeepSWE v1.1 コーディング性能)

Sol (max設定) 68.8%



Luna (max設定) 66.6%



Lunaも一定の基礎開発能力を有し、単純な自動化スクリプトには十分適応する。

高度な自律的作業 (Coding Agent Index)



Sol: 55 → 57 (上昇)



Luna: 43 → 41 (低下)

複雑なコード作業や自律型エージェントとしての能力は、Lunaでは後退している。

結論：Lunaの低価格化は圧倒的だが、高度な自動化パイプライン構築まで一律に改善したわけではない。複雑なシステム統合やデータ変換には、依然としてSol以上の推論モデルが必要不可欠である。

プロンプトキャッシュの威力：長文処理コストの劇的圧縮

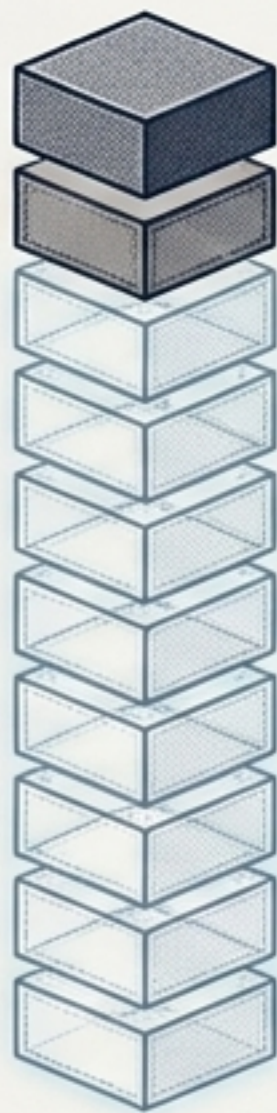
旧方式
(毎回ゼロから読み込み)



全トークンに
標準課金

旧方式
(毎回ゼロから読み込み)

GPT-6
キャッシュ機能



新規トークンの
み通常課金

キャッシュ適用：
読み出し料金
90%割引

GPT-6
キャッシュ機能

コア・メカニズム

- 条件を満たす入力を読み出しが90%オフ
- 30分以内の再利用で有効。推論の強さやツールを変更してもキャッシュを保持。

知財実務への適用例

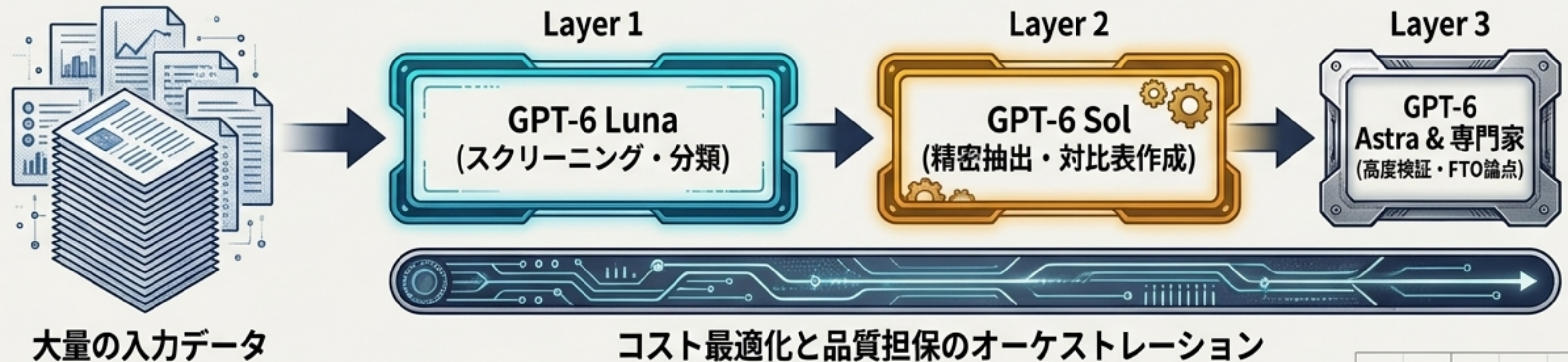
- 共通の評価基準（審査基準や特定技術の定義書）を先頭に固定
- 数千件の特許公報を順番に流し込むバッチ処理において、大幅なコスト破壊を実現。

The Fundamental Shift: 知財業務パラダイムの転換

【旧】 単一モデル依存 (The Single-Model Era)



【新】 マルチモデル・パイプライン (The Multi-Model Pipeline)



Task Allocation Matrix: 知財業務別の役割と検証ポイント

業務タスク (Task)	推奨モデル (Model)	人間が必ず確認すべき死角 (Blindspots)
一次スクリーニング・ 書誌情報抽出	Luna	⚠ 抽出漏れ、関連文献の見落とし、 「判断保留 (Refusal)」の不適切なスキップ
請求項分解・構成要件 対比表下案・技術動向	Sol	⚠ 構成要件の脱落、引用段落の不 正確さ、必須項目の網羅性欠如
複雑なFTO論点・ 拒絶理由対応	Astra	⚠ 解釈の妥当性、法的根拠、最終 的な専門家による総合的判断

Designing the Validation Pipeline: マルチモデル検証アーキテクチャ

フェーズ1:
網羅的フィルタリング

⚠ 判断保留フラグ

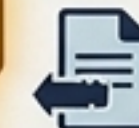
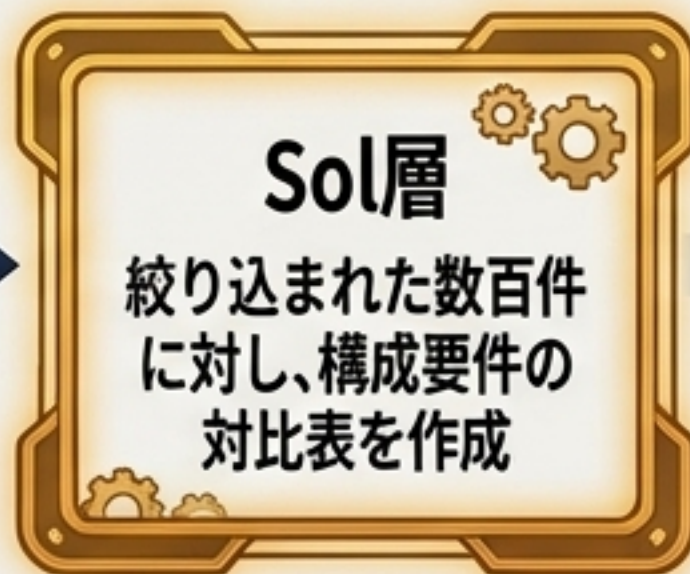
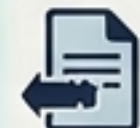
フェーズ3:
高度検証と最終判断

フェーズ2:
精密抽出と下案作成

数万件の
特許公報



数万件の特許公報



結論：大量のLunaによるスクリーニング → Solによる要件抽出 → 不足部分をAstraと専門家が補完。これがGPT-6時代の最適解。

Risks & Safety Guardrails: 導入における致命的な注意点



長文コンテキストの ペナルティ

The >272k Threshold

入力が27.2万トークンを超えると、リクエスト全体に
入力・キャッシュ2倍、出力1.5倍の特別料金が適用
される。文献の分割単位の
最適化が必須。



Lunaの過剰拒否リスク

Over-refusal in Safety

安全性評価での高い拒否率
が、正当な特許調査依頼
までブロックする可能性が
ある。「安全性の改善」と
「業務完遂能力」のトレード
オフを常に監視する仕組み
が必要。



Solのエラー残存

Remaining Severity Issues

内部評価で深刻度3以上の
問題行動は36%減少（66件
→42件）したものの、ゼ
ロではない。
外部利用での事故率を独自
に検証するフェーズが不
可欠。

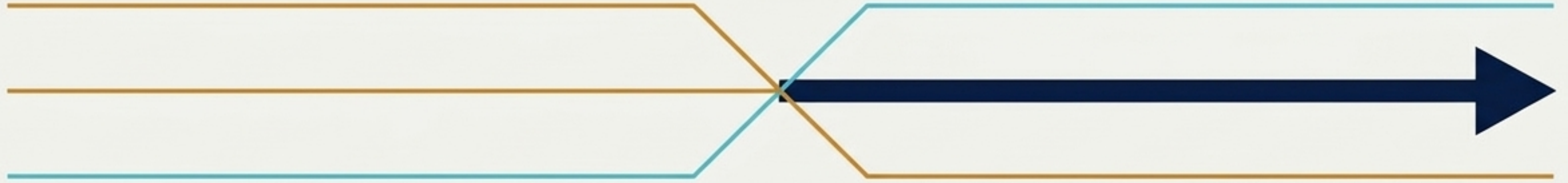
Rollout & Availability: 提供プランと展開状況

プラットフォーム / プラン	提供モデル	ステータス
ChatGPT Work / Codex (Plus, Pro, Business, Enterprise, Edu)	Sol / Luna 両方	✓ 提供済み
無料版 (Free / Go)	Luna のみ (デスクトップアプリ限定)	✓ 提供済み
API	Sol / Luna 両方	✓ 提供済み
GitHub Copilot	Sol (Pro+等) / Luna (Pro含む)	✓ 提供済み

注意喚起アラート

通常のChat版（無料ウェブ等）では発表時点（2026/9/23）で未提供。また、定額プランにおける利用枠（レートリミット）の早期消費に関する指摘も散見されるため、API単価と定額プランの体験は分けて評価すること。

プロンプト・エンジニアリングから、 パイプライン・オーケストレーションへ。



GPT-6 SolとLunaは、知能の魔法ではなく、圧倒的なコスト破壊と安全機構という強力な産業用ツールを我々にもたらした。

これからの知財専門家の役割は、単一のAIに完璧な回答を求めることではない。モデルの特性と死角を熟知し、Lunaの機動力、Solの抽出処理、そしてAstraと人間の高度な判断力を織り交ぜた『絶対的な信頼性を担保する検証パイプライン』を設計・統括することである。