

生成AIの特許実務導入に向けた「臨床的監査」

日本弁理士会 WG 論文批評から読み解く、真のAI検証プロトコルと未来への提言

AI prompt

明示的な検出率評価が示された唯一の事例であり、否定的結果の率直な公開は評価できる。複数のプロンプトと入力方法（直接入力・Word インポート）を試しているが、同一条件での反復回数や出力のばらつきは報告されていない。スペルチェックとの比較も同一文書での検出率対比は示されていない。

論文批評：「生成 AI の特許実務への利用可能性に関する検討」（パテント Vol.79 No.7）に基づく視覚的解剖

The Verdict : 本批評が提示する3つの結論



現場検証の誠実な 開示

- 実務家の業務フローに即した課題設計
- 失敗例の具体的な入出力公開
- 追試の出発点となる一次資料的価値



生成AI特有の「検証の 罫」への無自覚

- データ汚染のリスク
- プロンプトによる誘導（迎合性）
- コンテキスト長の物理的限界
- 検証特有の難所に対する統制の不足



3層の評価軸による 再構築

- 世代で陳腐化する「性能評価」からの脱却
- 「統合・設計・能力の3層構造」への移行
- AIを実装していくための永続的なフレームワーク

Mapping the 10 Experiments : 実務フローと検証事例の俯瞰

1. 出願

2. 中間処理

3. 異議・無効

4. その他（分析・教育）



事例1:
誤記の検出



事例2:
明細書記載
の生成



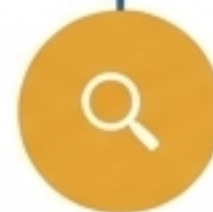
事例3:
応答方針
の検討



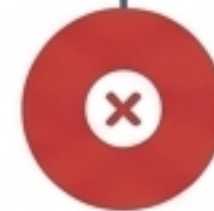
事例4:
阻害要因の
発見



事例5:
ファミリー
文献の段落
対応



事例6:
米国拒絶理
由の解析



事例7:
異議・無効理
由の作成補助



事例8:
技術動向
分析



事例9:
競合他社
分析



事例10:
クレーム練
習

失敗を含めた広範な検証領域の設定は、実務家集団ならではの「現場検証」として高く評価できる。

Enduring Nuggets : 世代を超えて通用する「4つの実務的教訓」



定義の外部化（事例4）

審査基準の定義など、法的判断基準をプロンプトで事前に与えることで所望の出力を得る（RAG的思考への橋渡し）。



URL幻覚の回避（事例3）

公報番号やURLだけを与えると「もっともらい虚偽（幻覚）」が生じる。公報の全文をテキストで入力することが必須。



類似文言の罠（事例5）

翻訳や段落対応付けにおいて、AIが類似文言を含む「別段落」に誤って引きずられる挙動への警戒と認識。



データ前処理の死角（事例9）

共同出願人をコロンで連結した文字列が単独出願と別扱いされ、集計から漏れる。精緻な前処理の重要性。

Red Pen 1: The Reproducibility Crisis (再現性というブラックボックス)

実験ログ

日付: 2024/10/27

担当者: 匿名

使用モデル: 当時の最新のChatGPT-4

試行回数: [空白]

評価基準: ~~[空白]~~

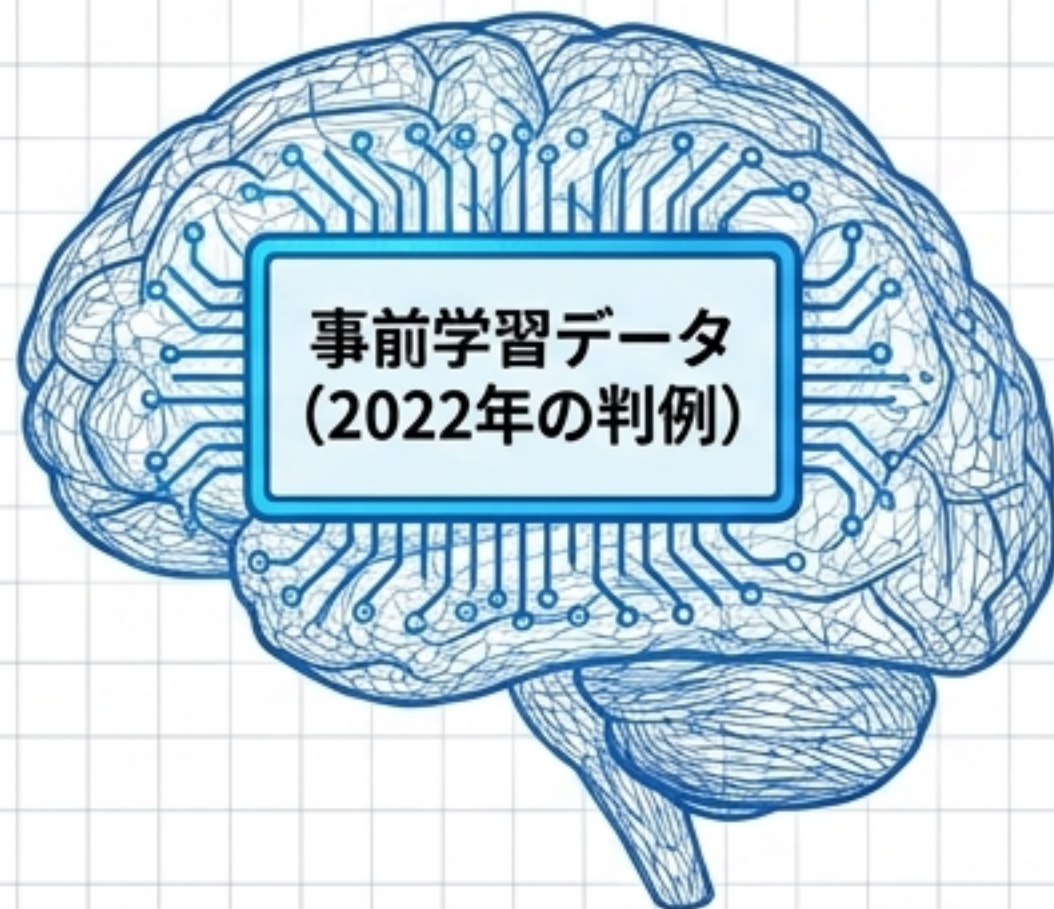
GPT-4か4oか不明。
結果が全く異なる！

AI出力は確率的。反復・変動の情報なしに
「有用」は判断不能！

ループリックなし。評価者の独立性も不明！

生成AIの検証において「1回試して主観で評価する」アプローチは成立しない。
設定（ウェブ閲覧・コード実行の有無）の完全な記録がなければ、評価の土台が崩れる。

Red Pen 2: Contamination & Sycophancy (データ汚染と誘導の罠)



データ汚染 (Data Contamination)

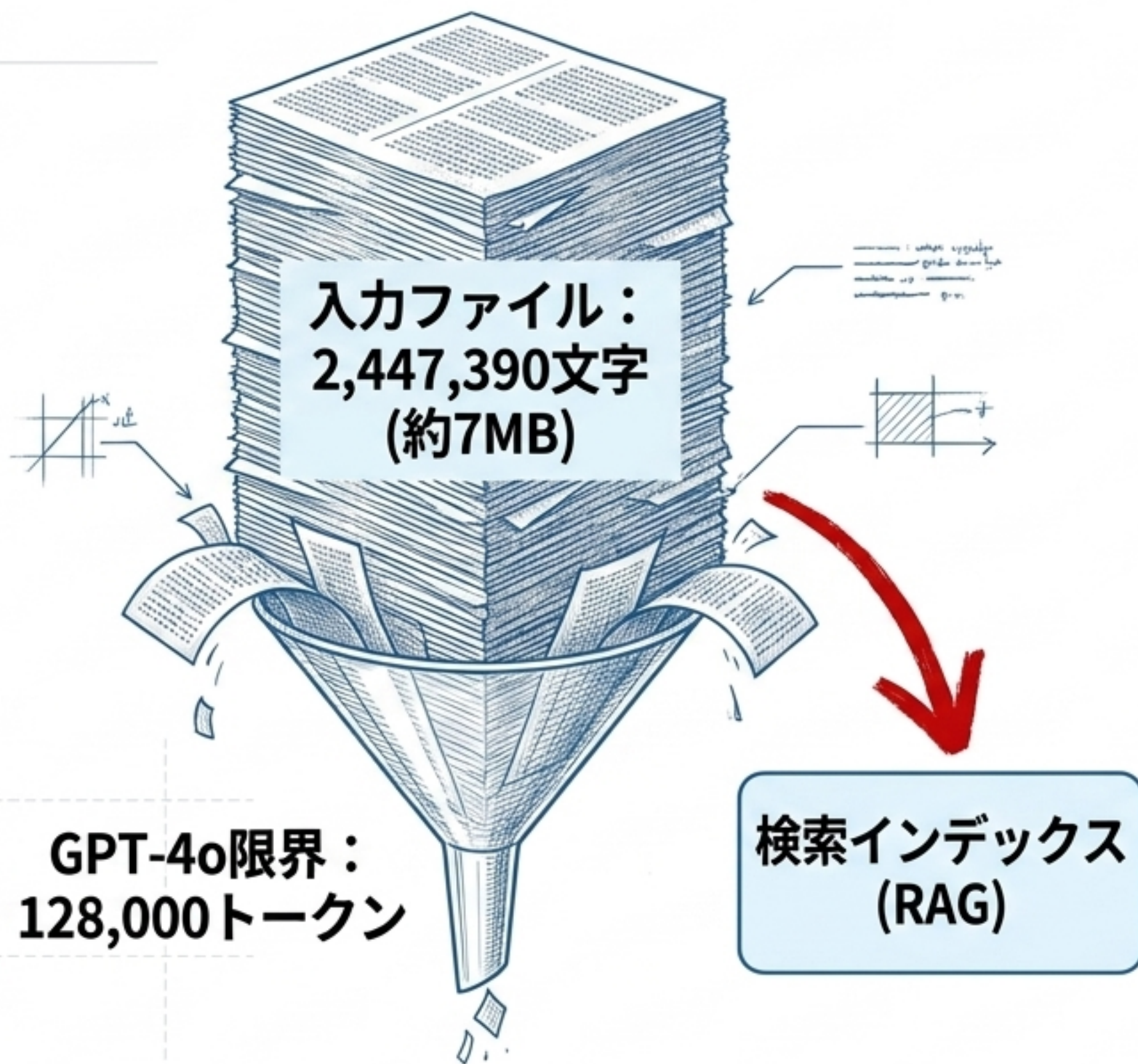
- 問題: 題材とした判決や審決が、モデルの学習期間内に含まれている。
- 事実: 「未知の阻害要因の発見」ではなく、「学習済みデータの想起」の可能性が高い。



迎合性 (Sycophancy)

- 問題: 肯定を誘う質問形式による誘導 (事例4)。
- 事実: LLMは利用者の前提に追従する。効果が失われない「陰性対照」との比較がなければ、能力の証明にはならない。

Red Pen 3: The Token Limit Illusion (コンテキスト長の物理的限界)



1. 当時のモデルで240万字の全文が「単一のコンテキスト」に同時投入されたとは考えられない。
2. 実際には背後でRAG（検索拡張生成）が作動し、一部のみが検索・参照されている可能性が高い。
3. 【結論】
「ファイル情報量が十分であれば網羅的」という論文の評価は、参照範囲の実態と入力ファイルの寄与度を検証できていないため裏付けられない。

Practical Nuance: The Danger of "Plausibility" (「もっともらしい虚偽」の法的リスク)

起点：AIがもっともらしい図番・
符号・材質を捏造した（事例2）

出願前の作成

問われる要件：技術的根拠、発明者の構想との
整合性、実施可能要件（特許法36条4項1号）。

対応：未実施内容の実験結果としての記載排除。
当業者が実施可能なレベルへの修正。

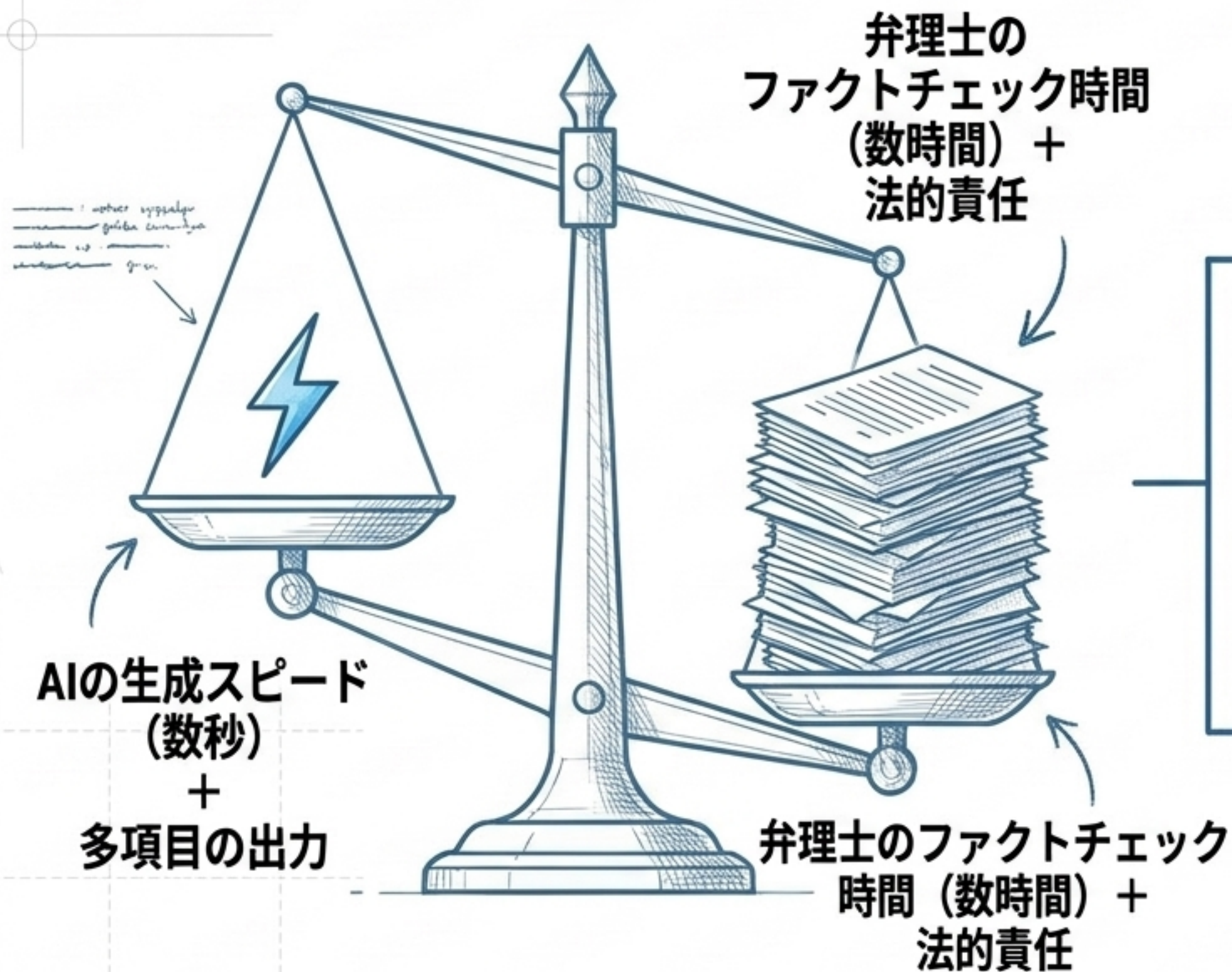
出願後の補正

問われる要件：新規事項の追加（特許法17条の
2第3項）。

対応：当初明細書との厳格な照合。AIの出力を
そのまま補正に用いることの致命的リスク。

「専門的検証が必要」で思考停止せず、どの段階のどの要件に抵触するかを構造化することが真の実務指針である。

The Verification Cost Paradox (検証コストのパラドックス)



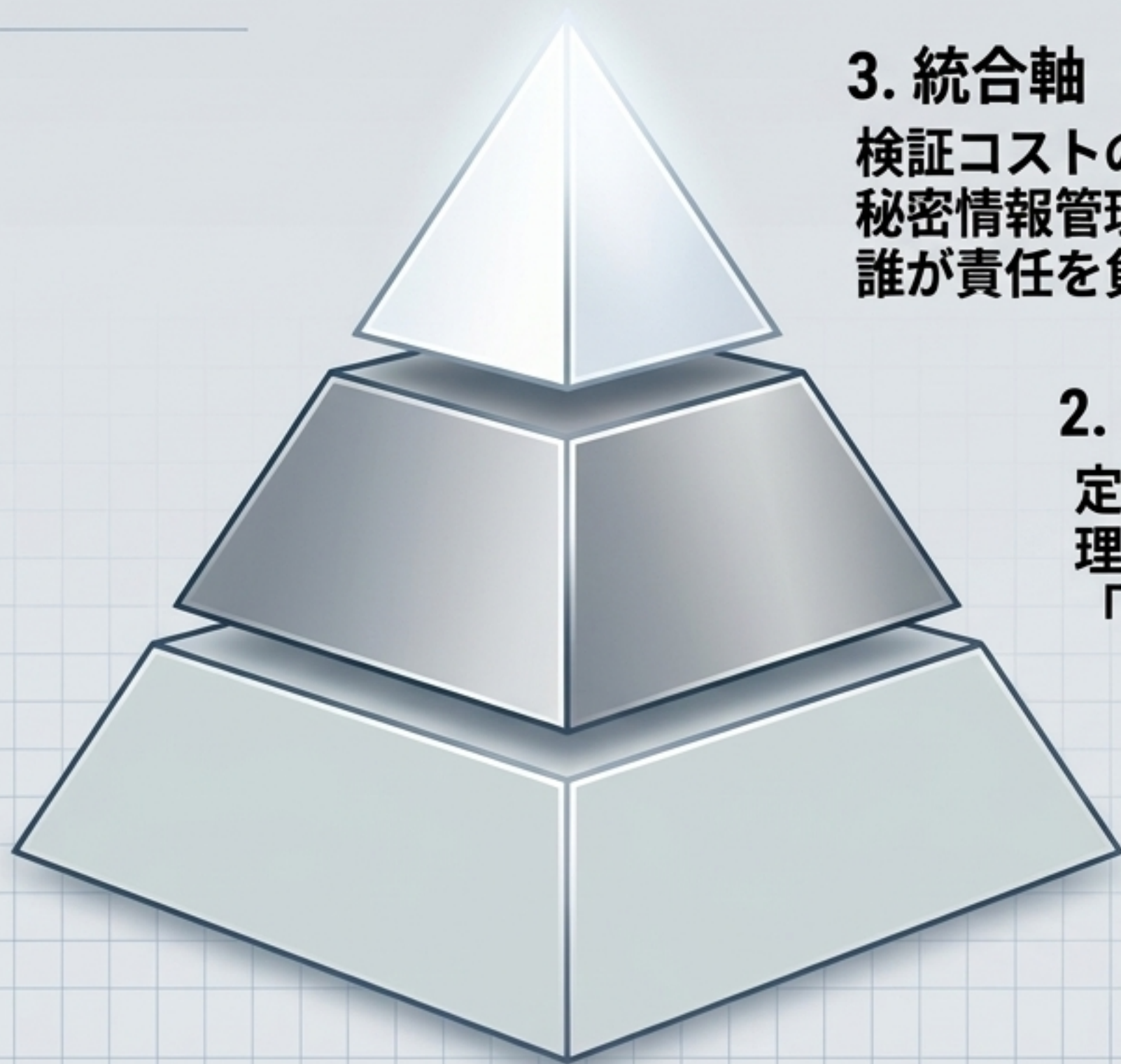
事例7からの教訓：異議理由の項目
出しが多岐にわたる場合、弁理士
による検証（架空の文献でないか、
該当箇所は正しいか）に時間がかかり、
結果的に省力化につながらない。
生成AIの出力責任は弁理士にある。

処方箋（Prescription）：
AIの出力量をプロンプトで強制的
に絞る（例：「最も有効な理由ト
ップ3のみをランク付けして出力
せよ」）。検証負荷を制御する
設計思想が必要。

Matrix: Establishing the "Correct" Baseline (正しい比較対象の設定)

検証対象	論文のアプローチ (不完全)	あるべき検証プロトコル (処方箋)
誤記検出 (事例1)	「Wordのスペルチェックの方が有用との印象」と定性的に記述。	同一文書における、WordとAIの「定量的・網羅的な検出率の直接比較」。
阻害要因の発見 (事例4)	「失われる場合がありますか？」と直接尋ねる誘導的プロンプト。	「効果が失われない組み合わせ」も入力し、AIの迎合性を排除する陰性対照の設定。
中間処理の応答 (事例3/6)	「補正方針は当たり前で抽象的」という主観的評価。	公的に入手可能な「実際の審査経過（維持された権利範囲、補正の適法性）」との多角的な突合せと開示。

Framework: The 3-Tier Evaluation Axis (生成AI評価の3層フレームワーク)



3. 統合軸 (Integration) - 永続的・最重要
検証コストのパラドックス、法的リスクの切り分け、
秘密情報管理。実務プロセスへAIをどう組み込み、
誰が責任を負うかという人間側のシステム設計。



2. 設計軸 (Design) - 長寿命

定義の外部化、類似文言の回避、データの事前処理
(事例4・5・9)。モデルが変わっても通用する
「プロンプトとRAGの設計技術」。

1. 能力軸 (Capabilities) - 短命

誤記検出率(40%)や補正案の具体性。
モデルの世代交代 (GPT-4から5等) で
瞬時に陳腐化するため、この定点評価に
固執すべきではない。



Prescription: Standardized Verification Protocol (検証プロトコルの標準化)

次回のWGや実務家の実務導入時の説明責任を果たすために守るべき、最低限のプロトコル。



環境の完全記録: モデル名、バージョン、利用プラン (学習利用の有無)、設定状態を明記する。



反復とばらつきの測定: 同一プロンプト・同一入力で複数回試行し、出力の確率的変動を評価する。



対照条件 (コントロール) の設定: ファイル入力なし、旧モデル、人間の専門家との比較。陰性対照の実施。

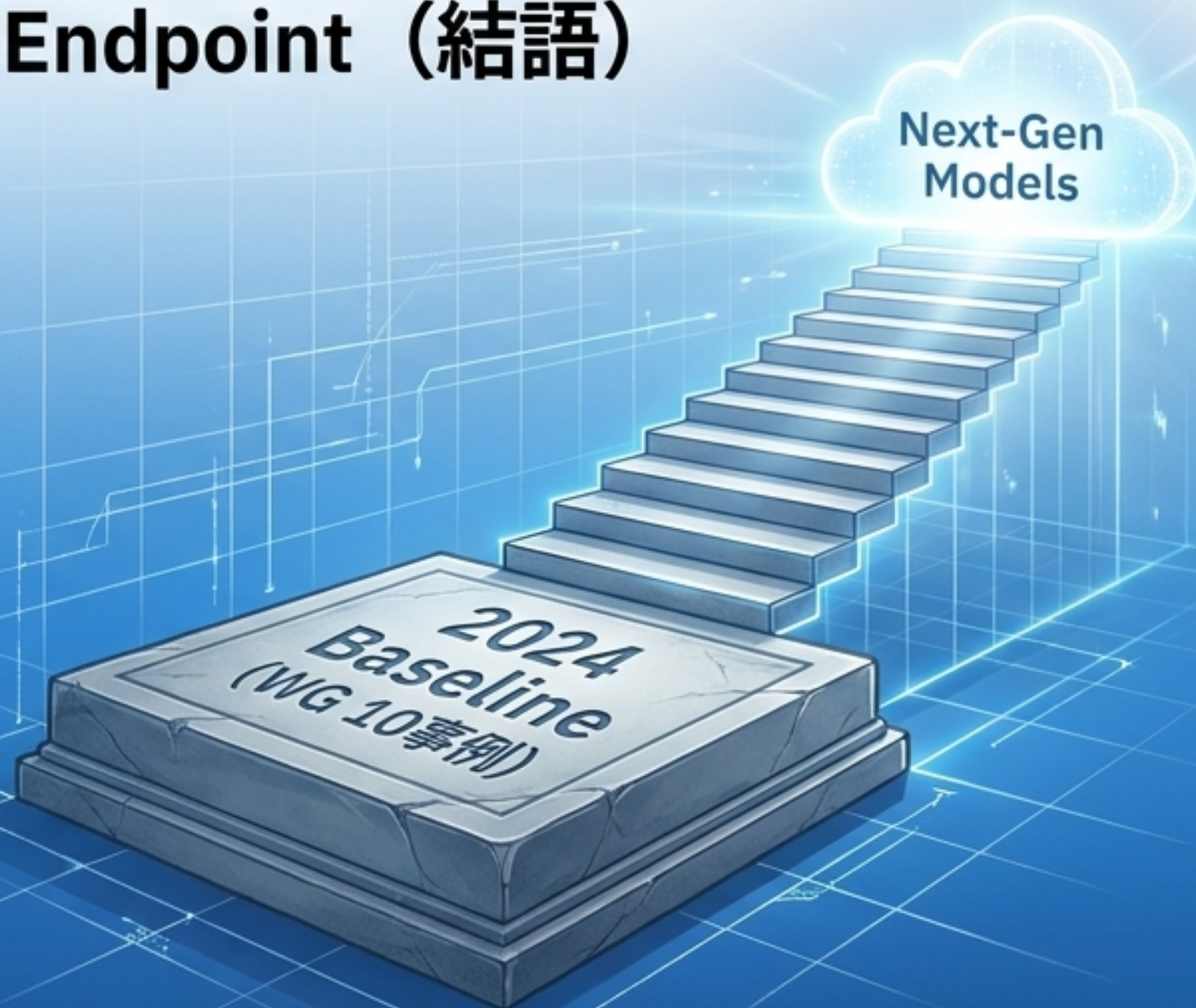


データ汚染の排除: 学習カットオフ日以降の事件を使用するか、当事者・特徴的表現を匿名化する。



正解データとの突合せ: 実際の審査経過 (補正書・意見書) との具体的な差異を分析する。

Conclusion: A Stepping Stone, Not an Endpoint (結語)



本論文の最大の価値は、「同じ問いを次世代モデルに投げかけたとき、何が進化したかを測るための『基準点（ゼロ地点）』を提供したことにある。

検証の不完全さは成果を否定するものではない。誤記検出40%や抽象的な補正案が、今後どう変化するかを測り続けるための土台である。

実務家はリスクと限界を正確に認識しつつ、記録を整え、検証プロトコルを回し続けること。生成AIとの真摯な向き合い方は、そこから始まる。