

INTELLECTUAL PROPERTY / AI GOVERNANCE

評 釈・改訂版

「生成 AI に基づく進歩性判断の モデル化評価」 評釈

法的整合性・実証設計・実務実装の観点から

対象文献

山崎 薫「生成 AI に基づく進歩性判断のモデル化評価」

『パテント』 Vol.79, No.6, 2026, pp.44-54

弁理士・弁護士・企業知財部・特許審査実務に携わる専門家向け

2026 年 8 月 18 日

GPT-5.6 Sol

EXECUTIVE ASSESSMENT

01 総合評価

結論

対象論文は、生成 AI が進歩性を正しく自動判断できることを実証した研究ではない。他方、請求項の構成対応、相違点、効果、根拠資料、結論を分節化し、LLM がどこで破綻し、どこで専門家介入を要するかを具体的に示した探索的研究として価値がある。現段階の最適な位置付けは「自動判断モデル」ではなく、専門家による論点抽出・証拠整理・反対仮説生成を支援する構造化ワークフローである。 [1]

対象論文は、進歩性判断を段階的なプロンプトに翻訳し、バッグと電気化学系という二つの事例で生成 AI の応答を検討する。プロンプトのテンプレート、著者による一部の訂正指示および結果の変化を記載し、構成対応・図面読解・検索の失敗点を示した点は有益である。ただし、実際の全入力、全中間出力、全対話ログ、検索クエリおよび採用 URL は開示されておらず、第三者による完全な追試には限界がある。 [1]

しかし、その「厳格ルール」は、日本法上の容易想到性判断を忠実に操作化したものではない。とりわけ、客観資料に効果の記載が見つからないことから「想定外のベネフィット」を認定し、進歩性肯定へ接続する構造は、検索不成功を法的非予測性へ読み替えるものであり、特許法 29 条 2 項および審査基準の枠組みと整合しない。 [2] [3] [4]

対象論文の実行ステップ 1 は、請求項と公報 A との共通項を、技術常識、論理的帰結および物理的可能性まで用いて広く抽出する段階であるのに対し、実行ステップ 2 は、これらの主観的補完を禁止して差分を抽出する。実行ステップ 5 は、客観的資料に効果の記載がなければ当該効果を「想定外のベネフィット」とし、その存在だけで進歩性を肯定する。さらに、厳格ルール 8 は上位概念とともに例示された下位概念の選択を原則として動機付けありとし、同 11 は複数差分の効果の予測可能性を明示的開示がない限り否定し、同 12 は動機付けを主引用発明の構成に起因する技術的課題の解決に限定する。 [1] (46–48 頁)

評価軸	評価	要旨
問題設定・着眼	A	暗黙知である判断工程を機械可読な手順へ分解した点は有益。
監査可能性	B+	プロンプトのテンプレート、一部の訂正指示・結果を開示。ただし、全ログ・全中間出力・検索記録がなく追試可能性に限界。
日本法との整合性	C	新規性の「記載されているに等しい事項」、動機付け、有利な効果、技術常識の扱いに重要なずれ。
実証設計・再現性	C-	2 事例、誘導的介入、モデル版・反復条件・評価基準の不足により、精度・一貫性は未検証。
専門家支援への可能性	B+	結論装置ではなく、クレームチャート、論点表、検索課題票としてなら実装価値が見込める。

注：A＝高く評価できる、B＝概ね評価できる、C＝重要な留保・修正を要する、D＝現状では妥当性を支持しにくい。＋／－は各段階内の相対的位置。

対象論文の実質的な貢献

- 請求項要素→引用発明との対応→相違点→相違点に起因する効果→根拠資料という連鎖を、出力させる設計として明示した。 [1]
- 人間の最終関与を残し、図面理解や構成対応の誤りを専門家が検証すべきことを、失敗事例から示した。 [1]
- 「根拠を引用できない」状態を出力上のフラグとして露出させる発想は有用である。ただし、その法的帰結は進歩性肯定ではなく、証拠不足・判定保留であるべきである。
- 一回の最終結論ではなく、中間結果をレビュー対象とする設計を提示し、その一部を記載した点は、知財実務における説明責任・再検証可能性と親和的である。

LEGAL ANALYSIS

02 法的フレームワークとの整合性

請求項解釈：明細書を「見ない」ことは厳格さではない

対象論文のプロンプトは、明細書中の技術的定義を参照せず、請求項文言を文字どおりに扱う方向へ強く振れている。特許庁の改訂概要は、令和 8 年 6 月改訂（同年 7 月 1 日以降に行われる審査に利用）による第 I 部第 2 章第 1 節の改訂点として、審査官が本願の特許出願の時または日を必ず確認する旨を明記したことを挙げる。改訂後の現行審査基準は、発明の認定を請求項の記載に基づいて行う一方、請求項に記載された用語の意味を解釈するに当たり、明細書・図面および出願時の技術常識を考慮し、書類全体を精読するものとしている。[7] [9]

最高裁リパーゼ判決は、進歩性判断の前提となる発明の要旨認定は、特段の事情がない限り特許請求の範囲の記載に基づくべきであり、発明の詳細な説明を参酌できるのは、請求項の技術的意義を一義的に明確に理解できない場合や一見して誤記である場合などに限られるとした。その上で、実施例が特定のリパーゼのみを用いていたことから、請求項上限定のない「リパーゼ」を当該リパーゼに限定することを否定した。他方、現行審査基準は、請求項に基づく発明認定を維持しつつ、用語の意義の解釈では明細書・図面および出願時の技術常識を考慮するとする。したがって、用語の意味・定義・技術的文脈を確認することと、実施例または下位概念を請求項に読み込んで発明を限定することとを区別すべきである。LLM には、通常意味、明示定義、広義・狭義の候補、解釈により結論が変わる点を併記させ、専門家承認後に先行技術対比へ進ませるのがよい。[9] [11]

新規性：「物理的に可能」と「記載されているに等しい」は異なる

対象論文の厳格ルール 3 前段は、引用文献の構成による使用態様が物理的に可能であり、かつ通常の使用態様として常識的であることを共通項認定の条件とする一方、「選択的に束ねる」等の動作表現については、構造上物理的に可能であれば共通項とする特則を置く。したがって、対象論文全体が「物理的に可能なら共通項」としているわけではない。ただし、後者の特則を新規性判断へそのまま用いる限り、範囲が広すぎる。刊行物に記載された発明には、明示的な記載事項だけでなく、その記載事項から出願時の技術常識を参酌することにより当業者が導き出せる「記載されているに等しい事項」も含まれるが、単なる物理的可能性だけでは足りない。日本法上は「明示的に記載」「記載されているに等しい事項」「開示なし」「判定不能」に区分するのが適切である。[1] [2]

バッグ事例（対象論文 49-51 頁）でも、「選択的に束ねることができる」ことが請求項の機能的限定を満たすのか、それとも特定の配置・関係・使用状態まで要求するのかを先に解釈しなければならない。図面から構成が視認できても、その図面が当該技術的關係を直接開示しているのか、観察者が事後的に用途を付与しているだけなのかは別問題である。[1] [2]

進歩性：効果探索を容易想到性判断の代替にしない

日本の進歩性判断は、主引用発明から請求項に係る発明へ当業者が容易に到達できたかを、動機付け、設計変更等、先行技術の単なる寄せ集め、阻害要因、有利な効果などを総合して論理付ける作業である。特許法 29 条 2 項の条文は当業者を「その発明の属する技術の分野における通常の知識を有する者」と表現し、審査基準はこれを、出願時の技術常識、通常の技術的手段および通常の創作能力を備える者として具体化している。動機付けには、技術分野の関連性、課題の共通性、作用・機能の共通性、引用発明中の示唆が含まれ得る。[2] [3] [5]

これに対し、対象論文の実行ステップ 5 は、相違点の効果に関する客観資料が見つからない場合に「想定外のベネフィット」を強制し、進歩性肯定へ接続する。検索エンジンで記載を発見できなかったことは、出願時の当業者が効果を予測できなかったことの証明ではない。検索式、コーパス、言語、検索順位、検索時点の限界を法的非予測性に転化しており、無知からの論証に陥る。反対に、似た効果の記載が検索で見つかったことも、それだけで予測可能性を証明しない。公開時点、技術分野、当業者の認識、請求項構成との因果関係および組合せの動機付けを別途検討する必要がある。[1] [2] [4]

判定ロジックの置換案

「効果記載を発見できない」→「現在の証拠集合では予測可能性を判定できない／追加調査が必要」。進歩性を肯定するのは、請求項に係る構成との因果的結び付き、適切な比較対象、出願時技術水準からの非予測性、効果の顕著性・クレーム全域性を検討した後である。[2] [4]

「想定外のベネフィット」より「引用発明と比較した有利な効果」

対象論文の用語は、効果が明細書に記載され、外部資料に同じ文言がなければ、直ちに進歩性を基礎付けるかの印象を与える。審査基準上の「引用発明と比較した有利な効果」は、独立した法定要件ではなく、進歩性を肯定する方向に働く総合評価の一要素である。そのうち、当業者が技術水準から予測できない効果であって、定性的に異質または同質でも際立って優れたものは、「予測できない顕著な効果」として有力な事情となり得る。なお、最高裁令和元年判決は、効果の法的位置付けに関する学説上の対立を明示的には解決していない。単に有利な効果があるにとどまる場合には、これを参酌しても容易想到性が十分に論理付けられれば、進歩性は否定され得る。[2] [4]

最高裁令和元年 8 月 27 日判決は、同じ効果を有する構造の異なる他の化合物が知られていたことだけから、直ちに予測できない顕著な効果を否定した原審の判断を是認せず、当該発明の構成が奏する効果を優先日当時に当業者が予測できたか、または予測できた範囲を超える顕著なものであったかを検討すべきとした。同判決は破棄差戻しであり、進歩性の有無を確定してはいない。LLM には、①請求項構成との因果関係、②引用発明との有利性、③定性的差異または定量的優位、④出願時の予測可能性、⑤請求項の範囲全体で奏するか、⑥当初明細書における根拠を分けて評価させるべきである。[4] [6]

動機付け・選択肢・複数相違点

厳格ルール 11 が複数差分について効果の明示的記載を過度に要求し、同 12 が動機付けを主引用発明の構成から生じ得る技術的課題の解決に限定する点には問題がある。容易想到性の論理付けは、引用文献に明記された課題だけでなく、当業者に自明な課題、設計変更、最適化、均等物置換などから構成され得る。また、相違点間の機能的・技術的関連性を確認し、関連する場合は一体的に、独立する場合は個別にも検討した上で、発明全体を評価すべきである。[2] [5] [12]

現行審査基準は以前から、当業者に自明な課題または容易に着想し得る課題も考慮し、各要素の有無だけでなく程度の差異も踏まえて総合評価するものとしている。令和 8 年 6 月改訂では、そのうち阻害要因等に関する記載が明確化された。したがって、課題の起点を主引用発明の構成に限定し、効果の明示的記載を機械的な分岐条件とする対象論文の設計は、現行基準との関係でも狭すぎる。なお、この改訂は誌面に記載された対象論文の日付（2025 年 10 月 30 日）後であるため、本段落は現行基準からの評価である。[1] [2] [5] [7]

また、上位概念と下位概念または例が併記されていれば、その選択が当然に公知・容易であるとするルール 8 も自動化できない。ピリミジン誘導体事件が扱ったように、化合物を表す式が非常に多数の選択肢を有する場合には、特定の選択肢を積極的または優先的に選択すべき事情があるかが重要であり、選択肢の列挙だけから具体的な組合せを抽出してはならない。[5]

技術常識と証拠時点

前述の当業者モデルを前提とすれば、実行ステップ 2 以降で技術常識の利用を全面的に禁ずる設計は、特許法 29 条 2 項および審査基準の枠組みと整合しない。他方、LLM の自由な常識補完は幻覚・時点混同を招く。解決策は全面禁止ではなく、技術分野、基準日、根拠資料、周知性の程度を記録し、証拠強度を区別した限定利用である。[2] [3]

さらに、原論文のフレームワークでは、AI が公開時期不明のウェブ資料を用いて進歩性否定へ転じ、その直後に著者が公開時期不明と記している。特許法 29 条 1 項・2 項は出願前の先行技術を基準とし、インターネット情報については公衆利用可能となった時点、改変可能性、内容の真正性を確認する必要がある。基準日前の版を確定できない資料は、先行技術としてではなく、追加調査の手掛かりにとどめるべきである。[1] [3] [10]

主要ルール別の修正方向

対象	現行設計の問題	推奨する出力・判断
請求項解釈	明細書定義を一律に排除	通常意味・明示定義・解釈分岐を提示し、実施例限定の読みみとは区別する。[9]
新規性	基本原則と動作表現の特則が併存し、後者では物理的可能性だけで共通項化	明示的に記載／記載されているに等しい／開示なし／判定不能の 4 区分とし、使用態様の付与を分離する。[1] [2]

対象	現行設計の問題	推奨する出力・判断
効果不発見	直ちに想定外・進歩性あり	証拠不足・追加調査。非予測性と検索性能を別評価する。 [2] [4]
動機付け	主引用発明の構成から生じ得る課題の解決へ限定	分野、課題、作用・機能、示唆、設計変更、阻害要因を双方向に検討する。 [2] [5]
選択発明	上位／下位概念の併記で選択容易	選択肢の数・具体性、積極的選択理由、組合せの示唆を検討する。 [5]
ウェブ資料	公開時期不明でも結論に使用	基準日前公開・版・真正性を確認できない限り、先行技術として採用しない。 [3] [10]

RESEARCH DESIGN

03 実証方法の評価

検証されているのは「能力」より「条件付き実行可能性」

対象論文は、同一の厳格プロンプトを与えれば一定の一貫性をもって進歩性判断が可能であるとの仮説を掲げるが、「一貫性」の定義、測定単位、閾値、正解基準が示されていない。プロンプト遵守、同一入力への反復安定性、法的正確性、審査結果の予測性は別の構成概念であり、切り分けて測定する必要がある。 [1] [8]

しかも、効果資料がない場合に進歩性肯定を強制する規則が結論の一部をハードコードしているため、観察される肯定結論は生成 AI の法的判断能力というより、規則追従能力を反映する。現状で実証されているのは、「一定の入力条件下で LLM が所定の出力様式を部分的に実行できること」と、「構成対応、図面理解、検索には人間の補正が必要なこと」である。 [1]

2 事例・対話介入・停止規則

二事例は異なる失敗モードを示す。バグ事例では、構成対応および図面理解の誤りが人間の教示によって訂正された後、異なる試行で取得資料が変わり、結論が進歩性肯定から否定へ反転した。AI はその理由を「検索戦略を変えた」と説明したが、検索クエリやログがないため原因は検証できない。これに対し、電気化学的なイオン回収システム事例では、著者が引用発明の流路構成による濃度低下を説明しても、AI は請求項所定の第 3 流路が明示されていないとして進歩性肯定を維持した。前者は画像理解および検索再現性の問題を、後者は機能的帰結と物理的構成の対応認定の硬直性を示す。 [1] (49–53 頁)

事例数は二つにとどまり、選定基準、技術分野の代表性、失敗を含む全試行数が不明である。同一会話内で構成関係を示唆する訂正を行えば入力条件が変わり、期待される結論の情報が混入し得る。回収システム事例では、著者は拒絶理由通知を意識しつつ濃度低下を生む構造の技術的説明を入力しているが、拒絶理由通知の事実自体を AI に伝えたとは本文から確認できない。補正後の出力を AI 単独の性能として数えることはできない。 [1]

また、期待する対応が得られるまで訂正を重ねる手続は、事後的停止と確認バイアスを招く。一方、これらの対話は、人間-AI 協働系における介入負担や、どの種類の補正が必要かを観察する資料としては有用である。AI 単独条件と協働条件を分け、訂正文を中立・定型化し、最大ターン数と停止規則を事前に固定すべきである。 [8]

モデル特定・検索・再現性

「ChatGPT Pro」は契約プラン名であり、モデルの特定にはならない。「Google Gemini」も製品群名である。対象論文は二製品のどちらが各事例・各出力を生成したかも区別していない。第三者が追試するためには、モデル ID またはスナップショット、実行日時、検索・画像解析機能、システム指示、生成パラメータ、会話履歴のない新規セッションか否か、入力ファイルのハッシュ、各出力と製品・モデルの対応、全対話ログを保存する必要がある。クラウド推論では、ローカル PC の CPU や RAM より、これらの条件の方が再現性に直結する。 [1] [8]

検索については、クエリ、言語、地域、対象コーパス、上位何件を評価したか、採用 URL、引用箇所、公開日確認が記録されていない。本文上、AI は「検索戦略を変えた」と説明しているが、その戦略・クエリ・ログは非開示である。したがって、異なる資料の取得と結論反転の因果を確認できず、検索性能と固定証拠パッケージに基づく法的推論性能を別モジュールとして評価する必要がある。 [1] [8]

正解ラベルと評価指標

イオン回収システム事例では、著者は、現実の拒絶理由通知において新規性が否定されていたことを AI の出力を評価する際の実務的な参照結果としている。しかし、同通知を事前に実験上の参照正解ラベルとして定義したとの記載はなく、出願番号、通知日、対象請求項版、通知書本文も示されていないため、独立した一次資料照合はできない。拒絶理由通知自体も特定時点の請求項に対する審査官の暫定的見解であって、確定的な法的判断ではない。したがって、この事例は AI の正誤を確定する検証というより、既存の審査判断と AI の構成対応が一致しなかった事例として位置付けるのが適切である。 [1]

評価は最終三値だけでなく、次の層に分けるべきである。

- 5. 構成要件対応：評価単位（要素／下位要素）、部分一致、引用箇所との証拠対応を事前定義し、micro 平均の precision・recall・F1 と案件単位 macro 平均を示す。
- 6. 相違点抽出・効果帰属：段階別 precision・recall・F1 に加え、差分→効果リンク単位の precision・recall・F1 を測る。
- 7. 引用資料：引用存在率、主張支持率、基準日前公開率、真正性確認率を分けて測る。
- 8. 結論：新規性段階と、新規性あり案件に対する進歩性段階を階層別に評価し、混同行列、クラス別 F1、macro-F1、accuracy、多クラス MCC および案件単位 bootstrap 信頼区間を示す。
- 9. 反復安定性・評価者一致：同一条件での生一致率と出力分布を示し、2 名なら Cohen の κ 、3 名以上なら Fleiss の κ または Krippendorff の α 、順番付きの評価なら重み付き係数を用い、クラス分布も併記する。
- 10. 協働負担：介入ターン数、所要時間、再作業、コスト。
- 11. 保留性能：証拠不足例に対する適切な保留率、不当な断定率、coverage および selective risk。

REDESIGN

04 知財実務に耐える再設計

最終結論ではなく「論点別の証拠状態」を出す

専門家向けシステムとしては、「進歩性あり／なし」を一足飛びに出力させるより、各判断要素の証拠状態と不確実性を構造化する方が安全である。主判定は「拒絶論理を構成可能」「現証拠では構成困難」「保留」の三つとし、「解釈分岐」「基準日・証拠適格性未確認」「追加調査要」を独立フラグとして付すべきである。

推奨ワークフロー

段階	AI に行わせる処理	専門家の承認点
0 条件固定	法域、優先日・出願日、請求項版、証拠コーパス、モデル版を記録。	秘密管理、利用環境、対象資料の確定。
1 請求項解釈	通常意味、明示定義、曖昧語、広狭二案、構成要件表を作成。	解釈と要素分割を承認。 [9]
2 証拠適格性	公開日、版、出所、ファミリー、引用箇所を抽出。	基準日前資料か、真正性は十分か。 [3] [10]
3 新規性	各要素を明示的に記載／記載されているに等しい／開示なし／判定不能で対応付け。	一文献性、「記載されているに等しい事項」、図面の読取り。 [2]
4 進歩性	主引用＋副引用・技術常識、動機付け、設計変更、阻害要因を両論併記。	最も強い拒絶側・出願人側の論理。 [2] [5]
5 有利な効果	因果関係、比較対象、顕著性、非予測性、全範囲性、当初根拠を点検。	効果の法的関連性と証拠強度。 [2] [4]
6 統合出力	論理チェーン、反証、欠落証拠、追加検索クエリ、証拠強度、不確実性の原因、未確認事項を提示。	結論、補正・主張方針、追加調査の要否。

検証プロトコルの最低要件

- 12.研究質問を分離する。構造化プロンプトの遵守率、同一条件での反復安定性、専門家ラベルとの正確性、人間-AI 協働の時間・品質を別仮説とする。
- 13.評価をモジュール化する。A＝構成対応、B＝検索、C＝固定証拠パッケージによる法的推論、D＝人間-AI の一連の工程（end-to-end）とし、検索失敗と推論失敗を分ける。
- 14.機械・電気・化学・バイオ・ソフトウェア等の技術分野を層化する。探索的試験で誤り率、出力分散および専門家間一致度を推定し、その結果を用いて、確認的試験の標本数を、事前に定めた主要評価項目、想定効果量、有意水準、検出力および多重性調整に基づいて決定する。同一事例の反復出力は独立事例として数えず、公開事件については学習データ記憶の影響を抑える。
- 15.3 名以上の知財専門家が、モデル出力を伏せ、可能な範囲で実事件の最終結果・識別情報もマスクした盲検方式で、構成対応、相違点、動機付け、効果、結論を独立採点する。合議または第三者裁定で参照正解ラベルを形成し、合議前の評価者間一致も保存する。
- 16.比較基準条件、対象論文のプロンプト、法的修正版、検索の有無、標準化介入の有無を要因計画で比較し、各条件を会話履歴のない新規セッションで反復する。
- 17.プロンプト、許容訂正、最大ターン、検索の締切時点、除外基準、主要評価項目である macro-F1、副次評価項目および統計分析を事前登録し、全ログと検索結果を保存する。〔8〕
- 18.実務効果は別試験で測る。専門家単独条件と専門家+AI 条件を参加者内無作為化・カウンターバランスで比較し、条件間で重複しない同等難度案件を割り当て、提示順・前条件の影響を管理する。時間、誤り、見落とし、参加専門家の確信度、再作業を測定し、参加者・案件の random effects を考慮して初めて「時短・コスト削減」を論じる。

実験デザインの例

要因	水準例	モジュール別評価指標例
モデル	固定モデル版 A / B	反復間一致率、段階別 F1
プロンプト	比較基準 / 対象論文 / 法的修正版	ルール遵守率、法的誤り率
証拠	外部検索なし / 固定コーパス / オンライン検索	固定コーパス：recall@k・nDCG@k（参照関連文献集合を事前作成） ／オンライン：precision@k・success@k・引用存在率・該当率・基準 日前公開率／固定証拠：推論精度
入力	テキストのみ / PDF のみ / 両方	図面・構成対応の改善幅
介入	1 ターン完結 / 中立定型介入	精度、介入回数、作業時間

PRACTICE IMPLICATIONS

05 実務上の使用場面とガバナンス

適する用途

- 19.初期クレームチャートと相違点候補の作成。
- 20.拒絶理由通知に対する論点一覧、引用箇所一覧、追加調査課題の作成。
- 21.出願人側・審査官側の双方の論理を生成する反対論生成型レビュー（red-team review）。
- 22.発明発掘時の効果仮説と、データ・比較例の不足箇所の洗い出し。
- 23.若手向け研修での、判断工程と証拠状態の可視化。

避けるべき用途

- 24.公開時期不明のウェブ情報を混在させたまま、最終的な進歩性結論を自動確定すること。
- 25.未公開発明・顧客情報を、契約、保存期間、学習利用、アクセス権が不明な環境へ投入すること。

26.引用箇所、検索条件、モデル版、会話履歴を保存せず、後から再現できない結論を意見書・鑑定へ転用すること。

27.AI の自信表現を、法的確信度または証拠の十分性と同一視すること。

ガバナンス上の最低ライン

知財実務への導入では、案件ごとに利用可能な環境、投入可能な情報、保存期間、学習利用の有無、アクセスログ、外部送信、出力レビュー責任者を定める必要がある。モデルの更新や検索結果の変動を前提に、入力・出力・参照証拠・人間の修正履歴を追跡できる状態を維持すべきである。生成 AI の評価とリスク管理についても、単発の成功例ではなく、反復試験、文書化、利用文脈に応じた継続評価が求められる。〔8〕

CONCLUSION

06 結語

対象論文は、進歩性の「自動判断モデル」を実証した研究ではない。しかし、判断工程を機械可読なチェックリストへ翻訳し、LLM の誤対応、図面理解の限界、検索依存性、専門家介入の必要性を具体的に露呈した点に、探索的研究としての本当の価値がある。〔1〕

今後の課題は、進歩性を効果記載の有無へ還元するのではなく、請求項解釈、引用発明の認定、動機付け、設計変更、阻害要因、有利な効果、証拠時点を、日本法の枠組みに沿って別々に評価し、最後に統合することである。結論の二値化よりも、根拠の強度、不確実性、追加調査の要否を出力させる方が、専門家の判断を強化し、誤用を防ぐ。〔2〕〔3〕〔4〕〔5〕

防御可能な結論表現

二つの著者介入型事例では、構造化プロンプトにより中間結果を分節して出力する設計と、その一部の結果が示された。一方、構成対応・図面理解・検索には顕著な変動があり、法的正確性と反復一貫性は未検証である。現段階では、専門家用チェックリスト／仮説生成・証拠整理支援としての可能性を示すにとどまる。

REFERENCES

07 参考文献

本文中の〔番号〕は以下の文献に対応する。URL の最終閲覧日は、特に記載のない限り 2026 年 8 月 18 日である。

- [1] 山崎 薫「生成 AI に基づく進歩性判断のモデル化評価」『パテント』Vol.79, No.6, 2026, pp.44–54。論文 PDF : <https://jpaa-patent.info/patent/ViewPdf/4836>。掲載号案内 : <https://jpaa-patent.info/patent?month%5Bmonth%5D=06&year%5Byear%5D=2026>。
- [2] 特許庁『特許・実用新案審査基準』第 III 部第 2 章「新規性・進歩性」（令和 8 年 6 月改訂後の現行版）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/tukujitu_kijun/ht/03_0200.html。
- [3] 特許法（昭和 34 年法律第 121 号）第 29 条第 1 項・第 2 項。e-Gov 法令検索。<https://laws.e-gov.go.jp/law/334AC0000000121>。
- [4] 最高裁判所第三小法廷令和元年 8 月 27 日判決・平成 30 年（行ヒ）第 69 号、裁判集民事 262 号 51 頁（ヒト結膜肥満細胞安定化剤事件最高裁判決）。裁判例詳細 : <https://www.courts.go.jp/hanrei/88888/detail7/index.html>。判決全文 : <https://www.courts.go.jp/assets/hanrei/hanrei-pdf-88888.pdf>。
- [5] 知的財産高等裁判所特別部平成 30 年 4 月 13 日判決、平成 28 年（行ケ）第 10182 号・第 10184 号（ピリミジン誘導体事件）。
https://www.courts.go.jp/ip/vc-files/ip/file/zen_28gk10182_10184.pdf。
- [6] 特許庁『特許・実用新案審査ハンドブック』第 III 部第 2 章 3202「ヒト結膜肥満細胞安定化剤事件最高裁判決」（現行版）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/handbook_shinsa/document/index/03.pdf#page=7。
- [7] 特許庁「『特許・実用新案審査基準』の改訂について（令和 8 年 6 月）」2026 年 6 月 25 日（改訂後基準は同年 7 月 1 日以降に行われる審査に利用）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/tukujitu_kijun/kaitei2/r8_shinsa_kijun_kaitei.html。
- [8] Autio, C. et al., Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1, National Institute of Standards and Technology, July 2024. NIST publication webpage updated Apr. 8, 2026. Publication page: <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>. DOI: <https://doi.org/10.6028/NIST.AI.600-1>。
- [9] 特許庁『特許・実用新案審査基準』第 I 部第 2 章第 1 節「本願発明の認定」（令和 8 年 6 月改訂後の現行版）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/tukujitu_kijun/ht/01_0200.html。
- [10] 特許庁『特許・実用新案審査ハンドブック』第 III 部第 2 章 3208～3210（電気通信回線を通じて公衆に利用可能となった発明の取扱い、公衆利用可能性および改変の疑義。2026 年 7 月 1 日更新の現行版）。
https://www.jpo.go.jp/system/laws/rule/guideline/patent/handbook_shinsa/document/index/03.pdf#page=13。
- [11] 最高裁判所第二小法廷平成 3 年 3 月 8 日判決・昭和 62 年（行ツ）第 3 号、民集 45 卷 3 号 123 頁（リパーゼ事件最高裁判決）。裁判例詳細 : <https://www.courts.go.jp/hanrei/52747/detail7/index.html>。判決全文 : <https://www.courts.go.jp/assets/hanrei/hanrei-pdf-52747.pdf>。
- [12] 知的財産高等裁判所平成 17 年 6 月 30 日判決・平成 17 年（行ケ）第 10068 号。判決全文 : <https://www.courts.go.jp/assets/hanrei/hanrei-pdf-9259.pdf>。

注：本評釈は、対象論文の法的・方法論的構造を検討するものであり、個別案件に対する法律意見ではない。