

AI時代の知財部門が採るべき業務設計：Can / Trust / Own の三軸による委任フレームワーク

Manus

はじめに

生成 AI および AI エージェントの発展により、知的財産業務の多くは自動化または高度な補助の対象となりつつある。特許調査、商標の類否判断、契約書の一次レビューなど、定型的な情報処理業務においては、AI が人間の処理能力を凌駕するスピードと精度を示し始めている。特に、IP ランドスケープの作成や特許文献の要約において、AI は劇的な効率化をもたらしている [2] [6]。

しかし、知財業務は企業の競争力そのものに関わる高度な判断を含む。未公開発明の取り扱い、営業秘密の管理、出願戦略の策定、他社権利リスクの評価、契約上の権利帰属の判断などは、単なる情報処理ではなく、ビジネス上のリスクとリターンを秤にかける戦略的意思決定である。そのため、AI 活用において「人間が最後に確認する」という抽象的な精神論だけでは、実務上のガバナンスとして不十分である [3] [5]。

本稿の中心命題は以下の通りである。

知財業務における AI 活用は、AI に判断を委ねることではなく、知財判断のプロセスを分解し、Can / Trust / Own の三軸で『任せられる範囲』を設計することである。

本稿では、AI に仕事を任せるための評価軸である「Can（能力）」「Trust（信頼・検証可能性）」「Own（責任と帰属）」の三軸を知財業務に適用し、AI時代の知財部門が採るべき具体的な業務設計のあり方を論じる。

1. 知財業務における AI 委任の現状と限界

現在、知財実務における AI 活用は急速に進展している。特許調査や IP ランドスケープ作成などの領域では、AI が膨大な文献を瞬時に処理し、劇的な効率化をもたらしている。一方で、特許明細書の完成品作成や高度な権利化戦略においては、AI の限界も明らかになっている [6]。

特に、AI の出力には「ハルシネーション（もっともらしい虚偽情報）」のリスクが内在しており、米国では AI が捏造した架空の判例をそのまま裁判所に提出した弁護士が制裁を受ける事件も発生している [5]。また、未公開発明を外部の AI に入力することによる新規性喪失や情報漏洩のリスクは、企業の存続を揺るがしかねない重大な問題である [5]。

このような背景から、日本弁理士会や法務省のガイドラインでも、AI をあくまで「道具」として位置づけ、最終的なファクトチェックと法的妥当性の担保は人間の専門家が行うべきであるという「ヒューマン・イン・ザ・ループ」の原則が強調されている [5] [6]。しかし、実務の現場では、どこまでを AI に任せ、どこから人間が介入すべきかの境界線が曖昧なままである。

2. Can / Trust / Own フレームワークの知財業務への適用

AI への業務委任を適切に設計するためには、業務を「Can」「Trust」「Own」の三軸で評価することが有効である [4]。

2.1 Can（能力・適合性）

知財業務における「Can」は、単なるテキスト生成能力ではなく、技術的・法的な文脈の理解力である。高い Can を示す領域としては、先行技術文献の要約、IP ランドスケープのためのデータ分類、商標の図形要素（ウィーン分類）の特定、契約書からの定型条項の抽出などが挙げられる。これらはパターン認識と情報整理であり、AI が得意とする分野である。一方で、低い Can を示す領域としては、独立クレームの作成、技術的な「飛躍」や「進歩性」の論理構築、事業戦略と連動した出願国の選定などがある。これらは、文脈に依存した高度な推論と、未来の事業展開を見据えた戦略的思考を要するため、現行の AI には遂行が困難である [3] [6]。

2.2 Trust（信頼性・検証可能性）

知財業務における「Trust」は、ハルシネーションのリスクと直結する。AI の出力結果の正しさを、人間が合理的なコストで検証・信頼できるかが問われる。高い Trust（検証容易）な領域としては、AI が抽出した検索キーワードの確認、翻訳の一次チェック、形式的な期限管理などがある。これらは元データとの照合が容易であり、誤りがあっても発見しやすい。対照的に、低い Trust（検証困難）な領域としては、「侵害リスクなし（FTO クリア）」という消極的事実の証明や、高度に専門的な技術分野における先行技術の網羅的調査などが挙げられる。AI が「見落とした」ものを人間が後から発見するのは極めて困難であり、AI の結論を盲信すること（自動化への過信）は致命的なリスクとなる [3] [5]。

2.3 Own（責任・権利帰属）

知財業務における「Own」は、法的責任とビジネス上の決断の所在である。その判断がもたらす結果（リスクとリターン）の責任を誰が負うべきか、また、成果物の権利は誰に帰属すべきかが問われる。AI に委ねられない Own としては、未公開発明の入力可否の判断（情報漏洩リスクの受容）、他社特許への抵触判断と事業継続の可否、最終的な特許請求の範囲の確定などが挙げられる。これらは、企業の競争力に直結する経営判断であり、弁理士の善管注意義務や企業のコンプライアンスに関わるため、人間（専門家・経営層）が完全に Own（所有・責任負担）しなければならない [5]。また、日本の裁判所（東京地裁令和 6 年 5 月 16 日判決）が示したように、発明者は自然人に限られ、AI は発明者になり得ないという権利帰属の観点からも、人間が主体となる必要がある [1]。

3. 三軸に基づく知財部門の業務設計

上記の三軸（Can/Trust/Own）を掛け合わせることで、知財業務を以下の 4 つのレベルに分類し、適切な委任設計を行うことができる [4]。

レベル	名称	評価（Can / Trust / Own）	該当業務の例	業務設計の指針
レベル 1	完全人間主導 （Human-Led）	Own が極めて重く、Trust が困難、Can も限定的	知財戦略の立案、侵害・非侵害の最終判断、独立クレームの権利範囲確定、重要契約の最終承認	AI は情報提供（インフォーム）にとどめ、人間が 100% の責任と決定権を持つ。
レベル 2	AI 支援・人間主導 （AI-Assisted, Human-Led）	Can は高いが、Own が重く、Trust が必須	従属クレームの作成、明細書初稿の作成、IP ランドスケープの分析と示唆出し、契約書のリスク条項抽出	AI がドラフトや分析結果を提示し、人間が必ずレビューし、文脈に沿って修正・承認する。
レベル 3	AI 主導・人間検証 （AI-Led, Human-Verified）	Can が高く、Trust が容易で、Own のリスクが中程度	先行技術文献の一次スクリーニング、商標の類似検索（一次）、定型的な社内知財相談への一次回答	AI が自律的に作業を進めるが、人間が最終的な「サニティチェック（健全性確認）」を行ってからリリースする。
レベル 4	自動化 （Automated）	Can が十分にあり、Trust が不要または極めて容易で、Own のリスクが低い	期限管理のアラート、公開情報のデータクレンジング、定型的な書誌事項の抽出	AI（システム）に完全に委任し、例外発生時のみ人間に通知する。

4. 結論：AI エージェント時代の知財専門家の役割

知財業務における AI 活用は、単なる「作業の自動化」から、AI エージェントによる「自律的なタスク遂行」へと進化しつつある。しかし、知財という企業のコア競争力を扱う領域において、最終的な判断（Own）を AI に委ねることはできない。

知財部門が採るべき業務設計は、AI の能力（Can）を最大限に引き出しつつ、検証可能性（Trust）を担保するプロセスを構築し、人間が責任を負うべき核心部分（Own）を明確に切り分けることである。

今後の知財専門家（弁理士や企業知財部員）に求められるのは、情報検索や定型文案の作成といった作業スキルではなく、AIが提示した分析結果やドラフトをビジネス戦略と照らし合わせて評価し、リスクをコントロールしながら最終的な決断を下す「判断力」と「オーケストレーション能力」である。AIを恐れるのではなく、Can / Trust / Own のフレームワークを用いて「任せられる範囲」を戦略的に設計することこそが、次世代の知財部門の競争力の源泉となる。

参考文献

- [1] 経済産業省 特許庁. "特許制度に関する検討課題について."
https://www.jpo.go.jp/resources/shingikai/sangyo-kouzou/shousai/tokkyo_shoi/document/52-shiryuu/01.pdf [2] 安藤俊幸. "AI特許調査ツールと生成系 AI の連携による高精度化検討." 情報の科学と技術, 2024.
- [3] Samuel W. Apicelli. "AI in Patent Prosecution: Judgment, Risk, and the Limits of Automation." Duane Morris LLP, Feb 23, 2026. <https://blogs.duanemorris.com/artificialintelligence/2026/02/23/ai-in-patent-prosecution-judgment-risk-and-the-limits-of-automation/> [4] Artha Learning Inc. "The Human-AI Delegation Framework for L&D." Feb 6, 2026. <https://arthalearning.com/the-human-ai-delegation-framework-for-ld/> [5] PatentRevenue. "生成 AI 出力の法的信頼性と実務リスク：専門家の確認義務と知財戦略の展望." Mar 5, 2026. <https://patent-revenue.iprich.jp/%E5%B0%82%E9%96%80%E5%AE%B6%E5%90%91%E3%81%91/4347/> [6] GrIP. "生成 AI 時代の知財調査実務 2026 記事調査レポート." Mar 2, 2026. <https://growing-ip.com/?p=1258>