

AI 導入の 8 段階（The Eight Levels of AI Adoption）を 知的財産部門の業務に当てはめる

— 知財実務家・知財部門管理職向け実践レポート（第 2 版） —

2026 年 7 月 20 日

Claude Fable 5

1. 要旨

Every 誌の Mike Taylor 氏による「The Eight Levels of AI Adoption」（2026 年 6 月公開）は、AI 活用を Level 1（Chatbot）から Level 8（Orchestrator）までの 8 段階に整理し、「higher is not always better（レベルが高いほど良いわけではない）」——最適レベルは、①AI が無介入で適切に処理できると信頼できる度合いと、②失敗時の影響の大きさという二軸を総合して決定する——という原則を核心に置く¹⁾。この考え方は、法的責任・秘密管理・証拠性が重い知財業務にそのまま当てはまり、業務類型ごとに最適レベルが異なる。

知財業務では、定型的・大量で、誤りを後工程で検出・訂正しやすい業務（一次スクリーニング、翻訳下訳、社内問合せ対応、発明提案書の下書き等）については、十分な情報管理と品質評価を前提として、AI への委譲を Level 3~5 まで高める余地がある。一方、明細書の最終確定、中間応答の主張・補正方針、FTO・侵害評価、契約締結、権利維持・放棄など、法的責任または不可逆性の高い判断については、AI の利用を情報整理・候補提示・検証補助に限定し、最終判断を人間が保持すべきである（本レポートの分析・提案）。

なお、この 8 段階は AI への委譲・自律性の程度を表すものであり、情報セキュリティや法的安全性の等級ではない。したがって、業務ごとの適正レベルは、AI の処理能力、失敗時の影響、可逆性、秘密性、検証可能性および人間による統制を総合して決定する必要がある。また、本評価は 2026 年 7 月時点のモデル性能、製品機能および法制度・ガイドラインを前提とするものであり、定期的な見直しを要する。

2. 原典フレームワークの確認（一次情報）

Every 誌の記事「The Eight Levels of AI Adoption」（著者 Mike Taylor、共著 Laura Entis・Claude、2026 年 6 月 2 日公開、7 月 5 日更新、2026 年 7 月参照）を直接確認した¹⁾。Mike Taylor 氏は Every の head of tech consulting で、『Prompt Engineering for Generative AI』（O'Reilly）の共著者である。フレームワークは Steve Yegge 氏の「Gas Town」投稿（コーディ

ングエージェント論）を土台に、Taylor 氏がクライアント向けに一般化・修正したものである。各レベルの名称と定義は表 1 のとおりである。

表 1 AI 導入の 8 段階（原典に基づく整理。例示ツール名は原典記載）

Level	名称（原文／訳）	原典の定義（例示ツールを含む）
1	Chatbot（チャットボット）	尋ねれば答える。ファイルやシステムに埋め込まれていない汎用 AI（ChatGPT, Claude, Gemini）
2	Copilot（コパイロット）	AI が自分のファイルの中に存在し、隣で一緒に作業する（Cursor, Claude in Excel, Gemini in Google Docs）
3	Agent（エージェント）	タスクを記述するとエージェントが段階ごとに実行し、次に進む前に承認を求める（Cowork, Codex）
4	Autopilot（オートパイロット）	承認をスキップし、エージェントが自力で完遂し、後で結果をレビューする（Lovable, Codex, Claude Code）
5	Workflows（ワークフロー）	エージェントの出力を専門品質化する仕組みを構築する（Compound engineering, Claude Workflows, Copilot AI Studio）
6	Assistant（アシスタント）	指示されなくてもバックグラウンドで能動的に動く（OpenClaw, Hermes Agent, Claude Managed Agents）
7	Multi-agent（マルチエージェント）	複数の長時間稼働エージェントを同時に管理する（Claude Managed Agents, OpenClaw, Codex Goals）
8	Orchestrator（オーケストレーター）	マネージャー役のエージェントがサブエージェントのチームを率いる（Gas Town, Paperclip, Symphony）

レベル間の移行について原典は、各レベルが上がるごとに「より多くの仕事を AI に委譲し、より多くの信頼を置く」とし、移行の一般条件は「品質が業務に足るか」「コスト（時間・エンジニアリング資源・トークン）が見合うか」であり、「安全に次のレベルに上がるには努力と実験、あるいはモデル能力の飛躍が必要」とする¹⁾。

「higher is not always better」の論拠として原典は、①最も洗練された AI ユーザーは複数レベルを同時に操り、目の前の課題ごとに最適なレベルを選ぶこと、②タスクに適したレベルは「介入なしで AI がうまくやると信頼できる度合い」と「失敗したときの影響の大きさ」の二軸で決まること、③ハイクラスなタスクでは低いレベルに留まって監督するか、同等品質を高レベルで得るために時間・資源・トークンを投じる覚悟が必要なこと、④ナレッジワーカーのスイートスポットは Level 1～4、エンジニアは Level 5～8 に多いこと（後者は不安定な新システムの足場=scaffolding を自ら作れるため）を挙げる¹⁾²⁾。なお、原典はこの二軸を判断の観

点として述べるものであり、数値的な計算式を定めるものではない。

（注記）原典ページは Level 1 の詳細まで無料公開であり、Level 2 以降の各節は Every の登録が必要なコンテンツである。本レポートの Level 2～8 の定義は、原典の一覧表（公開部分）、コンパニオン記事²⁾、および調査時に取得した原典本文に基づく要約であり、直接引用ではない。

3. 経理業務への翻訳例（参考動画・書き起こしで検証済み）

YouTube 番組「AI×BPO ラジオ」第 80 回「【完全保存版】あなたの会社の経理は今レベルいくつ？AI 導入の 8 段階」³⁾（株式会社管理のプロ⁴⁾運営、代表・松村氏と AI/DX 担当・田中氏が出演）の書き起こしにより、番組内容を確認した。番組は、X 上で「『AI 使ってます』のレベル感が違いすぎる、客観的な指標が欲しい」という投稿が話題になったことを契機に、上記 8 段階を日本の中小企業のバックオフィス（特に経理）向けに翻訳したものである。この物差しは「人の技量ではなく AI への任せ方の深さ」を測るため、会社単位でも部署単位でも適用可能とする。Level 1～4 は「個人が AI を使う」段階、Level 5～8 は「組織の仕組み」の段階であり、Level 4 と 5 の間に大きなキャズム（溝）があると整理する。原典の「レベルが高いほど偉いわけではない」「多くのホワイトカラーの仕事は Level 1～4 で十分に価値が出る」という留保も明示的に紹介し、適正レベルは「そのタスクで AI をどこまで信頼できるか」と「失敗した時のダメージの大きさ」の「掛け算」で選ぶ（番組の表現）と説明している³⁾。番組による各レベルの経理業務への割り当ては以下のとおりである³⁾。

- ・ **L1 チャットボット**：勘定科目の相談、インボイス制度の個別ケースの調べもの、支払延期依頼メールの下書き。「顧問税理士に聞くほどでもない小さな疑問を投げる」使い方から始まる。注意点は「答えの確認は全部人間の仕事」であり、特に税務は専門家確認前提（ハルシネーション対策）。
- ・ **L2 コパイロット**：会計ソフト画面内での仕訳提案、試算表ファイルを渡して前月比の異常検知。コピペの往復が消え「作業現場に AI が同席」する。L1～2 は導入コストほぼゼロで、世の中の「AI 使ってます」の大部分がここに集中。
- ・ **L3 エージェント**：AI が「答える存在から働く存在」になる。経費精算チェック（申請を規程と照合し、上限超過等の怪しいものだけ挙げて差し戻し確認）、請求書の山からのデータ読取→仕訳ドラフト→「登録していいですか」との確認。「実行は AI、判断は人間」の分業であり、「めちゃくちゃ頭のいい新人」の比喻が用いられる。

- ・ **L4 オートパイロット**：承認を手放し、人は結果レビューのみ。向くのは「ルールが明確で、間違えてもすぐ戻せる定型業務」＝入金消込（入金データと請求データの突合）。決算の最終数値や税務判断など失敗ダメージが大きい業務は L3 で止める（止めるべき）とする。
- ・ **L5 ワークフロー**：単発利用では出力品質にムラが出るため、毎回同じ品質が出るよう手順をパイプライン化する。典型は月次決算（銀行データ取込→仕訳ドラフト→規程チェック→異常あれば担当者に通知→レポート）。前提条件は業務の棚卸しと言語化（「業務が言葉になっていないとラインが引けない」）。
- ・ **L6 アシスタント**：L5 まで起点は人間だが、L6 は AI が指示を待たない。受信箱・共有フォルダを監視し、請求書が届いた瞬間に処理を開始、朝出社したら夜間に取込とドラフト作成が完了している。各社の常駐型 AI エージェント（Google Gemini 系、マネーフォワード「AI の同僚」コンセプト等）が例示される。勝手に動く範囲が広がるほど暴走リスクも増えるため、L5 のライン設計＋権限設計（ハーネス）がセットで必要。
- ・ **L7 マルチエージェント**：経費チェック担当 AI・売掛金督促担当 AI・月次レポート担当 AI が同時に連携稼働。事例として SHIFT 社（社内で「天才君」と呼ぶ AI エージェント 4,400 体、バックオフィスを半年で 110 人分 AI 化、営業事務 137 人→52 人）を「上場企業が全社で L7～8 をやってみせた例」として紹介（番組内の発言であり、SHIFT 社の一次資料とは未照合）。
- ・ **L8 オーケストレーター**：マネージャー AI が部下 AI チームに仕事を割り振り、人間はオーケストレーター AI に指示するだけ。残る人間の仕事は「何をやらせるかを決めること」と「成果が業務として正しいかのジャッジ」。中小企業が今すぐ L8 を目指すべきではなく、世の中の方向を知る「北極星」の位置づけとする。

番組はさらに、現在地診断の 3 質問（①AI に相談しているか→L1～2、②承認付きで作業を任せている業務があるか→L3～4、③人が指示しなくても動く仕組みがあるか→L5～6）と、登り方の 3 原則（①1 段ずつ上がる、②業務ごとにレベルを変える——入金消込 L4・経費チェック L3・税務判断は AI に委ねない、という「デコボコ」で問題ない、③キャズムを渡る前に業務の棚卸し・言語化）を示す。やってはいけないのは「整っていない業務にいきなり高レベルのツールを丸投げ」することであり、混乱を AI に渡すと混乱が高速で再生産されるとする³⁾。これらの論点は、知財への当てはめでもそのまま重要な示唆となる。

4. 知財業務における 8 段階の当てはめ（筆者による分析・提案）

以下は原典フレームワークを知財業務に当てはめた本レポート独自の分析であり、事実の記述ではなく提案である。各レベルにつき「知財での具体的な姿／ユースケース／必要な体制・スキル／リスク」を示す。

Level 1 — Chatbot（チャットボット）

- ・ **姿**：汎用生成 AI（ChatGPT／Claude／Gemini 等）への単発の質問。特許用語の解説、技術文献の要約、クレームの平易化、知財教育の教材下書き、社内知財ヘルプデスク。
- ・ **ユースケース**：明細書ドラフト前の技術理解、審査基準・審査ハンドブックの論点整理、発明者向け Q&A。なお、法令、審査基準、判例、審決および特許文献については、AI の回答を根拠とせず、公式データベース上の原文（該当箇所）を確認する運用とする。
- ・ **体制・スキル**：プロンプト設計の基礎。未公開情報を入力しない運用ルール。
- ・ **リスク**：ハルシネーション（存在しない特許番号・判例の生成）。未公開発明を外部生成 AI に入力した場合、利用規約、秘密保持義務、学習利用、ログ保存・再提供の条件によっては、守秘義務違反、営業秘密性の喪失、情報漏えい、さらに公開状態に至った場合の新規性への影響が生じ得る。

Level 2 — Copilot（コパイロット）

- ・ **姿**：知財担当者の作業環境（Word 上の明細書作成支援、特許検索 DB 内蔵 AI）に組み込まれ、隣で提案。
- ・ **ユースケース**：明細書作成支援ツール（過去出願、社内用語集および標準テンプレートを参照して請求項・明細書の候補文を提示するもの）、拒絶理由通知の論点整理、検索式作成支援、査読・可読性向上支援。
- ・ **体制・スキル**：入力内容が学習利用されない管理された環境（後述の閉域環境）の導入、社内文体・用語集の整備。レビューは全文確認に加え、数値・符号・従属関係・サポート要件・用語一貫性のチェック項目を定める。
- ・ **リスク**：出力の過信。担当者による人間の最終確認が前提。

Level 3 — Agent（エージェント）

- ・ **姿**：タスクを記述すると AI が段階的に実行し、次に進む前に承認を求める（実務上は全ステップではなく、重要なゲートに承認点を設計する）。

- ・ **ユースケース**：先行技術調査（技術要素抽出→検索式提案→スクリーニング→対比表を段階承認）、中間応答のゼロ稿作成（引用文献とクレームの相違点抽出→反論骨子提示）、無効資料調査の候補文献抽出。
- ・ **体制・スキル**：ワークフロー設計、承認ゲートの設定、サーチャー・弁理士のレビュー体制。
- ・ **リスク**：進歩性（複数文献の組合せ動機付け）反論など高度判断の精度不足。引用文献、技術内容および手続上の主張について、正確性、根拠資料との整合性および説明可能性を確保する必要がある（米国出願では Duty of Candor 等にも留意。第 6.4 節参照）。

Level 4 — Autopilot（オートパイロット）

- ・ **姿**：承認をスキップし、完遂後にまとめてレビュー。
- ・ **ユースケース**：大量特許の一次分類・タグ付け、SDI（定期監視）の初期スクリーニング、公開文献の翻訳下訳、パテントマップの母集団前処理。
- ・ **体制・スキル**：出力品質の検証（第 7 章参照）、精度 KPI のモニタリング。
- ・ **リスク**：見落とし（false negative。以下「見落とし」）が後工程に伝播する。FTO や無効資料調査では、Level 4 を唯一の調査手段または網羅性を保証する手段として用いることは適切でない（人手調査を補完する候補抽出や母集団形成には利用可能）。

Level 5 — Workflows（ワークフロー）

- ・ **姿**：エージェント出力を専門品質化する再現可能な仕組みの構築。入力、処理、検証、例外処理、承認および記録が再現可能な業務工程として設計されていることが本質であり、RAG・API 連携・MCP 等はこれを実現する技術要素になり得るが、それ自体が Level 5 を意味するわけではない。
- ・ **ユースケース**：発明発掘～出願依頼文の標準化パイプライン、IP ランドスケープのレポート自動生成テンプレート、社内ナレッジ（過去審査経過・成功反論パターン）と連携した RAG。
- ・ **体制・スキル**：知財 DX・情シス・法務の連携、プロンプト・テンプレートの資産化、品質評価・例外処理の組み込み。
- ・ **リスク**：仕組みの陳腐化、品質保証プロセスの形骸化、監査可能性の維持。

Level 6 — Assistant（アシスタント）

- ・ **姿**：指示なしでバックグラウンドで能動稼働。
- ・ **ユースケース**：他社特許の常時監視（競合の新規公開を自動検知しアラート）、ホワイトスペースの継続的モニタリング。年金・法定期限管理については、確立された期限管理システムを正本とし、AI は異常検知・予兆通知・二重チェックの補助として利用する（AI 通知のみを期限遵守の根拠としてはならない）。
- ・ **体制・スキル**：監視対象・閾値の設計、誤検知トリアージ、権限設計を含むガバナンス。
- ・ **リスク**：過検知・見逃し。権利行使判断や期限遵守の最終責任は人。

Level 7 — Multi-agent（マルチエージェント）

- ・ **姿**：複数の長時間稼働エージェントの同時管理。単なる並行実行ではなく、調査、分類、技術分析、市場分析および可視化を担当する複数のエージェントが、共有された目標・データ・評価基準の下で連携し、相互の成果を引き継ぐ構造を指す。
- ・ **ユースケース**：大規模 IP ランドスケープ（調査・分析・可視化・競合分析を役割分担して連携）、ポートフォリオ全体の棚卸し・整理。
- ・ **体制・スキル**：エージェント設計・オーケストレーション、競合・矛盾の処理設計、検証スキル、専任の DX 人材。
- ・ **リスク**：複合的なハルシネーションの検出困難、戦略判断の外部委任リスク。

Level 8 — Orchestrator（オーケストレーター）

- ・ **姿**：上位エージェントが目標を分解し、複数の下位エージェントに割り当て、成果を評価し、必要に応じて再計画する構造。
- ・ **ユースケース**：知財ライフサイクル全体（発明発掘→調査→ドラフト→中間→ポートフォリオ）の統合運用の実験。なお、Patlytics は特許ライフサイクルの広範な業務を対象とする統合 AI プラットフォームを標榜しているが、広い機能範囲を有することだけで本フレームワーク上の Level 8 に該当するとは限らず、該当性は上記の分解・割当・評価・再計画の機能を備えるかによって判断すべきである。
- ・ **体制・スキル**：高度なエンジニアリング資源、トークンコスト負担、厳格な検証体制。
- ・ **リスク**：2026 年時点では知財の高責任領域では実験段階。戦略・法的最終判断は人が担うべき。

5. 「高いレベルほど良いわけではない」の知財的含意と業務類型別の適正レベル

【重要な注意】本フレームワークの Level は、AI への委譲・自律性の程度を示すものであり、情報セキュリティや法的安全性の等級ではない。低い Level であっても外部の汎用 AI に秘密情報を入力する運用であれば高リスクとなり得る一方、適切な閉域環境、権限管理および承認点を備えた Level 3 の方が安全な場合もある。本章の「目安」は、閉域環境（入力・出力が学習利用・二次利用されず、アクセスが管理された環境。第 6.5 節参照）の利用と人間の最終確認を前提とした、委譲の深さの目安である。

知財業務の適正レベルは、原典の二軸——①AI への信頼度、②失敗時の影響の大きさ¹⁾——に、知財固有の要因（法的責任・代理人責任、秘密情報管理、証拠性〔出典・公開日の真正性と追跡可能性〕、ハルシネーション、判断業務か定型業務か、可逆か不可逆か）を重ねて総合的に決定すべきである。また、L1（チャットボット）は「人間が判断すること」を意味しないため、表 2 では「AI に任せやすい工程とその目安」と「人間が保持する判断」を分けて示す。契約締結判断のように、AI による委譲の対象とせず人間のみが行うべき事項（Human-only）は、レベルを付さずその旨を記す。

表 2 知財業務類型別の適正レベルの目安（筆者による 2026 年時点の提案）

業務類型	AI に任せやすい工程	目安	人間が保持する判断
社内問合せ・知財教育	回答案、教材案、FAQ 検索	L1～L5	公式見解の確定
先行技術調査	要素分解、検索式案、候補抽出	L2～L4	調査範囲の設計、特許性評価
翻訳・分類・要約	下訳、タグ付け、一次要約	L4～L5	重要箇所・訳語の確定
発明発掘・発明提案書	質問案、構成案、ゼロ稿	L2～L4	発明の把握、発明者への確認
明細書作成	初稿、整合性チェック	L2～L3	権利範囲・記載の最終確定（弁理士・担当者）
中間処理（拒絶理由対応）	引用箇所抽出、対比表、反論骨子案	L2～L3	補正・主張方針の決定
FTO・侵害予防調査	公報抽出、クレーム分解、対比表案	L2～L3	調査設計、侵害評価、事業判断（弁理士・弁護士等）
無効資料調査	候補抽出、時系列整理	L2～L4	証拠性・公開日・無効理由の評価

契約・ライセンス	条項抽出、差分比較、論点整理	L2～L3	交渉・受諾。締結判断は委譲対象外（Human-only、権限者）
年金・期限管理	補助通知、異常検知（正本は期限管理システム）	L4～L6	正本確認、法定期限遵守。権利維持・放棄判断は権限者
IP ランドスケープ	集計、分類、可視化、レポート案	L4～L7	経営的解釈、戦略提案

原則は、「AI が作業（情報処理）を担い、人間が判断を担う」HITL（Human-in-the-loop）構造を維持し、業務の不可逆性・法的責任が高いほど委譲を浅くすることである。これは番組が経理で示した「入金消込 L4・経費チェック L3・税務判断は委ねない」というデコボコ運用³⁾と同型であり、知財では「翻訳下訳 L4～5・先行技術調査 L2～4・FTO 評価や明細書確定は人間が保持」というデコボコが正解となる。効率化で浮いた時間を、入口（問いの設計）と出口（解釈・提案・責任ある判断）という高付加価値業務に振り向けることが、知財部門の価値を縮小させないための条件である。

6. 日本企業の知財部門の実情と政策動向（事実）

6.1 人員・分業構造

会社勤務（企業内）の弁理士は増加傾向にあり、特許庁の産業構造審議会弁理士制度小委員会資料によれば 2022 年末時点で約 3,000 人である⁸⁾（最新の内訳は日本弁理士会「会員分布状況」⁹⁾を参照）。日本企業には、「明細書作成」を外部特許事務所に、「発明発掘・出願レビュー」を社内知財部に、と役割分担する例が多く、調査（先行技術・無効・FTO）を外部（特許事務所・調査会社等）に外注する例も少なくない。

6.2 特許庁の政策

特許庁は「特許庁における人工知能（AI）技術の活用に向けたアクション・プラン」（令和 4～8 年度版、令和 7 年度改定版）において、実証事業の結果を踏まえ「先行技術調査②」（検索手法の高度化）を導入フェーズに移行させ、「生成 AI の特許審査業務への適用」を新設して生成 AI の活用を検討するとしている⁶⁾。また、特許・実用新案審査ハンドブック附属書「AI 関連技術に関する特許審査の事例」に令和 6 年 3 月に 10 事例を追加した⁷⁾。

6.3 日本弁理士会ガイドライン

日本弁理士会は「弁理士業務 AI 利活用ガイドライン」（令和 7 年 4 月）を公表した（β 版は

2025 年初頭に先行公開、同年 4 月 22 日の記者説明会で正式公表）⁵⁾。同ガイドラインは、生成 AI の生成物について「生成 AI から得られた生成物を活用する際には、弁理士としてその正確性を確認する必要がある、最終的には弁理士が責任をもって提供するべきである」と述べる（弁理士は、依頼者との委任・準委任その他の契約関係に応じ、善良な管理者としての注意義務を負い得る）。守秘義務については、弁理士法第 30 条を踏まえ、外部事業者が提供する生成 AI に秘密情報を入力する行為は「生成 AI 提供者という第三者に秘密情報を開示することになるため、守秘義務に違反するおそれ」があり、秘密保持契約がある場合は契約上の義務違反となるおそれがあるとする。また、入力情報が学習に利用されると他者への出力に使用され情報が漏えいする可能性を指摘し、「知財の分野においては、新規性喪失の観点からも十分な注意が必要」とする⁵⁾。すなわち、秘密保持措置が不十分な第三者プラットフォームへの入力、守秘義務違反や情報漏えいを招く可能性があり、その情報が公衆に利用可能な状態となった場合には新規性に影響するおそれがある——入力が直ちに新規性喪失を意味するわけではないが、そのリスク経路を遮断する管理が求められる。

6.4 米国実務上の留意（Duty of Candor 等）

米国では、AI 生成物を十分に検証せず、重要な虚偽情報または誤解を招く情報を USPTO に提出した場合、提出書類に関する署名・確認義務や Duty of Candor（誠実義務）との関係で問題となり得る。事情によっては制裁、手続上の不利益、または訴訟における Inequitable Conduct（不公正行為）の主張につながる可能性がある。Inequitable Conduct は単なる誤りや検証不足のみで直ちに成立するものではなく、一般に情報の重要性（materiality）や欺く意図（intent to deceive）が問題となるが、いずれにせよ人間による根拠資料との照合と最終確認が不可欠である。米国出願の詳細は、USPTO の AI 利用に関するガイダンスおよび関連規則の原文を確認されたい。

6.5 ツール選定のセキュリティ要件

本レポートで「閉域環境」とは、入力・出力データが基盤モデルの学習に利用されず（学習利用の禁止・オプトアウト）、目的外の二次利用が契約上禁止され、アクセスが管理された環境をいう。なお、「学習利用の禁止（オプトアウト）」と「ゼロデータ保持（入力・出力を原則保存しない）」は別概念であり、限定保持（不正利用監視等のための一定期間保存）や契約上の非利用と区別して確認する必要がある。ツール選定では、通信時・保存時の暗号化、学習利用の禁止、ログの保存期間とゼロデータ保持の適用範囲、事業者・再委託先によるアクセス条件、データ保存国・越境移転、テナント分離、アクセス権限と監査ログ、契約終了時の削除、

インシデント通知、モデル改善目的での二次利用の有無を個別に確認する。一般的な通信時・保存時の暗号化のみでは、事業者が処理のために平文データへアクセスできる場合がある点にも留意する。

6.6 企業・事務所の導入事例とツール動向

三井化学は、特許分析・新規用途探索・営業支援の3機能を備えた独自開発の生成AIチャット搭載プラットフォームを開発し、2024年度内に実証実験を完了、2025年度から社内での本格運用を開始すると発表した¹⁰⁾（実証において業務時間を80%削減できたとの報道がある〔同社側の報告値〕¹¹⁾）。旭化成は用途探索支援AIを活用し、知財部門が事業判断を支援していると報告されている¹²⁾。ある企業知財部はTokkyo.Aiの「生成AI Plus」により出願依頼文作成と簡易調査を従来約15時間から1~2時間に短縮したと報告されている（同社公表の事例値）¹³⁾。

IPTech弁理士法人は2025年4月1日から生成AIを本格導入し、使用ツールを「①暗号化処理がされており、サービス提供事業者に入力内容が解析できないこと、②入力されたデータが再学習の対象とならないこと、③上記①②が利用規約等で明示的に保証されていること」（同法人公表の要件・原文）の3要件を満たすものに限定している¹⁴⁾。

ツール動向としては、明細書作成支援（Tokkyo.Ai、Rowan Patents等）、中間処理支援（AISamurai ONE、パテント・インテグレーション「サマリヤ」等）、侵害分析・クレームチャート（Patlytics、XLSCOUT等）が挙げられる。Patlyticsは2026年4月8日にSignalFireがリードする4,000万ドルのSeries B（累計調達額約6,500万ドル）を発表し、Am Law 100の40%超が利用、顧客報告効果として「プロジェクト時間80%削減、クレームチャート1件あたり3万ドル超の節減、特許出願1件あたり最大15時間の回収」（いずれも同社公表値）を公表している¹⁵⁾。ただし進歩性反論など高度判断はAIの精度が発展途上であり、人間の補完が不可欠との指摘が一貫している。

7. 導入ロードマップと判断基準（筆者提案。原典の推奨手順ではない）

番組が示す登り方3原則（①1段ずつ上がる、②業務ごとにレベルを変える、③キャズムを渡る前に棚卸し・言語化）³⁾は知財部門にもそのまま適用でき、以下のロードマップの骨格をなす。

- ・ **第1段階（0~1か月）現在地の可視化**：業務を「定型度×処理量×失敗時の影響（不可逆性・法的責任）」の3軸で棚卸しし、業務類型ごとの現行レベルを診断する。診断には番組の3質問の知財版——①AIに相談しているか（L1~2）、②承認付きで調査・ドラフトを任せている業務があるか（L3~4）、③人が指示しなくても監視・処理が動く仕組みが

あるか（L5～6）——が使える。まず社内知財問合せ・教育、翻訳・分類の下処理、先行技術調査の一次スクリーニングといった低リスク業務を対象に選定する。

- ・ **第2段階（1～3 か月）基盤とルールの整備**：弁理士会ガイドライン⁵⁾に整合する社内ガイドラインを策定する（正確性確認・守秘義務・最終責任・秘密情報管理）。ツールは第 6.5 節の要件を個別に確認した閉域環境のものを選定し、未公開発明の外部汎用 AI への入力を原則禁止する。記録として、入力、出力、参照資料、利用モデル・バージョン、実行日時、担当者の修正内容、レビュー結果および承認記録を保存する（記録自体が未公開発明等を含み得るため、保存期間とアクセス権限も定める）。
- ・ **第3段階（3～6 か月）スモールスタートと PoC**：選定業務で Level 1～3 を試行する。品質検証は、単純なサンプリング監査に加え、正解付き評価セット、過去案件によるバックテスト、AI 非使用チームとの比較、高リスク案件の全件レビュー、技術分野別の精度評価、モデル更新時の再検証を組み合わせ、誤検知と見落としを別個に測定する。
- ・ **第4段階（6～12 か月）段階的拡大**：定型・大量業務（翻訳下訳、分類、SDI 一次監視）を Level 4～5 へ。ワークフロー化（テンプレート・RAG・過去審査経過連携。ただし Level 5 の要件は第 4 章参照）で品質を標準化する。高責任業務（明細書確定、進捗性反論の方針、FTO 評価、契約締結）は人間の判断として保持し、AI は情報整理・候補提示・検証補助に限定する。
- ・ **第5段階（12 か月～）高度化の検討**：他社監視を Level 6 へ（期限管理は期限管理システムを正本とし AI は補助）、IP ランドスケープの生成工程を Level 7 での連携処理へ。ただし解釈・戦略・法的最終判断は人が担う設計を堅持する。

表 3 業務別 KPI の例（筆者提案）

業務	主な KPI
文献スクリーニング	再現率、見落とし率、適合率
翻訳	用語誤り、数値誤り、欠落率
明細書支援	用語一貫性、従属関係誤り、サポート欠如の検出数
中間処理支援	引用箇所の正確性、根拠のない主張の数
社内 FAQ	正答率、一次資料提示率、エスカレーション率
ワークフロー全般	処理時間、レビュー時間、差し戻し率、例外率

レベルを上げてよい**閾値**：①対象業務の AI 出力品質が上記の検証（評価セット・バックテスト等）で目標値（例：一次スクリーニングの見落とし率が人手同等以下）を安定的に満たす、②

失敗が可逆で後工程での検出・訂正が容易、③記録と人間の最終確認が仕組みで担保されている——の3条件がそろったとき。

委譲を浅くする・据え置くべきシグナル：法的責任・不可逆性が高い、ハルシネーションが致命傷になる、出典・公開日の真正性と追跡可能性の説明責任が問われる、秘密情報の管理・漏えいリスクが管理しきれない場合。

8. 留意事項

- ・本レポートの評価は2026年7月時点のモデル性能、製品機能および法制度・ガイドラインを前提とする。原典（文献1）の参照日は2026年7月であり、Level 2～8の記述は登録コンテンツの要約（直接引用ではない）である。例示ツール名は原典公開時点（2026年6月）の原典記載による。
- ・参考動画（文献3）の内容は自動書き起こし（Scripsy.app）に基づき確認したものであり、固有名詞等に転記揺れがあり得るほか、動画概要欄の全文および画面表示の図表は未確認である。番組が言及するSHIFT社の事例数値は番組内の発言であり、SHIFT社の一次資料との照合は行っていない。
- ・第4章・第5章・第7章（知財業務への当てはめ、表2・表3、ロードマップ）は筆者の分析・提案であり、原典・番組の記述ではない。実際の適用は各社の体制・技術分野・リスク許容度に応じた調整を要する。本レポートは法的助言ではなく、個別案件は弁理士・弁護士に相談されたい。
- ・日本弁理士会ガイドラインは今後改訂の可能性がある。特許庁アクション・プランは令和8年度までの計画であり、検討・実証段階の項目を含む（確定した制度ではない）。
- ・企業導入事例・ツールの効果数値（時間削減率、コスト削減額等）は、いずれも各社・ベンダーによる自己申告の公表値であり、独立に検証されたものではない。

参考文献

- 1) Mike Taylor (with Laura Entis, Claude) , "The Eight Levels of AI Adoption", Every, 2026年6月2日公開・7月5日更新. <https://every.to/guides/the-eight-levels-of-ai-adoption>（最終アクセス：2026年7月20日。Level 2以降は要登録）
- 2) Every, "Where Do You Fall on the Eight Levels of AI Adoption?"（コンパニオン記事）. <https://every.to/p/where-do-you-fall-on-the-eight-levels-of-ai-adoption>（最終アクセス：2026年7月20日）

- 3) AI×BPO ラジオ 第 80 回「【完全保存版】あなたの会社の経理は今レベルいくつ？AI 導入の 8 段階」株式会社管理のプロ (YouTube) . <https://www.youtube.com/watch?v=KuO6lnUAAVE> (内容は Scripsy.app による自動書き起こしに基づき確認)
- 4) 株式会社管理のプロ 公式サイト . <https://www.kanri-pro.com/> (最終アクセス：2026 年 7 月 20 日)
- 5) 日本弁理士会「弁理士業務 AI 利活用ガイドライン」令和 7 年 4 月 . <https://www.jpaa.or.jp/cms/wp-content/uploads/2025/04/AIservices-guideline.pdf> (最終アクセス：2026 年 7 月 20 日)
- 6) 特許庁「特許庁における人工知能 (AI) 技術の活用に向けたアクション・プランの令和 7 年度改定版について」 . https://www.jpo.go.jp/system/laws/sesaku/ai_action_plan/ai_action_plan-fy2025.html (最終アクセス：2026 年 7 月 20 日)
- 7) 特許庁 特許・実用新案審査ハンドブック附属書「AI 関連技術に関する特許審査の事例」 (令和 6 年 3 月、事例追加) . (掲載 URL は本文書では未検証)
- 8) 特許庁「弁理士制度の現状と今後の課題」産業構造審議会知的財産分科会弁理士制度小委員会 資料 2、令和 6 年 1 月 29 日 . https://www.jpo.go.jp/resources/shingikai/sangyo-kouzou/shousai/benrishi_shoi/document/20-shiryoku/02.pdf (最終アクセス：2026 年 7 月 20 日)
- 9) 日本弁理士会「会員分布状況」 . <https://www.jpaa.or.jp/about-us/members/> (最終アクセス：2026 年 7 月 20 日)
- 10) 三井化学株式会社 ニュースリリース「三井化学、生成 AI を活用した特許チャットを開発」2024 年 12 月 25 日 . https://jp.mitsuichemicals.com/jp/release/2024/2024_1225/index.htm (最終アクセス：2026 年 7 月 20 日)
- 11) 日経クロステック「三井化学が特許探索に生成 AI チャット、業務時間を 80%削減へ」2025 年 1 月 . <https://xtech.nikkei.com/atcl/nxt/news/24/02031/> (最終アクセス：2026 年 7 月 20 日)
- 12) 旭化成株式会社 用途探索支援 AI の知財部門活用に関する公表資料 (知的財産報告書等) ・報道 . (掲載 URL は本文書では未検証)
- 13) Tokyo.Ai「生成 AI Plus」導入事例に関する公表情報 . (掲載 URL は本文書では未検証)
- 14) IPTech 弁理士法人「お知らせ」 (2025 年 3 月 28 日付・生成 AI 本格導入の告知を含む) . <https://iptech.jp/info> (一覧ページ。個別告知の URL は未検証。最終アクセス：2026 年 7 月 20 日)
- 15) Business Wire, "Patlytics Raises \$40 Million Series B to Expand the AI Platform Purpose-Built for IP Work", 2026 年 4 月 8 日 . <https://www.businesswire.com/news/home/20260408770722/en/> (最終アクセス：2026 年 7 月 20 日)