

AI時代の知財業務設計：Can / Trust / Own による「任せられる範囲」の再定義

AI活用の中心命題：抽象論からの脱却
「最後に人間が確認」という曖昧さから、責任と工程を特定した設計へ。

抽象的な「人間による最終確認」



知財特有の「二つの非対称性」

- ・誤りの不可逆性（事後修正不能）
- ・法的・契約的な責任帰属の固定（発明者認定・署名）

判断のための三軸評価（Can / Trust / Own）



Can（能力）
技術的な実行可能性：AIが必要な情報にアクセスし、求める精度の出力を生成できるか。



Trust（信頼）
検証なしの採用可能性：誤りの検出可能性とコスト、再現性、根拠の追跡可能性が十分か。



Own（責任）
結果責任の引き受け構造：法的・戦略的傾斜の最終責任を、人間が構造的に引き受けられるか。

業務設計のゾーニングと三つの自律性段階

ゾーンA：委任可能（Delegate） **HOOTL**
(Human-out-of-the-loop)
AIを主担当とし、人間はサンプリング検証



公開文献の要約や形式チェック



オムロンの「ALBERT」：
AIが生垣首、人間は
サンプリング検証

設計原則: Can: 高 / Trust: 高 / Own: 軽 → 人間は事後的にのみ関与

ゾーンB：協働必須（Augment） **HOTL**
(Human-on-the-loop)
AIは「選択肢生成・論点可視化」に限定し、最終判断は人間



FTO評価やクレーム設計



三井化学：
AIは選択肢生成・
論点可視化に限定

設計原則: Can: 高 / Trust: 中 / Own: 重 → 人間が監視し、例外時のみ介入

ゾーンC：人間専管（Reserve） **HITL**
(Human-in-the-loop)
AI出力を一次資料にせず、一般論の壁打ちに留める



出願・秘匿の最終意思決定や署名行為

設計原則: Own: 移譲不可 → 人間が連次介入、移譲不可の工程

設計を支える三つの仕組み



工程ごとの検証プロトコル化
サンプリング率、検証観点、根拠提示要件の具体化。



情報フローの分離設計
機微情報が、外部API、社内LLM、オンプレミスのどこに流れるかを厳格管理。



責任構造の文書化
AI担当工程と最終責任者(Own)を記録し、事後検証・紛争に備える。