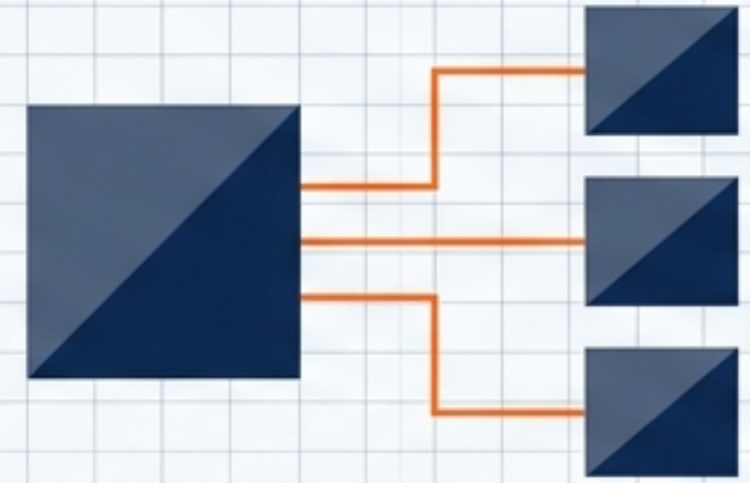




AI時代の知財部門の業務設計

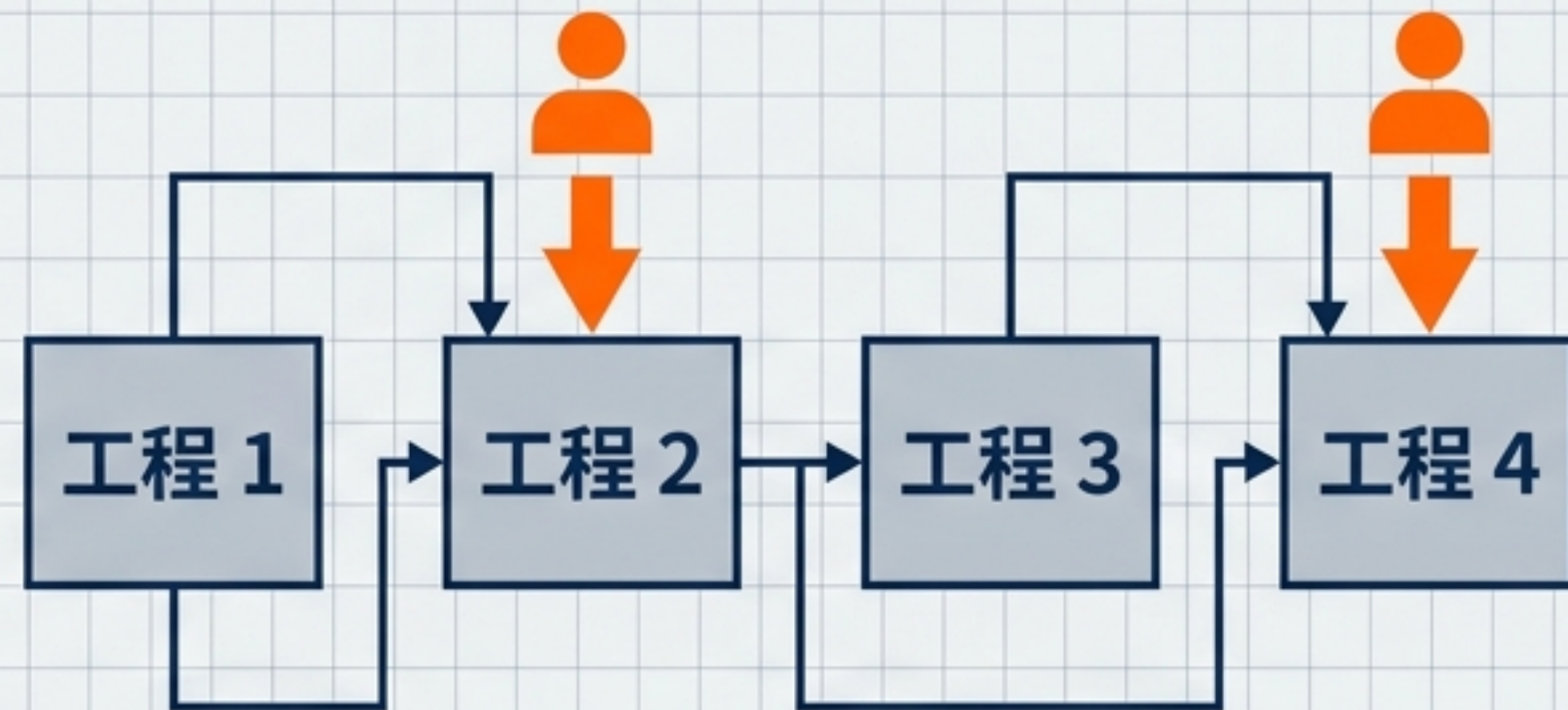
Can / Trust / Own による「任せられる範囲」の分解

2026年6月27日 / Claude Opus 4.8

抽象的な「人間による最終確認」の危険性



「最後に人間が確認する」という抽象論は、いつ・どの工程の・何を確認するのかを特定しないまま、責任の所在をブラックボックス化する。



知財判断のプロセスを細かい「工程」に分解し、意図的に設計しなければならない。判断を丸投げするのではなく、工程を切り分ける。

知財業務における2つの決定的非対称性

知財判断には、他業務にはない2つの厳格な制約が存在する。
確認すべきは成果物ではなく、工程ごとの「リスク構造」である。

誤りの不可逆性

- 未公開発明の取扱い、出願タイミング、新規性喪失の判断。
- 一度誤れば権利そのものが消滅し、事後修正が効かない。

Trust（信頼）の閾値が
極端に高くなる。

責任帰属の固定

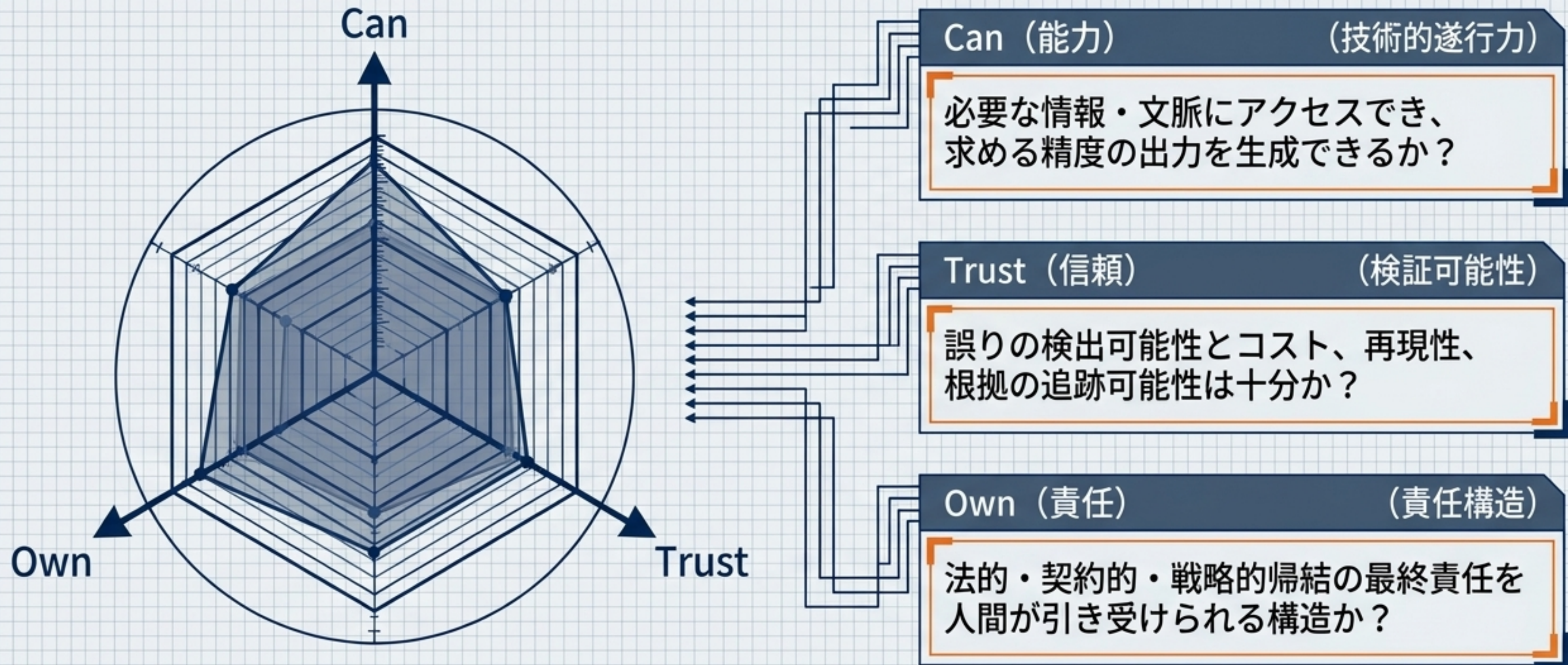
- 発明者認定、権利帰属、代理人としての署名。
- 法定・契約によりAIへの移譲が禁止されている領域。

Own（責任）が委任不可
（ゼロ評価）に張り付く。

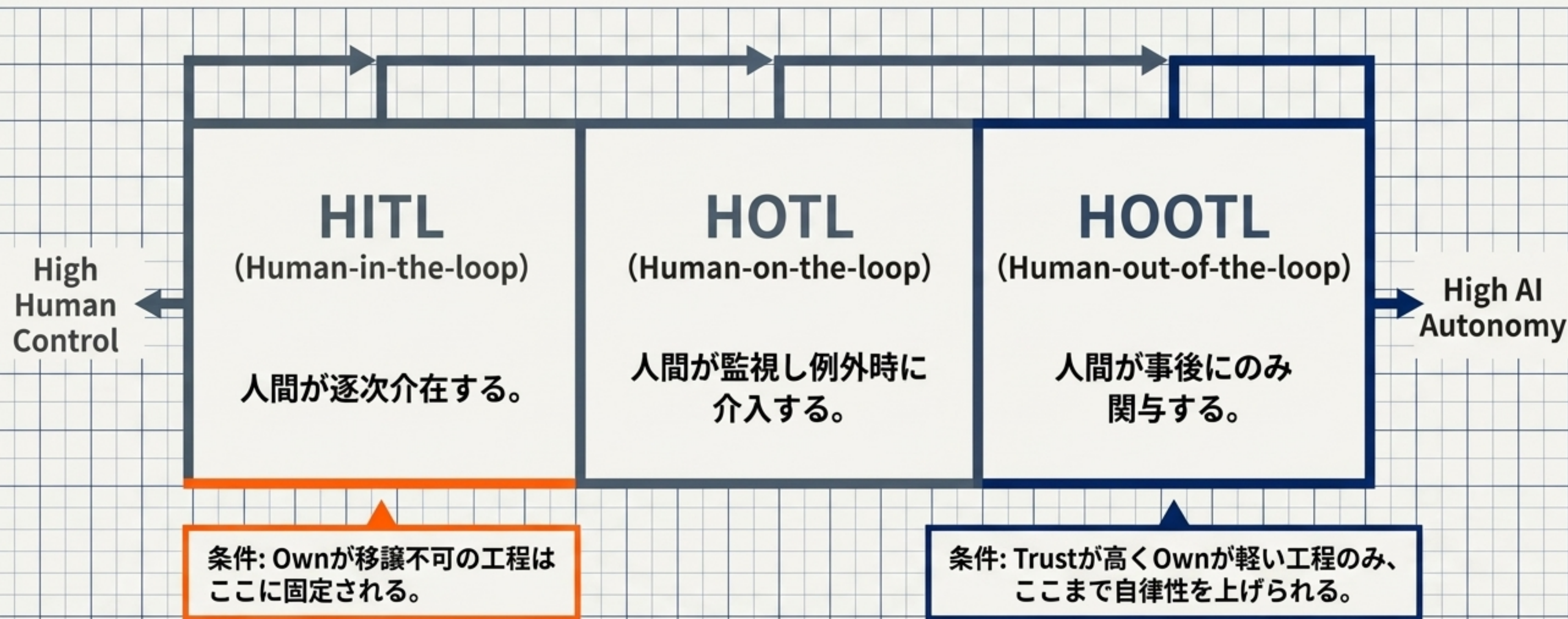
知財判断

「任せられる範囲」を分解する3つの評価軸

重要な原則：「Can（能力）が高い＝任せてよい」ではない。3軸は完全に独立して評価される。



評価軸と「自律性の段階区分」の接続



フロンティアAIの安全性議論（能力向上と監督維持の分離）と、知財実務の自律性設計は完全に軌を一にする。

業務設計マスターブループリント：3つの委任ゾーン

ゾーン (Zone)	三軸プロファイル (3-Axis Profile)	設計原則 (Design Principle)
Zone A: Delegate (委任可能)	Can 高・Trust 高・Own 軽	AIを主担当に据え、人間はサンプリング検証に回る
Zone B: Augment (協働必須)	Can 高・Trust 中・Own 重	AIの役割を結論ではなく「選択肢の生成と論点の可視化」に限定する
Zone C: Reserve (人間専管)	Own 移譲不可 (ゼロ)	AI出力を一次資料にしない。機微情報を切り離れた一般論の壁打ちに留める

Zone A: 委任可能 (Delegate)

Profile

A 委任可能 (Delegate)

Core Logic

検証コストが低く、誤りを事後的に修正でき、最終責任が人間に残る領域。

高度な補助ではなく「主担当」としてAIを据える。

Tasks & Application

対象工程

- ・ 公開特許文献の要約・分類、技術トレンドの一次抽出
- ・ 先行技術調査の検索式の初期生成、ヒット予備スクリーニング
- ・ 明細書ドラフト初稿、形式チェック、引用整合

Case Study (オムロン「ALBERT」)

Amazon Bedrock上のClaudeを活用。文献調査や情報整理においてAIを主担当とし、人間がサンプリング検証に回る合理的な設計実例。

Zone B: 協働必須 (Augment)

Profile

B

協働必須
(Augment)

Core Logic

競争力に直結し、誤り検出に
専門判断を要する領域。出力要
件として「根拠の追跡可能性」
を必ず組み込む。

Tasks & Application

対象工程

- FTO評価：抽出はAI、侵害該当性の「結論」は人間
- 新規性・進歩性の予備判断：論理の叩き台はAI、引用文献の「認定」は人間
- クレーム文言の戦略的設計：案出しはAI、権利範囲の「意思決定」は人間

Case Study（三井化学「OCSR」）

化学構造をAIが自律抽出し、その結果に基づく有効性判断を人間が行う。労働集約的工程と結論判断の完璧な切り分け。

Zone C: 人間専管 (Reserve)

Core Logic

Profile

法的・戦略的に責任が人間に固定される。CanやTrustが
どれほど高くても、設計上「委任」を完全に排除する。

Tasks & Application

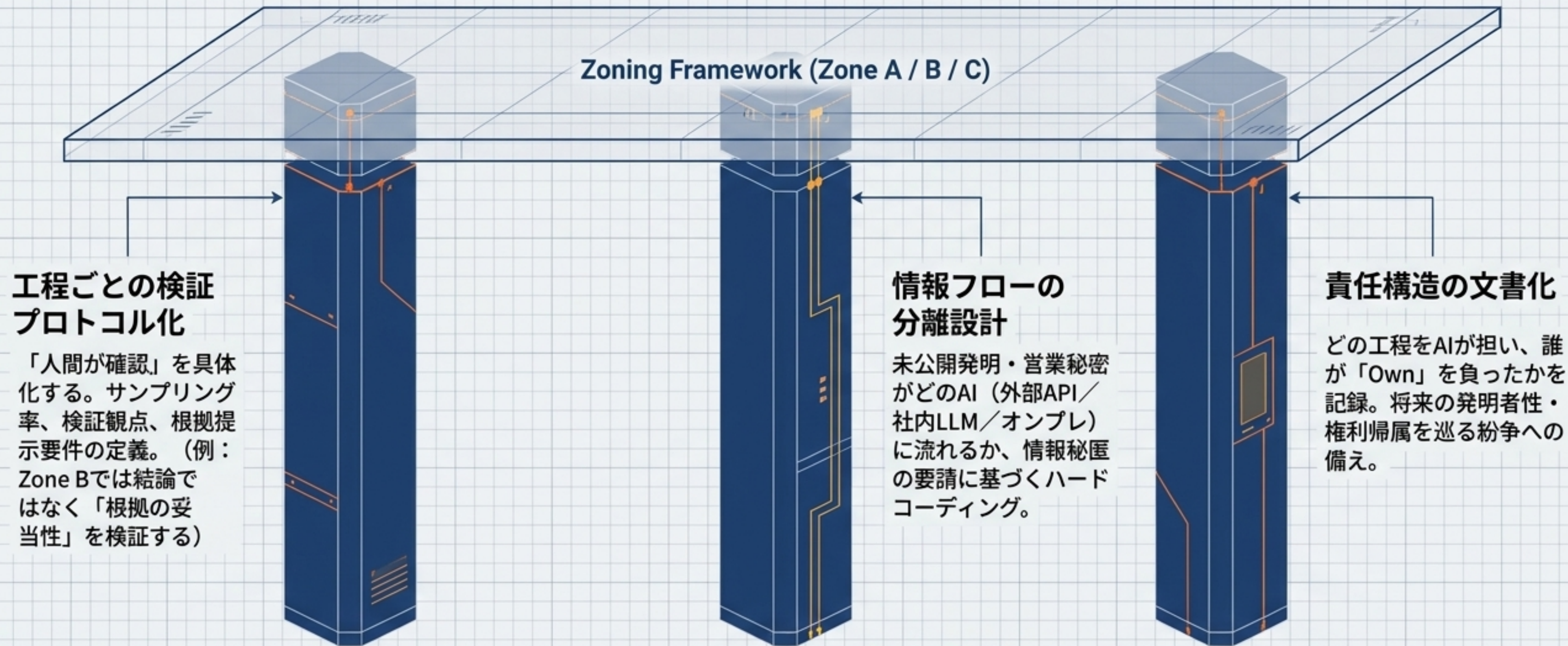
- 🔒 - 出願／非出願・営業秘密化の最終意思決定
- 🔒 - 発明者認定、権利帰属の契約判断
- 🔒 - 未公開発明の機微情報を含む判断
- 🔒 - 代理人・責任者としての署名行為

Design Rule

情報を外部AIに渡すこと自体がリスク。AIを利用する場合でも、
機微情報を完全に切り離れた「一般論の壁打ち」に限定する。

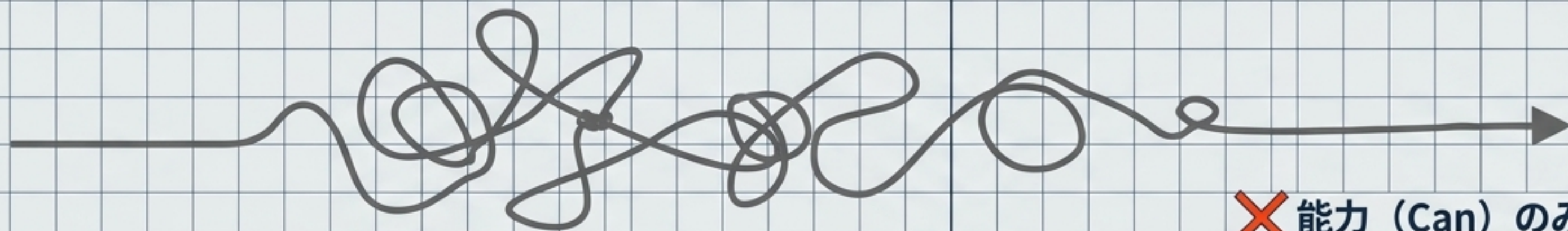
設計を支える3つの運用インフラ

ゾーニングは静的な分類ではなく、以下の運用メカニズムで支えなければ機能しない。



抽象論から、実装可能な「業務設計」へ

AIの能力（Can）の高さに引きずられて委任範囲を広げてはならない。



✕ 能力（Can）のみに依存した丸投げ



Can / Trust / Own に基づく精緻な範囲設計

知財部門の競争力は、「AIに何を任せるか」によって決まるのではない。

Can・Trust・Ownの分解を通じて、『任せられる範囲をどれだけ精緻に設計できるか』によって決まる。