

2025年 米国知財実務の変革

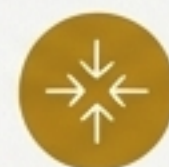
規律の収斂と司法の分断：生成AI時代のリスクと戦略



2025年：実験から実装、そして規律の年へ

2025年、米国の知財実務における生成AI利用は、無秩序な実験期を終え、「規律と実装」のフェーズに移行しました。法務部門の導入率は80%に達し、もはやAIは特殊なツールではなく、業務基盤そのものです。

この新たな環境は、2つの相反する力によって定義されます。



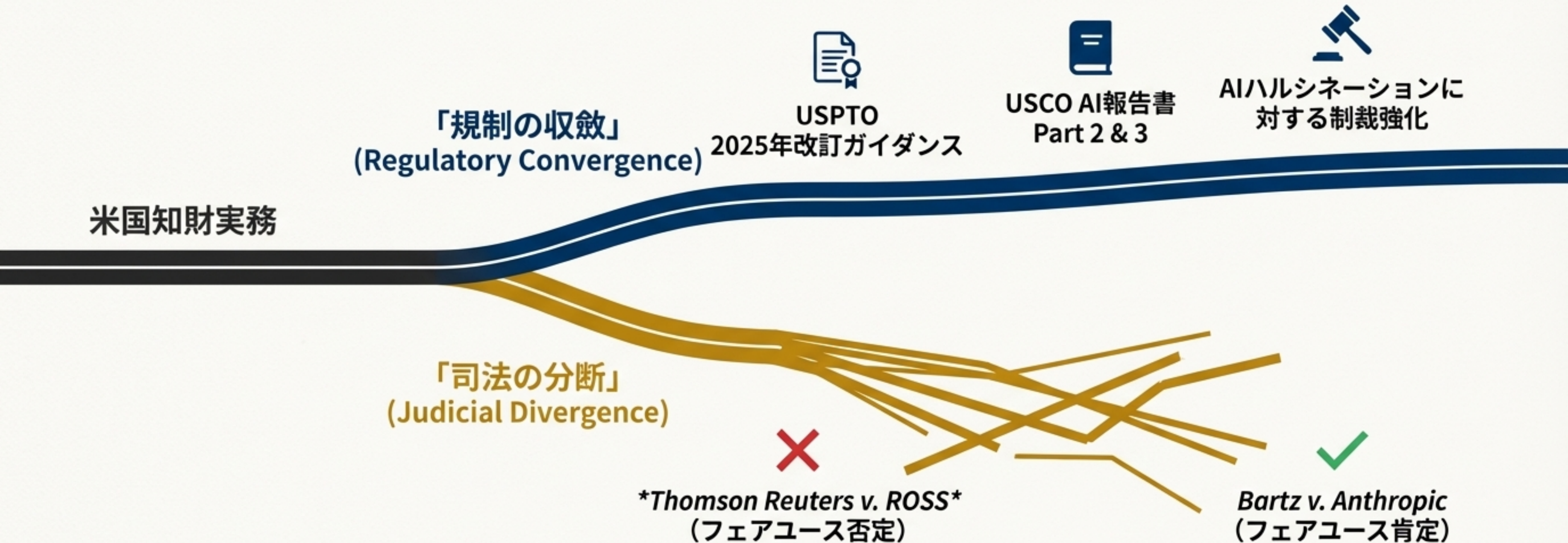
規制の収斂 (Convergence)：USPTOとUSCOは、「人間の創造性の優越」という原則に立ち返り、明確で予測可能なルールを整備しています。



司法の分断 (Divergence)：一方、連邦裁判所ではAI学習データのフェアユースを巡る判断が分裂し、深刻な法的・事業リスクを生み出しています。

本報告書は、この二極化するランドスケープを航海するための戦略的インサイトを提供します。

米国知財の2つの現実：収斂する規制、分断する司法



この2つの異なる現実を理解することが、2026年以降の戦略の鍵となります。

特許実務のパラダイムシフト： USPTO 2025年改訂ガイダンス

2025年11月28日、USPTOはAI支援発明に関する発明者シップガイダンスを抜本的に改訂。これにより、長らく実務家を悩ませてきた曖昧さが解消されました。



主要な変更点

- 撤回：2024年2月の旧ガイダンスを全面的に撤回。
- 核心：AIと人間の貢献度を比較する「**Pannu要素**」の適用を完全に廃止。
 - 理由：Pannuテストは本来、複数の「自然人」間の貢献度を測るものであり、自然人でないAIと関係に適用するのは論理的矛盾。AIを暗黙的に「共同発明者」として扱ってしまう問題を修正。
- 結果：審査の焦点は「人間が単独で、あるいは他の人間と共同で、発明の**着想（Conception）**を完了したか」という一点に集約されました。

発明者シップの新原則：「人間の着想」の絶対化



着想 (Conception)

1. AIは発明者になり得ない

AIシステムは法的に発明者や共同発明者にはなれないと明言。

2. 「着想」が全て

特許法の伝統的定義を厳格に適用。
「発明者の心の中で、完全かつ作用する発明の、明確かつ恒久的なアイデアが形成されること」が着想である。人間の精神的活動 (Mental Act) が介在しない限り、特許権は発生しない。

3. AIは単なる「ツール」

AIを「高度な実験機器」や「研究データベース」と同列のツールとして再定義。AIの出力が高度であっても、その有用性を認識し、特定の課題解決手段として採用・定義づけた自然人が発明者となる。

実務上の対応：厳格化するドキュメンテーションとグローバル戦略

特許取得のハードルは下がった一方、プロセス管理の重要性が増大。将来の無効審判や訴訟に備える必要があります。

アクションプラン



1. 発明届出書（Invention Disclosure Form）の再設計

- 発明者欄には必ず「自然人のみ」を記載。
- 人間の「課題設定」「指示（プロンプト）」「選択・修正」のプロセスを詳細に記録。
- AIとの対話ログや修正履歴を「着想の証拠」として保全。



2. グローバル出願における不整合リスクの管理

警告：一部の国（例：南アフリカ）でAIを発明者として認める可能性があるが、USPTOはAIを発明者として記載した外国出願に基づく優先権主張を受け入れない方針を示唆。

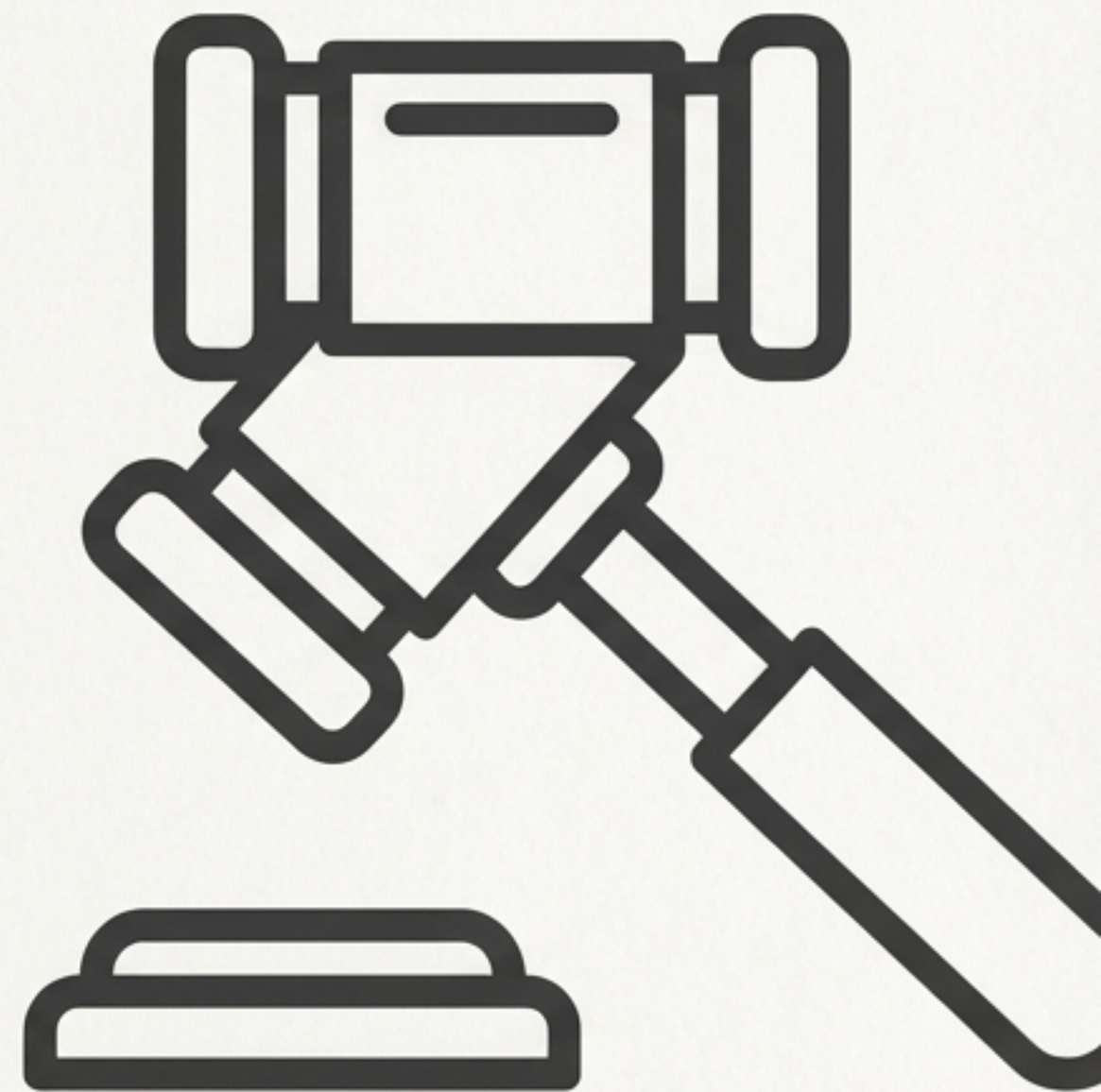
戦略：すべての法域で「自然人のみを発明者とする」戦略を徹底。不整合は米国での権利化不能や優先権喪失のリスクに直結する。

倫理の新基準：AIハルシネーションに対する司法の鉄槌

AIが捏造した架空判例の引用は、もはや「不注意」では済まされない重大な倫理違反です。連邦裁判所は、AI生成物の検証責任を怠った弁護士に極めて厳しい制裁を科しています。

ケーススタディ：Johnson v. Dunn (アラバマ州北部地裁)

- 事案： 経験豊富な弁護士3名が、ChatGPTが生成した5つの架空判例を含む申立書を提出。
- 判決： 裁判所はこれを「悪意に匹敵する (tantamount to bad faith)」行為と認定。
- 制裁：
 - 公的な譴責
 - 弁護士会への懲戒請求
 - 全クライアント・同僚・担当裁判官への制裁命令の送付
- 示唆： AIの使用自体ではなく、検証の欠如が「**キャリアを脅かす**」レベルのリスクとなることが明確化。連邦民事訴訟規則11条違反が確立されました。



フェアユース論争：USCOによる境界線の提示と司法の対立

USCOの行政的立場（2025年 AI報告書）

著作物性（Part 2）：

- 「**人間の著作権性**（Human Authorship）」を必須要件とし、AI生成物そのものの**著作権は否定**。
- プロンプトは「**アイデア**」に過ぎず、著作者とはみなされない。
- 人間による創造的な「**選択・配置**」がなされたハイブリッド作品のみ登録の道を残す。

AIトレーニング（Part 3）：

- 強制許諾制度の導入には慎重な姿勢。
- AI企業が主張する「AIトレーニングは常に**フェアユース**」という説に強い懸念を表明。特に「**市場代替**」のリスクを警告。



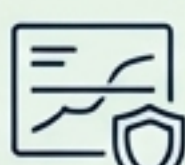
核心：



USCOの懸念は現実となり、連邦地裁レベルで判断が真っ二つに割れる事態に発展しました。



フェアユースの分岐点：ROSS事件 vs Bartz事件 判決比較

Thomson Reuters v. ROSS デラウェア地裁（2025年）	Bartz v. Anthropic カリフォルニア北部地裁（2025年）
 判決（Outcome） フェアユース否定 （原告勝訴）	 判決（Outcome） フェアユース肯定 （学習利用に関して）
 第1要素：利用の目的と性質 変革性なし・商業利用 競合するWestlawの要約を学習し、代替品を開発する行為。	 第1要素：利用の目的と性質 「劇的に」変革的 人間が本を読んで学ぶプロセスに類似しており、元の書籍とは全く異なる目的を持つと評価。
 第4要素：市場への影響 市場阻害あり Westlawの「直接的な競合相手」であり、市場シェアを奪う可能性が高いと判断。	 第4要素：市場への影響 直接競合なし AIシステムと原告の書籍は直接競合しない。 *ただし海賊版データの利用は否定。⚠

KEY TAKEAWAY

判決の決定的な違いは「**競合性**」にあります。

- **ROSS事件**: 「競合製品を作るための学習」はフェアユースが否定される。
- **Bartz事件**: 生成AIそのものが著作物の市場を直接代替するわけではないため、学習行為自体は変革的と認められる。

実務の構造転換：エージェンティックAIワークフローへの移行

2025年の最大の変化は、受動的な「対話型AI」から、目標達成のために自律的に行動する「エージェンティックAI」への移行です。

伝統的なプロセス



エージェンティックAI導入



人間の役割の変化

従来 (Before)

作業員 (Doer)

現在 (Now)

監督者 (Supervisor) - AIが生成した成果物をレビューし、戦略的判断に集中する。

特許実務AIツール比較：最適なユースケースの選定

市場には特許実務に特化した高度なAIツールが登場しており、自社のワークフローに最適なものを選択することが重要です。

ツール名	主要機能と特徴	最適なユースケース
Solve Intelligence 	ブラウザベースの「特許コパイロット」。発明届出書から明細書、図面まで一貫して生成。法域別カスタマイズ対応。	ドラフティングから中間処理まで一貫したAI支援を求める場合。
DeepIP 	MS Word完全統合プラグイン。既存フローを維持しつつクレーム推敲や実施例拡充を支援。SOC2認証を重視。	既存のWordベースの業務フローを維持しつつ生産性を向上させたい場合。
Rowan Patents 	エージェント型機能。用語の整合性チェック、図面番号の自動同期、実施例の自動拡張を自律的に実行。	複雑な明細書の整合性管理や形式的ミスの撲滅を重視する場合。
Pinch Patent Drawings 	CADデータからUSPTO/EPO準拠のベクター図面を数分で自動生成。図面作成コストを劇的に削減。	図面作成のアウトソーシングコストと時間を削減したい場合。

組織的リスク：「シャドーAI」の脅威とガバナンス戦略



問題定義：組織的な導入が進む一方、弁護士が個人で無料版ChatGPT等を業務利用する「シャドーAI (Shadow AI)」が深刻なセキュリティホールとなっています。

リスク：無料版ツールの多くは入力データを再学習に利用。未公開の特許情報やM&Aの機密情報が競合他社への出力として流出する危険性。

企業の対抗策：2つの柱



1. 安全な代替手段の提供 (Walled Garden)

OpenAI EnterpriseやMicrosoft Copilotなど、データが学習に利用されない安全な環境を全従業員に提供する。



2. ポリシー策定と教育

「どのデータなら入力して良いか」を明確にするデータ分類ポリシーを策定。シャドーAIの危険性を継続的に教育し、IT部門と法務部門が連携してガバナンスを徹底する。

戦略的要約：収斂と分断の世界で求められる3つの新資質

- **収斂の世界（特許・倫理）**：AIは発明者にはなれないが、最強の「発明支援ツール」として定着。ルールは明確化された。
- **分断の世界（著作権）**：AI学習データの法的リスクは依然として高い。特に競合製品を作るための利用（ROSS型）や海賊版データの利用（Anthropic型）は致命的リスク。

知財プロフェッショナルの新たな役割：



1. アーキビスト (The Archivist)

「人間による着想」を証明するための詳細な記録保持能力。



2. リスクマネージャー (The Risk Manager)

著作権の司法的判断の分裂がもたらす不確実性をナビゲートする能力。



3. AIスーパーバイザー (The AI Supervisor)

エージェンティックAIの出力を監督・検証し、戦略的な指示を与える能力。

2026年への展望：躊躇の時は終わった

実験的な導入期は終わり、今は「いかに安全に、効果的に使いこなすか」という最適化のフェーズにある。

もはやAI導入を躊躇する選択肢はない。

2026年に向けた唯一の戦略的選択肢は、

「倫理的かつ安全なAI活用基盤」

を構築することである。