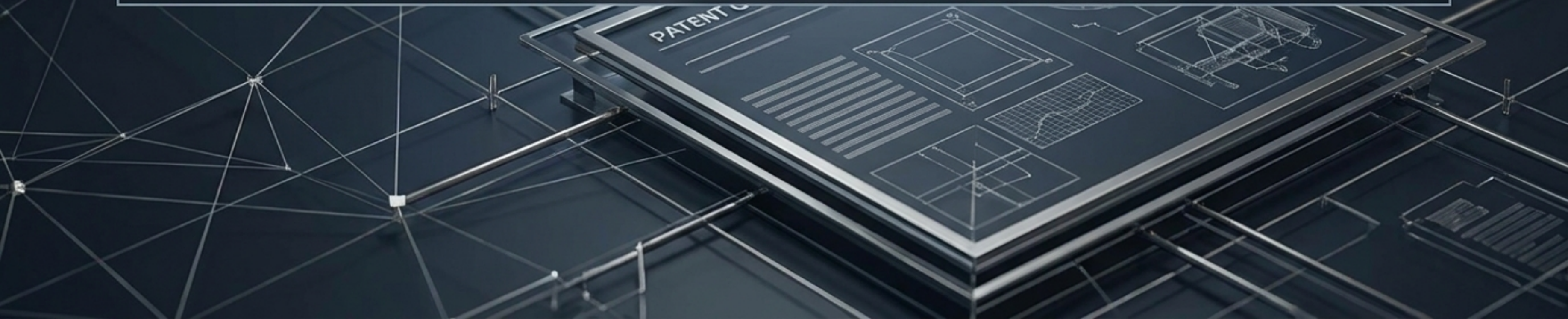


IP5特許審査におけるAI実装の 最前線と出願戦略の再構築

日米欧中韓の独自ツール動向・ガバナンス評価と、
次世代ドラフティング実務への提言





実務展開の拡大

全庁共通のトレンド。
AIはPoC（概念実証）を終え、
「先行技術探索・分類・翻訳
・手続文書作成」の最前線
ワークフローへ完全に実装。



ガバナンスの二極化

人間責任を明示し制度化を
急ぐJPO/EPO/USPTOと、
内部運用モデルが非公開で
ブラックボックス化する
CNIPA/KIPO。
スタンスの明確な分断。



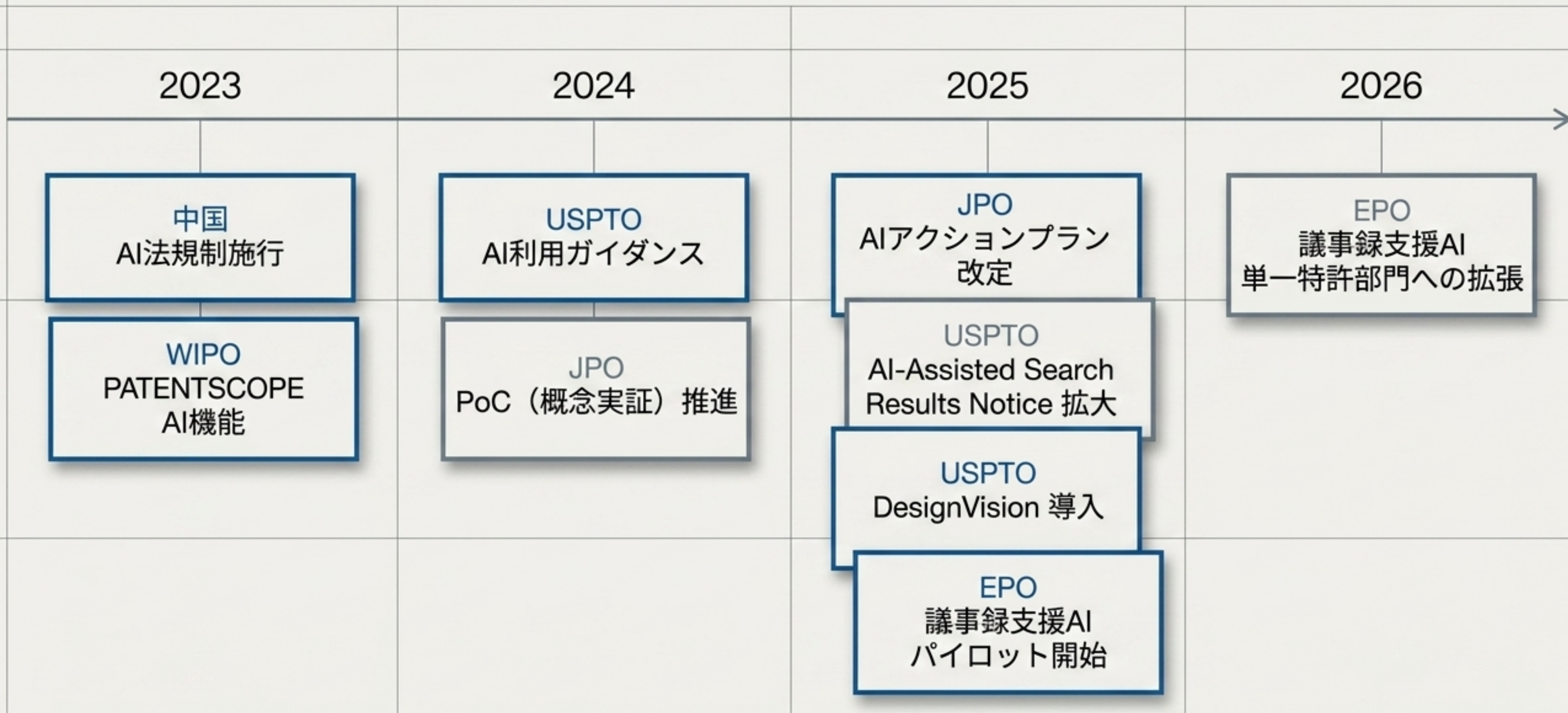
出願実務の転換

AI検索網を逆算した
「階層的クレーム設計」と、
出願前の「AI自己監査」が、
今後の特許成立率と権利の強さを
左右する不可避の要件へ。

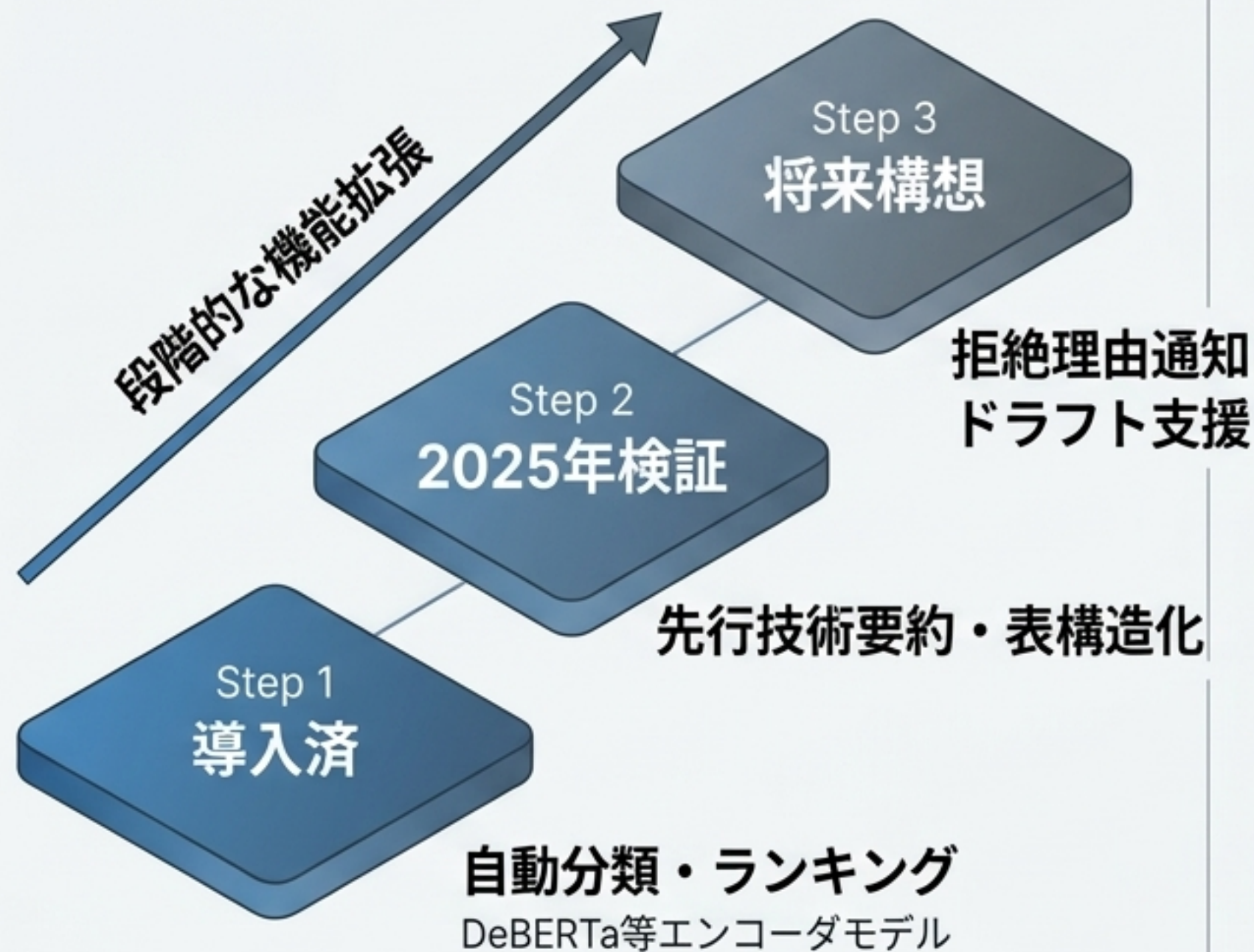
主要五庁 AI導入・ガバナンス Heatmap Matrix

	導入状況	主要機能	技術アプローチ	ガバナンス評価	主要課題
JPO (日本)	導入済+PoC継続	分類, 検索式支援 生成AI要約(検討)	ML, 生成AI, ハイブリッド	中〜高	LLM適合性, 誤引用抑制
USPTO (米国)	導入済多数	PE2E Search AI, DesignVision	CV, NLP, 生成AI試験環境	中	外部生成AIの 業務利用統制
EPO (欧州)	導入済+制度化 明確	Patent Translate, AI議事録支援	NMT, 生成AI補助	高	手続保障と透明 性の両立
CNIPA (中国)	一次資料詳細不 不均一	智能検索・審査 の示唆	ML, LLM活用の 可能性	中〜低	内部ツールの不 透明さ
KIPO (韓国)	一次資料詳細不 不均一	翻訳基盤(K2E- PAT), WIPO連携	NMT, 検索支援	中〜低	公開情報・効果 指標の不足

Implementation Timeline (2023–2026)



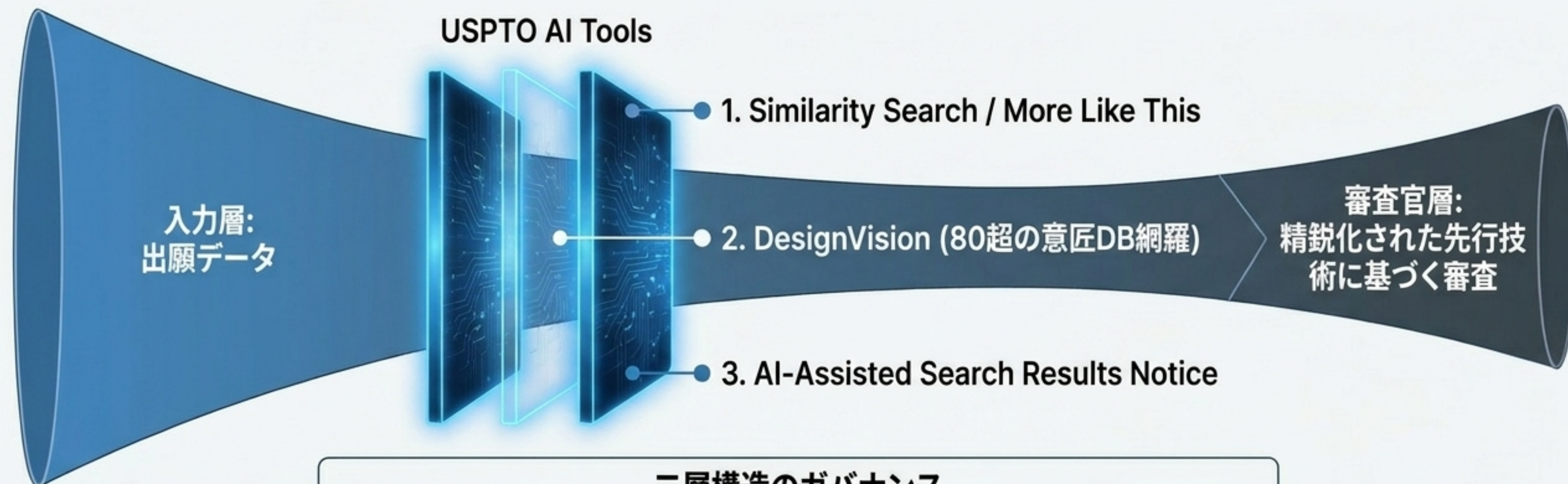
Spotlight - JPO (Japan) | 段階的かつ実証的なAI実装



厳格なハイブリッド評価

- 単なる生成AI導入ではなく、実務適合性を厳密に比較検証。
- クラウド型LLM vs オンプレミスモデル
- プロンプト vs ファインチューニング
- 戦略的意味: 初期探索の母集団構築を最速化し、見落としリスクを大幅低減。

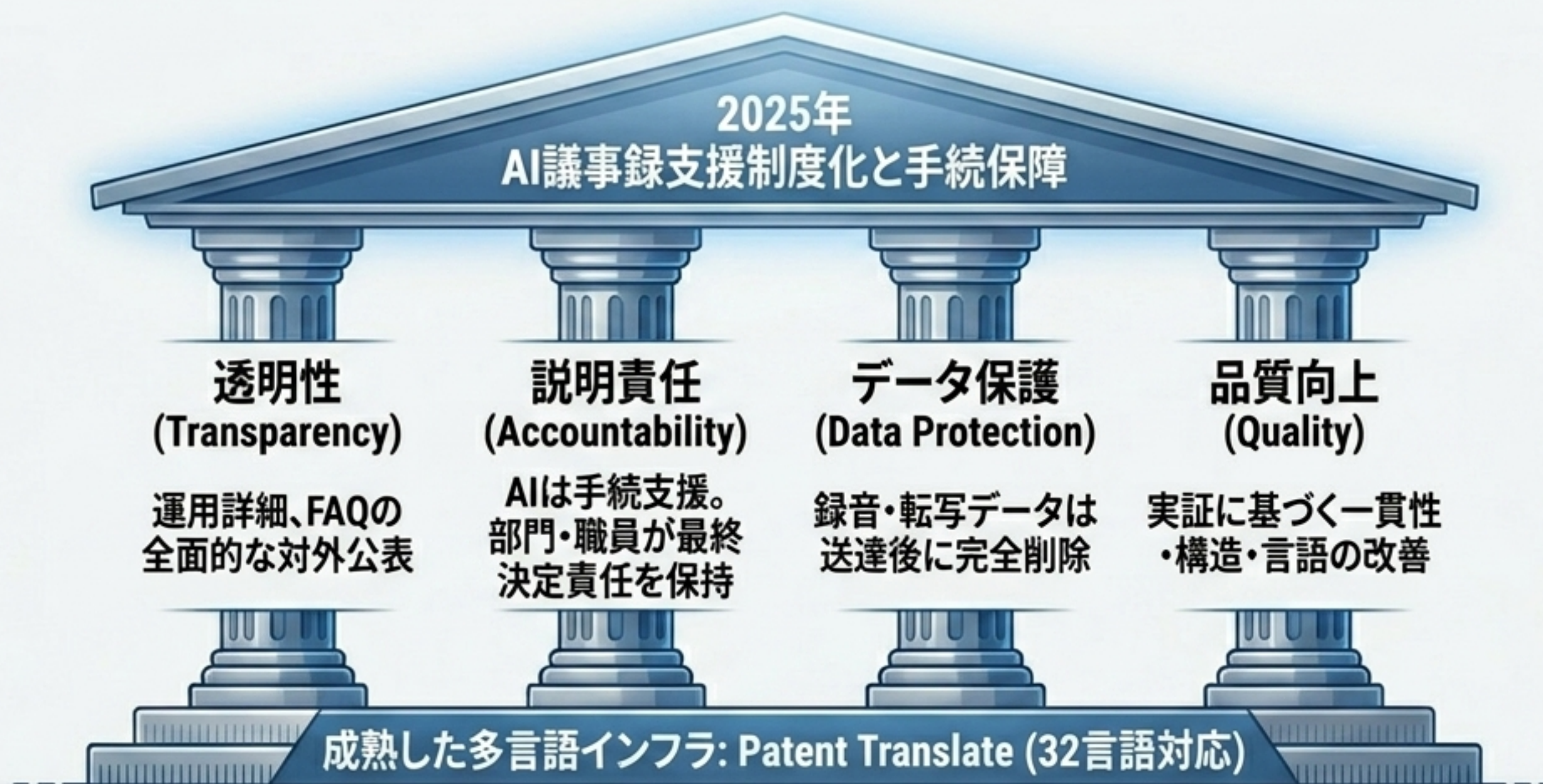
Spotlight - USPTO (US) | 審査ワークフローのフロントエンド変革



二層構造のガバナンス

積極推進	厳格統制
PE2E Search等の検索用機械学習モデルをフロントエンドへ統合	外部生成AIの自由利用は禁止。機密性を担保するAI Lab環境でのみ検証

Spotlight - EPO (Europe) | 透明性とガバナンスの世界標準



Key Takeaway: EPOのガバナンス設計は、他庁が追随すべき「世界標準モデル」として機能している。

Spotlight - CNIPA & KIPO | 巨大出願と「見えない」AIインフラ

公開確認済みの制度情報(水面上)

CNIPA: 2023年AI法規制、審査指針改訂。倫理・法令適合性による絞り込み(顔認証・差別的推論の制限)。

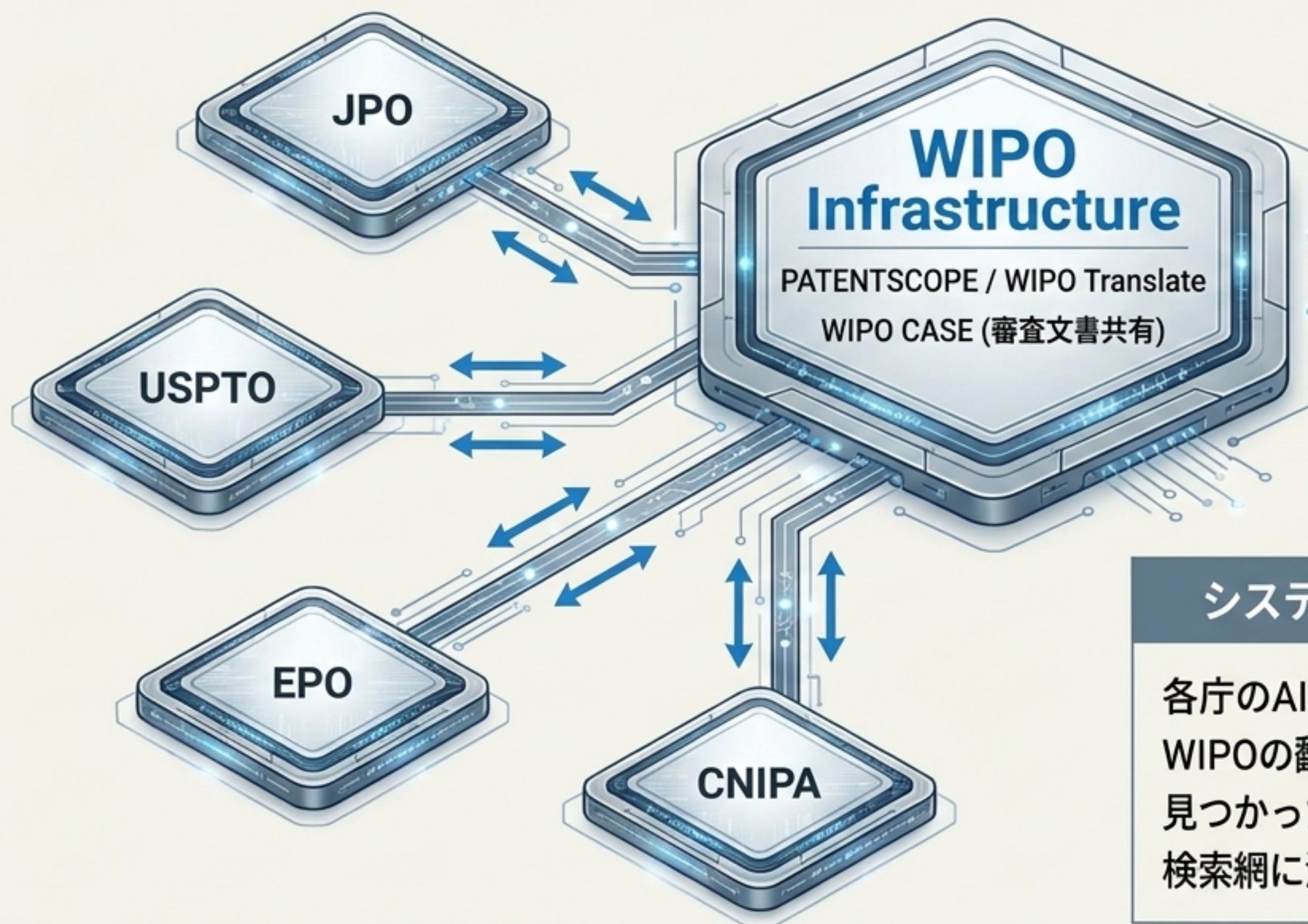
KIPO: 高度な翻訳基盤(K2E-PAT)と強固なWIPO連携。

ブラックボックス化された運用実態(水面下)

内部の「智能検索・智能審査」ツールの具体名やベンダー情報が著しく限定的。

巨大な出願件数进行处理するためのインフラは確実に稼働しており、出願人には見えない形で検索高度化が前提となる。

The Global Network | WIPOと国際協調インフラ



システムの波及効果（Network Effect）

各庁のAI導入は「国内の効率化」に留まらない。WIPOの翻訳・連携基盤を通じ、ある国のAI検索で見つかった先行技術が、即座に他国の審査官のAI検索網に波及するエコシステムが完成しつつある。

Legal & Ethical Divergence | 制度と倫理のスペクトラム

技術・手続重視
(Strict Technical/Procedural)

倫理・法規制重視
(Broad Ethical/Regulatory)



人間中心と義務。

AIは発明者になり得ず、代理人に対する
正確性確認と秘密保持の義務を徹底。

適正手続と透明性。

人間の最終責任と、データ削除による
手続き上の確実性を担保。

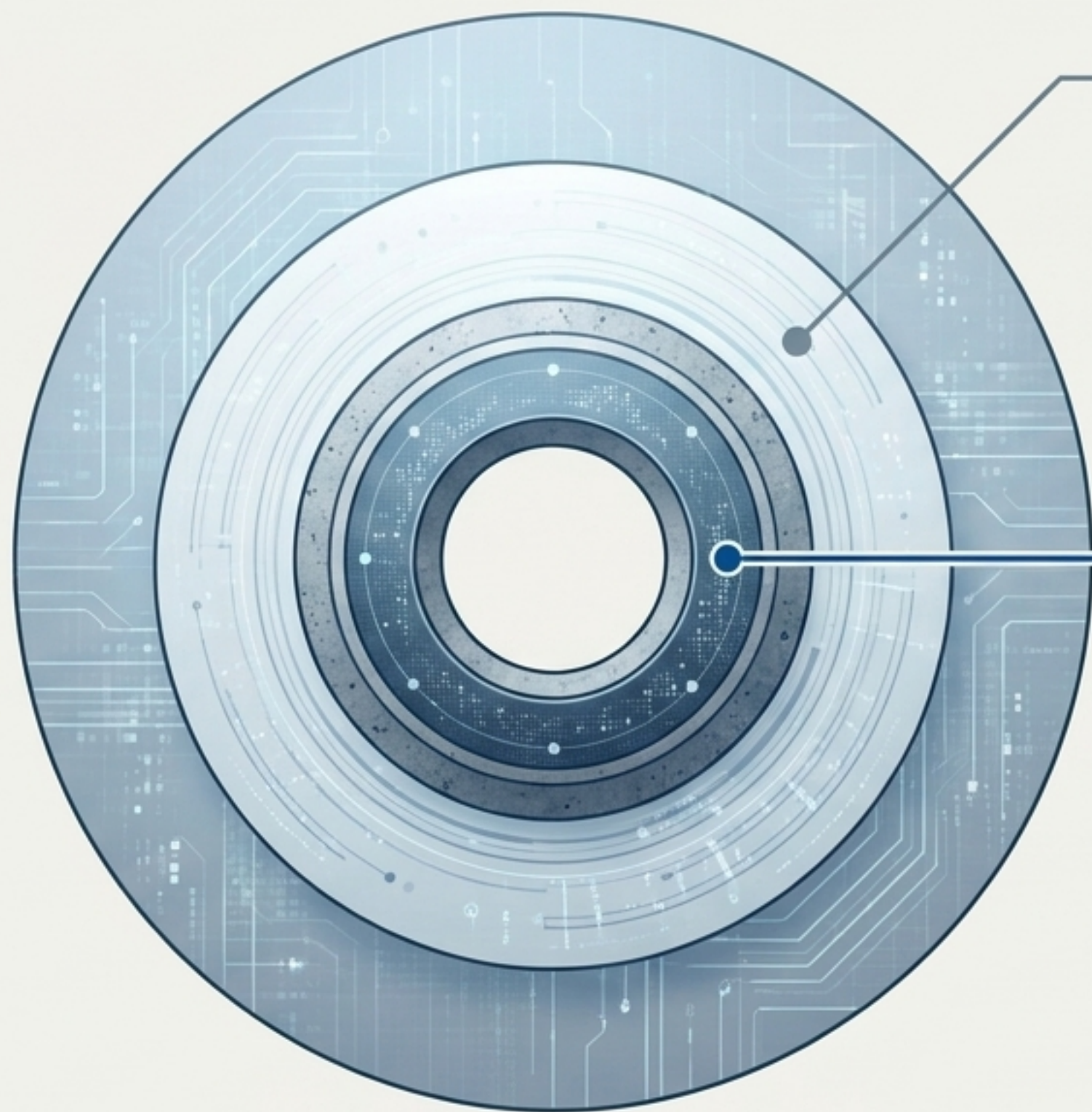
倫理と公序良俗。

技術性に加え、個人情報の無断取得や差別的
モデルなど倫理的逸脱を特許性から排除。

Strategic Insight:

「技術的特徴」の記述にとどまらず、「AI利用の適法性・倫理性・データ由来」
までを出願書類上で国別にコントロールする時代へ。

Strategic Imperative 1 | クレームの「階層的ドラフティング」



Layer 1: AIの守備範囲 (AI Search Capture Zone)

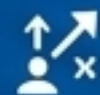
広い概念語、用途、データ前処理、モデル入出力



対策: AIによる後知恵拒絶 (Hindsight bias) の標的。網羅性は確保しつつも特許性の抛り所にはしない。

Layer 2: 人間の判断領域 (Human Differentiator Zone)

具体的パラメータ、微細な構成関係、制御条件、特有の作用効果



対策: AIが意味論的差異を見落としす層。従属項や明細書に厚く記述し、人間審査官に対する『技術的差別化』の絶対的防波堤とする。

Strategic Imperative 2 | 出願プロセス・体制の再構築



1. AI事前自己監査 (Pre-Filing Self-Audit)

各庁の多言語・画像検索AIを逆算し、同水準のAIを用いて自他庁文献・NPL・翻訳揺れまで先回りして潰す。



2. データ起源・秘密保持管理 (Data Origin)

外部LLMへの入力履歴、学習データの適法性をトラッキング。明細書で「適法なデータ取得」を記述可能にする。



3. 審査ナレッジの社内AI基盤化 (Knowledge Base)

庁側の要約・検索力向上に対抗するため、自社側の過去の引用文献・面接記録を社内AIでナレッジ化し対抗する。

AI実装は「実証」から「実務と法整備」のフェーズへ移行した。 各庁の"AIのクセ"を読み切る者だけが、次世代の強靱な特許網 を構築できる。

■ **Governance Matrix:**
制度化の進む欧米日と、
不透明な中韓。
審査インフラの透明度に
応じたりスクヘッジが必
須。

■ **Drafting Evolution:**
概念と具体の「階層的
ドラフティング」が、AI
検索網を突破する唯一の
最適解。

■ **Proactive Defense:**
庁のAI導入を待つのではなく、
出願人側が「事前自己
監査・データ由来管理」の
防御壁を構築すべきタイミ
ング。