

知財翻訳における生成AIの進化：2025年から2026年への構造転換

2025年：探索的導入と「人手検証」の限界



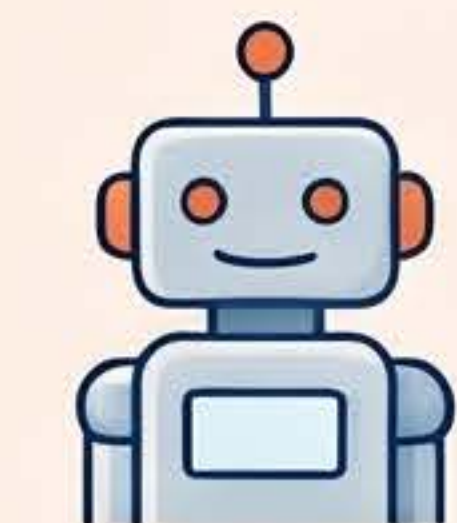
「手戻りの逆転現象」による負担増

AIの渡帳な誤訳を確認するため、最初から人手で訳す以上の工数が発生した。



セキュリティ懸念による導入の足踏み

未公開出願書類をクラウドへ送信するリスクが、大手製造業での導入障壁となった。



単一プロンプトによる一括翻訳

汎用LLMを用いた一括処理が中心で、特許特有の機微な文脈維持に課題があった。



2026年：自律検証型エージェントと実務インフラの確立



修正必要率 0.05%の達成

エージェント型AIが自律的に推論・修正を行うことで、驚異的な精度を実現した。



妥当性根拠の明示 (ホワイトボックス化)

AIが訳語を選択した論理的背景を示すことで、人間の検証負担を大幅に削減。



スタンドアロン型 ローカルLLM

物理的にネット隔離された環境で動作し、機密情報の漏洩リスクを完全に遮断。

2025年 実務モデル



検証手法

人手による全件
目視修正(MTPE)



セキュリティ

クラウドSaaS
への依存



人間の役割

誤訳を監す
「作業員」

2026年 実務モデル



検証手法

AIによる自動検
証と根拠提示



セキュリティ

ローカルLLMに
よる物理的遮断



人間の役割

プロセスを統制
する「責任者」