

知財AIの実力検証：PatentBenchが示す専用AIの優位性

知財調査の「壁」と汎用AIの限界



特許調査は退屈な業務

100件以上の特許公報を精査するには数日～数週間を要し、総務者でも重複文献を発見しずリスクが常に存在する。



汎用AIは特許実務に不向き

一般テキストで学習した汎用LLMは、特許特有の専門用語やクレーム意図を正確に理解せず、誤回答は深刻な法的リスクに連鎖しかわない。

公平で実践的な評価設計



PatentBenchとは？

知財実務に特化した全特許の特許クレーム・言語のAIの特許調査能力を体系的かつ客観的に評価するためにPatentBenchが誕生。



公平で実践的な評価設計

約840件の多様な特許クレーム・言語（英語60%、中国語22%）の最新特許クレームをテストデータとし、AIが特許調査をしない特許の課題で性能を測定する。



評価指標1：Xヒット率

約840件の多様な特許クレーム・言語（英語60%、中国語22%）の最新特許クレームをテストデータとし、AIが特許調査をしない特許の課題で性能を測定する。



評価指標1：Xヒット率

調査対象のうち、少なくとも1件の最重要先行文献（X文献）を発見できた特許の割合。「有用な文献を見つけ出す能力」を示す。



評価指標2：Xリコール率

正確となる全文文献のうち、AIが発見できた文献の割合。「どれだけ機能的に（漏れなく）扱えたか」を示す。

衝撃のベンチマーク結果：専用AI vs 汎用AI

専用AI「Patsnap Eureka」が汎用AIを圧倒

Patsnap Eureka（専用AI）

Xヒット率
（有効文献の発見率）

Xリコール率
（網羅率）

81%

36%

新規性調査において、Patsnap Eurekaは有効文献の発見率（ヒット率）と網羅性（リコール率）の両方で汎用LLMを大きく上回った。

汎用AI （ChatGPT-o3, DeepSeek-R1）

ChatGPT-o3

32%

11%

Xヒット率

Xリコール率

DeepSeek-R1

9%

3%

Xヒット率

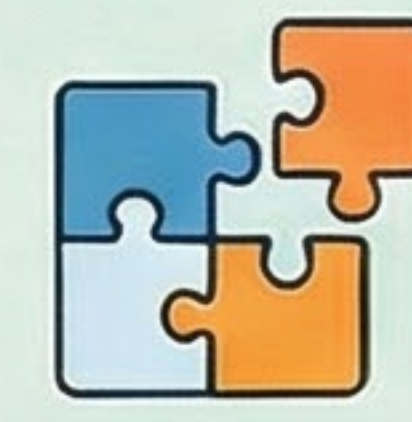
Xリコール率

高性能の秘密：Patsnap Eurekaの仕組み



技術的特徴の抽出

入力された発明内容からキーワードを抽出し、特許の土台を構築する。



クレーム要素の分解

特許クレームを構成要素ごとに分解し、特許すべきポイントを明確化する。



複合検索式の自動生成

キーワード、特許分類、質問の構成要素などを組み合わせた高度な検索クエリを生成する。



検索とスクリーニング

世界中の特許・論文DBから関連文献を精査し、関連性に基づき優先度を評価する。



クレーム対応付け（マッピング）

先行文献のどの部分が特許クレームのどの部分と対応するかを特定し、対応関係を構築する。



比較表・レポートの作成

調査結果を先行特許との対比表や特許図の形で可視化し、専門家の判断を支援する。

実務導入への道：可能性と課題



導入のメリット

調査効率の飛躍的向上、グローバルな多言語検索への対応、専門家が集約加価値集約へ集中できる環境の実現。



導入時の課題

AIの精度を信頼できない信頼性の問題（リコール率36%）、判断相場の透明可能性、機密情報を守るセキュリティ対策が不可欠。



人間との協働が鍵

AIは専門家を代替するのではなく、強力なアシスタント。AIが下調べを行い、人間が最終判断と動議立案を担うという協働が重要になる。