

生成 AI に基づく進歩性判断の モデル化評価

会員 山崎 薫



要 約

本稿は、生成 AI を用いて「進歩性」の判断過程を実験的に検討したものである。生成 AI に対し、「進歩性」判断に関する思考手順を言語化した「厳格ルール」および「実行ステップ」を入力したところ、AI は一定の一貫性をもってその規則性を再現した。一方で、技術常識、動機づけ、実質的同一性といった、明示的な定義や形式的処理のみでは十分に扱いにくい要素については、AI の判断に限界も確認された。これらの結果は、「進歩性」判断のうち、どの部分がフレームワークとして整理しやすく、どの部分で追加的な検討を要するのかを示すものである。かかるフレームワークが一つの判断手法として受け入れられれば、出願人や弁理士はフレームワークに従って請求項の書き方を決定することができる。請求項の記述や進歩性の確保にあたって AI は出願人や弁理士の最終判断を支援する可能性を秘める。本稿は、その点を実験で示す試みである。

目次

1. はじめに
2. 仮説
 2. 1 AI による規則性の抽出と従来との限界
 2. 2 生成 AI の活用による核心的仮説
3. 実験方法
 3. 1 手順
 3. 2 プロンプト
 3. 3 注意事項
 3. 4 使用環境
4. 実験結果（その 1）
 4. 1 【請求項 1】の記述
 4. 2 対応する作用効果
 4. 3 1 件の従来技術（特許文献）
 4. 4 新規性
 4. 5 進歩性
5. 実験結果（その 2）
 5. 1 【請求項 1】の記述
 5. 2 対応する作用効果
 5. 3 1 件の従来技術（特許文献）
 5. 4 新規性
6. 考察
 6. 1 AI の判断傾向とロジックの限界
 6. 2 実務導入への提言と今後の展望
7. おわりに

1. はじめに

特許要件「進歩性（非自明性）」の判断は、特許出願の成否を左右する重要な判断の一つである。その判断は、

請求項の記載、引用文献の記載、作用効果の把握、並びに動機づけや示唆の有無の検討など、複数の要素を段階的に整理しながら行われる。

「進歩性」は法的枠組みの下で判断されるものの、個別事案ごとに検討対象となる技術内容や引用文献の組合せは相違することから、その判断過程は相応に複雑である。とりわけ、どの範囲までを文献に明示された事項として把握し、どの範囲で技術常識や論理的理解を用いるのかという点は、出願人や弁理士が拒絶理由への対応方針を検討する上でも重要である。実務上は、反論を重視すべきか、補正を優先すべきか、あるいはその双方をどのように組み合わせるべきかの見通しを立てにくい場面が少なくない。

本稿は、このような「進歩性」判断の過程のうち、どの部分が言語化されたフレームワークとして整理可能であり、生成 AI がどの範囲までその実行補助を担い得るのかを検証するものである。具体的には、「進歩性」判断の手順を「厳格ルール」および「実行ステップ」として明文化し、生成 AI に入力することで、AI がどの範囲まで一貫して判断を実行できるかを観察した。

本稿の目的は、生成 AI に「進歩性」判断のフレームワークを実装することで、請求項の記述や進歩性の確保にあたって生成 AI はどのくらい出願人や弁理士の最終判断を支援し得るのかを検討する点にある。

2. 仮説

2. 1 AI による規則性の抽出と従来の限界

AI は、与えられた膨大な学習データから現実の事象に潜む規則性（モデル）を抽出する能力に長けている。この習性を「進歩性」判断に適用すれば、「進歩性あり」として権利化された特許や、「進歩性なし」として拒絶された公報群を学習することで、AI は「進歩性」判断に潜む暗黙の規則性を見出すはずである。

従来の機械学習を用いたアプローチでは、特定の判断モデルを構築するために、学習データ（特許文献）の選定や特徴量エンジニアリング、学習アルゴリズムの設計（モデルの「作り方」）に多大な労力が注がれてきた。しかしながら、技術分野の広大さや判断基準の解釈の柔軟性から、汎用性と判断の安定性とを両立させることは困難であった。

2. 2 生成 AI の活用による核心的仮説

本稿が提示する核心的な仮説は、新たに特化された学習モデルを構築せずとも、一般に提供される高性能な生成 AI（大規模言語モデル [LLM]）を用いることで、「進歩性」判断の一貫性を確保し、予測性の確度を高めるという点にある。生成 AI は、従来のモデルとは異なり、その判断ロジック（規則性）が外部から入力されるプロンプトによって定義され、実行されるという特性を持つ。

仮説：

特許法で規定される「進歩性」判断の規則性（例：動機づけの有無、示唆の程度）をフレームワークとして明確に言語化し、これを生成 AI に実装した場合、生成 AI は、そのフレームワークに則った判断を一定の一貫性をもって再現することが可能である。フレームワークは生成 AI に入力されるプロンプトとして記述される。プロンプトの記述に応じて生成 AI は進歩性の有無を判断する。かかるフレームワークは決められた【請求項】の書き方や反論の焦点を要求する。このフレームワークが1つの判断手法として受け入れられれば、出願人や弁理士はフレームワークに従って【請求項】の書き方や反論の焦点を決定することができる。

本稿は、現時点で生成 AI によって確保される判断の正確性を追求するのではなく、プロンプトで定義された「進歩性」判断の『一貫性』に着目する。この一貫性ある「進歩性」判断は、【請求項】の記述や進歩性の確保にあたって出願人や弁理士の最終判断を支援するものとする。

3. 実験方法

3. 1 手順

次の（手順1）～（手順4）に従うと、生成 AI は「新規性なし」「進歩性なし」「進歩性あり」のいずれかを回答する。

（手順1）対象の特許出願の〔テキストデータ〕、並びに、1 件の従来技術（特許文献）の〔.pdf データ〕および〔テキストデータ〕を用意する。

（手順2）プロンプトの指定された位置に、【請求項 1】のテキストデータ、【発明の詳細な説明】（明細書）のテキストデータ、並びに、特許文献のテキストデータを転記する。

（手順3）プロンプト全文をコピーし、生成 AI のメッセージボックスに貼り付けるとともに、特許文献の〔.pdf データ〕をアップロードする。（AI は文章も参照しつつ図面を理解するので、〔.pdf データ〕で図面を入力することが推奨される）

（手順4）プロンプトを送信する。

3. 2 プロンプト

-----<ここから、プロンプトのフォーマット>-----

あなたは、技術的知識を備える特許庁審査官として振る舞ってください。

以下の個々の【厳格ルール】を絶対的な制約として遵守し、忠実に【実行ステップ 1】～【実行ステップ 5】に従って「進歩性は肯定される」または「進歩性は否定される」を回答してください。

【実行ステップ 1】

以下に示す【請求項 1】の構成要素ごとに、添付の公報 A の記述（図面を含む）と比較し、＜共通項＞を抽出せよ。明確化を目的に、以下に、公報 A のテキストを転載する。【請求項 1】の解釈にあたって【請求項 1】に記載される技術的定義以外、＜本件発明＞の明細書（発明の詳細な説明）中に記載されていてもその技術的定義を考慮してはならない。

【厳格ルール】

- 最初に【請求項 1】の全体像を理解し、構成要素同士の位置関係を考慮しながら公報 A 中で【請求項 1】の各構成要素を対応付けよ。
- 【実質的同一性の確保】：『技術常識』や『論理的な帰結』に基づき、最大限に共通項を探し出す強い意志を反映し、構成要素単体に対して、公報 A 中で対応する構成要素が同一の技術的意義を果たす場合には、その構成要素を＜共通項＞とみなさなければならない。
- 【請求項 1】に記載される構成要素の動作や機能が、公報 A（図面を含む）において物理的に可能であり、かつ一般的な使用の態様として常識的と判断される場合、その動作または機能は、＜従来技術＞の開示目的に関わらず、＜共通項＞に含めなければならない。【請求項 1】の各構成要素について、公報 A における対応構成を探し出す際、「選択的に束ねる」「選択的に接続する」などの動作表現が含まれていても、公報 A の構造上その動作が物理的に可能であれば、同一の技術的意義を有するものとして＜共通項＞に含めること。

ここで全ての構成が＜共通項＞に該当したら処理を終了し、「新規性は否定される」と回答せよ。【実行ステップ 2】以降を無視する。＜共通項＞以外の構成が残っていれば、【実行ステップ 2】開始時点で、【実行ステップ 1】で広義に抽出した結果を全てリセットし、【実行ステップ 2】以降では＜共通項＞以外の構成についてのみ処理を行う。

【実行ステップ 2】

先に示した【請求項 1】の構成要素ごとに、公報 A の記述（図面を含む）と比較し、先の＜共通項＞を排除しながら、技術的定義の＜差分＞を抽出せよ。【請求項 1】の解釈にあたって、【請求項 1】に記載される技術的定義以外、＜本件発明＞の明細書（発明の詳細な説明）中に記載されていてもその技術的定義を考慮してはならない。

【厳格ルール】

4. 【主観的解釈の排除】：『技術常識』や『論理的な帰結』『内在』『容易』『周知』『当業者なら』といった主観的解釈は厳重に禁止する。

【実行ステップ 3】

抽出した各差分から生み出される利用のベネフィット（効果）を＜本件発明＞の出願書類（明細書）の記述から探し出し、＜差分＞ごとに記述せよ。以下の本文内に＜本件発明＞の明細書のテキストを転載する。

【厳格ルール】

4. 【主観的解釈の排除】：『技術常識』や『論理的な帰結』『内在』『容易』『周知』『当業者なら』といった主観的解釈は厳重に禁止する。
5. 【効果の一体性原則】：効果と差分構成とを不可分で評価し、独立に効果を推論してはならない。

【再確認指示】

【実行ステップ 4】 および 【実行ステップ 5】 それぞれの完了後に、次の検証を必ず行え：

- (a) この判断に使った根拠は公報 A または客観的資料に明記されているか？
- (b) 『技術常識』または『物理的に可能』といった主観的補完を使っていないか？
- (c) もし曖昧なら、「根拠不明」と明示し、想定外の効果として扱え。
- (d) 根拠の明示がない項目があれば、次ステップで【強制フラグ】を発動せよ。

【実行ステップ 4】

【公知性の徹底確認】：抽出された各＜差分＞を単体で捉え、Google 検索ツールを用いて、＜差分＞の技術的定義を示す客観的資料を徹底的に探索せよ。その客観的資料から、単体の＜差分＞の技術的定義から生み出される利用のベネフィット（効果）を探し出せ。

【厳格ルール】

6. 【客観的根拠の限定】：ある事象が自明または予測可能であると判断する根拠は、公報 A または検索可能な客観的資料（教科書、ハンドブック、ウェブサイト等）の明確な記述に限定する。

【実行ステップ 5】

【進歩性判断】：客観的資料に利用のベネフィット（効果）が掲載されていたら、それを「予測される範囲内の効果」として扱え。客観的資料（公報 A または検索結果）に裏付けられないベネフィットのみを『想定外のベネフィット』として特定し、『想定外のベネフィット』があれば「進歩性は肯定される」と回答し、それ以外なら「進歩性は否定される」と回答せよ。

【強制フラグ】

根拠を引用できないときは、必ず次の一文を出力せよ：

『※客観的資料に該当箇所なし：想定外のベネフィットとして扱う。』

このフラグが出力されない限り、進歩性を否定してはならない。

【実行ステップ 5】 の【厳格ルール】

4. 【主観的解釈の排除】：『技術常識』や『論理的な帰結』『内在』『容易』『周知』『当業者なら』といった主観的解釈は厳重に禁止する。
7. 【想定外の効果の定義】：【請求項 1】の出願書類（明細書）に記載される効果は、上記客観的根拠に示されない限り、すべて『想定外のベネフィット』として取り扱う。
8. 【選択肢の予測可能性】：＜従来技術＞（公報 A）に、＜本件発明＞の【請求項 1】の構成要素に対応する上位

概念と、その下位概念または具体例が併記されている場合には、下位概念または具体例を本件発明で選択することは、＜従来技術＞の開示に基づく公知の選択肢の採用とみなし、原則として進歩性判断の動機づけを充足するものとする。

9. 【効果の厳密な対応】：【請求項 1】の明細書に記載された効果（ベネフィット）が、公報 A に記載された効果と本質的に同一または同種である場合には、その効果は『想定外のベネフィット』として取り扱うことはできない。
10. 【同種効果の判定条件】：差分構成が同一の場合にのみ「同種」と扱う。異なる構成間で効果を類推してはならない。＜従来技術＞（公報 A）に、差分構成に対応する下位概念または具体例が記載されている場合には、下位概念または具体例は差分と同等の効果を奏するものとする。＜従来技術＞（公報 A）の文章に差分構成に対応する効果（ベネフィット）が記載されなくとも、＜従来技術＞（公報 A）の図面から、差分構成に対応する物理的構成が読み取れる場合には、その物理的構成は差分と同等の効果を奏するものとする。効果の相違が使用の態様のみの場合には、その技術的定義は差分と同等の効果を奏するものとする。
11. 【組み合わせる動機の成立条件】：複数の差分から得られる効果の予測可能性は、明示的な開示がない限り否定されなければならない。
12. 【動機づけの技術的課題への限定】：公報 A の開示は、公報 A の構成に起因する技術的課題の解決のみを動機づけとする。公報 A の構成からは生じ得ない新しい技術的課題に対する解決策は動機づけとは一切認められない。

▼【請求項 1】

（ここに請求項 1 を貼る）

▼【公報 A 全文（図面説明含む）】

（ここに公報 A のテキストを貼る）

▼【本件明細書（効果記載の根拠箇所）】

（ここに本件明細書の該当テキストを貼る）

【最終出力フォーマット】

- 1) 実行ステップ 1 の結果（要素対応表／＜共通項＞）
- 2) 実行ステップ 2 の結果（＜差分＞一覧）
- 3) 実行ステップ 3 の結果（差分→効果の対応）
- 4) 実行ステップ 4 の結果（差分ごとの客観的資料および引用）
- 5) 【強制フラグ】の発動有無（差分ごとに明示）
- 6) 結論：「進歩性は肯定される」または「進歩性は否定される」

-----<ここまで、プロンプトのフォーマット>-----

3. 3 注意事項

現段階では、一般に提供される生成 AI は、しばしば、【実行ステップ】内の制限事項および【厳格ルール】に従わない。【実行ステップ 1】で＜共通項＞の対応付けが不十分であると感じたら、それに応答するプロンプトで『技術常識に鑑みれば公報 A で…の構成は正しくは…と解釈されるべき』と入力する。【実行ステップ 1】で「主観的解釈は排除される」と指摘されたら、それに応答するプロンプトで『【実行ステップ 1】では構成の解釈は技術常識や論理的な帰結で補強されるべき』と入力する。生成 AI に技術的理解を促す。正しい理論に則って説得を試みる。逆に【実行ステップ 2】以降では根拠は客観的資料の記載に限定されることから、【実行ステップ 2】以降で判断の根拠が客観的資料から逸脱すると感じられたら、それに応答する形で『その…（動作や関係性、効果など）といった解釈は客観的資料に形成されていない』と入力する。客観的根拠の提示を促す。正しい理論に則って説得を試みる。

※必ず、【発明の詳細な説明】（明細書）のテキストデータ内に、【請求項 1】から生み出される『想定外のベネ

フィット』を記述すること。『想定外のベネフィット』は【請求項 1】と従来技術との＜差分＞（相違点）に基づき生じるものである。生成 AI は＜差分＞から生じるベネフィットの「予測可能性」に応じて「進歩性あり」「進歩性なし」を判断する。ここでの「予測」は、想像的な思考を排除し客観的資料に記載された「予測」に限定される。本実験では、生成 AI は進歩性の審査にあたって整理された特許文献のデータベースにはアクセスしていない。入力された 1 件目の特許文献に組み合わせられる 2 件目以降の特許文献は検索されていない。

3. 4 使用環境

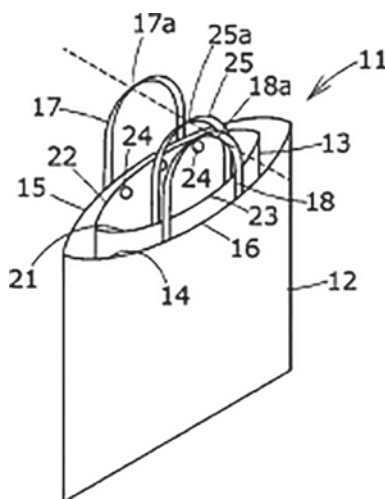
市販の Windows 11 Pro OS パーソナルコンピューター（CPU 2.8GHz + RAM16.0GB）を使用した。ウェブブラウザ Google Chrome から ChatGPT Pro および Google Gemini に個別にアクセスし対話形式で ChatGPT Pro および Google Gemini とやり取りした。対話には日本語が用いられた。

4. 実験結果（その 1）

4. 1 【請求項 1】の記述

【請求項 1】

出し入れの口（14）を形成し合わさると前記口を閉じる第 1 外縁および第 2 外縁を有する外袋体（12）と、
前記第 1 外縁に取り付けられる第 1 持ち手（17）と、
前記第 2 外縁に取り付けられる第 2 持ち手（18）と、
前記外袋体の内側に収容されて、出し入れの口（21）を形成し合わさると前記口を閉じる第 1 内縁および第 2 内縁を有する内袋体（13）と、
前記第 1 外縁に前記第 1 内縁を固定する固定具（24）と、
前記第 2 内縁に取り付けられて、前記第 1 持ち手または前記第 2 持ち手に選択的に束ねられる補助持ち手（25）とを備えることを特徴とするバッグ。



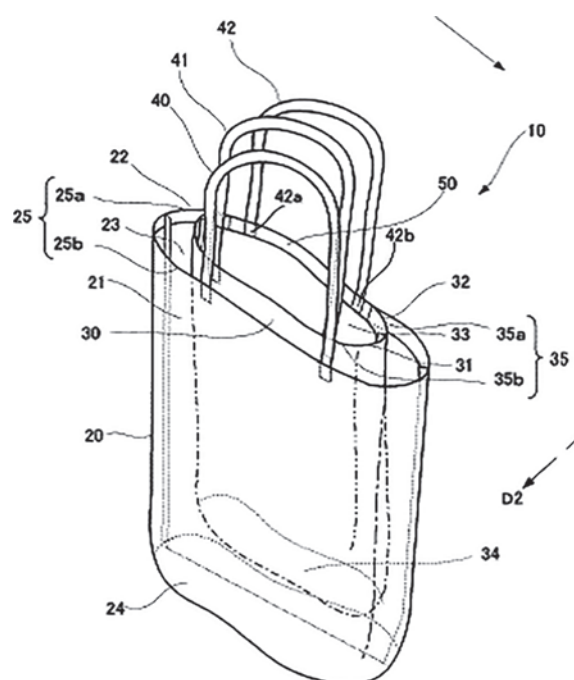
4. 2 対応する作用効果

第 1 持ち手、第 2 持ち手および補助持ち手が束ねられると、外袋体の口および内袋体の口は閉じられる。第 1 持ち手と第 2 持ち手とで外袋体は吊り下げられる。内袋体の第 1 内縁は外袋体の第 1 外縁に固定されることから、第 1 持ち手と補助持ち手とで内袋体は吊り下げられる。第 1 持ち手と補助持ち手とが束ねられたまま第 2 持ち手が第 1 持ち手から遠ざかると、外袋体の口は開かれる。内袋体の口は閉じられたまま維持される。外袋体に運搬物は進入することができる。このとき、第 1 持ち手と第 2 持ち手とで外袋体は吊り下げられる。第 1 持ち手と補助持ち手とで内袋体は吊り下げられる。第 2 持ち手と補助持ち手とが束ねられたまま第 2 持ち手が第 1 持ち手から遠ざかると、内袋体の口は開かれる。内袋体に運搬物は進入することができる。このとき、第 1 持ち手と第 2 持ち手とで外袋体は吊り下げられる。第 1 持ち手と補助持ち手とで内袋体は吊り下げられる。持ち手の束ね方だけで外袋体の口

と内袋体の口とは選択的に開かれる。

4. 3 1 件の従来技術（特許文献）

實用新案登録第 3190517 号「公報 A」



4. 4 新規性

最終出力フォーマットに従って生成 AI は「進歩性は肯定される」と結論づけた。この結論に至る過程で新規性は肯定されたものと理解される。しかしながら、最終出力フォーマットに従って【実行ステップ 1】の結果を確認したところ、共通項の抽出にあたって生成 AI は公報 A の構成を間違えていた。【請求項 1】に記載の「第 2 持ち手」に共通項として公報 A の「第 2 の把持部 41」が対応付けられた。「補助持ち手」に「第 3 の把持部 42」が対応づけられた。ここで、『「第 2 持ち手」には「第 3 の把持部 42」が対応づけられ「補助持ち手」には「第 2 の把持部 41」が対応づけられるべき』と入力したところ、生成 AI はその取り違いに理解を示した。生成 AI は、【請求項 1】に記載の「外袋体」「第 1 持ち手」「第 2 持ち手」「内袋体」に共通項を見出した。生成 AI は改めて【実行ステップ】を実行した。

それでも生成 AI は「進歩性は肯定される」と結論づけた。共通項の抽出にあたって「補助持ち手は第 1 持ち手または第 2 持ち手に選択的に束ねられる」構成は＜差分＞として認識された。確かに公報 A の文章中では「選択的に束ねられる」は示唆されていない。その一方で、図面を見ればだれでも公報 A から「選択的に束ねられる」を理解することができる。ここで、『公報 A の図 1 を見てほしい。第 1 の把持部 40 に第 2 の把持部 41 を束ねるか、第 3 の把持部 42 に第 2 の把持部 41 を束ねるか、ユーザーは選択できる。そういう物理的構造だよ』と入力したところ、生成 AI は「選択的に束ねられる」構成を共通項と判断した。生成 AI は改めて【実行ステップ】を実行した。その結果、最終的な回答として生成 AI は「新規性は否定される」と結論づけた。

4.5 進歩性

【請求項 1】の新規性は否定されたので、次のとおり補正（下線部分）に基づき新たに【請求項 1】を作成した。

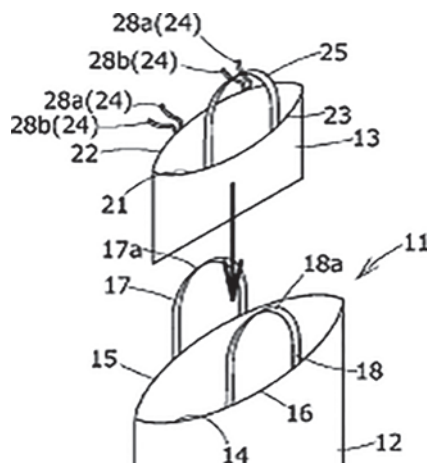
【請求項 1】

出し入れの口（14）を形成し合わさると前記口を閉じる第1外縁および第2外縁を有する外袋体（12）と、
前記第1外縁に取り付けられる第1持ち手（17）と、
前記第2外縁に取り付けられる第2持ち手（18）と、

前記外袋体の内側に収容されて、出し入れの口（21）を形成し合わさると前記口を閉じる第 1 内縁および第 2 内縁を有する内袋体（13）と、

前記第 1 内縁に固定されて、相互に結合されることで前記第 1 持ち手に引っかかる 1 対の紐状固定具（28a、28b）と、

前記第 2 内縁に取り付けられて、前記第 1 持ち手または前記第 2 持ち手に選択的に束ねられる補助持ち手（25）とを備えることを特徴とするバッグ。



新たな構成要件に応じて、次の通りに、対応する作用効果は追加された。

「1 対の紐状固定具が相互に結合されることで第 1 外縁に第 1 内縁は固定されることから、外袋体に改めて加工が施されなくても、外袋体に内袋体は連結されることができる。内袋体は 1 形態であっても様々な形態の外袋体に組み合わせられることができる。」

生成 AI は「紐状固定具」および「補助持ち手」に＜差分＞を見出した。「選択的に束ねられる」構成に関して再考を促したところ、生成 AI は唯一の＜差分＞として「紐状固定具」を理解した。【実行ステップ 1】で＜差分＞が確認されたことから、新規性は肯定され、生成 AI は【実行ステップ 2】以降に移行した。生成 AI は、インターネット上で、「紐状固定具」に対応する「作用効果」を読み取ることができる客観的資料を検索した。客観的資料は発見されなかったことから、『想定外のベネフィット』は認定された。最終的な回答として「進歩性は肯定される」は結論づけられた。

こうした手順をなんども繰り返したところ、対話に基づき生成 AI の理解漏れが是正されると、確実に最終的な結論『進歩性は肯定される』に行き着くことが確認された。

なんども繰り返すうちに、ある時点から最終的な結論に『進歩性は否定される』が出現した。【実行ステップ 4】を確認したところ、生成 AI はインターネットから客観的資料を探し当てていた。バッグ内の内袋らしきものに、持ち手に引っかかる紐状固定具が発見された。『なぜ、これまで、この客観的資料は発見されなかったのか』と、生成 AI にその理由を問い合わせたところ、検索の戦略を変えたということであった。

ウェブサイトなので公開時期は不明である。

5. 実験結果（その 2）

5. 1 【請求項 1】の記述

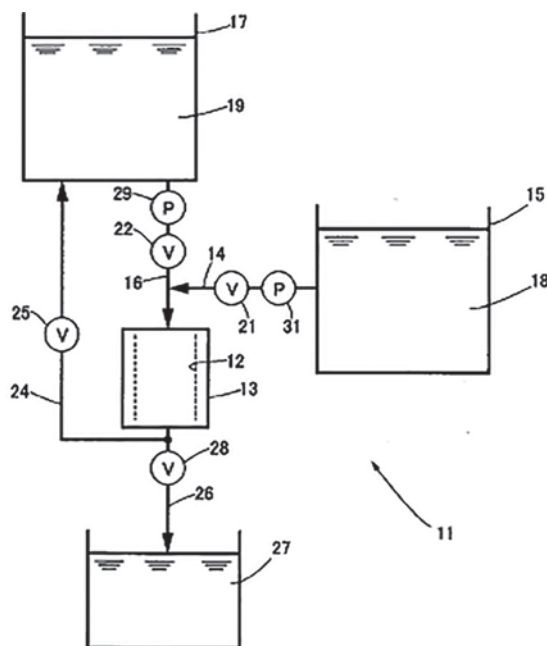
【請求項 1】

活物質を向き合わせてそれらの間に第 1 流路（12）を形成し、活物質の間で作用する直流電位に応じてイオンを吸着する正極および負極と、

決められた濃度以上でイオンを含有する液体（18）の通路を形成する第 2 流路（14）と、

第 2 流路に希釈液の通路を結合し、液体（18）に希釈液（19）を混合し、第 1 流路に向かって前記決められた濃度未満でイオンを含有する液体を導入する第 3 流路（16）と、

を備えることを特徴とする回収システム。

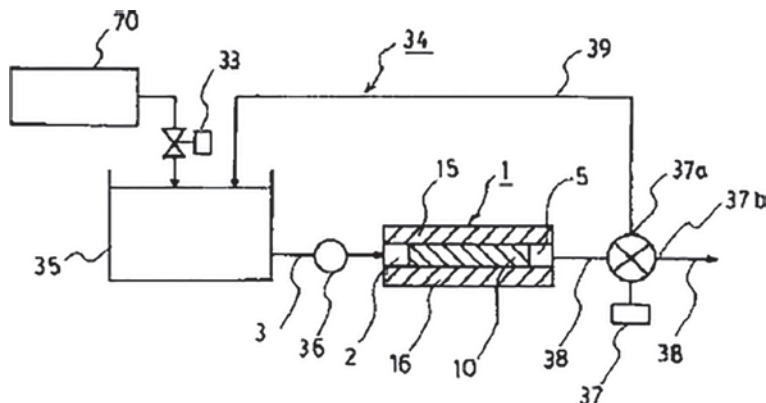


5. 2 対応する作用効果

第 1 流路ではイオンの濃度は決められた濃度未満に設定されることから、イオンの析出は良好に回避され（抑制され）ることができる。正極および負極からイオンは効率的に回収されることことができる。

5. 3 1 件の従来技術（特許文献）

特開 2001 - 191077 号公報「公報 A」



5. 4 新規性

最終出力フォーマットに従って生成 AI は「進歩性は肯定される」と結論づけた。生成 AI の理解は不十分であると考えられる。ここで、『公報 A では循環経路 34 から被処理液（希釈液）が貯留タンク 35 に流入するから、結果的に通液コンデンサー 1 に流入する液体でイオンの濃度は低下する』とプロンプトを入力し生成 AI の理解を促した。現実には【請求項 1】に対する拒絶理由通知で「新規性」は否定されたからだ。しかしながら、生成 AI は「進歩性は肯定される」の結論を維持した。

生成 AI の理屈は明快だった。公報 A には「第 2 流路に結合されて、決められた濃度未満でイオンを含有する液体を導入する第 3 流路」は開示されていないというものだ。なんども言い方を変えながら、説得を試みたものの、生成 AI は頑なに「進歩性は肯定される」の結論を維持し続けた。新規性は否定されると思いきや、新規性は肯定されてしまった。

6. 考察

本稿で実施した生成 AI による「進歩性」判断の検証は、AI が「進歩性あり」「進歩性なし」の判断において発揮できる能力と、依然として人間に依存せざるを得ない限界との両方を明確に示した。

6. 1 AI の判断傾向とロジックの限界

検証結果から、生成 AI は、【請求項 1】の構成要件と引用文献の構成要素とを照合する作業において 1 対 1 の対応関係を構築することで＜差分（相違点）＞を特定することが確認された。しかしながら、「バグ」の事例（項 4.5）では、「補助持ち手」に関する定義（「選択的に束ねられる」）に関して図面の参照を指示したように、生成 AI は人間との対話を通じて初めて文脈や図面の示唆を理解した。生成 AI の技術的理解は正確性に欠ける。生成 AI は人間の補助なしに単独では対応関係を構築することができない。最終的な判断にあたって人間は重要な役割を果たす。

「回収システム」の事例（項 5.4）では、確かに生成 AI は『技術常識』や『論理的な帰結』を駆使した上で物理的構成に固執した。生成 AI は、一貫して【実行ステップ 1】の制限事項および【厳格ルール】を遵守し続けた。判断は全くおぼれなかった。

6. 2 実務導入への提言と今後の展望

今回の検証で最も重要な知見は、生成 AI の判断は、プロンプトで設定されたガードレールに基づき、人間の介入に対しても高い一貫性を維持した点である。人間との対話を通じて技術的理解を補助する必要があるとはいえ、判断基準そのものに一貫性が確保されたことは特許出願および審査に大きな実務的貢献をもたらす可能性を秘めている。

（1）実務導入の根拠：無駄な作業の回避

発明の発掘や出願書類の準備にあたって、生成 AI が【請求項 1】の構成要件と引用文献との間で対応関係を構築することは極めて大きなメリットとなる。＜差分＞の特定にあたって生成 AI との対話を通じ発明の発掘や出願書類の準備に関わる者は発明の理解を深めることができる。生成 AI の指摘は引用文献の読解負担を軽減し時短を実現すると期待される。出願人は、事前に懸念される問題点を特定し、出願戦略を修正し、あるいは出願そのものを回避することができるかもしれない。生成 AI の支援の下で請求項の記述や進歩性の焦点を決定することができる。その結果、「無駄な作業」を大幅に削減し、特許の取得コストと時間とを最適化し得る。

ただし、相変わらず弁理士の関与は欠かせないだろう。【請求項 1】の書き方次第で【請求項 1】の構成要件は変化する。生成 AI は【請求項 1】の記載に基づき引用文献との対応関係を構築することから、【請求項 1】は的確に発明を定義しなければならないと考えられる。請求項の記述や進歩性の焦点の決定にあたって最終的な判断に弁理士が果たす役割はなおも重要であると考えられる。

（2）審査官の役割の高度化

AI は【請求項 1】と従来技術との＜差分＞の抽出や＜差分＞を示す客観的資料の検索を担当することが期待される。AI の指摘は引用文献の読解負担を軽減し時短を実現すると期待される。AI によって提起される不都合な一貫性に対して、審査官は、より厳密な法理と論理構成とで考察を加えることができる。

（3）今後の課題と研究の方向性

今回の実験は、実務上の「進歩性」判断を完全に反映するわけではない。1 つの仮説として、客観的資料に＜差分＞およびベネフィットの記述を要求することで、明確な客観性の確保が試みられた。「動機づけ」の成立に目的は欠かせないとの前提に立脚する。ベネフィットの記述は「動機づけ」を可視化する。記述の有無だからこそ予測可能なのである。生成 AI にはまだまだ判断の正確性は期待できないものの、生成 AI の活用に立脚した「進歩性」

判断の可能性は確認された。生成 AI が人間の判断や決定力に置き換わるとまで言えないまでも、作業負担の軽減や時短を実現する「支援ツール」として大いに活躍するであろうことは容易く想像される。

「進歩性」の判断といった非物理的な思考の領域では、判断過程の言語化が重要であると考えられる。人間の審査官の思考を追跡し、AI によって理解可能な表現にひとつひとつ思考を落とし込んでいく必要がある。フレームワークの蓄積に応じて出願人や弁理士の利便性は高まるに違いない。

7. おわりに

AI は、設定されたフレームワークに従って迷わず瞬時に判断を下す。よく言えば、常に一貫性を持ち、悪く言えば、柔軟性に欠ける。しかし、使い方次第で作業負担は軽減される。時短の効果は計り知れない。今回の検証を通じて率直にそう感じられる。

AI の活用が声高に叫ばれる中、思ったほどに AI の活用は進まない。巷には AI の専門家が溢れるにも拘わらず、である。なぜなら、様々な分野で言語化の壁が立ちだかるからだだろう。実は、AI の専門知識だけでは真に役立つ AI を作ることができない。業務に役立つ思考を持ち合わせていないからだ。「進歩性」判断でも同じことが言えるかもしれない。「進歩性」判断の思考を言語化できなければ、役立つ AI は成立し得ない。AI の出来具合は言語化の質で決まってしまう。だからこそ、現場の思考を知る実務家が AI の開発に関与しなければならない。私たちはいま、AI を通して、自らの思考と制度の形とを見つめ直す時代に立っている。本稿がその一助となれば幸いである。

なお、本稿で述べた見解は、所属組織の公式な立場を示すものではなく、筆者個人の学術的見解である。今後の議論と批判的検証とを通じ、AI が支える新しい「進歩性」判断の姿を模索していきたい。

(原稿受領 2025.10.30)