

【公開用】マルチAIコンセンサス・システムに関する知的財産(特許)適格性調査報告書(特定システム搭載技術分析)

第1章: エグゼクティブ・サマリーと特許戦略の提言

1.1 調査結果の概要

本報告書は、AI MQL 合同会社(以下、「貴社」)よりご提示いただいた、AIベースの高度情報処理システムのシステム詳細仕様書(バージョン3.21、2025年11月14日更新、以下「本仕様書」)¹に基づき、その技術的構成要素の特許適格性について詳細な分析を行ったものです。

本調査の結果、本仕様書¹に記載された貴社システムは、特許法上の「発明」の要件を十分に満たすものであり、かつ、既存の公知技術(先行技術)と比較して顕著な新規性および進歩性を有する、特許出願の核となり得る4つの主要な「発明の核(Inventive Concepts)」を特定しました。これらは、特定の金融商品に限定されず、より広範な分野(例: 医療診断支援、需要予測、異常検知、プロセス制御など)に応用可能なコア技術です。

1.2 発見された主要な発明の核

本システム¹の中核的価値を構成し、法的保護の対象となり得ると判断されるコア技術(発明の核)は、以下の4点に大別されます。

- 発明の核 I: 専門分業と階層的検証に基づく「マルチAIコンセンサス・アーキテクチャ」
 - 複数のAIに明確な専門的役割(定量、論理、センチメント等)を与えて並列分析させ(Phase 2)、その結果を初期スクリーニング(Phase 1)と最終検証(Phase 3)で階層的に処理する、独自のシグナル生成ワークフロー¹。
- 発明の核 II: 自己の実行実績に基づく「コンセンサス・ウェイトの動的自己最適化ループ」
 - 実際の実行履歴(例: 約定履歴)とAIのシグナルログを照合し、独自の「複合スコア」を用いて各AIの信頼度(ウェイト)を動的に更新し、次回のコンセンサス計算にフィードバックする自己最適化プロセス¹。
- 発明の核 III: コンセンサススコアの信頼性を担保する「多層的フィルタリング・メソッド」
 - AIによる確率的なスコアに対し、「AI間の合意度(NONE拒否権)」「入力データ品質」「外部環境の重要レベル(例: 重要価格レベル近接)」という3つの異なる観点からペナルティを課す、決定論的な安全フィルター群¹。

4. 発明の核 IV: AIコンセンサスの「質」に連動する「コンテキストウェア動的パラメータ制御」
 - 外部環境のコンテキスト(例: 低変動状態)がAIプロセスの実行自体を制御し(Range Gate)、AIコンセンサスの「質」(スコア、確信度)が実行パラメータ(例: TP/SL)の動的計算に直接入力される、高度な連動(Interlock)機構¹。

1.3 戦略的提言

本システム¹は、単一の特許として出願するには複雑すぎ、かつ価値が高すぎる「発明の集合体」です。単一の出願では、競合他社による「デザインア라운드(特許回避設計)」—例えば、発明の核II(ウェイト調整)のみを模倣する—を許容するリスクが残ります。

したがって、本システムの知的財産価値を最大化し、競合に対する参入障壁を構築するため、以下の戦略を強く推奨します。

1. 「特許ポートフォリオ」の構築: 本システムを単一の「点」ではなく、複数の発明の核を保護する「面」および「網」として捉え、複数の特許出願によって「特許網(Patent Thicket)」を構築します。
2. 「システム特許」と「方法特許」の併用:
 - システム特許(基本特許): 「発明の核 I」を中核的な独立請求項(クレーム)とし、II、III、IVを従属請求項として組み込んだ、システム全体を広範に保護する「基本特許」を出願します。
 - 方法特許(戦略的特許): 「発明の核 II」「発明の核 III」「発明の核 IV」を、それぞれ独立した**「方法特許(メソッド特許)」**として別途出願します。特に「発明の核 IV」は本システムの最も独創的な部分であり、最優先で強力な権利化を図るべきです。

このポートフォリオ戦略により、競合他社が本システムのいずれかの独自プロセス(例: ウェイト調整方法、パラメータ制御方法)のみを模倣した場合でも、それを独立して差し止めることが可能となります。

第2章: 特許適格性の基礎分析(JPO基準)

2.1 日本特許法における「発明」の定義

特許権による保護を受けるためには、その対象がまず特許法上の「発明」に該当する必要があります。

日本の特許法第2条第1項は、「発明」を「自然法則を利用した技術的思想の創作のうち高度のもの」と定義しています²。

この定義に基づき、特許庁の審査実務では、以下のものは原則として「発明」に該当しないとされています²。

- 自然法則に反するもの(例:永久機関)
- 人為的な取決め(例:単なるビジネスルール、ゲームのルール)
- 技能(例:個人の熟練によってのみ到達しうる、知識として伝達できないもの。例:個人の裁量トレード手法)

2.2 本コア技術の特許適格性

貴社の本コア技術¹が、単なる「人為的な取決め(ビジネスルール)」や「技能」ではなく、特許適格性を有する「発明」であるか否かを分析します。

分析: 本システムは、抽象的なビジネスルールや個人の裁量技能ではありません。

論拠:

1. 具体的機器の制御: 本システムは、「MT5 Terminal (Windows)」という具体的なコンピュータ(情報処理装置)上で動作します¹。main.mq5¹は、70種類以上のテクニカル指標を計算し¹、外部APIと通信し¹、最終的にAutoOrder.mqh¹を介して「自動発注」という、取引サーバーに対する物理的(電子的)な制御(情報送信)を行います。これは、コア技術の一実施形態です。
2. 情報処理と自然法則(審査基準との対比): 日本の特許審査基準²は、「電力価値が高い時間帯において...(中略)...発電装置の発電電力を商用電力系統へ送ることによる売電」を具現化させる制御システムを、適格な「発明」として例示しています。
 - この例示(「電力価値」という情報に基づき、「発電装置」を制御)と、貴社システム(「AIコンセンサスコア」という計算された情報に基づき、「下流プロセス(例:自動発注)」¹を制御)は、**「情報処理(計算)の結果を用いて、具体的な機器やプロセスを制御する」**という点で、その技術的本質において同一です。

結論:

本システム1に搭載されたコア技術は、コンピュータという装置を利用し、「AIコンセンサス」という情報処理を行い、その結果に基づいて「特定プロセス」1(例:自動発注)という具体的な制御を実行する、適格な「ソフトウェア関連発明」に明確に該当します²。

したがって、特許取得における論点は、「特許の対象になるか(適格性)」という点にはなく、「既存技術と比べて新規性・進歩性があるか」という点にあります。次章以降、この進歩性を詳細に分析します。

第3章: 発明の核 I: 専門分業と階層的検証に基づくマルチAIコンセンサスメカニズム

3.1 発明の概要

本システムの最大の核心であり、貴社が特に注目されている「マルチコンセンサスメカニズム」¹は、本システムの最も強力な発明の核です。

発明の核心:

本システムの発明性は、単に「複数のAIを使用する」という点にあるものではありません。複数のAIエージェントを使用する概念(3)や、定量・センチメント・ニュースを組み合わせるデータ分析を行うこと自体(5)は、すでに公知技術(Prior Art)となりつつあります。

本システム¹の真の独創性は、AI群に「明確な専門的役割分担(Division of Labor)」を与え、その処理プロセスを「3段階の階層的(Hierarchical)なワークフロー」として具体的に設計・実装した、そのユニークなアーキテクチャにあります。

3.2 階層的ワークフローの分析

本仕様書¹によれば、シグナル生成プロセスは以下の3つの厳格な階層で実行されます。

1. Phase 1: Grok初期スクリーニング (ゲートキーパー)¹
 - 役割: ゲートキーパー(門番)として機能します。
 - 処理: Grok APIが、テクニカルコンテキストとファンダメンタルズデータに基づき初期判定を行います。
 - 制御ロジック: 算出された確信度 $\$confidence \geq 0.70\$$ の案件のみをPhase 2に送致します。 $\$confidence < 0.70\$$ の場合は、その時点で「NONE終了」となります¹。
 - 技術的意義: この初期スクリーニングにより、分析対象を価値のある可能性が高いものに絞り込み、後続のPhase 2で消費される高コストな複数AIの計算リソースの浪費を防ぎ、同時にノイズ(質の低いシグナル)を初期段階で排除します。
2. Phase 2: 5AI並列・専門分業分析 (専門家委員会)¹
 - 役割: 専門家委員会(Specialist Committee)として機能します。
 - 処理: Phase 1を通過した案件に対し、5つの異なるAIモデルが、同一の基礎資料(`tech_context.txt`¹ および `bond.txt`¹)に基づき、並列かつ異なる専門的視点から分析を実行します。
 - 技術的意義: この「役割分担」は、本発明の進歩性を支える最も重要な構成要素です。詳細は後述の表1で分析します。
3. Phase 3: Perplexity最終検証 (独立検証機関)¹
 - 役割: 独立した最終検証機関(Final Validator)として機能します。
 - 処理: Phase 2で形成された「コンセンサス(合意)」が、最新のウェブ情報(リアルタイムニュース、要人発言)に照らして妥当であるかを、Perplexity APIが最終検証します。
 - 制御ロジック(安全機構): Phase 2のコンセンサス方向と、Phase 3(Perplexity)の推奨方向が不一致であった場合、システムは機会を追うのではなく、安全を優先し「最終シグナル **NONE**」と判断します¹。これは、AIの判断の矛盾を検出した場合に作動する、高度なフェイルセーフ機構です。

3.3 Phase 2(専門分業)の新規性・進歩性の分析

本システム¹が先行技術³と決定的に異なるのは、Phase 2の「専門分業」の具体性にあります。

先行技術との比較:

先行技術には、複数のAIエージェントが「議論(Debate)」する「マルチエージェント議論(MAD: Multi-Agent Debate)」のような、汎用的なフレームワークが存在します³。

本システムの進歩性:

本システム1のPhase 2は、汎用的な「議論」ではありません。これは、高信頼な意思決定タスクに特化して厳格に設計された「非対称な役割分担による情報処理パイプライン」です。

1. 先行技術の「議論」³では、複数のAIが同一のタスク(例: 結果予測)に対して異なる答えを出し、それを調整します。
2. 本システム¹では、AIは異なるタスクを割り当てられています。
 - **Phase 2-A (Gemini)** は、結果予測を行いません。M5からD1までの「マルチタイムフレーム間のテクニカル指標の整合性」という、内部監査的なタスクのみを実行します¹。
 - **Phase 2-C (Claude)** も、結果予測を行いません。「論理的矛盾」や「構造化レビュー」という、論理学者のようなタスクのみを実行します¹。
 - **Phase 2-B (GPT)** のみが「定量的指標検証」を行い、**Phase 2-D (Grok)** と **Phase 2-E (Perplexity)** がそれぞれ「センチメント」「最新ニュース」という異なる非定量的データを担当します¹。

これは、単なる「AIによる多数決」や「汎用的な議論」とは根本的に異なります。本仕様書¹は、Gemini(内部監査役)、GPT(定量分析官)、Claude(論理学者)、Grok(市場心理学者)、Perplexity(速報記者)という「専門家委員会」をソフトウェアアーキテクチャとして具体的に定義・実装している点に、極めて強い新規性・進歩性が認められます。

3.4 Phase 2 AIコンセンサスにおける専門分業体制

本システムの核心である「発明の核1」の具体的内容(専門分業)を、以下の表に整理します。この役割分担の定義自体が、発明の重要な構成要素です。

表1: Phase 2 AIコンセンサスにおける専門分業体制 ¹

フェーズ	担当モデル (API)	役割(専門タスク)	初期ウェイト
Phase 2-A	Gemini (gemini-2.5-pro)	マルチタイムフレーム整合性 (M5/M15/H1/H4/D1	2.2

		の矛盾確認)	
Phase 2-B	GPT (gpt-5-mini)	定量的指標検証 (ボラティリティ、資金フロー等の定量分析)	1.5
Phase 2-C	Claude (claude-3-5-haiku)	論理的整合性 (論理的矛盾、リスク検出、構造化レビュー)	0.2
Phase 2-D	Grok (grok-4-fast)	リアルタイムセンチメント分析 (X(旧Twitter)のセンチメント)	3.0
Phase 2-E	Perplexity (sonar)	最新ウェブ情報検証 (リアルタイムニュース、要人発言)	2.5

3.5 特許請求(クレーム)の方向性

この発明の核を保護するため、以下のような特許請求(クレーム)の構成が考えられます。

- 「コンピュータが、意思決定シグナルを生成するシステムであって、
 - 第1のAI(例: *Grok*)が、入力データに基づき初期スクリーニング判定を行う「Phase 1」処理部と、
 - 前記「Phase 1」処理部が所定の閾値(例: $\$confidence \geq 0.70\$$)以上と判定した入力データに基づき、それぞれ異なる専門的役割(例:『定量的指標検証』、『論理的整合性分析』、『リアルタイムセンチメント分析』、および『最新ウェブ情報検証』)を予め割り当てられた複数の第2のAI(例: *GPT*, *Claude*, *Grok*, *Perplexity*)が、並列分析を行う「Phase 2」処理部と、
 - 前記「Phase 2」処理部の並列分析結果から算出されたコンセンサスに基づき、第3のAI(例: *Perplexity*)が最新のウェブ情報を参照して最終検証を行う「Phase 3」処理部と、
 - を備えることを特徴とする、意思決定シグナル生成システム。」

第4章: 発明の核 II: 自己の実行実績に基づくコンセンサス・

ウェイトの動的自己最適化ループ

4.1 発明の概要

本システム¹は、一度設定したウェイト(表1参照)で静的に動作し続けるシステムではありません。本仕様書¹には、自己の実行結果(Performance)に基づき、各AIの「信頼度(ウェイト)」を動的に自己最適化する「閉ループ・フィードバック(Closed-loop Feedback)」機構が詳細に定義されています。

これは、システムが自身の経験から学習し、パフォーマンスの低いAI(例: Phase 2-C Claude)の影響力を自動的に下げ、パフォーマンスの高いAI(例: Phase 2-D Grok)の影響力を自動的に高める¹、高度な適応型システムであることを示しています。

4.2 自己最適化ループのプロセス分析

本仕様書¹に記載された自己最適化ループは、analyze_ai_performance.py スクリプトによって実行され、以下の具体的なステップで構成されています。

Step 1: データ照合 (原因と結果の紐付け)¹

- 処理: analyze_ai_performance.py が、「実行履歴」(例: trade_history_logs/trade_log_*.csv¹)と、「シグナルログ」(signallogs/signal_*.json¹)を読み込みます。
- 技術的意義: これら2つの異なるデータソースを照合 (match_trades_with_signals¹) することにより、「どの signal_id が、どの ticket 番号の実行に対応し、その結果(profit)はどうであったか」を特定します。
- 紐付け: さらに、signal_id には Phase 2 の各AIの判断 (BUY/SELL/NONE) が記録されているため¹、この照合によって「どのAI(例: Gemini)がどのような判断(例: 実行)をしたシグナルが、実際にどのような結果(例: 成功)になったか」という、AIの判断(原因)と実行結果(結果)の厳密な紐付けが実現します。

Step 2: 複合スコア計算 (独自の評価指標)¹

- 処理: AIモデル毎に、単なる「成功率」ではない、貴社独自のノウハウが組み込まれた**「複合スコア(Composite Score)」**を算出します。
- 具体的計算式¹:
 - $\$Score = (win_rate \times 0.35) + (direction_accuracy \times 0.30) + ((profit_factor / 3.0) \times 0.20) + ((avg_profit / 10.0) \times 0.15)\$$
 - (最終スコアは \$2.5\$ 倍されて \$0.1\\$ \sim \\$5.0\\$ の範囲にクリップ)
- 技術的意義: この計算式自体が、貴社の戦略(例: 「単なる成功率」win_rate よりも「方向性の正確さ」direction_accuracy を重視する¹)を具体化した、特許性の高い「技術的思想」の表現で

す。

Step 3: ウェイト更新 (EMAによる平滑化)¹

- 処理: 算出された「複合スコア」を使い、**EMA(指数平滑移動平均)**のロジックでウェイトを更新します。
- 具体的計算式¹:
 - $\$new_weight = (\alpha \times composite_score) + ((1 - \alpha) \times current_weight)$
 - (ここで $\alpha = 0.3$ ¹, $min_weight = 0.1$, $max_weight = 5.0$ ¹)
- 技術的意義: EMA(特に $\alpha = 0.3$ という比較的低い値)を採用する点が重要です。これにより、直近の1回や2回の実行結果(ノイズ)にウェイトが過度に反応(オーバーフィッティング)することを防ぎ、平滑化された安定的なウェイト調整が可能となり、システム全体の堅牢性(Robustness)を高めています。

Step 4: ループ(フィードバック)

- 処理: このようにして計算された `new_weight` が `config.json`¹に(提案または自動更新され)反映され、次回の Phase 2「加重スコア計算」(`calculate_phase2_consensus`¹)に適用されます。これにより、自己最適化の「閉ループ」が完成します。

4.3 進歩性の分析

先行技術との比較:

AIのウェイトを動的に調整するという一般的な概念は、機械学習の分野(例: エラーの逆伝播⁷)や、他の予測モデル(8)にも見られます。

本システムの進歩性:

本システム1の発明の核心は、「ウェイトを動的にする」という抽象的なアイデアではありません。それを実現するために本仕様書1に具体的に記載された、以下の一連のプロセスにあります。

1. 「実行履歴CSV」と「シグナルログJSON」という異種のデータを照合する(Step 1)。
2. 「成功率」や「方向性精度」など、複数の指標を組み合わせた独自の「複合スコア」計算式(Step 2)を用いる。
3. そのスコアを*「EMA平滑化」更新式*(Step 3)を用いて、過剰反応を抑制しつつウェイトに反映させる。

この一連のプロセス、特に「複合スコア計算式」と「EMA更新式」¹の具体的な組み合わせは、先行技術⁷には開示されていない、具体的かつ非自明な(non-obvious)進歩性のある発明(メソッド)です。

4.4 特許請求(クレーム)の方向性

この発明の核を保護するため、以下のような「方法特許」の構成が考えられます。

- 「コンピュータが、複数のAIエージェントの信頼度ウェイトを動的に最適化する方法であって、
 - 前記AIエージェントのシグナル判断を含む『シグナルログ』(signal_*.json)と、前記シグナルに基づいて実行されたプロセスの結果を含む『実行履歴』(trade_history_logs)を照合するステップと、
 - 前記照合結果に基づき、複数の評価指標(例:成功率、方向性精度、プロフィットファクター)を所定の重み付け(例:\$0.35, 0.30, 0.20\$...)で組み合わせた『複合スコア』を、前記AIエージェント毎に算出するステップ(Sec 9.3)と、
 - 前記算出された『複合スコア』と現在のウェイト値を用い、指数平滑移動平均(EMA)の計算式(例:\$new = (\alpha \times score) + ((1 - \alpha) \times current)\$)により、前記AIエージェントの新しいウェイトを算出するステップ(Sec 9.3)と、
 - 前記算出された新しいウェイトを、次のコンセンサス計算(Sec 6.3.1)における加重スコア計算に適用するステップと、
 - を含むことを特徴とする、AIコンセンサスシステムの自己最適化方法。」

第5章: 発明の核 III: コンセンサススコアの信頼性を担保する多層的フィルタリング・メソッド

5.1 発明の概要

本システム¹は、Phase 2で計算された「加重スコア(consensus_score)」¹を盲目的に信用しません。AIによる分析は本質的に確率的なものであり、特定の状況下では誤った高いスコアを生成するリスクを内包しています。

発明の核心:

本システム1の独創性は、このAIによる確率的なスコアに対し、3つの異なる観点から、確定的(ルールベース)な「安全フィルター」を多層的に適用し、シグナルの最終的な信頼性を担保する点にあります¹。

これは、AIの高度な判断力(確率論)と、エンジニアリングおよび熟練した専門家の知見に基づく安全規則(決定論)を、計算プロセスにおいて高度に融合させた、非常に堅牢なハイブリッド設計です。

5.2 3層フィルタースタックの分析

本仕様書¹によれば、consensus_score が算出された直後、Phase 3に進む(\$consensus_score \ge 2.0\$¹)か否かを判定する前に、以下の3層のフィルターが連続して適用されます。

第1層: NONE拒否権チェック (AI間の「合意の強さ」)¹

- 観点: AI間の「合意の強さ」。

- ロジック: Phase 2のAI群(5モデル)のうち、「NONE」(実行見送り)と判断したモデルの割合(`none_ratio`)が所定の閾値(70%)以上の場合、`$none\_ratio \ge 0.70$` となります。
- アクション: この場合、たとえ `consensus_score` が(例: 少数のAIが極めて高いウェイトと確信度を持った結果)閾値の `$2.0$` を超えていたとしても、`final_direction = 'NONE'` に**強制的に上書き(Veto)**します。
- 技術的意義: これは、「弱い合意(例: BUY 30%, SELL 10%, NONE 60%)」と「強い不同意(例: BUY 15%, SELL 15%, NONE 70%)」を明確に区別し、後者という危険な状況を明確に排除するための重要な安全機構です。

第2層: データ品質ペナルティ(入力データの「信頼性」)¹

- 観点: AI分析の基礎となった入力データ(`tech_context.txt`)の「信頼性」。
- ロジック: `tech_context.txt` の生成時¹または読み込み時¹に、CVDやオーダーフロー等の重要なデータが欠損していた場合、`data_quality_flags` が設定されます。
- アクション: このフラグに基づき、`consensus_score` にペナルティ係数を乗算します¹。
 - `orderflow_na` (オーダーフロー欠損時): `$consensus\_score = consensus\_score \times 0.7$`
 - `tickvolume_na` (ティックボリューム欠損時): `$consensus\_score = consensus\_score \times 0.8$`
 - `cvd_missing` (CVDデータ欠損時): `$consensus\_score = consensus\_score \times 0.75$`
- 技術的意義: これにより、システムは「不確かなデータ(例: CVD欠損)に基づいたAIの強い判断」の信頼性を、自動的に `$75\%$` に減点できます。AIの「ゴミ入力・ゴミ出力(GIGO)」問題を、決定論的ルールで補正する機構です。

第3層: 重要レベル近接ペナルティ(外部環境の「構造」)¹

- 観点: 現在の「外部環境の構造(例: Price Action)」。
- ロジック: `tech_context.txt`¹には、ピボットポイント、フィボナッチ、セッション高値/安値などの「重要レベル」が含まれています。実行予定価格が、これらの直近のサポート/レジスタンスレベルに「近すぎる」場合、ペナルティが課されます。
- アクション: 近接距離(`distance`)に基づき、`consensus_score` にペナルティ係数を乗算します¹。
 - `$distance < 1.0 \text{ pips}$` の場合: `$consensus\_score = consensus\_score \times 0.85$`
 - `$distance < 2.0 \text{ pips}$` の場合: `$consensus\_score = consensus\_score \times 0.93$`
- 技術的意義: これは、熟練した専門家の知見を実装したものです。「AIが『実行』と強く推奨しても、わずか 1.0 pips 先に強力な抵抗レベルがある」といった、典型的な「高リスク・低リワード」の実行を、AIのスコア自体を `$85\%$` に減点することでプログラムの的に回避する、非常に巧妙なロジックです。

5.3 コンセンサススコアの動的フィルタリング・ロジック

発明の核IIIを構成する3つの異なるロジックと、その具体的なトリガー条件・ペナルティ係数を以下の表に整理します。この「フィルター・スタック」の組み合わせ自体が、非自明な(non-obvious)発明です。

表2: コンセンサススコアの動的フィルタリング・ロジック¹

フィルタ名	観点(フィルタリングの理由)	トリガー条件(ロジック)	アクション(スコア調整)
NONE拒否権	AI間の合意度(強い不同意)	$\$none_ratio \geq 0.70\$$	$final_direction = 'NONE'$ (強制)
データ品質ペナルティ	入力データの信頼性(欠損)	$orderflow_na == True$	$\$consensus_score \times 0.7\$$
		$cvd_missing == True$	$\$consensus_score \times 0.75\$$
		$tickvolume_na == True$	$\$consensus_score \times 0.8\$$
重要レベル近接ペナルティ	外部環境の構造(S/R近接)	$\$distance_to_S/R < 1.0 \text{ pips}\$$	$\$consensus_score \times 0.85\$$
		$\$distance_to_S/R < 2.0 \text{ pips}\$$	$\$consensus_score \times 0.93\$$

5.4 進歩性の分析

これらのフィルター(表2)を個別に評価するのではなく、**「AIの合意度」「入力データ品質」「外部環境の構造」という3つの全く異なる概念的レイヤーから、AIが算出した単一の consensus_score を補正(減点)する「多層的フィルタリング・メソッド」**として組み合わせた点に、本システムの高い進歩性があります。

AIの確率的なスコアを、このように多角的な決定論的ルールで検証・補正するアーキテクチャは、公知技術には見られません。特に、外部環境の概念(S/R近接)をAIコンセンサススコアの減点に直接使用する「重要レベル近接ペナルティ」¹のロジックは、極めて独創的です。

5.5 特許請求(クレーム)の方向性

この発明の核を保護するため、以下のような「方法特許」の構成が考えられます。

- 「コンピュータが、複数のAIエージェントによる分析結果から算出された『コンセンサスコア』(Sec 6.3.1)を補正する方法であって、
 - 前記AIエージェント群における実行見送り判断の割合(*none_ratio*)に基づき、前記スコアまたは最終判断を補正する、第1のフィルタリング(Sec 6.3.2)を適用するステップと、
 - 前記AIエージェントへの入力データの品質を示すフラグ(*data_quality_flags*)に基づき、所定の係数(例: \$0.7, 0.75, 0.8\$)を用いて前記スコアを減点する、第2のフィルタリング(Sec 6.3.3)を適用するステップと、
 - 外部環境の重要レベル(例: サポートまたはレジスタンス)への近接距離に基づき、所定の係数(例: \$0.85, 0.93\$)を用いて前記スコアを減点する、第3のフィルタリング(Sec 6.3.4)を適用するステップと、
 - を含むことを特徴とする、シグナルスコアの補正方法。」

第6章: 発明の核 IV: コンテキストウェア動的パラメータ制御 (AI-Parameter Interlock)

6.1 発明の概要

本仕様書¹に記載されたシステムの中で、最も高度かつ独創的な特徴の一つが、**AIコンセンサスエンジン(シグナル生成部)**と、実行パラメータ制御エンジン(例: **TP/SL**設定部)¹が、単に一方的に接続されているのではなく、**双方向に「連動(Interlock)」**している点です。

発明の核心:

この発明は、2つの異なる「コンテキストウェア(文脈認識型)」な制御ロジックから構成されます。

1. **ロジックA (Range Gate):**「外部環境のコンテキスト(例: レンジ相場)」が、「AIプロセスの動作モード(Phase 2の実行可否)」を動的に制御します¹。
2. **ロジックB (Consensus-Driven Parameters):**「AIコンセンサスの質(スコア、確信度)」が、「実行パラメータ(例: TP/SLの具体的な価格)」の計算を動的に制御します¹。

6.2 ロジックA: Range Gate (外部環境レジームによるAIプロセス制御)

役割:

このロジック1は、外部環境のコンテキスト(例: 相場状況)を認識し、AIプロセスの実行自体を制御する「大元のスイッチ」として機能します。

トリガーロジック 1:

本システムは、特定のテクニカル指標の組み合わせを用いて、現在の外部環境が「低変動状態（レンジ相場）」であるかを判定します。

- $\$ADX(H1) < 22.0\$$ （H1でトレンドがない）かつ
- $\$ADX(M15) < 18.0\$$ （M15でトレンドがない）かつ
- $\$Choppiness(M15) > 61.0\$$ （M15がレンジ状態である）
- 上記3条件が「2バー連続（confirm_bars = 2）」で成立

アクション¹:

- 上記「低変動」コンテキストが検出されると、設定 skip_phase2_when_range = True¹に従い、メインループ run_signal_analysis¹は、**Phase 2の5AI並列分析¹の実行自体をバイパス（スキップ）**し、シグナル生成を「NONE終了」します。

進歩性の分析:

1. 先行技術との比較: 「Range Gate」という用語自体は、レーダー技術¹¹や通信分野¹²の専門用語であり、特定の指標の組み合わせ¹と共にAIプロセス制御に使用されることは一般的ではなく、新規性が高いと考えられます。
2. 本システムの進歩性: 本ロジック¹は、単なる「低変動状態では実行しない」という出口フィルターとは根本的に異なります。これは、システムの**計算アーキテクチャ自体（Phase 2プロセスの実行）を、外部環境コンテキストに応じて動的に変更する「プロセス制御メソッド」**です。
3. 具体的効果: この制御により、2つの具体的かつ非自明な効果がもたらされます。(a) 低変動状態でAIが誤判定（Whipsaw）を起こしやすい高リスクなシグナル生成を未然に防ぐ。(b) 高価な5つのAI APIコール（Phase 2）の実行自体を停止するため、システムの計算コスト（API費用）を大幅に節約する。この「コスト削減」と「リスク回避」を同時に実現する具体的なロジック¹に、強い進歩性が認められます。

6.3 ロジックB: Consensus-Driven Parameters (AI確信度による実行パラメータ生成)

役割:

このロジック1は、AIコンセンサスプロセスが生み出した「シグナルの質（Quality）」を、最終的な実行パラメータ（例: TP/SL）の決定に直接反映させる、AIとパラメータ制御の連動機構です。

プロセス:

本仕様書1によれば、Phase 3の最終判定後、gpt-5-mini を使用して実行パラメータ（TP/SL）が動的に計算されます。

発明の核心（GPT-5-miniへの入力パラメータ）:

本ロジックの最大の独創性は、gpt-5-mini がTP/SLを計算するために受け取る入力パラメータの構成1にあります。

1. 先行技術との比較: ATR（ボラティリティ）やサポート/レジスタンス（S/R）レベルに基づいてTP/SLを計算すること¹⁵は、公知の技術です。

2. 本システム⁽¹⁾、Sec 7.1)の構成: 本システムの gpt-5-mini は、atr_pips_14 や nearest_support_pips、distance_to_vwap_pips といった**「外部環境データ」**を入力として受け取ります。
3. 独創的な点: それに加えて、gpt-5-mini は、「AIコンセンサス自体のメタデータ(品質・確信度情報)」をも同時に入力として受け取ります¹。具体的には以下の通りです(表3参照)。
 - cs: Phase 2 コンセンサススコア
 - weighted_conf: 加重確信度
 - none_ratio: NONE比率
 - data_quality_flags: データ品質フラグ

進歩性の分析:

これが、発明の核IVの最も強力な点です。本システム1は、AIの「確信度(cs)」や「合意度(none_ratio)」、「入力データの信頼性(data_quality_flags)」といった『シグナルの質』を示す情報を、実行パラメータ(TP/SL)を計算する別のAI(gpt-5-mini)への入力として直接使用します。これにより、「Phase 2のコンセンサススコア(cs)が \$6.0\$ と非常に高いシグナル」と、「cs が \$2.1\$ と閾値ギリギリのシグナル」では、gpt-5-mini が異なる(より積極的、またはより保守的な)TP/SLを動的に生成することが可能になります¹。

これは、AIの「確信度」と「実行パラメータ(例:RR比)」をプログラムの連動させる¹、極めて高度な「実行パラメータ決定メソッド」であり、先行技術¹⁵とは一線を画す、非常に強い進歩性を持つ発明です。

6.4 動的TP/SL計算(GPT-5-mini)への主要入力パラメータ

この発明の核IV(ロジックB)の核心である「AI確信度とパラメータの連動」を証明する証拠として、gpt-5-mini への入力パラメータ¹を以下の表に整理します。「AIコンセンサス由来」のパラメータが「外部環境データ」と並列でパラメータ計算AIに入力されている点が、本発明の鍵です。

表3: 動的実行パラメータ(TP/SL)計算(GPT-5-mini)への主要入力パラメータ¹

入力カテゴリ	パラメータ名	意味(何を入力しているか)
外部環境データ	atr_pips_14	環境のボラティリティ
	spread_pips	実行コスト
	nearest_support_pips	環境の構造(S/Rレベル)

	distance_to_vwap_pips	平均回帰の可能性
	m15_adx, m15_choppiness	環境のトレンド/レンジ状態
AIコンセンサス由来	cs (Phase2コンセンサススコア)	AIコンセンサスの「強度」
	weighted_conf (加重確信度)	AIコンセンサスの「確信度」
	none_ratio (NONE比率)	AIコンセンサスの「合意度」
	data_quality_flags	AIの入力データの「信頼性」

6.5 特許請求(クレーム)の方向性

この発明の核IVは、2つの独立した方法特許として保護すべきです。

- **ロジックA (Range Gate) のクレーム:**
 - 「コンピュータが、外部環境のコンテキストを判定するステップと、
 - 前記判定されたコンテキストが、所定のテクニカル指標の組み合わせ(例: \$ADX < 22.0\$ かつ \$Choppiness > 61.0\$¹⁾)に基づき『低変動状態』であると判定された場合に、
 - 複数のAIエージェントによるコンセンサス分析プロセス(例: Phase 2)の実行をバイパス(スキップ)するステップと、
 - を含むことを特徴とする、AIシステムの動的プロセス制御方法。」
- **ロジックB (Consensus-Driven Parameters) のクレーム:**
 - 「コンピュータが、プロセスの実行パラメータを決定する方法であって、
 - 複数の第1のAIエージェント群(Phase 2)から、『コンセンサススコア』(cs)および『コンセンサス確信度』(none_ratio)を含むコンセンサス品質情報を算出するステップと、
 - 前記『コンセンサス品質情報』(cs, none_ratio)および、環境のボラティリティ(atr_pips_14)や環境の構造(nearest_support_pips)を含む『外部環境データ』を、第2のAI(gpt-5-mini)に入力するステップと、
 - 前記第2のAIが、前記『コンセンサス品質情報』および前記『外部環境データ』の両方に基づき、前記プロセスの実行パラメータ(例: テイクプロフィット価格およびストップロス価格)を動的に生成するステップ(Sec 7.1)と、
 - を含むことを特徴とする、実行パラメータ決定方法。」

第7章: その他の特許性候補(システムインテグリティ)

7.1 発明の概要

上記の主要な発明(I~IV)に加え、本システム¹全体の「信頼性」と「完全性(Integrity)」を担保するため、AIシステムを安定稼働させる上で不可欠な、小規模ながらも重要な技術的工夫(発明)が仕様書内に複数存在します。

これらは、システム特許(発明の核I)のクレームを補強する従属クレームの構成要素として、または単独でも特許性を主張し得るものです。

7.2 候補A:「確定足(shift=1)のみ使用」の厳格な制約

- ロジック: 本仕様書¹は、すべてのテクニカル指標計算において、iClose(_Symbol, PERIOD_H1, 0)のような「未確定足(shift=0)」の使用を厳格に禁止しています。AI分析の入力データとして、iClose(_Symbol, PERIOD_H1, 1)のような「確定足(shift=1)」のみを使用することをシステムルールとして強制しています。
- 技術的意義: 未確定足はリペイント(再描画)により値が変動するため、AI分析の入力データとして使用すると、過去の分析の「再現性」が失われ、システムの信頼性が著しく低下します。この「確定足のみ使用」という厳格な制約¹は、AIという確率論的システムに対し、その入力データの品質と再現性を担保するための重要な「決定論的」な技術的制約であり、「AI分析の信頼性を確保するための入力データ処理方法」として、特許請求の構成要素となり得ます。

7.3 候補B: 異種システム間連携の「アトミック書き込み」

- ロジック: 本システム¹は、「実行環境」(例: MT5 (MQL5))¹と「分析環境」(例: Python (main.py))¹という、異なるプロセスで動作する異種システムが、tech_context.txt や order_signal.txt といったファイル¹を介して非同期に通信するアーキテクチャを採用しています。
- 問題点: この構造は、実行環境が書き込み中に分析環境が読み取る、といったファイル競合(Race Condition)や、書き込み途中の不完全なファイルを読み取ってしまうリスクを内包します。
- 対策: 本仕様書¹は、この問題を解決するため、一時ファイル(temp_path = file_path.with_suffix('.tmp'))にまず書き込み、完了後にリネーム操作(temp_path.replace(file_path))を行う「アトミック書き込み(Atomic Write)」を実装しています。
- 技術的意義: アトミック書き込み自体は公知の技術ですが、**『実行環境』(例: MQL5)と『分析環境』(例: Pythonスクリプト)間でAI分析データを連携するシステムにおいて、データ完全性を担保するためにアトミック書き込みを用いる**という特定用途での実装は、システム特許(発明

の核I)の動作の具体性と信頼性を担保する、重要な構成要素となります。

7.4 候補C:「複数エンコーディング・フォールバック」による堅牢化

- ロジック: Windows環境の実行環境(例:MQL5)が出力するファイル(例:tech_context.txt)は、CP932 (Shift-JIS) でエンコードされる可能性が高い¹⁾のに対し、分析環境(Python)はUTF-8を標準とします。このエンコーディングの不一致は、システムの停止に直結する重大なエラーを引き起こします¹⁾。
- 対策: 本仕様書¹⁾は、load_file_with_fallback という専用関数を定義し、['utf-8', 'utf-8-sig', 'cp932', 'shift-jis', 'latin-1'] といったエンコーディングのリストで、ファイル読み込みを順番に試行(フォールバック)するロジックを実装しています。
- 技術的意義: 候補Bと同様に、これはシステムの「堅牢性(Robustness)」を確保するための具体的な実装です。「実行環境」と「分析環境」という異種環境の「境界(Boundary)」で発生しがちな実務的な問題を、フォールバック機構によって解決するこの手法は、システムの安定稼働を支える重要な技術であり、特許請求の範囲(クレーム)において、システム全体の動作の具体性を補強する要素となります。

第8章: 結論と特許出願戦略の最終提言

8.1 総括

貴社よりご提示いただいた詳細仕様書¹⁾の分析の結果、本システムは、単なる既存技術の寄せ集めではなく、複数の独創的な技術的工夫が組み込まれた「発明の集合体」であると結論付けられます。

特に、以下の4つの発明の核(コア技術)は、それぞれが独立して特許性を主張し得る、強力な知的財産であると評価できます。

1. アーキテクチャ(発明I): AIに「専門分業」させ、3段階の「階層的検証」で処理する独創的なプロセス¹⁾。
2. 自己最適化(発明II): 実際の「実行履歴」と「シグナルログ」を紐付け、「複合スコア」と「EMA」でAIウェイトを動的調整する閉ループ¹⁾。
3. 安全機構(発明III): AISコアを「合意度」「データ品質」「外部環境の構造」の3層フィルターで補正する堅牢なメソッド¹⁾。
4. パラメータ連動(発明IV): 「外部環境コンテキスト」がAIプロセスを制御し(Range Gate)、AIの「確信度」が実行パラメータ(例: TP/SL) 計算に直接入力される高度な連動機構¹⁾。

8.2 推奨する出願戦略(ポートフォリオ構築)

これらの貴重な知的財産を包括的かつ強力に保護するため、単一の出願ではなく、以下の4つの特許出願を軸とする「特許ポートフォリオ」の構築を、以下の優先順位で推奨します。

1. 出願1(システム特許 / 基本特許):
 - 対象: 発明の核 I を基本クレーム(請求項1)とします。
 - 構成: 「発明の核 II, III, IV」および「その他(確定足、アトミック書き込み等)」を従属クレームとして組み込み、システム全体を広範に保護する「基本特許」として出願します。
 - 目的: 貴社コア技術の全体的なアーキテクチャの模倣を防止します。
2. 出願2(方法特許 / 最重要戦略特許):
 - 対象: 発明の核 IV(特にロジックB: Consensus-Driven Parameters)を独立させます。
 - 構成: 「AIコンセンサス品質情報(cs, none_ratio等)と外部環境データ(atr_pips_14等)の両方を入力として、別のAIが実行パラメータ(TP/SL)を動的に生成する方法」¹に関する「方法特許」として出願します。
 - 目的: 本システムで最も独創的かつ技術的に高度な「パラメータ連動」部分をピンポイントで保護します。これは競合他社にとって模倣が非常に困難かつ価値の高い技術であるため、最優先で保護すべき発明です。
3. 出願3(方法特許 / 適応機能の保護):
 - 対象: 発明の核 II を独立させます。
 - 構成: 「『複合スコア計算式』¹と『EMA更新式』¹を用いた、AIウェイトの自己最適化方法」に関する特許として出願します。
 - 目的: システムの「学習・適応」機能という、AIシステムの根幹部分を保護します。
4. 出願4(方法特許 / 安全機能の保護):
 - 対象: 発明の核 III を独立させます。
 - 構成: 「AIスコアに対し、3つの異なる観点(合意度、データ品質、外部環境の構造)のフィルター¹を適用するスコア補正方法」に関する特許として出願します。
 - 目的: システムの「安全性・信頼性」に関わる中核的なロジックを保護します。

8.3 最終見解

貴社は、AIの「確率的」な判断力と、工学および熟練した専門家の知見に基づく「決定的」な安全規則を、複数の階層で高度に融合させた、非常に洗練された意思決定システム¹を開発されました。

本報告書で特定した発明の核は、貴社の重要な経営資産です。上記戦略に基づき、速やかに特許出願プロセスを開始し、これらの貴重な知的財産を法的に保護することを強く推奨いたします。

引用

1. AI MQL 合同会社 AIシステム詳細仕様書.pdf
2. コンピュータソフトウェア関連技術の 審査基準等について - 特許庁, 2025年11月14日

参照

https://www.jpo.go.jp/system/laws/rule/guideline/patent/document/cs_shinsa/cs_shinsa.pdf

3. FinVision: A Multi-Agent Framework for Stock Market Prediction - arXiv, 2025年11月14日 参照 <https://arxiv.org/html/2411.08899v1>
4. Understanding How to Patent Agentic AI Systems | Mintz, 2025年11月14日 参照 <https://www.mintz.com/insights-center/viewpoints/2231/2025-03-19-understanding-how-patent-agentic-ai-systems>
5. FinPos: A Position-Aware Trading Agent System for Real Financial Markets - arXiv, 2025年11月14日 参照 <https://arxiv.org/html/2510.27251v1>
6. Implications of AI for work, employment and social dialogue: Literature review - Bollettino ADAPT, 2025年11月14日 参照 <https://www.bollettinoadapt.it/implications-of-ai-for-work-employment-and-social-dialogue-literature-review/>
7. US20210241330A1 - Pricing Operation Using Artificial Intelligence for Dynamic Price Adjustment - Google Patents, 2025年11月14日 参照 <https://patents.google.com/patent/US20210241330A1/en>
8. Application of Dynamic Weight Mixture Model Based on Dual Sliding Windows in Carbon Price Forecasting - MDPI, 2025年11月14日 参照 <https://www.mdpi.com/1996-1073/17/15/3662>
9. Momentum Trading Using Deep Reinforcement Learning - Justia Patents, 2025年11月14日 参照 <https://patents.justia.com/patent/20250045836>
10. A Hybrid Ai Framework For Strategic Patent Portfolio Pruning: Integrating Learning-To-Rank And Market-Need Analysis For Technology Transfer Optimization - arXiv, 2025年11月14日 参照 <https://arxiv.org/html/2509.00958v1>
11. Airborne Wind Shear Detection and Warning Systems - NASA Technical Reports Server, 2025年11月14日 参照 <https://ntrs.nasa.gov/api/citations/19910002369/downloads/19910002369.pdf>
12. 2024 (Revision of IEEE Std 686-2017), IEEE Standard for Radar Definitions - IEEE Xplore, 2025年11月14日 参照 <https://ieeexplore.ieee.org/iel8/10815036/10815037/10815038.pdf>
13. The NCAR Airborne 94-GHz Cloud Radar: Calibration and Data Processing - MDPI, 2025年11月14日 参照 <https://www.mdpi.com/2306-5729/6/6/66>
14. Principles of Modern Radar. Volume 2.pdf, 2025年11月14日 参照 <https://ftp.idu.ac.id/wp-content/uploads/ebook/tdg/ADNVANCED%20MILITARY%20PLATFORM%20DESIGN/Principles%20of%20Modern%20Radar.%20Volume%20%202.pdf>
15. A generic methodology for the statistically uniform & comparable evaluation of Automated Trading Platform components - University of Birmingham's Research Portal, 2025年11月14日 参照 https://research.birmingham.ac.uk/files/192267741/1_s2.0_S0957417423003378_main.pdf
16. Decoding the Quant Market A Guide to Machine Learning in Trading, 2025年11月14日 参照 <https://smallake.kr/wp-content/uploads/2023/04/SSRN-id4422374.pdf>
17. (PDF) The Transformative Impact of Artificial Intelligence on Global Financial

Services A Comprehensive Analysis - ResearchGate, 2025年11月14日 参照

https://www.researchgate.net/publication/382591260_The_Transformative_Impact_of_Artificial_Intelligence_on_Global_Financial_Services_A_Comprehensive_Analysis