

人材の採用・育成・知識活用を支える AI/MAS 技術

AI/MAS technologies for supporting human resource recruitment, development, and knowledge management

阿部 雄大¹ 渡邊 佑典¹ 安島 聖¹ 佐藤 里恵¹ 村越 英樹¹ 林 久志^{1*}
Yuta Abe¹ Yusuke Watanabe¹ Akira Yasujima¹ Rie Sato¹ Hideki Murakoshi¹ Hisashi Hayashi^{1*}

¹ 東京都立産業技術大学院大学 Advanced Institute of Industrial Technology

*Corresponding author: Hisashi Hayashi, hayashi-hisashi@aait.ac.jp

Abstract This study addresses the social challenge of labor shortages in Japan's super-aging society by exploring approaches that leverage AI technologies and multi-agent simulations (MAS) to comprehensively support recruitment, development, and knowledge utilization. Due to the accelerating decline in the working-age population caused by rapid aging and low birth rates, Japan faces increasingly severe difficulties in knowledge transfer and securing human resources, exemplified by the so-called "2025 problem." Furthermore, delays in promoting digital transformation (DX), along with shortages of human resources, governance, and data infrastructure, have been identified as structural issues, indicating the need for comprehensive initiatives that go beyond simple IT utilization. This paper reports on three sub-themes undertaken by the Hayashi Project Team (PT) as part of the PBL-based education program. Chapter 2 focuses on the recruitment stage, discussing job-seeking support through AI interview training using LLMs. Chapter 3 addresses human resource development, examining the balance between short-term operational efficiency and long-term technical training. Chapter 4 focuses on retention, proposing the autonomous detection of knowledge gaps and their supplementation in conversational AI. Through these efforts, the project aims to establish a sustainable framework for human resource support spanning recruitment, development, and retention. It should be noted that this paper does not report completed research results but rather presents a research concept as a progress report.

Keywords AI; multi-agent simulation; recruitment; human resource development; knowledge management

1 はじめに

東京都立産業技術大学院大学（以下、AIIT とする）では、専門職大学院として、PBL（Project Based Learning）型教育を導入している。本稿では、PBL 型教育の一環として林 PT（Project Team）で取り組んでいる「人材の採用・育成・知識活用を支える AI/マルチエージェントシミュレーション (MAS) 技術」について報告する。

近年、日本社会は急速な高齢化と少子化の進行により、労働力人口の減少が大きな社会課題となっている。図 1 に示す国立社会保障・人口問題研究所による年齢別人口推計（平成 29 年推計）[1] では、2000 年から 2025 年にかけて労働力人口が 17%減少し、2050 年にはさらに 26%減少すると報告されている。この推計には、新型コロナの影響や出生率の低下などの要因は反映されておらず、実際の減少ペースはより深刻化する可能性が指摘されている [2,3]。さらに、「2025 年問題」と呼ばれる団塊世代の大量退職は、産業現場における技術・知識の継承断絶や中長期的な人材確保の困難化を引き起こすと懸念されており、人材に関する課題は早急な対応を要する領域である [3]。

労働力不足への対応として企業や自治体でデジタル化や DX（デジタルトランスフォーメーション）、AI の活用が推進されているが、図 2 に示すとおり日本の投資水準は国際的に見て低く、推進の遅れが課題となっている [4,5]。その背景には、人材不足に加えて、データ基盤やガバナンスの未整備といった構造的要因があることが指摘されており [4]、単なる「IT 活用」や「足りない人材を補う」といった短期的な対策にとどまらず、人材の入口から成長、そして定着に至るまでを一貫して支援する仕組みの構築が不可欠であると考えられる。

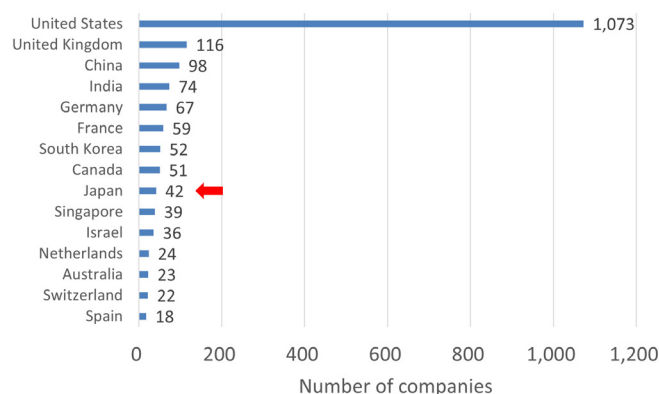


図 2 新たに資金調達を受けた AI 企業数（国別・2024 年）

以上の背景を踏まえ、本プロジェクトでは、超高齢社会における労働力不足に対して、人材の採用・育成・知識活用を支える AI/MAS 技術にてアプローチする。

本稿の構成は次のとおりである。第 2 章から 4 章で各小テーマについて記述する。第 2 章では、人材の入口である人材採用に着目し、LLM を活用した AI 面接訓練による就職支援について論じる。続く第 3 章では、人材の成長に焦点を当て、短期的業務効率化と長期的技術者育成の両立について述べる。第 4 章では、人材の定着に対して知識活用に着目し、対話型 AI における知識ギャップ検出と補完の自律制御を取り上げる。最後に第 5 章では、本研究のまとめと今後の展望を述べる。

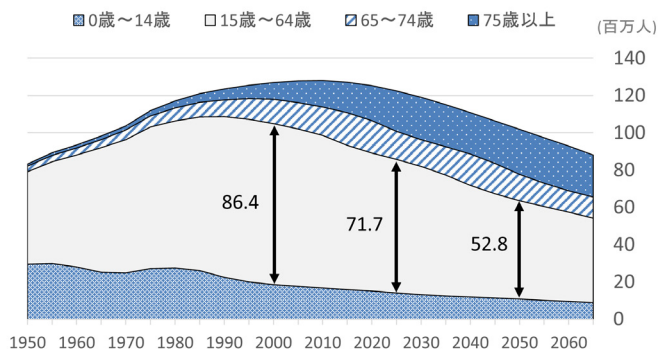


図 1 年齢別人口推計の推移（平成 29 年推計）

（注 1）2016 年以降は、将来推計人口は出生中位（死亡中位）推計による。

（注 2）2015 年までは総務省「国勢調査」（年齢不詳をあん分した人口）による。

2 LLM を活用した AI 面接訓練による就職支援

背景

2022 年の調査によると、面接を経験した 23 年卒の大学生・大学院生を対象とし、面接が得意であるかを対面式の面接とライブ形式の WEB 面接で質問したところ、対面式の面接の場合は「得意（どちらかといえば含む）」は 53.2%、「苦手（どちらかといえば含む）」は 46.8% となり、ライブ形式の WEB 面接の場合は「得意（どちらかといえば含む）」は 53.4%、「苦手（どちらかといえば含む）」は 46.6% となった [6]。22 年卒を対象とした質問では、対面式の面接の場合は「得意（どちらかといえば含む）」は 50.9%、「苦手（どちらかといえば含む）」は 49.1% となり、ライブ形式の WEB 面接の場合は「得意（どちらかといえば含む）」は 54.2%、「苦手（どちらかといえば含む）」は 45.8% となった [6]。この結果から、対面式の面接とライブ形式の WEB 面接の両方で「得意（どちらかといえば含む）」の割合は「苦手（どちらかといえば含む）」の割合を上回ったが、どちらの形式の面接でも半数近くは苦手であると回答した。加えて、同調査内で対面式の面接やライブ形式の WEB 面接が苦手と回答した学生にどのような点が苦手か複数回答で質問したところ、「緊張する（62.4%）」「言葉に詰まる（48.7%）」「自分に自信がない（46.9%）」「自分の意見や考えを上手く伝えられない（43.9%）」「自分の良さを上手く表現できない（42%）」が上位となった [6]。この調査結果から、半数近くの大学生・大学院生は面接が苦手であり、その主な要因として「言葉に詰まる」・「自分の意見や考えを上手く伝えられない」・「自分の良さを上手く表現できない」といった言語的な表現力やアピール力の課題が挙げられている。

関連研究

LLM を活用した AI 面接訓練の関連研究は、対話型フィードバック、LLM を使用した評価の 2 つに大別される。

対話型フィードバックに関しては、LLM によってフィードバックをチャット形式で進行させることで学習体験を向上させる手法が提案されている [7]。

LLM を使用した評価に関しては、主に LLM を利用した文章生成や文章評価 [8-11] と、LLM-as-a-Judge [12-15] である。文章生成や文章評価は、AI 面接訓練の際に回答を評価するため参考になっている。LLM-as-a-Judge に関連する論文は、LLM による評価作業という点において、就活生の発話内容を LLM で評価する本研究と類似点が見られる。

目的

本研究の目的は、言語的な表現力やアピール力に不安を抱き、面接が苦手である就活生を対象として、言語的な側面から面接練習サービスを提供し、発言内に含まれる長所や経験といった細かな言語的要素を評価することで、表現力や自己アピール能力を最大化させることである。

問題定義

対話型フィードバックを活用した面接練習の先行研究において、延々と修正点が出続けてしまう「修正ループ」状態があると言及

されている [7]。この「修正ループ」によって、ユーザーは自己否定感や焦燥感を感じ、学習に対するモチベーションが低下したとされている。また、「修正ループ」による自己否定感や焦燥感によって、学習の目的が「AI に承認されるための回答を考える」ことに偏ってしまい、実際の面接での最良の答えを見つけることがおろそかになる可能性があると言及されている [7]。そのため本研究では、少ない回数の「修正ループ」にすることを課題とした。

解決手法

LLM を利用して就活生の発話内容を解析し、言語的に長所や経験をアピール出来ているかを定量的・定性的にフィードバックする。ユーザーはフィードバックを受けた後に再度挑戦し、フィードバックを受ける。この繰り返しによって、少しずつ言語的な表現方法が上達する。フィードバックの手法として、対話型フィードバックを採用する。これは、対話形式による建築的なフィードバックや、フィードバックへの質問とそれへの回答の流れで構成される対話によって、ユーザーの理解を深めて学習体験を向上させることを目的としている。この時、「修正ループ」が発生しないよう、発言の段階評価や点数化・課題や問題の段階的な評価によってゴールを設定できるようにし、ループを脱却出来るようにする。

また、発言の段階評価や点数化・課題や問題の段階的な評価を出力する際に、LLM の精度を LLM 自身に評価させる「LLM-as-a-Judge」という仕組みで用いられる手法を取り入れることで、評価基準の安定化や評価理由をより明確化させる。これによって、ユーザーは「何故この修正が必要であるのか」が簡単に分かるようになる。

支援は図 3 のように 5 段階で構成されており、各段階はそれぞれ以下のようになっている。

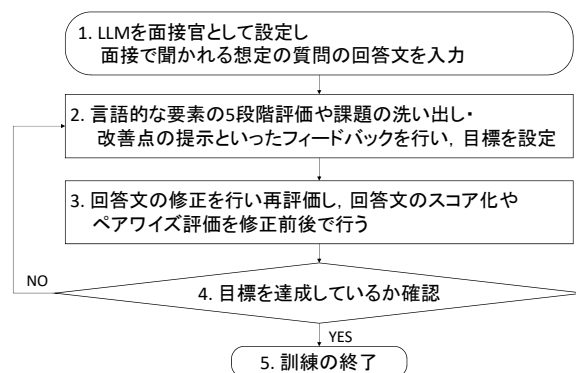


図3 考案中の支援の流れ

1 段階目は、面接で聞かれると想定される質問の回答文を入力するところから始まる。質問には予め設定されているものの他に、訓練したい質問を自分で設定することもできる。ユーザーはその質問に対する回答を入力する。入力音声入力と文字での入力を想定する。

2 段階目は、入力された回答文に対して LLM を利用し、言語的な要素の 5 段階評価や課題の洗い出し・改善点の提示などを行っ

たのちに対話型フィードバックを行う。この時、この訓練における目標を設定する。これは終了の条件に利用される。

3段階目は、ユーザーは2段階目でのフィードバックを元に、回答文の修正を行う。LLM は回答文のスコア化やペアワイズ評価を修正前後で行う。

4段階目は、修正された回答文が設定された目標に達しているかを判断する。設定した目標に達していない時、未達事項の重要度を考慮して改善の必要があると判断した場合は、2段階目に戻りフィードバックと修正を行い、訓練を続ける。修正を繰り返す間に新たな課題が発生した時、課題の重要度が高いと判断した場合は、これも修正対象とし目標に追加する。

目標に達していると判定された場合、5段階目の訓練の終了に移行する。また、設定した目標に到達している事のほかに、設定した目標に達していなくても未達事項の重要度が低い場合や、直近の文章の評価や課題が大きく変化せず頭打ちになったとユーザーが判断した場合は、ユーザーが自主的に訓練の終了に移行できるよう設計する。

評価軸

評価軸は① 言語的な要素の5段階評価と② 文章の課題の解決数を訓練開始時点と終了時点で比較し、長所や経験をアピールする能力を訓練出来たかを確認する。また、少ないループ回数で言語的な表現方法の上達や、より多い・より重大な課題を解決できることがベストであるとし、③ 修正ループの回数をそれに対応する評価軸とする。

①の「言語的な要素の5段階評価」は、支援内容の2段階目の部分で記載した通り LLM によって判定を行う。この時、LLM-as-a-judge の仕組みを応用し精度の確保する。また、②の「課題の解決数」については、少ないループ回数でより多い・より重大な課題を解決できることがベストであるとし、大まかな算出方法を「解決した課題の重みの総和を修正ループの回数で割る」とする。

まとめと今後の研究計画

本節では、AI 面接を活用した面接練習に着目し、修正ループへの対策と LLM-as-a-Judge の手法の活用、対話型フィードバックを組み込んだ AI 面接訓練を取り上げた。

今後は、回答の段階評価や点数化・課題や問題の段階的な評価を行うプロンプトを改良し、AI エージェント同士での面接練習シミュレーションなども活用し、提案手法の有効性の定量的検証に取り組む。

3 短期的業務効率化と長期的技術者育成の両立

背景

日本では今、熟練技術者の大量退職と若手の人材不足が同時に進行している。現場では「即戦力の活用」が求められるが、その一方で「新人の育成」は後回しにされがちである [16]。技術者チームでも、「目の前の成果を優先すれば育成が進まない」「育成に力を入れると業務効率が落ちる」というトレードオフが日常的に起きている。

関連研究

関連研究ではクラウドソーシングが挙げられ、作業者・タスクの異質性を前提に、品質・コスト・効率の調和を図る方法論が体系化されている [17-22]。代表例は、補完的スキルによるチーム形成・割当最適化、オンライン到着タスクへの逐次割当、および階層的スキル木による適合度最適化からなる一連の研究群である [18-22]。また、経験学習の体系化、組織内知識伝播のモデル化、手戻りリスクを考慮した優先規則、ならびに複合タスクの視点は、能力ー需要ー資源の適合という設計原理を補強する [23-26]。本研究はこれらの知見を接続し、タスク割当とスキル成長を同一枠組で評価する。

目的

本研究の目的は、こうしたジレンマに対して、タスクの割り当て戦略によって、短期成果と長期育成のバランスがどう変わるのかを、シミュレーションで定量的に評価・可視化することである。あわせて、現実の運用に資する意思決定指針（どの状況でどの戦略が望ましいか）を抽出することを目指す。

問題定義

ここで、本研究における2つの指標を定義する。まず「短期成果」である、これはたとえば「今、このステップで何件タスクを処理できたか」といった仕事の速さや効率である。一方でもう一つの「長期育成」は、特に新人技術者がどれだけ成長したかという、人材の伸び具合を示す指標となる。この2つは往々にして反発し合い、即戦力にばかりタスクを振れば、当然早く終わるが、新人が育たず将来的な生産性が伸びない。逆に新人に多く任せれば育成は進むものの、目の前の仕事が滞ってしまう。そこで本研究では、「どのような割り当て戦略を取れば、短期・長期の両方を高められるか？」を探ることが問題設定となる。

解決手法

この問いに対して、マルチエージェントシミュレーションの構築を行う。まず、技術者エージェントを設定する。このうち何名かは完全な新人で、スキル値が最低値からのスタートとする。各タスクは、ネットワーク／データベース／ソフトウェア／セキュリティの4カテゴリで構成され、毎ステップごとに複数タスクが生成される。このタスク群を、以下の3つの戦略で割り当てて、一定ステップ数でその効果を比較した。1つ目は、best 戦略。スキル値が最も高いエージェントに優先的に割り当てる。これは即戦力重視型となる。2つ目は、random 戦略。完全にランダムにタスクを配分するベースライン比較用となる。3つ目が、新人や中堅を優先戦略。スキルが平均以下のエージェント、新人や中堅を優先する育成重視型となる。この3つの方針が、チーム全体の処理効率や新人の成長にどう影響するかを可視化で示す。(1) タスク処理時間: ステップ毎のエージェントがタスクを実行する時間を可視化 (2) スキルの成長率: 各エージェントが初期からどれだけ成長したかの平均を可視化。これらを戦略間で比較し、短期・長期の折衷点を可視化することで、運用上の意思決定のための知見を提示する。以下の図4は今回のシミュレーションのフローチャートとなる。

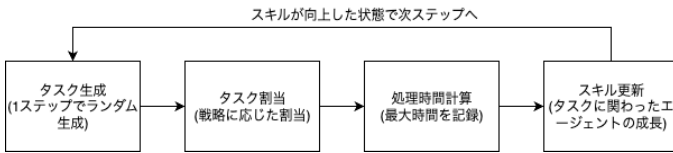


図4 フローチャート

- タスク生成：各ステップ s でタスク需要ベクトル $\mathbf{v}_s = (v_{NW}, v_{DB}, v_{SW}, v_{SEC})$ を確率モデルに基づき生成する (v は各カテゴリの件数である)。
- 割り当て：所与の戦略 (best/random/新人や中堅を優先戦略) に従い、タスクをエージェントへ割り当てる。ステップ s の「エージェント × カテゴリ」割当の集合を Assignments_s とする。
- 処理時間計算：割当ペア $(i, k) \in \text{Assignments}_s$ に対し処理時間

$$T(i, k) = \frac{\text{difficulty}_k}{\text{skill}_{i,k}} \quad (1)$$

を計算する ($\text{difficulty}_k > 0$: カテゴリ k の基礎難易度, $\text{skill}_{i,k} \geq 0$: エージェント i の当該スキル)。

- 記録：各ステップ s のボトルネック時間 (並行実行下の最大所要時間) を

$$T_{\max}(s) = \max_{(i,k) \in \text{Assignments}_s} T(i, k) \quad (2)$$

とし、完了件数・未処理件数、割り当てログを保存する。

- スキル成長：タスク完了後、担当カテゴリ k のスキルを

$$\text{skill}_{i,k} \leftarrow \min\{\text{skill}_{i,k} + \text{learn_rate}_k, \text{skill_cap}\} \quad (3)$$

に更新し、累積学習効果を反映させる ($\text{learn_rate}_k > 0$: カテゴリ k の学習レート (1件あたりの増分), skill_cap : スキル上限値)。ここで $\min\{\cdot\}$ は上限が決まっていることを意味する。

評価軸

評価は次の軸で行う。(1) 短期効率: 各ステップの処理ボトルネック時間 (式2) 及びその推移統計。(2) 長期育成: 期間終了時のカテゴリ別平均スキル上昇 (式3の累積値)。

まとめと今後の研究計画

今回の予備実験で、新人や中堅を優先すると、長期的に全体のスキルアップに繋がるが、短期的効率は悪く best 戦略だと、短期的効率は良いが、長期的に全体のスキルアップはできないことがわかった。今後は、シミュレーション設定を現実に近づけた上で「短期的効率」と「長期的スキルアップ」をバランスよく両立するためのアルゴリズムの新規開発に取り組む。

- 短期効率：best が最良である。処理時間の小ささと安定性で優位に立つ。
- 学習の広がり：random は効率の不安定性と引き換えに、平均スキルの上昇が広範に生じる。
- 育成投資の回収：新人や中堅の優先戦略は初期では遅いが、中盤以降に処理時間が通減し、弱点カテゴリの底上げが顕著である。

4 対話型 AI における知識ギャップ検出と補完の自律制御

背景

日本は急速な超高齢社会に突入しており、熟練技術者の退職に伴う知識や技能の消失が深刻な課題となっている。この状況において、企業や組織では LLM を活用したナレッジ共有や自動応答システムの導入が進められている。しかし、現場適用にはいくつかの課題が残されている。第1に、機密性の高い社内知識の取り扱いが難しく、クラウド依存型 AI の利用が制限される点であり、国内ガイドラインも機密データの取り扱い・偽情報対策 (RAG 等の活用を含む) を明確に求めている [27]。第2に、LLM が出力する誤情報 (ハルシネーション) により、誤答や不完全な情報が業務に混入する危険性がある点であり、公的ガイドラインおよび産業応用の研究でも検出・抑制の必要性が繰り返し指摘されている [27,28]。

関連研究

対話型 AI の信頼性向上や知識補完に関する研究は、主に RAG (Retrieval Augmented Generation) による知識補強、ハルシネーション抑制と信頼性評価、ユーザーへの情報欠落提示の3つに大別される。

RAG による知識補強に関しては、Web 分野の主要国際会議である WWW 2025 において、HTML 構造を保持することで従来のプレーンテキスト抽出に比べ、より正確な情報伝達を可能とする HtmlRAG が提案されている [29]。この手法は HTML クリーニングとブロックツリー型プルーニングを組み合わせ、構造化情報を活用することで RAG の性能向上を実現している。同様に、NeurIPS 2024 でもドキュメント構造を考慮した RAG 最適化手法が報告されており、情報欠落の抑制につながる事が示されている [30]。また、情報検索分野の国際会議である SIGIR-AP 2024 においては、検索文書の精度と更新性を同時に考慮する最適化が提案されており、情報鮮度を評価する視点が強調されている [31]。

LLM の信頼性評価に関しては、出力に不確実性推定を導入することでハルシネーションを早期検出する試みが報告されている [32]。さらに、複数回の再生成による一貫性評価 [33]、複数モデル間での応答比較と統合による信頼性向上 [34] が有効であることが示されている。NeurIPS 2024 の別研究では、外部知識ベースを用いたファクトチェック枠組みが提案され、応答の事実性を補強する有効性が報告されている [35]。

一方で、ユーザーに不足情報を明示し、補完を促す研究は限定的である。NAACL 2024 の産業応用研究では、実運用において利用者が不足部分を理解し補う仕組みの重要性が指摘されている [36]。また、経営学領域の研究では、暗黙知の欠落が知識継承や組織学習の阻害要因となることが報告されている [37]。産業保健分野においても、高齢労働者の知識伝承が業務効率や安全性に直結する課題として論じられている [38]。

これら先行研究の知見を踏まえると、本研究が目指す「知識ギャップの自動検出」と「非 IT 人材にも理解可能な自然言語による補完要請文の生成」は、既存研究には見られない独自のアプローチであり、学術的にも実務的にも意義を有する。

目的

本研究の目的は、超高齢社会における知識継承の断絶や人材不足という社会的課題に対し、LLM を活用して知識活用を支援する枠組みを構築することである。背景で述べたように、現場では熟練者の退職による知識消失、クラウド依存や機密性に起因する導入制約、そしてハルシネーション混入の危険性が課題となっている。また、関連研究の調査からも、既存のアプローチは主として「正確な回答生成」や「事実照合」に焦点を当てており、知識の欠落を利用者に伝え補完を促す仕組みは十分に検討されていないことが明らかとなった。

そこで本研究では、正答精度の追求にとどまらず、不足部分を明確化して知識継承を促進する新しい対話型 AI の枠組みを提案し、社内有機者や退職者の断片的知識を有効に活用できる仕組みを目指す。

問題定義

対話型 AI は、応答において根拠が不明瞭な部分や一貫性を欠く部分、あるいは社内存在しない情報に依拠する部分を含むことが多い。これらはユーザーにとって「知識ギャップ」となり、回答の信頼性や実務適用性を損なう要因となる。既存研究では、出力の再生成やファクトチェックにより精度を高める手法が提案されているが、「どの要素が不足しているのか」を明示し、利用者が補完行動に移れる形で提示する仕組みは不十分である。また、AI の導入やカスタマイズには高度な IT 知識を要し、教育・運用コストが高いことも課題である。経済産業省の DX 関連レポートや直近の調査では、この要因として人材・ガバナンス・データ基盤の不足がボトルネックとなることが示されている [4,39]。

したがって、本研究では応答内の知識ギャップを体系的に検出し、その属性（社内知識不足・外部情報不足・曖昧表現など）を分類した上で、利用者に補完を促す要請文として提示することを課題とする。

解決手法

本研究の提案手法は、対話型 AI の応答から不足要素を抽出し、行動可能な要請文へ接続するまでを① 信頼度・整合性の多角的判定と② 知識属性の判定に基づく補完アクション分岐の2段階で構成する。第1段階では応答の妥当性を複眼的に評価して「知識ギャップ候補」を抽出し、第2段階では不足の性質を判別し、依頼先・必要資料・期限・提出形式を含む要請文へと落とし込む。

まず第1段階の信頼度・整合性の多角的判定では、応答を主張単位に分解し、自己信頼度の手掛かり（曖昧表現や対数確率の校正値）、RAG ソースとの整合性（内容一致の被覆と参照文書の更新日による鮮度）、再生成による一貫性（同一クエリに対する複数出力の揺らぎ）、外部ファクト検証（公開データベースや外部 LLM による照合）、類似クエリ履歴との整合（過去の承認応答との矛盾や抜け漏れ）を定量化する。これらを総合して総合信頼度を次式で定義し、所定の運用しきい値を下回る主張を「根拠が弱い／欠落している」とみなしてギャップ候補とする。

$$R = \alpha S_{\text{self}} + \beta S_{\text{rag}} + \gamma S_{\text{cons}} + \delta S_{\text{fact}} + \varepsilon S_{\text{hist}}, \quad \sum \alpha_i = 1 \quad (4)$$

ここで、 S_{self} は自己信頼度の手掛かり、 S_{rag} は参照文書との一致

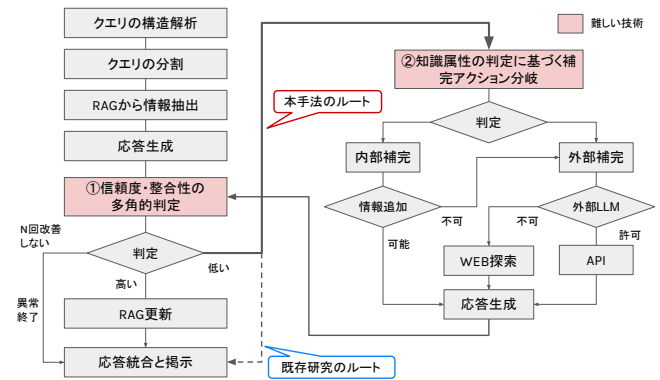


図5 考案中の ChatAgent のフレームワーク

度と鮮度を統合した指標、 S_{cons} は再生成一貫性、 S_{fact} は外部照合の支持度、 S_{hist} は履歴整合である。各スコアは $[0,1]$ に正規化し、重み $\alpha, \beta, \gamma, \delta, \varepsilon$ はドメインや運用目的に応じて校正する。

つづく第2段階の知識属性の判定と補完アクション分岐では、抽出したギャップ候補を「社内資料不足」「外部情報不足」「曖昧表現（定義・前提の不足）」に分類する。分類結果に応じて、社内資料不足であれば該当部署・担当（例：保全部門、品質保証部）に対する資料アップロードや版数・章節の特定依頼を、外部情報不足であれば公開データベース照会や外部 LLM の限定的参照を促す。曖昧表現の場合は、対象機種・適用範囲・前提条件などのクエリ再定義を提案する。生成される要請文は次の行動をアシスタントすると共に、可視化 UI 上で根拠リンク・版数・更新日とともに提示してトレーサビリティを担保する。分類後の社内知識に関する不足については、どの部署・担当に照会すべきか、あるいはどの社内データベースを参照すべきかを明示する仕組み、外部情報に関しては、参照可能な公開データベースや外部 LLM への照会を限定的に案内する仕組みの構築が求められる。利用者が具体的な補完アクションに直結できることが期待されるが、社内の担当部署やデータベースとの対応付けを動的に判定する仕組みは、組織固有の構造や知識管理の実態に依存するため実装が難しく、本研究における今後の重要な課題となる。

以上の2段階を統合するために、図5に示す自律制御型 ChatAgent のフレームワークを提案する。機能は、クエリの構造解析と分割 (*plan*)、RAG からの情報抽出 (*retrieve*)、応答生成 (*synthesize*)、信頼度・整合性の多角的判定 (*assess*; 式(4)に基づく R の算出)、知識属性の判定に基づく補完アクション分岐 (*classify & request*)、RAG 更新、および最終応答の統合提示を順次実行する。RAG 更新は判定スコアが既定値を上回る場合に限定し、不確実な情報による誤更新を防ぐ。図5に示すように、既存研究のルートが応答精度の向上に留まり、不十分な判定時の対処を欠いていたのに対し、本研究のルートは判定を挟んだ RAG 更新を導入することで、最終応答の信頼性向上と知識コーパスの持続的な進化を可能にする。

評価軸

評価は、最終応答の正確性と信頼性を中心に行う。外部ベンチマークは TruthfulQA を利用予定である。また、内部指標として Coverage（一致率）、Freshness（情報鮮度）、Self-report（自己信頼度）、Support strength（検索スコア）を用いる。さらに、応答の揺らぎや矛盾率を補助指標として設定し、定量的に評価する。

まとめと今後の研究計画

本節では、人材の定着に対して知識活用に着目し、対話型 AI における知識ギャップ検出と補完の自律制御を取り上げた。本研究により、応答の信頼度や知識属性に基づく制御の枠組みを提示し、知識継承を支援する仕組みの基盤を示すことができた。今後は、社内文書や現場の点検業務といった具体的タスクへの適用を想定し、自律制御型チャットエージェントの実装を進めるとともに、最終応答の正確性と信頼性の評価や提案手法の有効性の定量的検証に取り組む。

5 まとめ

本研究は、超高齢社会における労働力不足という社会的課題に対して、人材の採用・育成・知識活用を支援する AI/MAS 技術の活用をテーマとした。各小テーマごとに、人材の入口から成長、そして定着に至るまでを一貫して構想整理し、その有効性と今後の方向性を示した。今後は、それぞれの提案手法の実装を進め、各サブテーマの提案手法の有効性を検証していく。

第 2 章では、人材の入口である人材採用に着目し、LLM を活用した AI 面接訓練による就職支援について論じた。今後は回答の段階評価や点数化・課題や問題の段階的な評価を行うプロンプトを改良し、支援のテストを行う予定である。

第 3 章では、人材の成長に焦点を当て、短期的業務効率化と長期的技術者育成の両立について述べた。今後はよりリアルなシミュレーションとして OJT 要素を取り入れ、スキルの高いエージェントと低いエージェントが協力することにより、タスクをこなすことに加えて、スキルの成長を促すことを目指していく。

第 4 章では、人材の定着に対して知識活用に着目し、対話型 AI における知識ギャップ検出と補完の自律制御を取り上げた。今後は、社内文書や現場の点検業務といった具体的タスクへの適用を想定し、最終応答の正確性と信頼性の評価、提案手法の有効性を検証する。

参考文献

- 国立社会保障・人口問題研究所. 日本の将来推計人口 (平成 29 年推計). 国立社会保障・人口問題研究所; 2017 Jul. Available: https://www.ipss.go.jp/pp-zenkoku/j/zenkoku2017/pp29_ReportALL.pdf
- 統計情報課. 令和 6 年 (2024 年) 人口動態統計月報年計 (概況). 厚生労働省; 2025. Available: <https://www.mhlw.go.jp/toukei/saikin/hw/jinkou/geppo/nengai24/index.html>
- 内閣府. 令和 5 年版 高齢社会白書 (全体版). 内閣府 (Cabinet Office, Government of Japan); 2023. Available: https://www8.cao.go.jp/kourei/whitepaper/w-2023/zenbun/05pdf_index.html
- デジタルトランスフォーメーションに向けた研究会. DX レポート～IT システム「2025 年の崖」の克服と DX の本格的な展開～. 経済産業省; 2018 Sep. Available: https://www.meti.go.jp/policy/it_policy/dx/20180907_03.pdf
- Stanford University, Institute for Human-Centered Artificial Intelligence (HAI). Artificial Intelligence Index Report 2025. Stanford HAI; 2025. Available: https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf
- 株式会社マイナビ. マイナビ 2023 年卒学生就職モニター調査 6 月の活動状況. 株式会社マイナビ HR リサーチ統括部; 2022. Available: <https://career-research.mynavi.jp/wp-content/uploads/2022/07/s-monitor-23-6-001.pdf>
- Daryanto T, Ding X, Wilhelm LT, Stil S, Knutsen KM, Rho EH. Conversate: Supporting reflective learning in interview practice through interactive simulation and dialogic feedback. Association for Computing Machinery (ACM). 2025;9: 1–32. doi:10.1145/3701188
- Nakamoto S, Okamoto Y, Nakakouchi T, Shimada K. Towards Human-Level Evaluation: Assessing the Potential of GPT-4 in Automated Evaluation and

- Feedback Generation on Japanese Essays. Proceedings of the 16th IIAI International Congress on Advanced Applied Informatics (IIAI-AAI 2024). 2024. pp. 156–161. doi:10.1109/IIAI-AAI63651.2024.00039
- 中辻, 松本, 鈴木, 野本, 佐藤. 人間の意図を容易かつ正確に反映する LLM エージェントを用いた創造的文書作成支援技術. 人工知能学会全国大会論文集 第 39 回 (2025). 2025. pp. 3J6GS501–3J6GS501. doi:10.11517/pjsai.JSAI2025.0_3J6GS501
- Li R, Marrese-Taylor E, Matsuo Y. Evaluating Japanese Language Proficiency in Large Language Models through Definition Modeling Techniques. 人工知能学会全国大会論文集 第 38 回 (2024). 2024. pp. 3Q5IS2b02–3Q5IS2b02. doi:10.11517/pjsai.JSAI2024.0_3Q5IS2b02
- Shankar S, Zamfirescu-Pereira J, Hartmann B, Parameswaran A, Arawjo I. Who Validates the Validators? Aligning LLM-Assisted Evaluation of LLM Outputs with Human Preferences. Proceedings of the 37th Annual ACM Symposium on User Interface Software and Technology. Association for Computing Machinery (ACM); 2024. pp. 1–14. doi:10.1145/3654777.3676450
- Gu J, Jiang X, Shi Z, Tan H, Zhai X, Xu C, et al. A survey on LLM-as-a-judge. arXiv preprint arXiv:2411.15594. 2024. doi:10.48550/arXiv.2411.15594
- Gebreegziabher SA, Chiang C, Wang Z, Ashktorab Z, Brachman M, Geyer W, et al. MetricMate: An Interactive Tool for Generating Evaluation Criteria for LLM-as-a-Judge Workflow. Proceedings of the 4th Annual Symposium on Human-Computer Interaction for Work. 2025. pp. 1–18. doi:10.1145/3729176.3729199
- Szymanski A, Ziems N, Eicher-Miller HA, Li TJ-J, Jiang M, Metoyer RA. Limitations of the LLM-as-a-Judge Approach for Evaluating LLM Outputs in Expert Knowledge Tasks. Proceedings of the 30th International Conference on Intelligent User Interfaces. 2025. p. 952–966. doi:10.1145/3708359.3712091
- Wang R, Guo J, Gao C, Fan G, Chong CY, Xia X. Can LLMs Replace Human Evaluators? An Empirical Study of LLM-as-a-Judge in Software Engineering. Association for Computing Machinery (ACM). 2025;2: 1955–1977. doi:10.1145/3728963
- Nakatani H. Population aging in Japan: policy transformation, sustainable development goals, universal health coverage, and social determinates of health. Global Health & Medicine. 2019;1: 3–10. doi:10.35772/ghm.2019.01011
- Hettiachchi D, Gonçalves J, Kostakos V. A Survey on Task Assignment in Crowdsourcing. ACM Computing Surveys. 2022. Available: <https://dl.acm.org/doi/10.1145/3494522>
- Ho C-J, Jabbari S, Vaughan JW. Adaptive Task Assignment for Crowdsourced Classification. Proceedings of the 30th International Conference on Machine Learning (ICML 2013). 2013. Available: <https://proceedings.mlr.press/v28/ho13.html>
- Wang D, Ding W. A Hierarchical Pattern Learning Framework for Forecasting Extreme Weather Events. Proceedings of the IEEE International Conference on Data Mining (ICDM 2015). Institute of Electrical and Electronics Engineers (IEEE); 2015. Available: <https://ieeexplore.ieee.org/document/7373429>
- Mavridis P, David Gross-Amblard ZM. Using Hierarchical Skills for Optimized Task Assignment in Knowledge-Intensive Crowdsourcing. Proceedings of the 25th International Conference on World Wide Web (WWW '16). Montréal, Québec, Canada: Association for Computing Machinery (ACM); 2016. p. 843–853. doi:10.1145/2872427.2883070
- Sepehr Assadi SJ Justin Hsu. Online Assignment of Heterogeneous Tasks in Crowdsourcing Markets. Proceedings of the AAAI Conference on Human Computation and Crowdsourcing (HCOMP 2015). 2015. Available: <https://ojs.aaai.org/index.php/HCOMP/article/view/13236>
- Tang F. Optimal Complex Task Assignment in Service Crowdsourcing. Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence (IJCAI-20). 2020. p. 1563–1569. Available: <https://www.ijcai.org/proceedings/2020/0217.pdf>
- Biskup D. A State-of-the-Art Review on Scheduling with Learning Effects. European Journal of Operational Research. 2008. Available: https://www.researchgate.net/publication/222866154_A_State-of-the-Art_Review_on_Scheduling_with_Learning_Effects
- 橋本, 藤原, 鈴木, 奥田, 伊勢, 塩谷. 組織における知識伝播過程のマルチエージェントシミュレーション. 電気学会論文誌 C (電子・情報・システム部門誌). 2013. Available: https://www.jstage.jst.go.jp/article/ieejc/133/9/133_1770/_article/-char/ja/
- 満行, 大和, 稗方, モーザー, 磯沼, 岡田, et al. システム開発プロジェクトにおける手戻りリスクを考慮したタスク優先ルール設計に関する研究. 日本機械学会論文集. 2016. Available: https://www.jstage.jst.go.jp/article/transjsme/82/835/82_15-00474/_article/-char/ja/
- Uçar B, Aykanat C, Kaya K, İkinci M. Task Assignment in Heterogeneous Computing Systems. Journal of Parallel and Distributed Computing. 2006;66: 32–46. doi:10.1016/j.jpdc.2005.06.014

27. 産業サイバーセキュリティセンター. テキスト生成 AI の導入・運用ガイドライン. Information-technology Promotion Agency (IPA); 2024. Available: https://www.ipa.go.jp/jinzai/ics/core_human_resource/final_project/2024/f55m8k000003spo-att/f55m8k000003svn.pdf
28. Wang S, Wang X, Mei J, Xie Y, Chen S-Q, Xiong W. Developing a Reliable, Fast, General-Purpose Hallucination Detection and Mitigation Service. Proceedings of the 2025 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 3: Industry Track). Albuquerque, New Mexico: Association for Computational Linguistics (ACL); 2025. pp. 971–978. doi:10.18653/v1/2025.naacl-industry.72
29. Tan Z, others. HtmlRAG: Unlocking the Power of Structured Information in Retrieval-Augmented Generation. Proceedings of the ACM Web Conference 2025 (WWW '25). Sydney, Australia: Association for Computing Machinery (ACM); 2025. doi:10.1145/3589334.3645678
30. Bo X, Zhang Z, Dai Q, Feng X, Wang L, Li R, et al. Reflective Multi-Agent Collaboration based on Large Language Models. Proceedings of the 37th Annual Conference on Neural Information Processing Systems (NeurIPS 2024). New Orleans, USA: Curran Associates, Inc.; 2024. Available: <https://openreview.net/forum?id=wWiAR5mqXq>
31. Diaz F, Drozdov A, Kim TE, Salemi A, Zamani H. Retrieval-Enhanced Machine Learning: Synthesis and Opportunities. Proceedings of the 2024 Annual International ACM SIGIR Conference on Research and Development in Information Retrieval in the Asia Pacific Region (SIGIR-AP 2024). Tokyo, Japan: Association for Computing Machinery (ACM); 2024. pp. 299–302. doi:10.1145/3673791.3698439
32. Xiao Y, Wang WY. Quantifying Uncertainties in Natural Language Processing Tasks. Proceedings of the AAAI Conference on Artificial Intelligence. 2019. pp. 7322–7329. doi:10.1609/aaai.v33i01.33017322
33. Park J, Kim G, Kang J. Consistency Training with Virtual Adversarial Discrete Perturbation. Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL 2022). Association for Computational Linguistics (ACL); 2022. pp. 5646–5656. Available: <https://aclanthology.org/2022.naacl-main.414.pdf>
34. Manakul P, Liusie A, Gales M. SelfCheckGPT: Zero-Resource Black-Box Hallucination Detection for Generative Large Language Models. Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing. Singapore: Association for Computational Linguistics (ACL); 2023. pp. 9004–9017. doi:10.18653/v1/2023.emnlp-main.557
35. Wei J, Yang C, Song X, Lu Y, Hu N, Huang J, et al. Long-form factuality in large language models. Proceedings of the 37th Annual Conference on Neural Information Processing Systems (NeurIPS 2024). 2024. Available: <https://neurips.cc/virtual/2024/poster/96675>
36. Zhao Y, Singh P, Bhathena H, Ramos B, Joshi A, Gadiyaram S, et al. Optimizing LLM-Based Retrieval Augmented Generation Pipelines in the Financial Domain. Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Industry Track). Mexico City, Mexico: Association for Computational Linguistics (ACL); 2024. pp. 279–294. doi:10.18653/v1/2024.naacl-industry.23
37. 三輪. IT・AI の進歩による仕事と働き方の変化——知識労働・感情労働・定型労働のマネジメントの展望——. 日本経営学会誌. 2020;44: 72–81. doi:10.24472/keiejournal.44.0_72
38. Matsumoto Y, Kaneita Y, Itani O, Otsuka Y. Development and validation of the Work Style Reform Scale. Industrial Health. 2023;61: 462–474. doi:10.2486/indhealth.2022-0090
39. 経済産業省 商務情報政策局 情報技術利用促進課, 独立行政法人情報処理推進機構. デジタルトランスフォーメーション調査 2024 の分析. 経済産業省; 2024 May. Available: https://www.meti.go.jp/policy/it_policy/investment/keiei_meigara/dx-bunseki_2024.pdf



Open Access This article is licensed under CC BY 4.0. To view a copy of this license, visit <https://creativecommons.org/licenses/by/4.0/>