

プレスリリース

AI は使えている。 では、その使い方で人は育っているか

－「見えない AI 誤用」の作業定義と 4 つの構造仮説を公開

問い・事実確認・判断・結果確認のどこで、必要な過程が省かれているのか。AI 利用の良し悪しを、操作の巧拙だけでなく「一つの仕事の設計」から検討する

組織行動科学®の研究・教育開発を行うリクエスト株式会社（本社：東京都新宿区、代表取締役：甲畑 智康）は、AI 時代の新たな実践研究テーマとして、「見えない AI 誤用」の作業定義、4 つの構造仮説、今後の検証方針を公表します。

今回着目するのは、AI が誤った情報を出すことだけではありません。

AI の出力に目立った誤りがなく、成果物も完成している。それでも、その仕事に必要な問い直し、事実確認、判断、結果確認が十分に行われていない可能性です。

当社はこの問題を、AI 利用者の知識や操作技術だけでなく、仕事の目的、情報へのアクセス、権限、上司の関わり、評価方法、結果から学ぶ機会を含む「仕事設計」の観点から検討します。

本発表は、AI 誤用の発生率や、AI による能力低下を実証した調査結果ではありません。企業支援の中で生じた問題意識を、既公表の実践研究と国内外の関連研究に照らし、今後、具体的な仕事事例から確かめるための研究課題として提示するものです。

AI の回答が正しいことと、その仕事が適切に進んでいることは同じではない

AI が作成した提案書に、事実の誤りがなかったとします。

それでも、その提案が相手の実際の課題に合っているか、判断に必要な条件が含まれているか、相手が使った後に何が起きたかは、別に確認する必要があります。

また、完成した提案書の品質が高いことだけでは、担当者がその仕事から何を学び、条件が変わった次の仕事で何を確かめられるようになったかまでは分かりません。

ここで区別したいのは、「AI が適切な出力をしたか」「相手に役立つ仕事になったか」「人と組織に必要な経験が残ったか」です。

当社の既公表資料でも、成果物の正確さに加えて、その出力を求めた仕事の捉え方や、受け手が使った後の結果まで確認することを提案してきました。今回の発表は、その問題意識を AI の利用過程に沿って具体化するものです。 [1]

「見えない AI 誤用」とは何か

本発表では、「見えない AI 誤用」を、現段階の作業定義として次のように整理します。

AI を用いる仕事において、その仕事の目的・リスク・維持すべき技能に照らして必要な問い直し、根拠の確認、判断、結果確認が省略または形骸化し、それを補う人や仕組みも十分に確保されていない利用状態。

「見えない」とは、発見できないという意味ではありません。**出力の見た目や、納品・処理の完了だけでは把握しにくい**という意味です。

また、ここでいう「誤用」は、悪意ある利用、規則違反、利用者個人の能力不足と同義ではありません。本人が適切に取り組もうとしていても、仕事の条件によって必要な確認や判断が行いにくくなっている可能性を含めて検討します。

一方、AI に問いの候補を出させること、情報を検索・照合させること、選択肢を比較させること、定型的な処理を自動化すること自体は、誤用とは扱いません。別の担当者や検証済みの仕組みによって必要な機能が確保されている場合も、一人の担当者がすべてを行わないという理由だけで誤用とはしません。

問題にするのは、AI に任せた量ではなく、その仕事に必要な機能が確保されているかです。

なお、生成 AI の誤った出力そのものは技術上のリスクであり、それだけで利用者の誤用を意味するわけではありません。出力の誤り、利用方法、仕事条件、実際に生じた不利益は、分けて確認します。

既存の AI リスク研究を、日常の仕事へ接続する

AI への過度な依存や、人と AI の不適切な関わりは、既に研究・ガイドラインで扱われています。

米国 NIST の生成 AI 向けリスク管理資料は、人と AI の関係を独立したリスク領域として取り上げ、自動化された出力を過度に信頼する「自動化バイアス」や過剰な依存を指摘しています。同時に、AI を必要以上に避けることで、有益な利用機会を失う可能性も扱っています。[2]

日本の総務省・経済産業省「AI 事業者ガイドライン（第 1.2 版）」でも、自動化バイアスなど、AI に過度に依存するリスクへの注意と対策が示されています。[3]

したがって、本発表は、これまで知られていなかったリスクを初めて発見したと主張するものではありません。

当社が今回提示するのは、これらの論点を、「一つの仕事」の問い・事実確認・判断・結果確認に沿って捉え直し、成果と学習の両面から検討するための実務枠組みです。

「見えない AI 誤用」の 4 つの構造仮説

以下の 4 つは、今後の観察と検証のために置く初期仮説です。実証済みの分類ではなく、互いに重なる場合もあります。すべての AI 利用上の問題を網羅するものでもありません。

1. 問いを任せすぎる

何を解決するかを、現実には照らさない

AI に問題の整理や問いの候補を出してもらうことは、有用な使い方になり得ます。

問題として検討するのは、そこで示された問題設定を、顧客の状況、現場の記録、仕事の目的に照らさず、そのまま確定してしまう状態です。

もっともらしい問いに対して適切な答えが得られても、その問いが実際に向き合うべき問題とずれていれば、必要な仕事にはつながりません。

確認する問い：その問題設定は、誰の、どの事実に基づいているか。何が分かれば、問いを変える必要があるか。

2. 事実確認を省く

AI の説明を、そのまま「確認済み」にする

AI が説明したことと、根拠を確かめたことを取り違える状態です。

例えば、数値の出典が示されていても、その資料に当該数値が存在するか、対象期間や集計条件が合っているか、自分たちの仕事に適用できるかは、別に確かめる必要があります。

AI による検索や照合を排除するわけではありません。問題は、確認方法として AI を使ったかどうかではなく、**説明から元の記録や現実の状況へたどれるか**です。同じ AI に「正しいか」と尋ね直すことや、複数の AI の回答が一致することだけでは、根拠の確認を終えたとは限りません。

確認する問い：何を、どの資料・記録・相手・現場で確かめたか。確認済みの事実と、推測や未確認事項は分かれているか。

3. 判断を任せすぎる

案を選ぶが、必要な基準や条件を確かめない

AI が示した複数案から、人が一つを選ぶ。選択することも判断の一部です。

ただし、案を選んだという事実だけで、必要な判断が十分に行われたとは限りません。

目的に合っているか、守るべき条件を満たすか、他の選択肢を検討する必要はないか、誰に影響するか、どの条件なら採用を見送るか。こうした確認がないまま、AI の推奨や文章の説得力だけで採用する状態を検討します。

必要なのは、すべての処理で人が承認ボタンを押すことではありません。**何を自動処理に任せ、何を人や組織が引き受け、どの条件で止めるかが定まっていることです。**

確認する問い：採用・不採用を分ける条件は何か。その条件を誰が定め、例外が起きたときに誰が判断するか。

4. 結果確認をしない

アウトプットの完成だけで、仕事を閉じる

文章、提案書、分析、回答、マニュアルなどが完成したことだけで、仕事を完了扱いにする状態です。

その仕事の目的が、相手の判断や行動を支えることであれば、相手が使えたか、後工程が進んだか、どこで迷いや手戻りが生じたかも確認対象になります。

すべての成果物を、担当者が一件ずつ長期間追跡することを求めるものではありません。重要度に応じた確認時点、担当者、指標、代表事例の振り返りなどを設け、結果を次の仕事へ戻せるかを見ます。

確認する問い：完成後の何を、誰が、いつ確かめるか。その結果を、次の判断基準や手順へどう戻すか。

問題を「本人が考えていない」で終わらせない

AI への指示を工夫する。出力を比較する。表現を修正する。違和感のある箇所を直す。

これらも、仕事の中で行われる思考です。本研究では、AI を使っているという理由だけで、「本人は考えていない」と評価しません。

検討したいのは、**何を考えているかと、どこまでを仕事の範囲として認識しているか**です。

提案書を分かりやすくすることに十分取り組んでいても、相手が何を判断するための提案書なのかは、確認対象に入っていないかもしれません。

また、依頼する側が目的や背景を伝えず、完成物だけを求めている可能性もあります。確認に必要な情報へアクセスできない、判断する権限がない、結果を確認する担当が決まっていないといった条件も、検討対象です。

そこで本研究では、次の仮説を置きます。

必要な確認や判断が省かれる背景には、個人の AI リテラシーだけでなく、仕事の目的の共有、情報、権限、評価、上司の関わり、結果確認の仕組みが関係しているのではないかと。

個人の知識や技能を軽視するのではなく、それらを実際に使える仕事条件も併せて確かめます。

外部研究が示しているのは、AI の効果が一方向ではないこと

本研究は、「AI を使うと人が育たなくなる」という前提から出発しません。関連研究には、生産性や学習への便益を示すものと、利用条件によって学習上の不利益が生じ得ることを示すものの両方があります。

実務では、生産性の向上と学習を示唆する結果がある

Brynjolfsson、Li、Raymond による、顧客対応者 5,172 人への生成 AI 支援の段階的導入を分析した研究では、時間当たりの問題解決件数で測った生産性が平均 15% 向上しました。AI の提案が利用できないシステム停止時の分析から、学習の持続を示唆する結果も報告されています。

ただし、一企業・特定職務での研究です。停止時の分析にも推定上の不確実性があり、あらゆる AI 利用が学習につながることを示すものではありません。[4]

学習場面では、支援の設計によって結果が異なる

Bastani らによる、トルコの高校で約 1,000 人を対象に行われた数学学習の無作為化比較実験では、一般的な対話型 AI に近い支援を使った群は、練習中の成績が高まる一方、AI を使えない試験では対照群を下回りました。

これに対し、答えを直接与えるのではなく学習を支えるよう設計した支援では、その負の影響が大きく軽減されました。これは高校数学という限定された条件の結果であり、企業内の仕事へそのまま当てはめることはできません。 [5]

批判的思考は、「なくなるか」だけでなく「何に向かうか」を見る必要がある

Lee らは、知識労働者 319 人から 936 件の生成 AI 利用事例を収集し、AI への信頼が高いことと、自己申告上の批判的思考の少なさとの関連を報告しました。

同時に、AI 利用に伴い、思考の対象が情報の検証、出力の統合、仕事全体の管理へ移ることも示しています。自己報告に基づく研究であり、AI が長期的な思考能力の低下を引き起こしたと証明するものではありません。 [6]

人を関与させれば、自動的に最適になるわけでもない

Vaccaro らによる 106 の実験研究の系統的レビュー・メタ分析では、人と AI の組合せは、平均すると「人単独または AI 単独のうち、成績のよい方」を下回りました。対象は 2020 年から 2023 年上半期に公表された研究であり、現在の AI 全般を評価した結果ではありません。

それでも、人が関与しているという形式だけではなく、どのように組み合わせるかを検証する必要があることを考える材料になります。 [7]

2026 年の職業技能に関する研究でも、成果と理解を分けて測っている

Shen と Tamkin による公開プレプリントでは、Python 利用経験者 52 人を対象に、未経験のライブラリを使う課題を AI 支援の有無で比較しました。AI 利用群の理解テストの得点が低い一方、平均的な作業時間の短縮は統計的に有意ではありませんでした。

これは AI 開発企業の研究者らによる短時間の限定課題であり、査読済み研究とは区別して扱います。また、長期的な技能変化や、他職種への一般化は確認されていません。 [8]

これらの研究は、今回の「4 つの構造仮説」を直接実証したものではありません。当社は、対象・課題・利用方法・測定指標を区別しながら、仕事の成果と人の学習を別々に確かめるための背景資料として参照しています。

今後の検証では、「成果」「学習」「組織への蓄積」「責任・負担」を分ける

検討の単位は、個人の AI 利用回数ではなく、具体的な「一つの仕事」です。

対象業務ごとに、何を達成する仕事なのか、どの程度のリスクがあるか、誰にどの技能を残す必要があるかを確認したうえで、次の 4 つを分けて見ます。

確認する対象	主に確かめること
仕事の成果	品質、処理時間、誤り、手戻り、相手や後工程への影響
本人に残る学習	類似場面で必要な事実を確認できるか。条件の違いを捉え、理由を持って対応できるか
組織への蓄積	経験が他者にも使える記録や判断基準となり、実際に再利用・更新されるか
責任と運用負担	判断や確認の担当が明確か。必要な停止・相談ができるか。確認や記録が過剰な負担になっていないか

学習を確認する場合も、業務上の安全を損なってまで AI を外すことはしません。安全な事例演習や、条件を変えた課題、実務での継続的な観察などを、目的に応じて検討します。

また、AI を使う前から存在した問題、AI 利用によって表面化・増幅した問題、AI によって軽減された問題を区別します。前後の違いだけで因果関係を断定せず、仕事内容、経験、上司の支援、利用した AI や設定の違いも確認します。

不利益が疑われる事例だけではなく、**AI によって成果と学習の両方が高まった事例や、4 つの仮説に当てはまらない事例**も検討対象とします。

記録がないことだけで「考えていなかった」とは判定しません。確認できないものは不明として扱い、本人の説明と実際の仕事の記録を照合します。事例の収集・利用では関係者の同意と情報管理に配慮し、本枠組みを個人の能力順位や処遇を決める尺度としては用いません。

すべての仕事を、学習課題に戻すわけではない

本研究が目指すのは、AI 利用を減らすことでも、人の作業や確認を無条件に増やすことでもありません。

既に習熟している定型作業の省力化や、表現の整形など、本人に新しい判断経験を残すことが主目的ではない仕事も、検討上は区別します。

また、必要な技能を別の仕事や教育機会で維持できている場合まで、「その場で学ばなかった」という理由で誤用とはしません。

重要なのは、**どの技能を、誰が、どの仕事や機会**で身につけ、維持する必要があるかが明らかになっていることです。

作業時間が短くなったことと、必要な判断経験が失われたことは同じではありません。AI への委任と、必要な機能の喪失を分けて確かめます。

「判断デザイン」「アート思考」との接続

当社はこれまで、「仕事が完了すること」と「人と組織に判断経験が形成されること」を区別し、AI を含む仕事の中に、事実確認、比較、選択、実行、結果確認をどう組み込むかを検討してきました。[9]

また、アート思考については、判断の前提となる見方・問い・価値仮説をつくり直す働きとして、当社の実務上の定義を提示しています。これは学術上の統一定義や、効果が確立した方法としての主張ではありません。[10]

今回の発表は、これらを別のテーマへ置き換えるものではありません。

何を問うかを見直すことと、必要な事実を確かめ、条件に照らして判断し、結果から学ぶこと。その接続が、AI を使う仕事のどこで途切れているのかを確かめるための研究課題です。

なお、当社が既公表資料で示している「980 社・33.8 万人」は、研究・教育開発の背景となる既存のデータ基盤であり、今回新たに実施した AI 誤用調査の標本数ではありません。本発表では、AI 誤用の発生率を算出していません。[11]

リクエスト株式会社 代表取締役 甲畑智康 コメント

企業の皆様と関わる中で、AI が使われているかどうかよりも、「AI を使って進めたことで、本来必要だった確認まで済んだことになっていないか」と感じるがあります。

AI に案を出してもらうことも、文章を書いてもらうことも、それ自体が問題なのではありません。本人も、その人なりに考えています。

だからこそ、本人に「もっと考えて」と求める前に、何を良くする仕事なのか、どこまでを確認する仕事なのかを、依頼する側も一緒に見直す必要があると考えています。

AI に任せることと、必要な確認や判断が失われることは同じではありません。何を任せ、何を確かめ、どの経験を人に残し、その結果を誰が次の仕事へ戻すのか。

今回は、その条件を実際の仕事から確かめるための問いを公表します。

本発表の位置づけと限界

「見えない AI 誤用」は、当社が本発表で用いる実務上の研究概念です。4 つの構造仮説を含め、現時点で信頼性・妥当性が確立した診断尺度ではありません。

本発表は、関連する一次研究、公的ガイドライン、当社の既公表資料を選択的に参照した、概念整理と検証方針の提示です。網羅的な系統的レビュー、新たな大規模調査、提案した仕事設計の効果検証ではありません。

「ほとんどの企業が AI を誤用している」「AI を使うと人の思考力が低下する」「AI を使わない方が人材育成に優れている」といった一般的な結論は示していません。

今後は、必要な確認や判断が確保される条件と、確保されにくい条件の双方を確かめ、作業定義や仮説そのものも見直していきます。

AI を使うかどうかだけではない。

**AI を仕事のどこに置き、何を確かめ、
どの経験を残すかを設計する。**

主な参照研究・公開資料

以下は、今回の論点整理に用いた主な資料です。個々の研究が当社の作業定義や4つの構造仮説を承認・実証していることを意味しません。

本文の[番号]は、この参考文献一覧へ移動します。各資料の「原文を開く」から参照先へ進めます。

[4] Brynjolfsson, E., Li, D., & Raymond, L. (2025)

Generative AI at Work.

The Quarterly Journal of Economics, 140(2), 889–942.

参照点：実務における生産性と学習を区別して検討するための研究。

[原文を開く](#)

[5] Bastani, H., et al. (2025)

Generative AI without guardrails can harm learning: Evidence from high school mathematics.

PNAS, 122(26), e2422633122.

参照点：AI 利用中の成績と、AI を使わない場面での学習成果、支援設計による違い。著者所属に関する訂正も確認しています。

[原文を開く](#) | [訂正記事](#)

[6] Lee, H.-P., et al. (2025)

The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers.

CHI 2025.

参照点：自己報告に基づく批判的思考の実行と、その対象の変化。

[原文を開く](#)

[7] Vaccaro, M., Almaatouq, A., & Malone, T. (2024)

When combinations of humans and AI are useful: A systematic review and meta-analysis.

Nature Human Behaviour, 8, 2293–2303.

参照点：人と AI の組合せの効果を、比較条件と課題に即して評価する必要性。

[原文を開く](#)

[8] Shen, J. H., & Tamkin, A. (2026)

How AI Impacts Skill Formation.

arXiv:2601.20245 (v2) . 公開プレプリント。

参照点：新しい職業技能を学ぶ短時間課題で、作業完了と理解を分けて測る試み。

[原文を開く](#)

[2] NIST (2024)

Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile.

NIST AI 600-1.

参照点：人と AI の関係、過度な依存、利用結果の評価に関するリスク管理。

[原文を開く](#) | [資料本文 \(PDF\)](#)

[3] 総務省・経済産業省 (2026)

AI 事業者ガイドライン (第 1.2 版)

2026 年 3 月 31 日公表。

参照点：自動化バイアスなど、AI への過度な依存に関する留意事項。

[原文を開く](#) | [本編 \(PDF\)](#)

関連する当社の公開資料

[1] リクエスト株式会社 (2026 年 9 月 9 日)

AI で早くつくれた。その先の仕事は、よくなったか。－「アート思考」と「判断デザイン」を一つの仕事でつなぐ理由を 22 ページで公開

PR TIMES | No.241

[原文を開く](#)

[9] リクエスト株式会社 (2026 年 8 月 31 日)

AI 時代、仕事は進む。では、人と組織に「判断経験」は残っているか－980 社・33.8 万人の既公表分析等を統合し、新たな知見を 18 ページで公表

PR TIMES | No.225

[原文を開く](#)

[10] リクエスト株式会社 (2026 年 9 月 2 日)

AI 時代、「アート思考」を曖昧なままにしない－「判断前提生成」として定義した図解 11 ページを公開

PR TIMES | No.229

[原文を開く](#)

[11] リクエスト株式会社 (2026 年 9 月 10 日)

AI 時代、「考える力」はどう育てるのか－「時間量による育成」から「判断経験を設計する育成」へ。図解 20 枚と研究レビュー38 ページを公開

PR TIMES | No.244

[原文を開く](#)

参照情報の確認基準日：2026 年 9 月 12 日

リクエスト株式会社について

リクエスト株式会社は、組織行動科学®を基盤に、組織で働く人の行動、判断、経験形成、組織学習に関する研究・教育開発と実践支援を行っています。

会社名：リクエスト株式会社

代表者：代表取締役 甲畑 智康

所在地：東京都新宿区新宿 3 丁目 4 番 8 号。 [12]

本研究テーマに関するお問い合わせ・取材のご相談は、当社公式サイトのお問い合わせ窓口で受け付けています。 [13]

[会社概要 \(公式サイト\)](#) | [お問い合わせ窓口](#)