

内なる他者との対話：生成AIを用いた統計演習における学習支援の実践報告 －AI支援型リフレクションモデルによるメタ認知と学習感情の変容－

小野原 彩香^A

要 旨：本研究は、初年次教育の統計演習において生成AI（ChatGPT, Gemini等）を「内なる他者」として位置づけ、メタ認知と学習感情への影響を検討した実践報告である。Mercier & Sperber (2011)の理性の相互作用仮説に基づき、理解不足を言語化させる「AI支援型リフレクションモデル」を提案した。モデルは理論層と操作化層からなる二層構造とし、操作化層として5段階構造化プロンプトと裏取りルールを授業に実装した。全4回の演習(N=90)でAI支援指数(5項目)とメタ認知指数(6項目)は高水準を維持した。Plutchikの8基本感情では驚き・喜びが高く恐れは低く、心理的安全性の高い学習環境が示唆された。自由記述の探索的分析では、課題特性(コード実装・プロンプト設計等)への言及がある学習者ほど両指数が高い傾向が見られた。AIが擬似的な対話相手となることで、個人学習でも社会的議論に近い省察が促され得る。

キーワード：生成AI, メタ認知, 学習感情, 内なる他者, 初年次教育

Dialogue with the Inner Other: A Practical Report on Learning Support Using Generative AI in Statistics Exercises: Transformation of Metacognition and Learning Emotions through an AI-Assisted Reflection Model

Ayaka ONOHARA^A

Abstract: This practical report examines the use of generative AI (e.g., ChatGPT, Gemini) as an “inner other” to support reflective, self-regulated learning in first-year statistics exercises. Grounded in Mercier and Sperber’s argumentative theory of reasoning and the notion of epistemic vigilance, we propose an AI-Assisted Reflection Model (AIRM) with a two-layer structure: (1) a conceptual layer positioning AI as a non-evaluative dialogue partner that recreates an argumentative context within individual learning, and (2) an operational layer that implements structured dialogue and verification routines. Five-stage structured prompts and an explicit “fact-checking” rule were applied across four exercise sessions in a large introductory course (243 students; 90 valid post-session responses). Students completed an AI Support Index (5 items), a Metacognition Index (6 items), and an emotion profile based on Plutchik’s eight basic emotions. Across sessions, both indices remained high (around 4/5), while emotions showed high surprise and joy and low fear (0–3 scale), suggesting a supportive learning climate. Exploratory analyses of open-ended responses indicated that students who referenced task-specific strategies (e.g., coding, prompt design, data selection) tended to report higher indices. A small longitudinal subset (n=3) showed a qualitative shift from procedural help-seeking to more conceptual questioning with repeated use. Limitations include self-selection and low response rates. We discuss the implications for designing reflective dialogue with generative AI in first-year and remedial statistics education.

Keywords: generative AI, metacognition, learning emotions, inner other, first-year education

1 はじめに

1.1 背景と問題の所在

初年次教育において、学習者が自らの理解度を客観的に把握し、学習行動を調整する「メタ認知」の育成は重要な

課題である。特に統計教育のような積み上げ型の科目では、つまづきを早期に発見し、再学習につなげるプロセスが不可欠となる。

近年、生成AI (Generative AI) の教育利用に関する研究が急速に蓄積されつつある。Kasneci et al. (2023)は、大規模言語モデル (LLM) の教育利用について、個別化されたフィードバックや対話的な学習支援の可能性を指摘し

A：立教大学 社会情報教育研究センター / Center for Statistics and Information, Rikkyo University

た。また、Wang & Fan (2025) による51研究のメタ分析では、ChatGPTを用いた教育介入が学習成果 ($g=0.87$) および高次思考 ($g=0.46$) に中〜大程度の正の効果を持つことが示されている。

しかしながら、これらの先行研究の多くは、AIを「知識提供者」または「解答生成ツール」として位置づけており、学習者の内面的な省察 (Reflection) を促すメカニズムについては十分な検討がなされていない。Tankelevitch et al. (2024) は、生成AIの利用がメタ認知的なモニタリングと制御を要求することを指摘し、AIシステムへの自己評価プロンプトの統合を提案しているが、具体的な教育実践への応用は今後の課題とされている。

さらに、Fan et al. (2025) は、構造化されていないAI利用が「メタ認知的怠惰 (metacognitive laziness)」を招き、省察や自己評価への関与を低下させる危険性を報告している。すなわち、AIを単に便利な道具として使うだけでは、深い学習には結びつかない可能性がある。

本研究は、この課題に対し、進化認知科学の知見を援用することで新たな視座を提供する。具体的には、AIを「内なる他者」(評価なき対話相手) として位置づけ、対話を通じて省察が生起するメカニズムを理論的に説明し、生成AIを「内なる他者」として配置した授業実践において、当該メカニズムが観察されるかを検証することを試みる。

1.2 理論的枠組み：進化認知科学と「内なる他者」

本研究の理論的基盤は、進化認知科学における二つの重要な概念、「理性の相互作用仮説 (The Argumentative Theory of Reasoning)」と「認識的警戒 (Epistemic Vigilance)」にある。

第一に、Mercier & Sperber (2011) によれば、人間の理性は、孤独な真理探究のためではなく、集団内での議論において他者を説得し、自己の正当性を主張するために進化したとされる。したがって、学習者が深い理解 (概念変容) に到達するためには、一方的な講義聴講ではなく、他者からの反論や問いかけに応答する「対論的な文脈」が不可欠である。

第二に、しかしながら、現実の初年次教育現場において、この「対論」は容易ではない。Sperberら (2010) が提唱する「認識的警戒」の観点からは、人間は誤った情報を掴まないよう他者の発言を警戒する一方で、自らの評判 (Reputation) を守るために、「無知だと思われたくない」「変な質問をして評価を下げたくない」という防衛本能を働かせる。特にリメディアル教育の対象となる学習者は、過去の学習体験からこの社会的コストを過大に見積もる傾向があり、教員や友人への質問行動が抑制されやすい (Help-seekingの回避)。

ここに、生成AIを「内なる他者」として導入する意義がある。AIは感情や社会的評価を持たないため、学習者は対人関係に伴う社会的コストを支払うことなく、初歩的な質問や反論の想定を繰り返すことができる。理論的基盤で

ある二つの概念を踏まえ、本研究では、他者からの反論や問いを想定しながら、自らの理解を説明・正当化・修正する対論的プロセスを社会的議論に近い省察と位置づける。そして生成AIは、この対論的文脈を評価リスクなしに個人学習内で再現する装置として機能する。

本稿では、この仕組みを自己調整学習支援のために体系化した枠組みを「AI支援型リフレクションモデル (AI-Assisted Reflection Model)」と呼び、(i) 理論層 (生成AIを内なる他者として配置し、対論的文脈を個人学習内に再現することで、理解の不足の言語化・説明・正当化を促す) と、(ii) 操作化層 (理論層を授業実践に落とし込むための対話設計・運用ルール・振り返り記録の手続き) からなる二層構造として定義する。

なお本研究では、操作化層の中核として、対話を段階化して省察を誘発する「5段階の構造化プロンプト (Five-Stage Structured Prompts)」を採用し (表1)、併せて誤情報を前提とした裏取り行動を明示的に組み込むことで認識的警戒を喚起した。

1.3 研究の目的

本研究の目的は、初年次科目における統計演習において、生成AIを用いた対話的振り返りを導入し、以下の点を検証することである。(1) メタ認知の変化 (または維持) : AIによるフィードバックが学習者のメタ認知 (理解の整理・不足点の自覚) に与える影響を明らかにする。(2) 学習感情の分析 : 学習過程における感情の変化 (AIに対する信頼、驚き、喜び等) と学習環境の関連を分析する。(3) 課題特性との関連 : 課題の性質 (コード実装、データ探索等) によってAI活用の焦点がどのように異なるかを探索的に検討する。

2 方法

2.1 対象授業と参加者

立教大学の初年次科目「景気・格差問題と統計情報」(2025年度秋学期) 受講生 (243名) を対象とした。分析対象は、以下の全4回の演習完了後に実施したアンケートの有効回答計90件である。第5回 ($N=26$) : 地理情報と移動平均 (季節調整の理解)、第7回 ($N=36$) : 人口統計 (少子高齢化・従属人口指数)、第9回 ($N=18$) : 労働統計 (失業率・求人倍率)、第11回 ($N=10$) : 賃金統計 (賃金カーブ・格差)。

2.2 実践手順：AI支援型リフレクションモデル

各回 (第5, 7, 9, 11回) の演習は100分で構成された。前半30分では、公的統計の取得と加工 (Python/Excelによる可視化) について、教員が自らの操作画面を提示しながらレクチャーを行った。続く後半70分は、学生が各自で課題に取り組む演習時間とし、教員は個別の質問やコードの不具合に対応する形で進行した。

本研究では、1.2で定義したAI支援型リフレクションモデルのうち、授業内で観察可能な「操作化層」を、(1) 対話の構造化、(2) 振り返りの記録と共有、(3) 倫理的配慮と裏

表1 AI支援型リフレクションモデルの操作化としての5段階構造化プロンプト (Five-Stage Structured Prompts) の構成と意図

段階	問かけの例(第11回賃金統計の場合)	教育的・認知的意図
1. 困難点の特定	「今日の演習で難しかった点は〇〇である(例:賃金分布の歪み)」	メタ認知モニタリング 自身の理解不足を言語化・客観化させる。
2. 概念の深掘り	「〇〇がよく分かりません。別の具体例を使って説明してください」	個別最適化された足場かけ／精緻化 個人のレベルに合わせた解説を引き出す。
3. 社会的接続	「この統計は社会問題のどのような理解に役立つと考えるか」	学習の転移 (Transfer) 抽象的なデータを現実社会と結びつける。
4. 感情の省察	「〇〇のグラフに驚いた。その理由は～である」	認知的感情の活用 「驚き」を知的探究心へと昇華させる。
5. 次なる学習	「今日理解したことを要約し、次に学ぶべきトピックを3つ提案してください」	自己調整学習 (SRL) 学習プロセスを総括し、次回の目標を立てる。

取り(認知的警戒)の三手続きとして実装した。なお「生成AIとの対話および振り返り」は、課題提出後に実施するよう指示した。これは、演習での試行錯誤を終えた直後に、自らの思考プロセスを客観視させることを意図している。

(1) 対話の構造化：5段階の構造化プロンプト (Five-Stage Structured Prompts)

自由な対話では学習者が単に正解や評価を求めてしまう懸念があったため、本実践では「AI支援型リフレクションモデル」の操作化層として、意図的に設計された「5段階の構造化プロンプト (Five-Stage Structured Prompts)」を導入した。学習者には、以下の5段階に対応する自己記述文を生成AIに入力するよう指示した(表1参照)。各段階において、学習者はあらかじめ提示された問いに基づき、自身の理解状況・疑問・感情を文章化し、それを生成AIに提示する。その後、AIの応答を参照しながら省察を深める。このプロセスは、単なる知識獲得ではなく、理解の不足の言語化、概念理解の深掘り、学習の転移、認知的感情の言語化、次の学習方略の生成を促すよう設計されている。

1. 困難点の特定：

学習者は、「今日の演習で難しかった点は〇〇である(例：賃金分布の歪み)」という自己記述を生成AIに入力する。この過程を通じて、学習者は自身の理解不足を言語化・客観化する(意図：メタ認知モニタリング)。

2. 概念の深掘り：

「〇〇がよく分かりません。別の具体例を使って説明してください」と促すことで、抽象的な統計概念を具体的な文脈へ落とし込み、個人の理解段階に合わせた解説を引き出す(意図：個別最適化された足場かけ／精緻化)。

3. 社会的接続：

学習者は、「この統計は社会問題のどのような理解に役立つと考えるか」という自己記述を生成AIに入力する。この過程を通じて、抽象的なデータを現実社会の文脈へ接続し、学習内容の転移 (Transfer) を促す。

4. 感情の省察：

学習者は、「〇〇のグラフに驚いた。その理由は～である」という自己記述を生成AIに入力する。この過程を通じて、認知的感情(驚き)を言語化し、その背景にある理解や前提を省察する(意図：認知的感情の活用)。これは本研究の特徴であり、認知的感情を知的探究心へと接続させることを狙う。

5. 次なる学習：

「今日理解したことを要約し、次に学ぶべきトピックを3つ提案してください」と促すことで、学習プロセスを総括し、次回の学習目標を立てる(意図：自己調整学習 [SRL])。

また、プロンプト入力が苦手な学生に対しては、「中学生にも分かるように」「身近な例(売上や気温)に置き換えて」といった具体的なヒント (Scaffolding) を提示し、プロンプトエンジニアリングのスキル格差が学習効果に影響しないよう配慮した。

(2) 振り返りの記録と共有

対話終了後、直ちにWebアンケート (Google Forms) にてAI支援指数およびメタ認知指数を測定した。加えて任意課題として「AIの回答で印象に残った点」をLMS (学習管理システム) のディスカッションボード等に投稿させ、クラス全体での知見共有 (Collaborative Learning) を図った。

(3) 倫理的配慮とハルシネーション対策(認知的警戒の喚起)

生成AIの使用に際しては、毎回「AIの回答をそのまま転用せず、必ず出典や根拠を確認すること(裏取り)」を徹底した。これは誤情報(ハルシネーション)への対策であると同時に、Sperberら(2010)のいう「認知的警戒 (Epistemic Vigilance)」を喚起し、情報を批判的に検証する態度を養うための教育的介入でもある。

2.3 調査項目

(1) AI支援指数(5項目)¹⁾：「生成AIの問かけで理解が整理できた」「不足点が明確になった」などを問う(5件法：1=全くそう思わない～5=強くそう思う)。(2) メタ認知指数(6項目)²⁾：「今日の学習目標に対する達成度を説明できる」「次回までに復習すべき点が明確になった」「良い質問ができた」などを問う(5件法)。(3) 感情プロファイル：Plutchik (1980)の感情の輪に基づく8感情(期待、驚き、喜び、信頼、怒り、嫌悪、悲しみ、恐れ)の強度(4件法：0=感じなかった～3=強く感じた)。(4) 自由記述：AI活用の良かった点、改善点、AIへの追加質問内容。

2.4 分析手法

データ分析にはPython (pandas, matplotlib, seaborn) を使用した。分析コードはGitHubで公開している (<https://github.com/aonoa68/jade2025-genai-learning>)。

3 結果

3.1 AI支援指数とメタ認知指数の推移

本節では、研究目的(1)である「AIによるフィードバックが学習者のメタ認知に与える影響」について、AI支援指数およびメタ認知指数の推移から検討する。

各セッションにおけるAI支援指数およびメタ認知指数の平均値を表2に示す。全体を通じて、AI支援指数は4.0

表2 セッション別 AI支援指数・メタ認知指数の平均値

Session	内容	N	AI支援 M (SD)	メタ認知 M (SD)
第5回	地理・移動平均	26	4.23 (0.91)	3.83 (0.89)
第7回	人口統計	36	4.06 (0.90)	3.62 (0.89)
第9回	労働統計	18	4.19 (0.91)	3.83 (0.61)
第11回	賃金統計	10	3.85 (0.94)	3.60 (0.99)

表3 学習過程における感情プロフィール
(平均強度：0-3, 4件法)

Session	期待	驚き	喜び	信頼	怒り	嫌悪	悲しみ	恐れ
5	1.92	2.04	1.76	1.58	0.69	0.77	0.52	0.76
7	1.81	2.11	2.03	1.75	0.64	0.58	0.56	0.72
9	1.78	1.83	1.89	1.33	0.28	0.11	0.22	0.39
11	1.60	1.70	2.00	1.44	0.22	0.67	0.22	0.22

前後(5点満点)の高い値を維持しており、AIが学習の補助として有効に機能したことが示唆される。第5回および第9回では高い評価が得られた一方、第11回(賃金統計)では数値の微減が見られた。第11回の微減は「AIの回答が正しいか不安」という記述と整合しており、むしろ批判的思考の発達を示す指標として解釈できる。なお、回答数が後半で減少しており、この傾向の解釈には留意が必要である。

3.2 学習感情の推移

授業中に感じた感情の強度平均を表3に示す。Kruskal-Wallis検定の結果、いずれの感情においてもセッション間で有意差は認められなかった(全て $p > .05$)。これは、全4回を通じて感情プロファイルが安定していたことを示している。具体的には、「驚き(Surprise)」と「喜び(Joy)」が全回を通じて高く(平均1.7~2.1, 4件法：0-3)、肯定的な情動が維持されていることがわかった。一方で、「恐れ(Fear)」や「悲しみ(Sadness)」などのネガティブな感情が低水準(平均0.8以下)に留まったことから、心理的安全性の高い学習環境が構築されたことが示唆される。「信頼(Trust)」は概ね中程度(1.3~1.8)で推移しており、AIを過信せず、道具として冷静に利用している様子が見ええる。

表4 縦断追跡可能な回答者

ID	参加回	参加数
Case1	5, 7	2
Case2	5, 9	2
Case3	5, 7, 9, 11	4

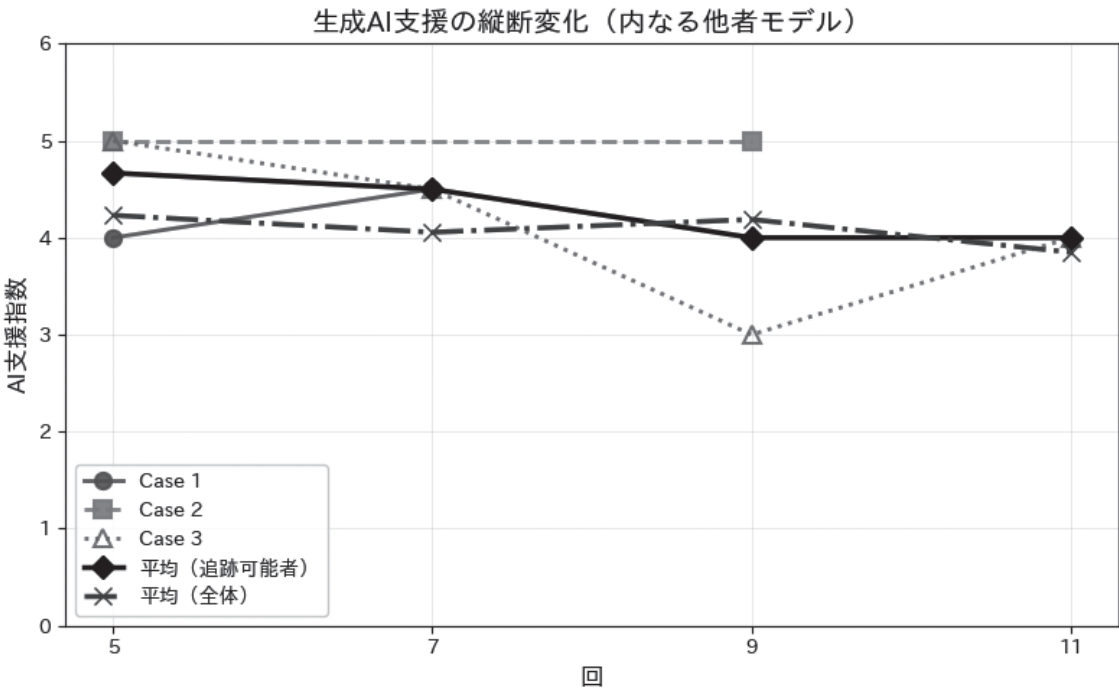


図1 AI支援指数の縦断変化

注：追跡可能者(N=3)と全体平均(各回N=26, 36, 18, 10)を示す。

表5 課題特性への言及と学習指数の関連

Session	課題特性	N	AI支援M	メタ認知M
7	コード実装・あり	5	4.50	4.20
7	コード実装・なし	31	3.98	3.53
9	プロンプト設計・あり	3	4.33	4.00
9	プロンプト設計・なし	15	4.15	3.80
11	データ選択・あり	1	5.00	5.00
11	データ選択・なし	9	3.72	3.44

注：第11回はN=1のため参考値。

3.3 探索的事例研究：縦断的变化

本実践におけるアンケートは匿名・任意であったが、IDを用いて複数回回答を紐付けられた学習者が3名存在した(表4)。図1に、縦断追跡可能な3名の個人内変動とその平均を示す。サンプル数は極めて限定的であり、本研究では回別集計を主とし、縦断分析は補助的に位置づけた。

全4セッションに回答した学生(ID: Case 3)の自由記述内容の変化に着目すると、AIへの依存の質が変わっていく様子が読み取れる。第5回(地理情報)および第7回(人口統計)の段階では、「エラーが出たコードの修正」や「グラフ描画の方法」といった手続き的知識(Procedural Knowledge)に関する問いが中心であった。しかし、第9回(労働統計)以降、「失業率と求人倍率の乖離についてAIと議論した」「なぜこの指標が重要なのかを尋ねた」といった概念的知識(Conceptual Knowledge)に関わる記述が見られるようになった。

これは、初期にはAIを単なる「作業代行ツール」として捉えていたが、対話を重ねる中で、自らの解釈の妥当性を検証する「壁打ち相手」へと認識を変化させたことを示唆している。

3.4 課題特性とAI活用の関連

自由記述の内容から、課題特性に関連する語句が含まれる回答を抽出し、指数との関連を探索的に検討した(表5)。

第7回(人口統計)では、コード実装に関する言及(「コード」「エラー」「修正」等)がある群(N=5)は、ない群(N=31)と比較してAI支援指数が高く(M=4.50 vs M=3.98)、メタ認知指数も高い傾向を示した(M=4.20 vs M=3.53)。第9回(労働統計)では、プロンプト設計に関する言及(「指示」「具体的」「形式」等)がある群(N=3)において、両指数がやや高い傾向がみられた(AI支援：M=4.33 vs M=4.15、メタ認知：M=4.00 vs M=3.80)。

4 考察

4.1 進化認知科学的視点：「内なる他者」としてのAIの特質

3.1および3.4の結果から、AI支援型リフレクションモデルを用いた学習環境において、メタ認知指数が一貫して中～高水準を示した。これらの結果は、AIとの対話が学習者のメタ認知を促進した可能性を示唆している。

では、なぜAIとの対話がこのような効果を持ったのだろうか。本研究の結果は、AIが「内なる他者」として機能したことに起因すると考えられる。Mercier & Sperber (2011)の「理性の相互作用仮説」によれば、人間の推論能力は対論的な文脈において活性化される。しかし、従来の対面指導では、教員や友人に対する「無知を晒すことへの恥じらい」や「評価への不安(社会的コスト)」が、学習者の問いを抑制してしまう傾向があった。これに対し、本実践における生成AIは、感情や社会的評価を持たない「内なる他者」として機能した。実際に、自由記述には心理的安全性の高さを裏付ける声が多数寄せられている。

「先生に質問するのは勇気がいるが、AIならどんな初歩的なことでも気兼ねなく聞けるので、理解が曖昧なまま進むことがなかった」(第7回回答)「自分のペースで、納得いくまで何度も同じことを聞いたのが良かった」(第9回回答)

感情分析(表3)においても「恐れ」が極めて低く維持されたことは、Sperberら(2010)のいう「認識的警戒」のうち、対人関係に由来する防衛本能が解除されたことを示唆する。AIは、心理的安全性を担保した状態で「議論」や「正当化」のプロセスを個人内で遂行させる「内なる他者」として機能したと言える。

4.2 課題特性とAI活用の焦点

表5に示した探索的分析の結果は、課題の性質に応じてAI活用の焦点が変化することを示唆している。第7回(人口統計)では、学習者の関心はコード実装に向けられており、自由記述にも「エラー箇所を教えてもらえて助かった」「グラフの日本語化の方法を聞いた」といった、手続き的知識の補完に関する記述が集中した。一方、第9回(労働統計)以降は、プロンプト設計(AIへの指示の出し方)そのものが学習課題となった。

「AIに指示を出すことが難しかったが、細かく指定することで精度の高いグラフを作成できた」(第9回回答)「求人倍率と失業率の関係について、納得いくまで議論した」(第9回回答)

このように、学習が進むにつれて、AIとの対話は「正解(コード)を得るプロセス」から「問いを精緻化し、概念理解を深めるプロセス」へと質的に変化していることが確認された。意識的に課題に応じた問いを立てられた学習者ほど、高い学習効果を得られた可能性がある。

4.3 AIリテラシーと「誤り」の教育的活用

第11回(賃金統計)におけるスコアの微減や、「AIの回答が正しいか不安だった」という記述は、一見ネガティブな結果に見えるが、教育的には重要な意味を持つ。

生成AI特有の「もっともらしい嘘（ハルシネーション）」に遭遇することは、学習者が「情報を鵜呑みにしない」という批判的思考（クリティカル・シンキング）を働かせる絶好の機会となるからである。実際に、第11回の自由記述では、AIに対する認識的警戒が、健全な批判精神へと転化している様子が見て取れる。

「AIがもっともらしい嘘をつくことがあったので、教科書や元のデータと照らし合わせる癖がついた」(第11回答)
「AIの説明が少しずれている気がしたので、指摘して再回答させた」(第11回答)

「内なる他者」としてのAIは、時に誤るからこそ、学習者を「教わる立場」から「評価・検証する立場」へと反転させる触媒となり得る。AIが間違えるかもしれないという前提が、学習者の認識的警戒を適切に喚起し、教科書との照合（ファクトチェック）という深い学習行動を促した。したがって、今後の授業設計においては、AIの誤りを前提とした「ファクトチェック演習」を意図的に組み込むことが、メタ認知をさらに深化させる鍵となると考えられる。

4.4 本研究の限界と今後の課題

本実践には限界がある。

第一に、縦断追跡可能なサンプルが3名に限られたため、「内なる他者」としてのAIが学習者の認知発達に与える影響を個人内変動として十分に検討することはできなかった。

第二に、回答数が後半で減少しており、脱落バイアスの影響を排除できない。特に後半回の結果解釈には留意が必要である。

第三に、本研究ではAI支援型リフレクションモデルの操作化として5段階の構造化プロンプト(Five-Stage Structured Prompts)を採用したが、自由対話や他の段階構成(段階数・順序・問いの文言)との比較検証は行っていない。したがって、観察された結果を5段階構造そのものの効果として同定することはできず、当該構造化の最適性・必要性は今後の検討課題である。今後は、構造化の有無や各段階を操作した比較(dismantling design)により、モデル要素の寄与を検証する必要がある。

第四に、課題特性分析は記述的な傾向の報告にとどまる。今後は、課題特性の定義・抽出手続き(コーディング規則や一致度など)を明確化したうえで、学習成果との関連を統計的に検証する必要がある。

第五に、本実践におけるアンケートは任意回答であり、各回の回答率は約4~15%(受講者243名中10~36名)に留まった。任意回答の性質上、課題に成功しAIに満足した学習者ほど回答しやすい傾向があると考えられる(自己選択バイアス)。したがって、本研究で得られた高いAI支援指数やポジティブな感情プロファイルは、母集団全体を代表するものではなく、「AIをうまく活用できた学習者」の特徴を反映している可能性がある。今後は、授業内での

全員回答形式や、つまづきを経験した学習者へのフォローアップインタビューなど、非回答者の声を拾う工夫が求められる。

以上を踏まえ、今後の課題として、(1)学生IDの一貫した取得方法の確立と縦断データの蓄積、(2)課題特性と学習成果の関連の統計的検証、(3)「良い問い」を立てるための足場かけ(Scaffolding)の授業設計への組み込み、(4)構造化プロンプトの要素分解による比較検証、が挙げられる。

5 おわり

本研究では、初年次統計教育において生成AIを導入し、その効果を定量的・定性的に分析した。その結果、AIとの対話は学習者のメタ認知を促し、肯定的な学習感情(特に驚きと喜び)を維持する効果があることが示された。また、事例分析からは、AIへの問いが手続き的知識から概念的知識へと深化するプロセスが観察された。さらに、課題特性への言及と学習指数の関連から、AI活用の焦点を意識化することの重要性が示唆された。

AIを「答えを出す機械」ではなく「思考を深めるための他者」として再定義することは、これからの教育における重要な視点となるだろう。今後は、サンプル数を拡大し縦断的なデータを蓄積するとともに、AIへの「問いの質」が学習成果に与える影響を詳細に分析することが課題である。

注

1) AI支援指数の構成項目(5項目):

- 生成AIの説明は、移動平均や季節調整の理解に役立った。
- 生成AIの問いかけや返答で、自分の理解が整理できた。
- 生成AIの説明で、人口統計や労働力データの扱い方を理解できた。
- 生成AIに質問することで、自分の理解の不足点が明確になった。
- 生成AIの説明により、賃金データの分布や特徴が理解しやすくなった。

2) メタ認知指数の構成項目(6項目):

- 今日の学習目標に対して、自分の達成度を説明できる。
- 次回までに何を復習・再学習すべきか明確になった。
- 生成AIに良い問いを投げる工夫ができた。
- 今日の課題内容(増加率・従属人口指数・出生動向など)を自分の言葉で説明できる。
- 今日扱った指標(失業率・無業率・求人倍率など)を自分の言葉で説明できる。
- 今日扱った指標(平均・中央値・分布形状など)を自分の言葉で説明できる。

引用資料

Fan, Y., et al. (2025). Metacognitive laziness in self-regulated learning with generative AI. *British Journal of Educational Technology*. (in press)

- Kasneci, E., Seßler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., ... & Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274.
- Mercier, H., & Sperber, D. (2011). Why do humans reason? Arguments for an argumentative theory. *Behavioral and Brain Sciences*, 34(2), 57–74.
- Plutchik, R. (1980). *Emotion: A Psychoevolutionary Synthesis*. Harper & Row.
- Sperber, D., Clément, F., Heintz, C., Mascaró, O., Mercier, H., Origgi, G., & Wilson, D. (2010). Epistemic vigilance. *Mind & Language*, 25(4), 359–393.
- Tankelevitch, L., Kewenig, V., Simkute, A., Scott, A., Sarkar, A., Sellen, A., & Rintel, S. (2024). The metacognitive demands and opportunities of generative AI. *Proceedings of the CHI Conference on Human Factors in Computing Systems (CHI '24)*, Article 607.
- Wang, X., & Fan, Y. (2025). The effect of ChatGPT on students' learning performance, learning perception, and higher-order thinking: insights from a meta-analysis. *Humanities and Social Sciences Communications*, 12, 287.

受付日 2025年12月28日, 受理日 2026年2月24日

J-STAGE早期公開日: 2026年4月1日