

# 高等教育機関における生成 AI 運用方策の 総合的分析と評価パラダイムの再構築

注：Google Gemini Deep Research により作成  
(内容の成否については確認されたい)

2026/07/17

駒水 孝裕 (名古屋大学)

## 1. 序論：生成 AI がもたらす高等教育へのパラダイムシフトと評価の妥当性の危機

生成 AI (Generative AI) の大規模な普及は、高等教育機関の根幹をなす「学習評価 (Assessment)」の前提に不可逆的なパラダイムシフトをもたらした。これまで大学教育における評価は、主に提出されたレポートや論文といった「成果物 (Product)」を通じて、学生の認知プロセスや学習到達度を間接的に推測することに依存してきた。しかし、大規模言語モデル (LLM) が人間と同等以上の流暢な学術的文章を生成し、複雑なコーディングやデータ分析を瞬時に遂行できるようになった現在、成果物の質が必ずしも学生個人の学力や思考力を反映しなくなっている。東京大学のガイドラインにおいて指摘されている通り、生成 AI は評価の「妥当性 (Validity)」を根本から揺るがしており、提出された課題は「学生個人の能力」ではなく「学生と AI の協働能力」を示すものへと変質している<sup>1</sup>。

評価本来の役割は、学習者が自身の状態を把握して次の一步を決める「学習の前進 (形成的評価)」と、教育機関が社会に対して学習の質を保証する「質の保証 (総括的評価)」に大別される。教育学的研究 (Black & Wiliam, 1998) によれば、形成的評価は数ある教育的介入の中でも最大級の効果量 ( $d = 0.4 \sim 0.7$ ) を持ち、特に学力格差の縮小に寄与するとされる<sup>1</sup>。しかし、生成 AI が学生の思考プロセスを代替してしまう環境下では、この形成的評価の基盤となる「学習の証拠」を正確に引き出すことが極めて困難になる。

本分析は、国内のトップ大学 (東京大学、京都大学、大阪大学、早稲田大学、慶應義塾大学など) および海外の重要大学 (ハーバード大学、スタンフォード大学、オックスフォード大学、ケンブリッジ大学、シンガポール国立大学など) における生成 AI の運用方策に関する膨大な一次情報およびガイドラインを統合したものである。単なるツールの利用可否を超え、「授業の設計・教育手法」「学生の評価方法の構造的再設計」「レポート課題における生成 AI の具体的な運用」に焦点を当て、次世代の教育機関に求められる持続可能かつ高度な評価パラダイムの実態を解き明かす。

分析から導き出される最大の潮流は、「AI 利用の検知と摘発 (Policing)」から「評価システムそのものの再設計 (Redesign)」への完全な移行である。AI 検知ツールの技術的限界が世界のトップ大学で共通認識となる中、教育機関は評価手法を根本的に見直し、人間の認知能力と AI の演算能力をいかに統合、あるいは切り離して測定するかという新たな学問的枠組みの構築

を迫られている。

## 2. 国内外トップ大学における運用ポリシーのガバナンスと全体傾向

世界の主要大学における生成 AI の運用ポリシーは、大学全体での一律の禁止（Blanket Ban）を避け、科目担当教員の裁量と科目ごとの特性に応じた個別最適化（Decentralization）を原則としている。しかし、そのデフォルトスタンス（明示的な指示がない場合の初期設定）や、コンプライアンスの要求水準においては、各大学の教育理念や法的背景に基づく明確な差異が観察される。

| 大学名        | デフォルトスタンス（明示的指示がない場合の扱い）        | ガバナンスの特徴と特記事項                                                                                  |
|------------|---------------------------------|------------------------------------------------------------------------------------------------|
| ハーバード大学    | 教員によるコースごとのポリシー設定を義務化           | 一律禁止を退け、教員が「最大限の制限」「全面的な推奨」「条件付きの許可」の3パターンのポリシーから選択しシラバスに明記する <sup>2</sup> 。                   |
| スタンフォード大学  | 明示的な許可がない限り「他者からの支援」と同等とみなし不正行為 | 学部・大学院はデフォルト禁止だが、経営大学院（GSB）においては「持ち帰り課題での AI 禁止」を教員に禁じるという極めて革新的な方針をとる <sup>5</sup> 。          |
| オックスフォード大学 | 事前の書面による明示的な許可がない限り不正行為         | サマティブ評価（成績評価）での無断使用は退学を含む嚴重処罰の対象。一方で、全学生・教職員に ChatGPT Edu ライセンスを無償提供しリテラシー教育を推進 <sup>7</sup> 。 |
| ケンブリッジ大学   | 評価の概要に明記されていない限り学業不正            | 形成的評価や個人学習での利用は認めるが、提出物への無断使用は剽窃とみな                                                            |

|                                |                             |                                                                                                  |
|--------------------------------|-----------------------------|--------------------------------------------------------------------------------------------------|
|                                |                             | す。AI 検知ツールの使用は証拠として不十分であるため教員に非推奨としている <sup>9</sup> 。                                            |
| シンガポール国立大学 (NUS)               | 持ち帰り課題での利用は原則許可（適切な開示が条件）   | 学生の作品の質のみで評価し、AI 使用の有無だけで減点しない。ただし、無断使用や偽の引用は剽窃として厳格に処罰する <sup>11</sup> 。                        |
| 国内トップ大学<br><br>（東大、京大、早大、慶大など） | 教員の指示を最優先。独力での課題における無断使用は厳罰 | 早稲田大学は無断利用を「無期停学・全科目不可」の対象とし <sup>13</sup> 、京都大学は精神的依存やハルシネーション（幻覚）への警戒を強く喚起している <sup>15</sup> 。 |

## 2.1. 組織内でのポリシー細分化と教育理念の反映

同一の大学内であっても、学部や学位の性質によってポリシーが細分化されている事例がスタンフォード大学に見られる。同校の Office of Community Standards (OCS) は、教員からの明確な許可がない状態での生成 AI 利用を「他者からの無許可の支援」とみなし、デフォルトで学業不正 (Academic Misconduct) として扱う方針を全学に示している<sup>6</sup>。しかし、スタンフォード大学経営大学院 (GSB) ではこれとは対照的に、教員に対して「持ち帰り課題や試験において、学生の AI 利用を禁止してはならない」という特例規則を適用している<sup>5</sup>。これは、現代のビジネス環境においてインターネットや AI ツールを駆使した課題解決が必須であり、実際の職務環境に即した評価を行うべきだという実学重視の理念に基づいている。一方、スタンフォード大学医学部では、患者のプライバシー保護と診断の正確性が人命に直結するため、パブリックな AI ツールへの保護対象保健情報 (PHI) の入力や、臨床文書作成への無断使用を厳しく禁じる独自のポリシーを展開している<sup>17</sup>。

## 2.2. アカデミック・インテグリティの再定義と懲戒プロセス

生成 AI の介在は「剽窃 (Plagiarism)」の定義を著しく複雑化させている。シンガポール国立大学のガイドラインが哲学的に指摘するように、AI は伝統的な意味での「著者」ではないため、AI から出力されたアイデアを自身のものとして発表しても、AI 自身の人格権や著作権を侵害するわけではない。しかし、自身で作成していないものを自身の成果として提出する行為は、自己の学習プロセスを偽るものであり、学問的誠実性 (Academic Integrity) に対する重大な裏切りとみなされる<sup>18</sup>。

早稲田大学や慶應義塾大学は、独力で取り組むことが求められる課題において、AI の出力をそ

のまま、あるいは僅かな改変のみで提出する行為を「代作」と認定し、無期停学等の重い懲戒処分の対象とすることを明文化している<sup>13</sup>。スタンフォード大学の懲戒プロセスにおいては、懸念（Concern）が報告された場合、予備調査を経て聴聞パネル（Hearing Panel）が招集され、学生の責任の有無と制裁が決定される厳密な法的・行政的手続きが整備されている<sup>6</sup>。

## 2.3. データプライバシーと著作権の法的制約

パブリックな生成 AI ツール（無料版の ChatGPT 等）に機密情報、学生の個人情報、あるいは未発表の研究データを入力することは、各国の情報保護法制に抵触する深刻なリスクを伴う。日本の文部科学省のガイドラインにおいても、個人情報を含むプロンプトを入力する場合、AI サービスの提供者がそのデータを機械学習に利用するか否かを十分に確認することが求められている<sup>22</sup>。また、他人の著作物を AI に入力して生成を行う行為について、学校の授業の範囲内（著作権法第 35 条）であれば無許諾での利用が例外的に認められる場合があるものの、それを広く一般公開するウェブサイトに掲載したり、外部コンテストに提出したりする場合は著作権侵害となり得ることが文部科学省によって警告されている<sup>23</sup>。

このプライバシーとセキュリティの課題に対するトップ大学の共通の解決策は、入力データが AI モデルの再学習に利用されない「エンタープライズ版（クローズド環境）」のライセンスを大学側が包括契約し、構成員に提供することである。ハーバード大学は「Harvard AI Sandbox」を整備し<sup>2</sup>、オックスフォード大学は英国の大学として初めて全教職員・学生に対して「ChatGPT Edu」および「Microsoft 365 Copilot」「Google Gemini」等のセキュアなライセンスを無償提供し、安全な実験環境を保証している<sup>8</sup>。

## 3. 授業のやり方・教育手法の変革と AI リテラシーの育成

AI 時代における教育機関の役割は、AI を単に制限することから、倫理的かつ効果的な活用法を指導する「AI リテラシー教育」の提供へとシフトしている。教員自身が AI を教育の「副操縦士（Copilot）」として活用する動きも加速している。

### 3.1. 教育手法への AI 統合に関する 5 つの原則

シンガポール国立大学は、生成 AI を授業や評価に統合する際の道標として、以下の「5 つの重要な教育学的原則（5 Key Pedagogical Principles）」を掲げている<sup>25</sup>。

1. **学習目標と真正の評価（Authentic Assessment）の優先:** AI ツールの導入ありきではなく、その課題がどのような知識やスキルを評価するためのものかを先に定義する。AI の活用が学習目標の達成を促進するか、阻害するかを文脈に応じて見極める。
2. **公平性、多様性、アクセシビリティの確保:** AI モデルに内在するバイアス（偏見）を認識し、特定の学生集団が不利益を被らないよう配慮する。また、有料ツールへのアクセスの有無が成績格差を生まないよう、大学が平等なアクセス環境を保障する。
3. **高次スキル（Higher-Order Skills）の開発への注力:** 情報の暗記や基礎的な問題解決といった低次の認知タスクは AI に委ね、人間の思考力、批判的分析、複雑な統合力といった高次スキルの育成に授業の時間を割く。
4. **学問的誠実性と人間の介在価値の維持:** AI は優れたアシスタントになり得るが、評価の目的や基準の設計、学生への全人的な個別フィードバックといった本質的な役割は、人

間の教員が担い続けなければならない。

5. **実験と継続的な評価の受容:** AI 技術は急速に進化するため、小規模なパイロットテストから始め、学生や教員からのフィードバックを基にアプローチを反復的（イテレーティブ）に修正していく姿勢を持つ。

これらの原則は、AI を排除するのではなく、AI の能力を所与のものとした上で人間の認知能力の限界を押し広げるためのカリキュラム再構築を求めている。

### 3.2. 学生と教員の AI リテラシーのギャップ

オックスフォード大学出版局（OUP）が英国の 13～18 歳の学生を対象に実施した調査「Teaching the AI-Native Generation」は、学生の AI 利用の実態と不安を浮き彫りにしている。調査によると、回答者の 80%が学校の課題で AI を使用している一方で、「AI が生成した情報の正確性を判別できる自信がある」と答えた学生は半数以下（47%）に留まった<sup>26</sup>。さらに、学生の 60%が「AI は独自の思考よりもコピーを助長する」と懸念し、51%が「バイアスやステレオタイプを強化する」と危惧している。学生の約半数は、教員に対して「AI コンテンツの信頼性の判断方法を教えてほしい」「明確なルールの提示をしてほしい」と求めており、教員側のリテラシー不足が教育現場のボトルネックとなっている現状が示唆されている<sup>26</sup>。

### 3.3. 教員の副操縦士としての AI 活用とメタプロンプティング

学生への指導を充実させるためには、まず教員自身が AI を高度に活用するスキルを身につける必要がある。オックスフォード大学や大阪大学の教員向けガイダンスでは、教員が教育活動を最適化するための AI 利用法が具体的に推奨されている。例えば、言語的・文化的背景が異なる学生や神経多様性（ニューロダイバーシティ）を持つ学生向けに、学術的な専門性を落とさずに文章の構文を平易に書き換えたり、視覚的なメタファーを提案させたりする「教材の差異化（Differentiating materials）」に AI は極めて有効である<sup>27</sup>。

また、大阪大学の全学教育推進機構では、授業内で行うグループワークの設計（例えば「情動知性」に関する 15 分間の学生間ワークショップのタイムラインや発問の構築）や、学生のアンケート結果の定性分析などに AI を活用する事例が共有されている<sup>28</sup>。

これらの高度な活用を支えるのが「プロンプトエンジニアリング」である。ケンブリッジ大学の教員向け指針では、効果的なプロンプトの構成要素として「明確性（Clarity）」「文脈の提供（Context）」「文体とトーンの指定（Desired style and tone）」「制約の明示

（Constraints and boundaries）」の徹底を求めている<sup>30</sup>。さらにオックスフォード大学では、複雑なタスクを AI に依頼する際、AI 自身に最適なプロンプトを書かせる「メタプロンプティング（Metaprompting）」の技術が推奨されている<sup>27</sup>。これにより、教員は AI を単なる検索エンジンとしてではなく、高度な対話型アシスタントとして教育設計に組み込むことが可能となる。

## 4. 学生の評価方法の構造的再設計

AI 時代において、提出された成果物の品質のみで学生の能力を測る評価パラダイムは崩壊した。これに対する世界のトップ大学の回答は、AI 検知ツールの放棄と、評価フォーマットの構造的な多層化である。

## 4.1. AI 検知ツールの放棄と技術的境界の認識

初期の段階では、多くの教育機関が Turnitin などの AI 作成テキスト検知ソフトウェアに依存し、不正を摘発しようと試みた。しかし現在、これらのツールの使用は世界のトップ大学において強く非推奨とされている。シンガポール国立大学（NUS）、南洋理工大学（NTU）、シンガポールマネジメント大学（SMU）の実情を取材した報道や、ハーバード大学、ケンブリッジ大学の公式指針が共通して指摘するように、現行の AI 検知ツールは極めて信頼性が低い<sup>4</sup>。

AI 検知アルゴリズムは、特定の単語の出現頻度や構文の予測可能性（Perplexity）に基づいて判定を行うが、学生がパラフレーズ（言い換え）ツールを使用したり、人間と AI の文章を混在させたりすることで容易に回避可能である<sup>12</sup>。さらに深刻な問題として、非ネイティブスピーカーが書いた文章が不当に AI と判定される「偽陽性（False Positive）」のリスクが高い。ケンブリッジ大学は「AI 検知ツールは学業不正の調査を裏付ける確固たる証拠には一切ならない」と断言し、これに依存した不正調査を退けている<sup>10</sup>。ハーバード大学においても、特定の文法構造や句読点の使用、シラバス外の文献の引用といった表面的な指標のみで AI 使用を断定することは不適切であると警告されている<sup>4</sup>。不正を「検知」しようとする警察的なアプローチではなく、課題そのものを AI に代替されにくい形式へ「再設計」することが唯一の解決策とされている。

## 4.2. Two-Lane Approach（二軌道アプローチ）による評価の分離

評価の再設計において標準モデルとなりつつあるのが、東京大学の生成 AI ガイドライン等で提唱されている「Two-Lane Approach」である<sup>1</sup>。これは、プログラム全体の中に「確証のための評価」と「学習のための評価」という相反する二つの評価レーンを意図的かつバランス良く配置する戦略である。

| 評価レーン                                   | 概念的定義と目的                                                       | 実施される環境的制約                               | 具体的な評価手法の実践例                                                           |
|-----------------------------------------|----------------------------------------------------------------|------------------------------------------|------------------------------------------------------------------------|
| <b>Lane 1: Secure</b><br><br>（確証のための評価） | 学生が AI の補助なしに独力で基礎的知識とスキルを満たしているかを確実に「確証」する。認知能力の最低保証ラインを設定する。 | 物理的に管理された状況、対面、ネットワーク遮断環境、あるいは教員の直接の監督下。 | 対面での筆記試験（手書きの Blue Book 等）、口頭試問（Viva）、ライブでの実演デモンストラーション <sup>1</sup> 。 |
| <b>Lane 2: Open</b><br><br>（学習のための評価）   | AI を使いこなす力、出力の真偽を疑                                             | AI を含むあらゆるリソース、インター                      | 長期的なリサーチプロジェクト、AI との協働によるレポー                                           |

|    |                                          |                      |                                          |
|----|------------------------------------------|----------------------|------------------------------------------|
| 価) | う批判的評価能力、膨大な情報を統合する力など、次世代のコンピテンシーを育成する。 | ネット、外部ツールの使用を全面的に許可。 | ト作成、複数回のマイルストーンを伴うポートフォリオ <sup>1</sup> 。 |
|----|------------------------------------------|----------------------|------------------------------------------|

ハーバード大学においても、AI で容易に完了できてしまう「ハイリスクな課題（持ち帰りのエッセイやプロブレムセット）」への依存を減らし、教室内の対面試験や口頭試問といった「ローリスク（AI 耐性の高い）課題」の比重を高めるハイブリッドな評価戦略への移行が強く推奨されている<sup>4</sup>。

### 4.3. AI アセスメントスケール（AIAS）による透明性の確保

個々の授業レベルにおいて、教員が学生に対して AI の利用可能範囲を明確に示すための指標として「AI アセスメントスケール（The AI Assessment Scale: AIAS）」の活用が広まっている。大阪大学等の資料で示されるこの指標は、AI の関与度を以下の 5 段階で定義し、学生と教員間の認識のズレを防ぐ<sup>1</sup>。

| レベル     | 呼称                                | 許容される AI の利用範囲と評価の主眼                                               |
|---------|-----------------------------------|--------------------------------------------------------------------|
| Level 1 | NO AI<br><br>(AI 使用不可)            | AI の支援を一切受けず、学生自身の知識とスキルのみで課題を遂行する。基礎力の証明に主眼が置かれる。                 |
| Level 2 | AI PLANNING<br><br>(構想段階のみ可)      | アイデア出し、ブレインストーミング、初期リサーチ、アウトライン作成等の構想段階に限定して AI を使用可。最終的な執筆は自力で行う。 |
| Level 3 | AI COLLABORATION<br><br>(AI との協働) | 草稿の作成や推敲段階で AI の支援を受けることを認める。ただし、AI の提案を批判的に評価し、最終的な編集と意思決定は学生自身が  |

|                |                                       |                                                                                 |
|----------------|---------------------------------------|---------------------------------------------------------------------------------|
|                |                                       | 行わなければならない。                                                                     |
| <b>Level 4</b> | <b>FULL AI</b><br><br>(全面活用)          | 成果物全体を通して AI を自由に駆使してよい。適切なプロンプトエンジニアリングを用いて最高品質の成果物を生み出す「使いこなしの質」そのものが評価対象となる。 |
| <b>Level 5</b> | <b>AI EXPLORATION</b><br><br>(AI の探究) | 生成 AI の限界を突破するような創造的活用や、特定分野における新しい AI の応用アプローチの開発自体を探究課題とする。                   |

教員はシラバスや課題の要項において、このレベルを明記することで、何が「学業不正」であり何が「適切な活用」であるかの境界線を確定させることができる。

#### 4.4. プロセス評価の徹底と「AI 耐性 (AI-resilient)」の強化

成果物が AI によって一瞬で生成される以上、評価の対象はその成果物に至るまでの「過程 (プロセス)」へと移行せざるを得ない。東京大学や大阪大学のガイドラインでは、Google Docs の編集履歴 (バージョン履歴) や下書きの提出を求め、段階的な推敲の痕跡を評価対象とすることが推奨されている<sup>1</sup>。単なる履歴だけでなく、「なぜその出力結果を採用したのか」「どのようにプロンプトを修正したのか」を学習ジャーナルとして言語化させるメタ認知の評価が重視されている<sup>1</sup>。

また、課題そのものを AI が自動生成しにくい形式へと変容させる工夫も求められる。大阪大学の教育学習支援部の指針では、以下の具体的な戦略が提案されている<sup>29</sup>。

- **最新・未学習データの活用:** LLM の学習データに含まれていない最新の出来事や、その日の授業内で発生した固有のディスカッション内容に基づいた記述を求める。
- **マルチモーダルな課題:** テキストのレポートではなく、概念マップ (コンセプトマップ)、動画作成、インフォグラフィックス、またはポスター発表などを要求する<sup>1</sup>。
- **一次情報と個人的経験の統合:** 抽象的な理論の解説だけでなく、学生自身の個人的な体験やフィールドワークから得られた一次情報と理論を統合して論じさせる<sup>1</sup>。

ハーバード大学では、「源泉に基づく批判 (Source-Anchored Critique)」として、コース独自の指定文献をベースにして、あえて AI に生成させたテキストの誤りやバイアスを学生に監査・修正させるという逆転のアプローチが実践されている<sup>4</sup>。これにより、AI の出力品質が成

績向上に直結するのではなく、AI の限界を指摘する批判的思考力こそが評価の対象となる。

## 4.5. 対面評価と口頭試問 (Viva) の再評価

最も古典的でありながら、AI 時代において最も強固な評価手法として再評価されているのが、対面での口頭試問 (Oral Exam / Viva) である。ハーバード大学は 2025 年秋学期より、Canvas (学習管理システム) 上で「Respondus」というブラウザのロックダウンツール (試験中のインターネットアクセスや他アプリの起動を強制遮断するツール) を導入し、教室内の座席で行うクローズド試験における不正を物理的に排除する体制を強化している<sup>2</sup>。

さらに強力な抑止力として、提出されたレポートや成果物に対して事後的な口頭チェックを行う手法が推奨されている。ハーバード大学や大阪大学の指針によれば、教員は学生に対し「この専門用語や理論を自身の言葉で説明できるか?」「なぜこの文献を引用源として選んだのか?」「この結論に至った論理的ステップを解説してほしい」といった問いを対面で投げかける<sup>4</sup>。もし学生が AI に過度に依存し、内容を内面化していない場合、この口頭試問で即座に破綻する。教員が AI 利用を疑う事案が発生した場合も、検知ツールに頼るのではなく、面談を通じてこの「理解度の欠如」を証明することが、大学の懲戒委員会 (Honor Council 等) に対する最も確実な根拠となるとハーバード大学は指導している<sup>4</sup>。

## 5. レポート課題における生成 AI の具体的な運用とアカデミック・インテグリティ

AIAS の Level 3 や 4 において生成 AI の利用が許可された場合、レポート課題の執筆における要件は、「使わないこと」から「透明性を持って開示すること」へと完全にシフトする。

### 5.1. 「どのように使ったか」の透明化と規範の醸成

スタンフォード大学経営大学院 (GSB) が提唱する画期的なアプローチは、学生に対して「AI を使ったかどうか (Whether)」を問うのではなく、すべての学生に対して「AI をどのように使ったか (How)」を定常的に報告させることである<sup>5</sup>。これは、AI 利用に関する罪悪感やスティグマを取り除き、透明性のある利用をクラスの標準 (規範) として定着させるための心理的アプローチである。

シンガポール国立大学 (NUS) 歴史学科は、この透明性を実務レベルで徹底させるため、全てのレポート提出時に「AI 利用開示表 (Disclosure Table)」の添付を義務付けている<sup>18</sup>。この表には、以下の項目の詳細な記載が求められる。

1. AI ツールの名称 (例: ChatGPT (OpenAI))
2. 入力したプロンプト (一語一句そのまま)
3. AI からの出力結果 (生成されたテキストの全文)
4. レポートへの具体的な適用方法 (出力をどのように解釈したか、論文のどの部分の参考にしたか、あるいはどのように修正して用いたか)

これにより、教員は AI の関与度合いを正確に把握でき、学生も自身の執筆プロセスを客観化することができる。

### 5.2. 引用 (Citation) のルールの確立

ケンブリッジ大学の CELTA/DELTA（英語教授法資格）向けガイダンスが示すように、生成 AI の出力は学術的慣行に従って厳密に引用（Cite）されなければならない<sup>32</sup>。APA スタイルなどの標準的なフォーマットにおいては、AI モデルの開発者（例：OpenAI）、公開年、モデルの名称（例：ChatGPT）、およびアクセス元の URL を参考文献リストに明記することが国際的な標準となりつつある。本文中（インテキストサイテーション）においても、AI の出力をそのまま、あるいは近似した形で用いた箇所は引用符で括り、どのプロンプトから得られた出力かを明示する必要がある。これを行わずに AI の生成物を自身の文章として織り込む行為は、NUS や早稲田大学の規定が明確に定義する通り「意図的な剽窃」として取り扱われる<sup>12</sup>。

### 5.3. 翻訳に関する厳格な規定

国際的な大学において特に議論となるのが「翻訳ツールとしての AI 利用」である。NUS の歴史学科のガイドラインでは、母語で書いた論文やノート、あるいは一次史料を AI に入力して英語に翻訳し、それを「そのまま（as is）」提出する行為は、学問的誠実性に反すると明確に禁じている<sup>18</sup>。AI による全自動翻訳を経た文章は、学生自身の言語的表現力や解釈の努力を反映したものではないため、大量のテキストを翻訳した場合は、学生自身が意味を深く解釈し、自身の語彙で「リフレーズ（再構築）」することが求められる。一次史料の翻訳に AI を使用した場合も、必ずその利用を明示的に引用しなければならない。

### 5.4. ハルシネーション（幻覚）と最終的な品質責任

生成 AI がもっともらしい嘘を出力する「ハルシネーション（幻覚）」に対する責任の所在は、一貫して「学生本人」にある。慶應義塾大学の基本 7 原則が「ファクトチェックの徹底」を筆頭に掲げているように、AI が出力した架空の文献、不正確な年号、あるいは論理的な矛盾をそのままレポートに記載した場合、それは単なるツールの不具合ではなく、学生自身の学問的怠慢として厳しく評価される<sup>33</sup>。

スタンフォード大学医学部のポリシーでも「生成されたコンテンツに欠陥やバイアスがあっても、提出する学生が一切の責任を負う。情報の検証と品質の判断は学生に期待されている」と明記されている<sup>17</sup>。つまり、情報源の一次確認（Verification）と批判的吟味は、AI の出力をコピー＆ペーストする能力よりもはるかに重要であり、AI 時代において最も要求される中核的なアカデミックスキルとして位置付けられているのである。

## 6. 教員による評価業務（採点・フィードバック）への AI 適用とガイドライン

AI 利用の議論は学生側にとどまらない。教員が学生のレポートを採点し、フィードバックを行うプロセスにおいて生成 AI をどこまで活用できるかについても、厳格なガイドラインが整備されつつある。ここには、業務効率化の計り知れない恩恵と、学生の知的財産権・個人情報保護を巡るジレンマが共存している。

### 6.1. 学生の知的財産（IP）の無断入力 of 禁止とプライバシー保護

ケンブリッジ大学やハーバード大学の指針において、教員に対して最も厳しく制限されている

のが、「学生が提出した課題やレポートを、学生の事前の明示的な同意なくパブリックな AI ツール（無料版 ChatGPT 等）に入力すること」である<sup>2</sup>。学生の著作物は学生自身の知的財産（Intellectual Property）であり、それを外部のサーバーに送信し、AI 企業の機械学習モデルの訓練データとして利用されうる状態に置くことは、深刻なプライバシー侵害および著作権侵害となる。ケンブリッジ大学は「いかなる審査官・教員も、提出された学生の作品を LLM にアップロードして分析やフィードバックの生成を行ってはならない」と明確に禁じている<sup>10</sup>。採点に AI を活用する場合、大阪大学の教員向けガイダンスが示すように、大学が契約しているセキュアな環境を利用するか、データ学習をオプトアウト（Chat history & training をオフ）に設定した上で、学生の氏名などの個人を特定できる情報を完全に匿名化する手続きが必須となる<sup>22</sup>。

## 6.2. 「ヒューマン・イン・ザ・ループ（Human-in-the-loop）」の原則とルーブリック評価

シンガポール国立大学の教育ポリシーは、AI を採点やフィードバックの「アシスタント」として利用することを推奨しつつも、AI の出力のみで成績評価を決定する「完全自動化」を固く禁じている<sup>11</sup>。客観式の選択問題を除き、記述式レポート等の評価においては、必ず「人間の教員による最終確認と修正（Human-in-the-loop）」を経なければならない<sup>11</sup>。

実際の運用例として、大阪大学の資料では、教員がルーブリック（評価基準となる観点と尺度）をプロンプトとして AI に事前入力し、一次的な添削案や誤字脱字のチェックを AI に行わせる手法が紹介されている<sup>28</sup>。例えば「序論・本論・結論の構成」「文字数」「客観的データの有無」といった観点を指定することで、AI はレポートの形式的要件を瞬時に監査できる。しかし、内容の独創性や論理の深さの評価においては、AI の出力はあくまで教員の「副操縦士（Copilot）」としての示唆に留まり、最終的な成績評価の責任は教員個人が負うことが強調されている<sup>11</sup>。

さらに、東京大学の講義内での実証例によれば、汎用的な「ChatGPT」と、特定の講義内容に関する専門知識を学習させた「熟達者 AI（Expert AI）」を用いて同一のレポートを採点させた場合、ChatGPT が主に文章の構成や書き方を中心に表面的な指摘を行うのに対し、熟達者 AI は専門知識に基づいた具体的な改善点や追加すべき事例を提示するなど、明確な質の差異が生じることが確認されている<sup>35</sup>。これは、教員が AI を評価補助として用いる場合、汎用ツールをそのまま使うのではなく、講義の文脈に特化したカスタムプロンプトや独自モデルを構築する能力が求められることを示唆している。

## 7. 結論：生成 AI 時代に求められる教育機関の持続可能な評価パラダイム

国内のトップ大学および海外の重要大学における生成 AI の運用方策を総合的に分析した結果、高等教育における学習評価のシステムは未曾有の構造的転換の只中にあることが明確となった。

第一に、「検知から再設計へ」の移行である。AI による不正をソフトウェアで検出するイタチ

ごっこは、技術的限界により既に終焉を迎えている。これに代わり、教育機関は評価手法そのものを、AI と協働することを前提とした「学習のための評価 (Lane 2) 」と、AI から完全に隔離された対面環境での「確証のための評価 (Lane 1) 」という二軌道へと解体し、再構築している。

第二に、「成果物からプロセスへ」の評価基軸の転換である。完成された最終成果物のみで学生の認知能力を測ることはもはや不可能であり、構想、推敲、情報の取捨選択、そして自身の意思決定に対するリフレクションという「学習プロセス」の可視化と評価が、今後の大学教育の主戦場となる。

第三に、「透明性と責任の規範化」である。一律の禁止令を敷くのではなく、学生自身に AI の利用方法を申告させる (Disclosure Table の導入等) ことで、AI を学術的ツールとして透明性をもって扱う文化を醸成している。同時に、AI の出力に対するファクトチェックと最終的な品質責任は人間 (学生・教員双方) にあることを、アカデミック・インテグリティの新たな中核として再定義している。

生成 AI は大学教育から「テキストを生成する」という物理的労力を奪ったが、その結果として、高等教育機関は「なぜ書くのか」「AI が提示した膨大な情報の中で、何を独自に思考し、どのように人間的価値を見出すのか」という、本質的な人間教育の次元へと回帰を迫られている。各大学が模索するこれらの運用方策は、単なる技術的対応ではなく、AI の高度な演算能力と人間の批判的思考力をいかに融合させ、次世代の知の枠組みを構築するかという、高等教育の存在意義を懸けた歴史的挑戦であると言える。

## Works cited

1. 学習評価 - 教育における生成 AI 活用推進リーダープログラム - 東京大学,  
[https://genai.t.u-tokyo.ac.jp/assets/materials/02\\_01\\_%E6%B4%BB%E7%94%A8%E3%81%AB%E3%81%A4%E3%81%84%E3%81%A6\\_06\\_%E5%AD%A6%E7%BF%92%E8%A9%95%E4%BE%A1.pdf](https://genai.t.u-tokyo.ac.jp/assets/materials/02_01_%E6%B4%BB%E7%94%A8%E3%81%AB%E3%81%A4%E3%81%84%E3%81%A6_06_%E5%AD%A6%E7%BF%92%E8%A9%95%E4%BE%A1.pdf)
2. Generative AI Guidance - Office of Undergraduate Education - Harvard University, <https://oue.fas.harvard.edu/faculty-resources/generative-ai-guidance/>
3. AI Code of Conduct | metaLAB (at) Harvard, <https://aicodeofconduct.mlml.io/>
4. Guidance for Faculty on Addressing AI-Related Academic Integrity Issues, <https://bokcenter.harvard.edu/guidance-faculty-addressing-ai-related-academic-integrity-issues>
5. Course Policies on Generative AI Use - Teaching and Learning Hub - Stanford University, <https://tlhub.stanford.edu/docs/course-policies-on-generative-ai-use/>
6. The Cardinal Rules of AI: Stanford University's Acceptable Artificial Intelligence Usage Policies - LLF National Law Firm, <https://www.studentdisciplinedefense.com/the-cardinal-rules-of-ai-stanford-university-s-acceptable-artificial-intelligence-usage-policies>

7. Guidance on safe and responsible use of GenAI - University of Oxford, <https://www.ox.ac.uk/students/life/it/genai-tools/guidance-on-safe-and-responsible-use-of-genai>
8. Generative AI at Oxford, <https://www.ox.ac.uk/gen-ai>
9. Using Generative AI - Blended Learning Service - University of Cambridge, <https://blendedlearning.cam.ac.uk/artificial-intelligence-and-education/using-generative-ai>
10. Generative AI and Assessment - Blended Learning Service - University of Cambridge, <https://blendedlearning.cam.ac.uk/artificial-intelligence-and-education/generative-ai-and-assessment>
11. Policy for Use of AI in Teaching and Learning - Singapore - NUS CTLT, <https://ctl.t.nus.edu.sg/wp-content/uploads/2026/04/Policy-for-Use-of-AI-in-Teaching-and-Learning.pdf>
12. Trying to catch students using AI a 'lost cause', Singapore university professors say - CNA, <https://www.channelnewsasia.com/singapore/nus-ntu-ai-cheating-plagiarism-misuse-university-5220781>
13. 生成 AI の利用に関する取り扱いと注意事項, <https://www.waseda.jp/fcom/soc/assets/uploads/2025/10/f56972858058641a7fb12ad5ed89d104.pdf>
14. レポート等作成における生成 AI ツールの利用に関する注意事項, <https://www.waseda.jp/fedu/edu/assets/uploads/2023/10/4ac144ff034de119d7a6cdc5e12e64f3.pdf>
15. 京都大学の教育・学修における生成 AI 利用に関するガイドライン - OEI, <https://oei.kyoto-u.ac.jp/use-of-generative-ai/ai-guideline/>
16. AI の基本方針・ガイドライン | 京都芸術大学, <https://www.kyoto-art.ac.jp/info/ai-policy/>
17. 3.32: Generative Artificial Intelligence (AI) Policy | MD Program - Stanford Medicine, <https://med.stanford.edu/md/mdhandbook/section-3-md-requirements-procedures/3-32--generative-artificial-intelligence--ai--policy.html>
18. Department of History, NUS Guiding Principles on Generative AI Use AY2025-2026, [https://fass.nus.edu.sg/hist/wp-content/uploads/sites/7/2025/08/HY\\_Guide-to-Use-of-AI-Tools\\_Aug2025.pdf](https://fass.nus.edu.sg/hist/wp-content/uploads/sites/7/2025/08/HY_Guide-to-Use-of-AI-Tools_Aug2025.pdf)
19. ChatGPT 等生成 AI の利用について - 慶應義塾大学 塾生サイト, <https://www.students.keio.ac.jp/com/class/registration/chatgpt.html>
20. レポート・論文は AI で書いて OK? - 日本の大学における生成 AI ガイドライン | データサイエンス百景, <https://ds100.jp/report/r-26008/>
21. BCA Guidance & Recommendations - Stanford Office of Community Standards, <https://communitystandards.stanford.edu/policies-guidance%23policies-guidance-links/bca-guidance-recommendations>
22. 初等中等教育段階における 生成 AI の利活用に関するガイドライン - 文部科学省, [https://www.mext.go.jp/content/20241226-mxt\\_shuukyo02-000030823\\_001.pdf](https://www.mext.go.jp/content/20241226-mxt_shuukyo02-000030823_001.pdf)

23. 初等中等教育段階における生成 AI の利用に関する暫定的なガイドライン,  
[https://www.mext.go.jp/content/20230718-mtx\\_syoto02-000031167\\_011.pdf](https://www.mext.go.jp/content/20230718-mtx_syoto02-000031167_011.pdf)
24. Teach with Generative AI - Harvard University,  
<https://www.harvard.edu/ai/teaching-resources/>
25. Assessment | NTU Singapore,  
<https://www.ntu.edu.sg/education/inspire/teaching-learning-assessment-with-genai/assessment>
26. Teaching the AI-native generation - Oxford University Press,  
<https://corp.oup.com/spotlights/teaching-the-ai-native-generation/>
27. Getting started with AI for educators | AI Competency Centre,  
<https://oerc.ox.ac.uk/ai-centre/ai-guides/getting-started-with-ai-for-educators>
28. 生成 AI による授業づくり支援 - 大阪大学 全学教育推進機構 教育学習支援部,  
[https://www.tlsc.osaka-u.ac.jp/project/generative\\_ai/support\\_ai.html](https://www.tlsc.osaka-u.ac.jp/project/generative_ai/support_ai.html)
29. 評価における生成 AI の影響 - 大阪大学 全学教育推進機構 教育学習支援部,  
[https://www.tlsc.osaka-u.ac.jp/project/generative\\_ai/assessment\\_ai.html](https://www.tlsc.osaka-u.ac.jp/project/generative_ai/assessment_ai.html)
30. Getting started with AI in the Classroom - Cambridge community,  
<https://www.cambridge-community.org.uk/professional-development/gswaic/>
31. Policies & Guidelines | NTU Singapore,  
<https://www.ntu.edu.sg/education/inspire/teaching-learning-assessment-with-genai/assessment/policies-guidelines>
32. Cambridge English Teaching Qualifications – advice on the use of generative AI in assessed work, <https://www.cambridgeenglish.org/Images/745146-guidance-for-use-of-generative-ai-tools-on-celta-and-delta.pdf>
33. 生成 AI 利用ガイドライン | 大阪学院大学 - OGU, <https://www.ogu.ac.jp/genai/>
34. 慶應義塾における生成 AI の利用ガイドライン,  
[https://www.sfc.itc.keio.ac.jp/ja/software\\_ai\\_guideline.html](https://www.sfc.itc.keio.ac.jp/ja/software_ai_guideline.html)
35. 【現役東大生が検証】 東大教授の生成 AI 「熟達者 AI」 でレポートを効率化する方法 - note, <https://note.com/shinchao/n/n9e3c47c9d04d>