

教育機関における生成AI運用： 評価手法・授業形態・ポリシーの国際比較と実践的方策

名古屋大学 数理・データ科学・人工知能教育研究センター

准教授 駒水 孝裕

2026.7.27

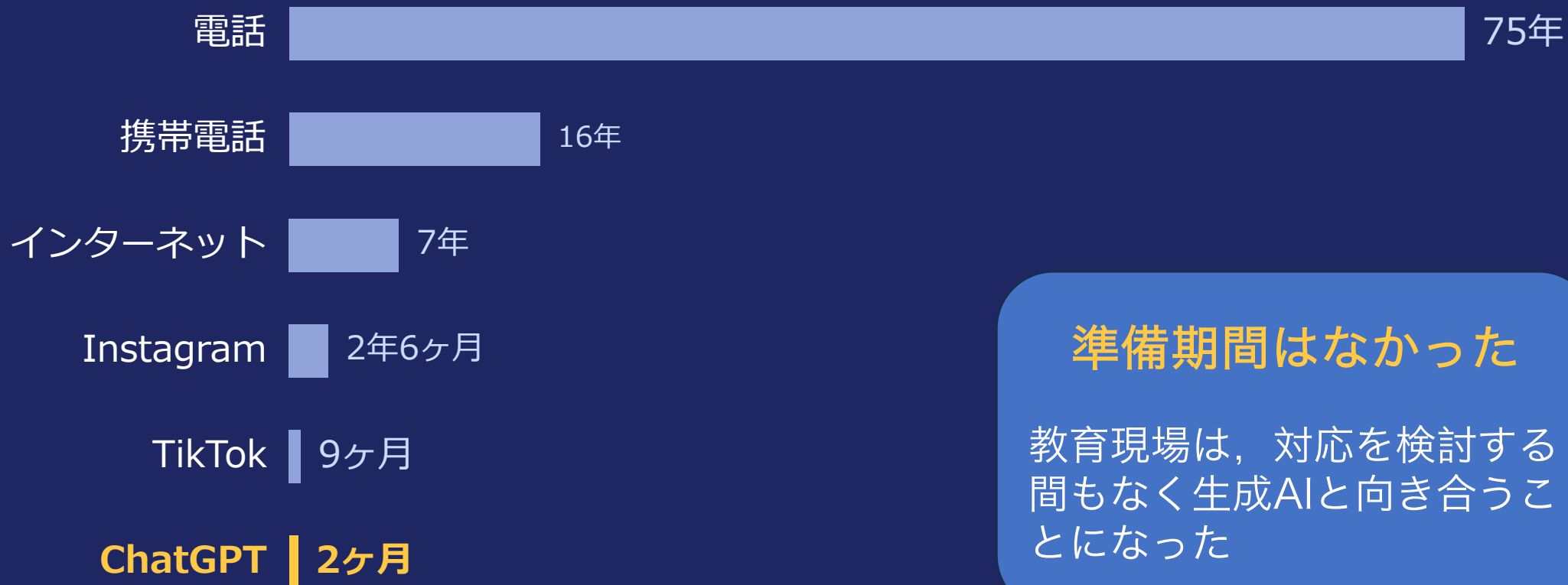


- Web資料置き場（左記 QRコード）：<https://taka-coma.pro/domestic/>
- 本スライド：https://taka-coma.pro/pdfs/ai_education_slide_20260727.pdf
- 調査まとめ資料1：https://taka-coma.pro/pdfs/ai_education_report_20260717.pdf
- 調査まとめ資料2：https://taka-coma.pro/pdfs/ai_guideline_report_20260414.pdf

（本スライドの参考文献は資料を参照されたい）

史上最速で普及したテクノロジー

ユーザー1億人到達までに要した期間



準備期間はなかった

教育現場は、対応を検討する間もなく生成AIと向き合うことになった

生成AIは「すでに社会インフラの規模」

9億人

ChatGPT 週間アクティブユーザー
(2026年2月・OpenAI発表)

9人に1人

世界人口のうち、毎週ChatGPTを
利用している計算

25億件/日

1日に処理されるプロンプト数

史上最速

月間利用者10億人に到達したアプリ
(2026年6月・Sensor Tower推計)

学生世代を中心に日常化した利用を前提に
評価・授業・ポリシーを設計する必要がある

AI時代の教育に問われる3つの課題

Q1

AIの活用可能範囲（境界線，運用ポリシー）を
どう設定するか？

Q2

AIを前提として，授業や学びのプロセスを
どう変革するか？

Q3

「成果物のみ」に頼らない，新たな評価方法を
どう再構築するか？

+

おまけ：教員が生成AIを使う時に気をつけるべきこと

本講演での内容

Q1

AIの活用可能範囲（境界線・運用ポリシー）をどう設定するか？

- 一律禁止から、科目ごとの個別最適化へ
- 米国7大学に共通する5つの基本方針

Q2

AIを前提に、授業や学びのプロセスをどう変革するか？

- AIリテラシー教育
- 教育の「副操縦士」としての活用
- NUSの5つの教育学的原則

Q3

「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

- 「検知」から「再設計」へ
- Two-Lane方式と
AIAS（AI関与度5段階）で透明性を担保

+

おまけ：教員が生成AIを使う時に気をつけるべきこと

- 学生の知的財産・プライバシーの保護
- 「ヒューマン・イン・ザ・ループ」原則

— 問いかけたいこと —

あなたの授業では、生成AIをどう位置づけ、学びと評価をどう再設計しますか？

Q1 AIの活用可能範囲（境界線, 運用ポリシー）をどう設定するか？

A1 一律禁止 → 科目ごとに個別最適化

大学名	デフォルトスタンス	ガバナンスの特徴・特記事項
ハーバード大学	コースごとのポリシー設定を教員に義務化	「最大限の制限」「全面的な推奨」「条件付きの許可」の3類型から教員が選択しシラバスに明記.
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」と同等の不正行為	原則デフォルト禁止だが、経営大学院（GSB）は持ち帰り課題でのAI禁止を教員に禁じる.
オックスフォード大学	書面による明示的な許可がない限り不正行為	成績評価での無断使用は退学を含む厳重処罰. 全構成員にChatGPT Eduを無償提供.
ケンブリッジ大学	評価の概要に明記がない限り学業不正	形成的評価・個人学習は容認. AI検知ツールは証拠不十分として教員に非推奨.
シンガポール国立大学	持ち帰り課題は原則許可（適切な開示が条件）	作品の質のみで評価しAI使用の有無で減点しない. 無断使用・偽の引用は剽窃として厳罰.
国内トップ大学 （東大・京大ほか）	教員の指示を最優先, 独力課題での無断使用は厳罰	早大は無断利用を「無期停学・全科目不可」の対象に. 京大は依存やハルシネーションへ警戒喚起.

Q1 AIの活用可能範囲（境界線，運用ポリシー）をどう設定するか？

米国7大学のAIガイドラインに共通する5つの基本方針

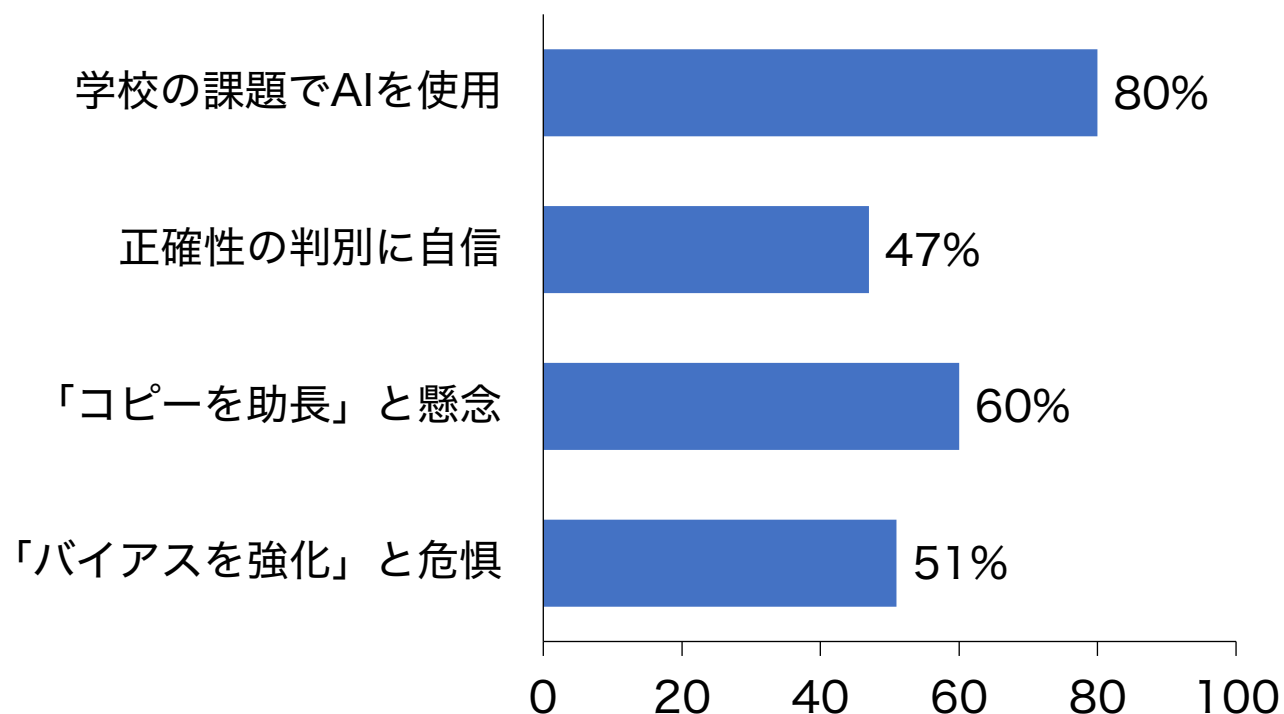
- 1 責任ある探求（禁止ではなく歓迎）**
生成AIを頭ごなしに禁止せず，生産性向上・イノベーションの機会として歓迎。
倫理と学問的誠実性を守る利用を推奨。
- 2 データ保護と「承認済みツール」**
機密情報の公開AIへの入力を厳禁。
大学がデータ保護契約を結んだ安全なツール（Microsoft Copilot等）の利用を推奨。
- 3 出力の検証は人間が最終責任**
ハルシネーションや偏見を前提に，正確性を批判的に検証する責任は常に利用者。
高度なツールより「基礎知識」の習得を優先。
- 4 透明性と教員の裁量（現場ルール）**
AI利用の事実を開示する義務。
許可範囲は一律ではなく，教員が個別に設定しシラバス等で明示する。
- 5 AIアシスタントの慎重な運用**
録音・要約ツールは事前通知と「拒否権」の保障が必須。
ただし障害のある学生への合理的配慮としての利用は明確に保護。

Q2 AIを前提として、授業や学びのプロセスをどう変革するか？

A2 AI リテラシー教育と教育の「副操縦士」として活用

オックスフォード大学出版局調査 “Teaching the AI-Native Generation” (英国の13～18歳の学生対象)

学生のAI利用の実態と不安



学生が教員に求めていること

- **48%** AIコンテンツの信頼性の判断方法を教えてほしい
- **51%** 利用の時期・方法の明確なルールがほしい
- **約1/3** (学生から見て) 教員自身がAI活用に自信がないと感じる
- **多数** 授業へのAIリソースの統合を希望

調査が示す示唆

教員側のAIリテラシー不足が
教育現場のボトルネックに

Q2 AIを前提として、授業や学びのプロセスをどう変革するか？

A2 AI リテラシー教育と教育の「副操縦士」として活用

学生への指導の充実には、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



教材の差異化

オックスフォード大
大阪大

- 言語・文化背景や神経多様性に応じて構文を平易に書き換え
- 視覚的メタファーを提案
- 学術的な専門性は落とさない



授業設計と定性分析

大阪大

- グループワークの設計支援
例：「情動知性」15分ワークショップのタイムラインと発問の構築
- 学生アンケート定性分析



プロンプト活用

ケンブリッジ大
オックスフォード大

明確性

文脈の提供

文体とトーン

制約の明示

メタプロンプティング

AI自身に最適なプロンプトを書かせる手法

Q2 AIを前提として、授業や学びのプロセスをどう変革するか？

シンガポール国立大学が掲げる「5つの重要な教育学的原則」

- 1 学習目標と真正の評価（Authentic Assessment）の優先
AI導入ありきではなく、課題が何を評価するのかを先に定義し、AI活用が学習目標の達成を促すか阻むかを文脈で見極める
- 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を認識し、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環境を保障する
- 3 高次スキル（Higher-Order Skills）の開発への注力
暗記や基礎的な問題解決など低次のタスクはAIに委ね、批判的分析・複雑な統合力など高次スキルの育成に授業時間を割く
- 4 学問的誠実性と人間の介在価値の維持
評価の目的・基準の設計や、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る
- 5 実験と継続的な評価の受容
AIの継続的な進化に対応するため、小規模なパイロットテストから始め、学生・教員のフィードバックを基にアプローチを反復的に修正していく

Q3 「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

A3 「検知」から「再設計」へ — 評価再設計の5本柱

AI検知ツールの放棄と技術的限界の認識

- 1 Turnitin等の検知は信頼性が低く偽陽性リスクも高い。
ハーバード・ケンブリッジ等は使用を強く非推奨し、課題の「再設計」を唯一の解とする。

Two-Lane Approach（二軌道アプローチ）

- 2 東大が提唱する標準モデル。「確証（Secure）」と「学習（Open）」の2レーンを配置し、独力の基礎力とAI協働力を分けて測る。

AIアセスメントスケール（AIAS）による透明性

- 3 AIの関与度をLevel 1～5で明示。シラバスで許容範囲を確定し、「学業不正」と「適切な活用」の境界を可視化する。

プロセス評価の徹底とAI耐性の強化

- 4 成果物から「過程」へ。編集履歴・下書き・学習ジャーナルでメタ認知を評価し、最新データや一次経験の統合で課題のAI耐性を高める。

対面評価・口頭試問（Viva）の再評価

- 5 最古典かつ最強の手法として再評価。提出物への事後の口頭チェックで理解の内面化を確認する。ロックダウン試験も併用。

Q3 「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

標準モデル「Two-Lane Approach」 — 評価を2つのレーンに分離

観点	Lane 1 Secure（確証のための評価）	Lane 2 Open（学習のための評価）
目的	AIの補助なしに基礎的知識・スキルを独力で満たすことを確証。認知能力の最低保証ラインを設定する。	AIを使いこなす力，出力の真偽を疑う批判的評価力，情報統合力など次世代コンピテンシーを育成する。
環境	物理的に管理された対面環境。ネットワーク遮断，教員の直接監督下で実施。	AIを含むあらゆるリソース・インターネット・外部ツールの使用を全面的に許可。
手法例	対面筆記試験（手書きBlue Book），口頭試問（Viva），ライブ実演デモ。	長期リサーチPJ，AIとの協働レポート，複数マイルストーンのポートフォリオ。

Q3 「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

AIアセスメントスケール（AIAS） — AI関与度を5段階で明示

- 1 Level 1 | NO AI (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を置く。
- 2 Level 2 | AI PLANNING (構想段階のみ可)

アイデア出し・初期リサーチ・アウトライン作成など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。
- 3 Level 3 | AI COLLABORATION (AIとの協働)

草稿作成や推敲でAIの支援を受けられる。
ただし、提案を批判的に評価し、最終編集と意思決定は学生自身が行う。
- 4 Level 4 | FULL AI (全面活用)

成果物全体を通してAIを自由に駆使可。
プロンプトを駆使した「使いこなしの質」そのものが評価対象となる。
- 5 Level 5 | AI EXPLORATION (AIの探究)

生成AIの限界を突破する創造的活用や、新しいAI応用アプローチの開発自体を探究課題とする。

Q3 「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

レポート課題での運用 — 透明性とアカデミック・インテグリティ

「どのように使ったか」の透明化・開示

- 1 スタンフォードGSBは「使ったか」でなく「どう使ったか」を報告させ、透明な利用を規範化。NUS歴史学科はAI利用開示表（ツール名／プロンプト／出力／適用方法）の添付を義務化。

引用（Citation）ルールの確立

- 2 AI出力はAPA等の書式で厳密に引用（開発者・公開年・モデル名・URL）。本文でも引用符とプロンプトを明示。無断使用は「意図的な剽窃」として扱う。

翻訳に関する厳格な規定

- 3 母語原稿や一次史料をAIで英訳し「そのまま」提出するのは学問的誠実性に反する。自身の語彙で「リフレーズ」し、翻訳利用も明示的に引用する。

ハルシネーションと最終的な品質責任

- 4 AIの幻覚・誤情報の責任は一貫して学生本人。慶應は「ファクトチェックの徹底」を筆頭に掲げ、スタンフォード医学部も欠陥・バイアスの責任は提出者と明記。一次確認と批判的吟味がAI時代の中核スキル。

+ おまけ：教員が生成AIを使う時に気をつけるべきこと

1 学生の知的財産の無断入力 of 禁止とプライバシー保護

原則 — 学生の成果物の無断入力を禁止

- 1 ケンブリッジ大学・ハーバード大学の指針は、学生が提出した課題・レポートを本人の明示的な同意なくパブリックなAIツール（無料版ChatGPT等）へ入力することを教員に最も厳しく制限している。

理由 — 知的財産権とプライバシーの侵害

- 2 学生の著作物は学生自身の知的財産であり、外部サーバーへ送信しAI企業の学習データに利用され得る状態に置くことは重大なプライバシー・著作権侵害にあたる。ケンブリッジは「学生の作品をLLMにアップロードして分析・フィードバックを生成してはならない」と明言。

採点にAIを使う場合の必須手続き

- 3 大阪大学の教員向け指針では、大学が契約するセキュアな環境を用いるか、データ学習をオプトアウト（Chat history & training をオフ）に設定し、かつ学生氏名など個人特定情報を完全に匿名化することが必須となる。

+ おまけ：教員が生成AIを使う時に気をつけるべきこと

2 「ヒューマン・イン・ザ・ループ」原則とループリック評価

原則 — 完全自動化の禁止 (Human-in-the-loop)

- 1 NUSは採点・フィードバックへのAI活用を「アシスタント」として推奨する一方、AIの出力のみで成績を決める「完全自動化」を固く禁止。客観式問題を除き、記述式評価は必ず人間の教員による最終確認・修正を経なければならない。

活用例 — ループリックによる一次チェック

- 2 大阪大学では教員が評価観点（構成／文字数／客観的データの有無等）をループリックとしてAIにプロンプト入力し、形式的要件の一次添削や誤字脱字チェックを行う。内容の独創性や論理の深さはAIの「副操縦士（Copilot）」的示唆にとどめ、最終的な成績評価の責任は教員が負う。

示唆 — 汎用ツールより文脈特化AIを

- 3 東京大学の実証では、汎用ChatGPTが文章構成中心の表面的な指摘にとどまるのに対し、講義内容を学習させた「熟達者AI」は専門知識に基づく具体的な改善点を提示。教員には講義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。

さいごに

Q1

AIの活用可能範囲（境界線・運用ポリシー）をどう設定するか？

- 一律禁止から、科目ごとの個別最適化へ
- 米国7大学に共通する5つの基本方針

Q2

AIを前提に、授業や学びのプロセスをどう変革するか？

- AIリテラシー教育
- 教育の「副操縦士」としての活用
- NUSの5つの教育学的原則

Q3

「成果物のみ」に頼らない、新たな評価方法をどう再構築するか？

- 「検知」から「再設計」へ
- Two-Lane方式と
AIAS（AI関与度5段階）で透明性を担保

+

おまけ：教員が生成AIを使う時に気をつけるべきこと

- 学生の知的財産・プライバシーの保護
- 「ヒューマン・イン・ザ・ループ」原則

— 最後に、問いかけたいこと —

あなたの授業では、生成AIをどう位置づけ、学びと評価をどう再設計しますか？