

学習者中心の倫理的AI利用評価フレームワーク (Vivian Tzu-I Chiang, 日本語訳：志賀, 喜多)

WH	カテゴリ	とても良い	良い	標準	改善を推奨	改善が必要
理由	学習目標の振り返り	学習目標との整合性を頻繁に振り返り、AIを使った作業の利点や逃した学習機会を継続的に評価する。	学習目標との整合性を振り返り、AIを使った作業の利点と逃した学習機会を評価する。	学習目標との整合性やAIを使った作業の利点について振り返りが限られている。	学習目的に反する、あるいは学習目的から逸脱したAIツールの使用を行っている。	学習目標を無視し、AIツールを無関係な目的で使用したり、意味のある考察をせずに単にタスクを完了する手段として使用したりする。
理由	創造性と効率の最適化	AIツールの機能を一貫して 振り返り、創造性を刺激し （例：AIが出発点を提供したり、人間が改良するための斬新なアイデアを生み出したりする）、 効率を最適化する （例：反復的な作業の自動化、データの要約、反復的な改良のためのフィードバックの提供）ための使用を評価する。	AIツールの機能を振り返り、創造性を刺激する可能性（例：AIが出発点を提供したり、人間が改良するための斬新なアイデアを生み出したりする）や、効率を最適化する可能性（例：タスクの自動化、データの要約、反復的な改良のためのフィードバックの提供）を考慮する。	AIツールの機能を利用する際、これらの機能が創造性の最大化（例：AIが出発点を提供したり、人間が改良するための新しいアイデアを生成したりすること）や効率の向上（例：タスクの自動化、データの要約、反復的な改善のためのフィードバックの提供）にどのように役立つかについての振り返りが最小限にとどまっている。	創造性（例えば、AIが出発点を提供したり、人間が改良するための斬新なアイデアを生み出したりする）や効率性（例えば、タスクの自動化、データの要約、反復的な改良のためのフィードバックの提供）を最大化するために、これらの機能の可能性について明確な考察を行うことなく、AIツールの機能を使用する。	AIツールの機能と、創造性（例えば、AIが出発点を提供したり、人間が改良するための斬新なアイデアを生み出したりすること）や効率性（例えば、タスクの自動化、データの要約、反復的な改良のためのフィードバックの提供）との関連性について考察していない。
手段	精度の検証	AIツールからの情報を信頼できる情報源と照合し、適切な引用を提供する。	AIツールからの情報に一貫性があるかチェックし、引用を含める。	精度を確認せず、引用も提供せずに、AIツールに完全に依存している。	AIツールからの不正確な情報を、疑問や適切な引用なしに受け入れる。	正確性の検証や引用を行わず、すべての情報を正確なものとして受け入れる。
手段	個人情報の保護	AIツールのプライバシーポリシーを定期的に確認し、個人情報や特定可能な情報の入力を受け、データ・プライバシーの重要性を理解していることを示す。	データ保護方針が明確なAIツールを使用し、個人情報の入力を避ける。データプライバシーの重要性をある程度理解している。	データプライバシーへの影響を理解せずにAIツールを使用し、個人情報を入力する可能性がある。	プライバシーのリスクや個人情報の入力を考慮せずに個人データを共有する。	個人情報保護を無視してAIツールを使用し、一貫して個人情報を入力する。
相手	エンゲージメントの維持	AIツールのインタラクティブな機能に積極的に参加し、AIだけに頼るのではなく、AIと協力してタスクをこなす。	AIツールの利用可能なインタラクティブ要素に関与し、同僚としてAIをタスクに関与させる。	インタラクティブな機能を探索したり、AIと協力したりすることなく、受動的にAIツールを使用する。	AIツールの使用中に注意が散漫になりやすく、積極的に貢献したり、プロセスに関与することなく、AIのアウトプットを受け入れる。	AIツールへの関心や関わりを示さず、AIが関与することなくすべての作業を行うことを許可する。
目的	AIリテラシーの向上	ファインチューニング、ハルシネーション（もっともらしい誤情報）の発見、事実確認技術など、AIリテラシー・スキルを一貫して向上させる。	ファインチューニング、ハルシネーション（もっともらしい誤情報）の発見、事実確認技術など、基本的なAIリテラシーに習熟していることを示す。	ある程度のAIリテラシーを身につけるが、深みや一貫性が不足している。	基本的なAIリテラシー・スキルと理解に課題がある。	AIリテラシーのスキルや理解に欠け、批判的な評価や理解をすることなく、すべてのAI出力を受動的に受け入れる。
目的	倫理的使用の促進	AIツールに見られる 倫理的な懸念や偏見を積極的に特定し 、対処する。間違いを率直に認め、 間違いから学び、改善点を積極的に模索し 、是正措置を実施する。AIの倫理的利用について積極的に指導し、コミュニティや組織内の 倫理意識向上に大きく貢献する 。倫理的なAIの実践を推進し、責任と透明性の文化を醸成することにより、一貫したインパクトのある取り組みを示す。	AIツールに見られる倫理的な懸念や偏りを指摘し、非倫理的なAIの使用に関する議論に参加する。個人的な過ちや改善すべき点を積極的に認め、対処する。AIの倫理的な使用について他者に支援や指導を提供し、身近な状況（教室やチームなど）における倫理的な意識向上に貢献する。	時折、AIツールに見られる倫理的な懸念や偏りを報告するが、非倫理的なAIの使用に関する議論には一貫して関与するとは限らない。間違いや改善すべき点を認める意欲がある。AIの倫理的な使用について、時折、他者を支援するが、倫理的な意識向上への一貫した貢献は不足している。	AIツールに見られる倫理的な懸念や偏みを報告することが稀で、非倫理的なAIの使用に関する議論に参加することが少ない。間違いや改善点を認めることに消極的または躊躇する。AIの倫理的な使用に関して、他者への支援は最小限にとどめ、倫理啓発活動への貢献はまれである。	AIツールに見られる倫理的な懸念や偏りを報告せず、非倫理的なAIの使用に関する議論を避ける。個人的な過ちや改善点を認めず、対処しない。AIの倫理的な使用について他者に支援を提供せず、倫理的な意識向上に貢献する努力を一切しない。