

「じゃない方」の AI倫理

教育におけるAI活用のリスク・懸念を普遍的な倫理で考える

麗澤大学 中園 長新

AI倫理 (AIの文脈に特化した倫理)

- AIが普及した社会の秩序を変えていくということ
- どのようにしたらよい社会, よいビジネス, よい暮らしにつながるのか, そうしたことをともに追求するプロセス

AI倫理の必要性に異論はないけど…

それって「AI特化」じゃないとダメなのかな?

AI特化「じゃない方」の“普遍的な倫理”でどこまでカバーできる?

文部科学省「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 (Ver.2.0)」における「学校現場において留意すべき代表的なリスクや懸念の例」を対象に、AI倫理と普遍的な倫理について検討

※倫理的課題はリスク・懸念と一体になっているわけではない
※ポジティブな文脈での倫理の検討も重要

(AIに人格があるかのように誤認するリスク)

生成AIは流暢な文章やコンテンツを生成することが可能であり, また人間のコミュニケーションと遜色ないスピードで反応するレベルに到達している。児童生徒が, 人間のように振る舞うAIに触れることで, AIに人格があるかのように誤認するリスクがある。

じゃない方

- 「ごっこ遊び」で動物や植物を擬人化
→ 人間以外の存在にかりそめの人格を与えて交流
「人間とは違う存在なのだ」という潜在意識
- 人間とAIにも, この関係性を援用することができることが期待

(資質・能力の育成に悪影響を与えるリスク)

学習活動の目的や育成したい資質・能力を十分に意識しないままに, 安易に生成AIを児童生徒の学習活動に導入することで, AIに依存したり, AIの答えを鵜呑みにしたりするなど, 目的に即した必要な学習過程が省略されてしまい, 資質・能力の育成に繋がらないリスクがある。

じゃない方

- 教育において教材・教具はその特性を十分に理解して活用しなければならない
- AIの責任ではなく, AIを活用する際の教材開発・授業研究の不十分さが原因

(バイアスの存在とそれによる公平性の欠如)

生成AIは既存の情報に基づいて回答を作るため, その答えを鵜呑みにする状況が続くと, 既存の情報に含まれる偏見を増幅し, 不公平及び差別的な出力が継続・拡大する可能性がある。生成AIサービスを利用する人間側にも, 流暢な出力を見ると正しいと感じてしまう流暢性バイアスや, 人間の判断や意思決定において自動化されたシステムや技術に過度に依存してしまう自動化バイアス等の様々なバイアスが存在している。

じゃない方

- 「AIは常に正しい」という先入観
- 人間同士の情報交換においては批判的思考 (critical thinking) を重視
- AIに対しても批判的思考を持つことが必要

(機密情報や個人情報に関するリスク)

生成AIサービスでは, 入力された機密情報や個人情報が, 生成AIの機械学習に利用されることがあり, 他の情報と統計的に結びついた上で, また, 正確又は不正確な内容で, 生成AIサービスから出力されるリスクがある。

じゃない方

- 技術的にはAI特有のリスクだが…
- 利用規約の同意等, 利用対象を理解した上での利用や, むやみに機密情報・個人情報を渡さないといったことは情報倫理の範疇

(著作権に関するリスク)

生成AIにおいては, 既存の著作物と類似した生成物が生成される可能性があり, そのような生成物の利用の態様によっては著作権侵害が生じるリスクがある。

AI倫理

- 生成AIは既存の著作物を元に回答を生成
- 既存の著作物の模倣によるリスクはAI特有
- ただし, その土台は普遍的な情報倫理に依拠

(外部サービスの利用に起因するリスク)

生成AIサービスはその利用形態も多様であり, 利用に当たってはサービス提供者の定める利用規約に基づくことが求められる。その際, 現在は無償のサービスであったとしても将来的に有料のサービスになる価格の変動リスク, サービス停止等の提供条件の変動リスク, 日本の法令が適用されないリスクや係争時における管轄裁判権が日本国外になるリスクがあるほか, 技術やサービスの進展が早いことから利用規約が頻繁に変更されるリスクも考えられる。

じゃない方

- 外部サービス起因のリスクはAIに限定されない
- 企業等によるサービスはすべて, 突然の有料化やサービス停止がリスクとして存在

AI倫理が必要とされる項目においても, 実際に適用する倫理は普遍的であったり, AI倫理ではなく既存の倫理だったりするんだね!
AI利活用は教育を大きく変化させ, 必要に応じて特別視しながら適切な利活用を検討していかねばならないけど, 従来の倫理観で十分に検討できる, これまでの世界の延長線上に位置づけることも可能であることが示唆されたよ!

あつ, でも, 普遍的な倫理で対応可能だからといって, AI倫理を全否定しているわけじゃないから! リスクや懸念を大枠で捉えて, 対応の方向性等を検討するのであれば, 普遍的な倫理で対応可能であるけど, 個別の具体的事象に深入りし, 特にその技術的背景を意識しようとするんだったら, そこにはAI特有の倫理が表出するんじゃないかな! 知らんけど。

