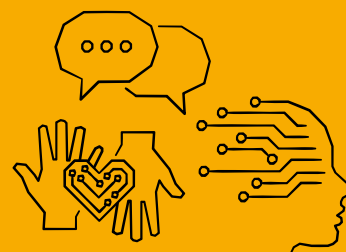




レゴ® エデュケーションのAI原則



レゴ® グループは、子どもたちの権利の尊重と保護に全力で取り組んでいます。デジタル体験のデザインにおいては、ユーザー、児童・生徒、子どもたちを開発プロセスの中心に据えた「安全第一設計」アプローチを取っています。この姿勢から生まれたのが、すべての製品イノベーションの指針となる当社のAI原則です。

1. AI機能が子どもたちの創造力・主体性・理解力を損なうことがあってはならない。

AI機能は、一人ひとりの子どもの学習者・ビルダーとしての積極的な役割をサポートするものであるべきです。AIが、子どもが考えるべきことを代行し、本来促進すべき学習そのものを阻害するのではなく、子どもの思考を支える役割を果たす必要があります。

2. AIに人間のような知性・感情・社会性があるという考えを助長しない。

子どもたちには、AIの本当の姿を知る権利と必要性があります。AIは、人間が考えたことに応答する有能な道具であり、AI自体が考えや気持ちはもちろん主体性をもつものではありません。子どもたちが自信をもってAIを使い、好奇心をもってその働きに疑問を抱くためにこの知識が必要です。

3. 当社が生成AIを使う場合には常にそれを開示する。

当社の製品の中で、AIが何らかのコンテンツ（テキスト、画像、音声等）を生成または修正したり、それらに影響を与えたりする場合には、AIの使用を教員と児童・生徒の双方に明確かつ年齢に応じた形で開示します。子どもたちはそれがAIを使用したものであると知らない限り、AIを使用したコンテンツを客観的・批判的に判断する力をつけられません。法的要件を



満たすためではなく、あくまで教育的な意図に基づく開示です。人間が作ったコンテンツと、機械が生成したコンテンツの違いを理解することこそが、AIリテラシーの基本スキルです。

4. すべてのAI処理・学習・推論は、子どもたちのデバイス上で完結する。

リスクにさらされやすい子どもたちのことを第一に考えるならば、デバイス上での処理が、データのプライバシー保護の観点から設計上の最強の形です。児童・生徒のデータがネットワーク上を行き来する必要性や、クラウドでのデータ保持リスクを排除すると同時に、ネット接続が限られやすい教室環境でも機能を確保します。

5. 児童・生徒のデータや画像をはじめとするあらゆる入力を、外部AIモデルに保存・共有せず、その学習に使用しない。

児童・生徒が行った作業は、児童・生徒自身や教員・保護者に帰属します。すべての処理がデバイス上で完結するため(項目4参照)、第三者の学習が禁じられるのみならず、不可能な構造になっています。

6. 製品は最高のアクセシビリティ基準を満たし、学びのユニバーサルデザイン:UDL原則に準ずる。

AI機能は、すべての子どもたちに開かれたものであるべきと考え、学校の教室に存在しうる、あらゆる身体的・感覚的・認知的・言語的多様性を考慮して設計されています。AIが創造力・主体性・理解力を高めることができるならば(項目1参照)、子どもたちの能力や障がい、ニューロダイバシティ(神経多様性)、地理的な場所・文化・言語等の違いにかかわらず、すべての子どもたちを包括する設計であるべきです。

7. 基礎となるあらゆるAIモデルのデータ出所を明確に文書化する。

学習データの出所、収集やラベリングの方法を説明できないAIモデルを、子ども向けの製品に使うことはありません。

8. 子どもたちに、構造化されていないもしくはフィルタリングされていないLLM (Large Language Model: 大規模言語モデル) を直接扱わせない。

生成AIシステムを「より安全」にすることは可能ですが、安全を保証できるものはほとんどありません。小学1年生～中学3年生の子どもたちが対象となることを踏まえると、脆弱性をつく行為(ジェイルブレイク)とそれに伴う不適切な出力のリスクを看過できません。この年齢層の子どもたちは、オープンエンドの生成システムではなく、特定の教育目的に即して設計・構造化された制限付きのAI出力を扱うのが適切です。

9. 当社の製品や機能は、GPAI (汎用AI) やAIコンパニオンに分類されるものであってはならない。

当社の製品は、特定の教育目的に即して設計された、限定的・タスク特化型のAIを使用しています。汎用的システムでもなければコンパニオン(相談相手)でもありません。法規制に準ずる戦略と設計理論の両面(項目1・2・7参照)に即した分類です。

10. 児童・生徒と教員はともに、あらゆるAI機能に対して実質的な管理権をもつ必要がある。

児童・生徒と教員が、すべてのAI機能に対し、有効化・無効化・設定・上書きする権限を持つ必要があります。人間が明示的に監督または行動をしない限り、AI機能が有効化されることはありません。

