



Edição Especial

VII Simpósio de Licenciaturas em Ciências Exatas e em Computação
Universidade Federal do Paraná – Pontal do Paraná (PR), 2025

DESAFIOS E PERSPECTIVAS DA REGULAÇÃO DA INTELIGÊNCIA ARTIFICIAL NA EDUCAÇÃO: UMA ANÁLISE COMPARATIVA ENTRE O BRASIL E A UNIÃO EUROPEIA

CHALLENGES AND PERSPECTIVES OF ARTIFICIAL INTELLIGENCE REGULATION IN EDUCATION: A COMPARATIVE ANALYSIS BETWEEN BRAZIL AND THE EUROPEAN UNION

Alberto Malheiros Junior¹
Claudia Leticia Filla²
Eliana Lisboa³

Resumo

O advento da Inteligência Artificial (IA) impactou as sociedades contemporâneas, tanto por seu potencial transformador quanto pelas complexas questões éticas que sua utilização suscita. Este artigo analisa o Projeto de Lei nº 2.338/2023, proposto no Brasil, que busca regulamentar o uso da IA com foco em ética, transparência e responsabilização. A análise concentra-se nas implicações dessa regulamentação para a educação, com ênfase na formação docente e nos desafios éticos, estabelecendo um diálogo comparativo com o *AI Act* da União Europeia, aprovado em 2024. O estudo baseia-se na análise documental e evidencia que, enquanto a UE prioriza uma abordagem regulatória robusta, o Brasil foca na capacitação e inclusão digital. Ambos os contextos, no entanto, carecem de detalhamento quanto à implementação prática e financiamento das políticas públicas. A regulamentação da IA na educação exige equilíbrio entre inovação tecnológica e equidade social, além de investimentos em infraestrutura e formação docente.

¹ Mestrando em Educação, UFPR.

² Mestranda em Educação, UFPR.

³ Doutora em Ciências da Educação, UMINHO Portugal.

REPPE: Revista do Programa de Pós-Graduação em Ensino

Universidade Estadual do Norte do Paraná, Cornélio Procopio (PR), v. 9, n. 2, p. 291-311, 2025

ISSN: 2526-9542



Palavras-chave: Inteligência artificial; Educação; Regulamentação; Ética; Brasil; União Europeia.

Abstract

The advent of Artificial Intelligence (AI) has impacted contemporary societies, both due to its transformative potential and the complex ethical questions its use raises. This article analyzes Brazil's Bill No. 2,338/2023, which seeks to regulate the use of AI with a focus on ethics, transparency, and accountability. The analysis concentrates on the implications of this regulation for education, emphasizing teacher training and ethical challenges, while establishing a comparative dialogue with the European Union's AI Act, approved in 2024. The study is based on documental analysis and shows that while the EU prioritizes a robust regulatory approach, Brazil focuses on digital capacity-building and inclusion. Both contexts, however, lack detail regarding practical implementation and funding for public policies. The regulation of AI in education requires a balance between technological innovation and social equity, as well as investments in infrastructure and teacher training.

Keywords: Artificial intelligence; Education; Regulation; Ethics; Brazil; European Union.

Introdução

A crescente influência da Inteligência Artificial (IA), impulsionada por seu rápido avanço tecnológico e seus multifacetados impactos sociais, econômicos e éticos, demanda urgente regulamentação. Neste cenário, autores como Barroso e Mello (2024) alertam para o potencial da IA de redefinir estruturas sociais, operando sobre bases algorítmicas opacas que podem reproduzir preconceitos e ampliar desigualdades, conforme reforça o editorial da Sociedade Brasileira para o Progresso da Ciência (SBPC) de abril de 2024, escrito por Renato Janine Ribeiro.

No campo educacional, essa necessidade de regulamentação se torna particularmente sensível, dada a natureza formativa do ambiente, em que subjetividades, conhecimentos e a própria cidadania são construídos, como indica o Guia para a IA generativa na Educação e na Pesquisa, produzido pela Unesco no final de 2023 (UNESCO, 2023). A IA tem se consolidado como fator de reconfiguração em diferentes campos da educação, conforme argumentam algumas análises recentes. Estudos como os de Alves (2023), Espírito Santo *et al.* (2023) e Unesco (2023), evidenciam que ela já atinge diversos aspectos da educação, desde algoritmos de personalização do ensino até sistemas de vigilância que monitoram o comportamento de alunos. Considerando esses aspectos, é apresentado como desafio a desconstrução de práticas que acentuam as desigualdades de acesso à tecnologia

digital (como em escolas com menor infraestrutura digital), desigualdade profissional e epistêmica (comprometimento da liberdade pedagógica), ao substituir o julgamento humano por decisões automatizadas, bem como desigualdades curriculares e culturais quando não valoriza a cultura local, impondo padrões globais como os ocidentais ou eurocêntricos. Diante desse quadro, torna-se necessária a implementação de diretrizes que orientem o uso ético, holístico e contextualizado dessa tecnologia no sistema educacional.

A regulamentação, portanto, deve garantir que o uso da IA seja orientado por princípios éticos, como o respeito à privacidade dos estudantes, à autonomia docente, à diversidade cultural e ao combate a preconceitos que possam ser reproduzidos por sistemas algorítmicos treinados com dados enviesados (UNESCO, 2023), ecoando a advertência de Paulo Freire (1996) sobre a natureza não neutra da educação, que inevitavelmente forma valores.

Essa temática consolidou-se ganhando consistência teórica e metodológica em 2022 com o lançamento do *ChatGPT*, tanto por seu potencial transformador quanto pelas complexas questões éticas que sua utilização suscita. Este artigo, que se debruça sobre o campo educacional e adota a concepção de IA apresentada na publicação *Intelligence Unleashed* de Luckin; Holmes e Forcier (2016) e na perspectiva de Santaella (2025). Para os primeiros autores, a IA consiste em sistemas computacionais criados para simular capacidades tipicamente humanas de interação com o mundo (como ver e ouvir) e de comportamento inteligente, que envolvem analisar informações e agir de forma lógica para atingir um fim determinado (Luckin; Holmes; Forcier, 2016).

Em complemento Santaella (2025), chama a atenção para o fato de que a Inteligência Artificial não deve ser vista como um monólito, mas sim como um "campo heteróclito" que se manifesta em diferentes formas. Em seguida ela distingue principalmente entre a IA preditiva, que analisa grandes volumes de dados históricos para identificar tendências e padrões, e a IA generativa, que se dedica à criação de conteúdo novo e original (como imagens, textos e outras mídias) a partir do aprendizado de padrões de dados existentes. Santaella (2025) enfatiza que a IA, em suas diversas manifestações, é uma tecnologia que amplia as capacidades humanas, mas que, inerentemente, carrega ambiguidades e riscos, tornando a ética um aspecto indissociável de seu desenvolvimento e aplicação.

Diante desse panorama, propõe-se uma análise comparativa entre duas iniciativas legislativas relevantes: o Projeto de Lei nº 2.338/2023, no Brasil, e o Regulamento (UE) 2021/0105 — conhecido como *AI Act* — aprovado em 2024 pela União Europeia e o Plano de Ação para Educação Digital, que complementa o *AI Act* e aborda a integração de tecnologias digitais na educação, ainda que não se concentre diretamente na regulamentação da IA no campo educacional.

Encaminhamentos metodológicos

A presente pesquisa adotou uma abordagem qualitativa de natureza exploratória, fundamentada na análise documental comparativa. O objetivo foi examinar dois marcos regulatórios contemporâneos cruciais para a governança da Inteligência Artificial (IA): o Projeto de Lei nº 2.338/2023, em tramitação no Congresso Nacional brasileiro, e o Regulamento (UE) 2021/0105, amplamente reconhecido como *AI Act*, aprovado pelo Parlamento Europeu em 2024.

A opção pela abordagem qualitativa se baseia na compreensão de que ela permite captar a complexidade e a historicidade dos fenômenos, sendo particularmente adequada para investigar legislações que emergem de um emaranhado de disputas discursivas e campos de interesse diversos (Minayo, 2016). A análise documental, por sua vez, permite acessar os sentidos produzidos e registrados nos textos oficiais. Compreendemos esses documentos como construções sociais que expressam valores, intenções e orientações políticas (Fávero; Centenaro, 2019).

Após a coleta dos documentos legais, procedeu-se à análise documental conforme preceitos metodológicos de Ludke e André (2022). A categorização dos dados levou em consideração os seguintes elementos presentes nos textos analisados: princípios orientadores da regulamentação, mecanismos de responsabilização, critérios de risco associados da IA e aspectos diretamente relacionados à educação.

A análise dos dados foi realizada por meio da leitura crítica dos documentos, que foi sistematicamente articulada com literatura especializada contemporânea. Para isso, utilizamos como suporte artigos de autores como Santaella (2025), Barroso e Mello (2024), Azambuja e Silva (2024), Alves (2023), Drullis (2023), Espírito Santo et

al. (2023). Essa triangulação de fontes permitiu uma compreensão mais rica das nuances e implicações das propostas regulatórias em questão.

Análise e discussão

A análise dos dados será desenvolvida em dois momentos complementares e interdependentes. No primeiro, será realizada uma análise comparativa do *corpus* documental, constituído pelo *Regulamento da União Europeia sobre Inteligência Artificial (AI Act – Regulamento UE 2021/0105)* e pelo *Projeto de Lei nº 2338/2023*, atualmente em tramitação no Congresso Nacional do Brasil. Essa etapa tem como propósito identificar convergências e divergências entre os dois marcos normativos, com foco nos critérios de classificação de risco, mecanismos de responsabilização, princípios orientadores e implicações diretas para o campo educacional.

No segundo momento, os dados resultantes dessa análise documental serão submetidos a um processo de triangulação com a literatura especializada, de modo a articular as evidências empíricas aos referenciais teóricos de autores que discutem os impactos regulatórios da inteligência artificial. A triangulação analítica buscará aprofundar a compreensão sobre os desafios e potenciais da regulação da IA, especialmente no que se refere à proteção de direitos fundamentais, à governança algorítmica e à sua inserção em contextos educativos, conferindo maior densidade interpretativa à pesquisa.

Comparação entre os marcos regulatórios

Com base na análise dos documentos oficiais selecionados, emergiram quatro categorias analíticas que, a priori, revelaram-se pertinentes para a compreensão das diretrizes éticas e regulatórias relativas à implementação da inteligência artificial no campo da educação: (i) **princípios orientadores da regulamentação**; (ii) **mecanismos de responsabilização**; (iii) **critérios de risco associados à IA**; e (iv) **aspectos diretamente relacionados à educação**. A seguir, cada uma dessas categorias é apresentada e discutida à luz do *corpus* documental analisado e da literatura especializada que fundamenta este estudo.

i) Princípios orientadores da regulamentação

As fontes evidenciam um consenso sobre a necessidade de regular a IA, impulsionado pela rápida popularização de tecnologias como o ChatGPT (Drullis, 2023). Diversos países, incluindo o Brasil, os Estados Unidos, a China e a União Europeia, estão elaborando propostas para tal regulação (Drullis, 2023; Barroso; Melo, 2024). No Brasil, o debate se concentra em projetos como o PL 2338-2023, que visa estabelecer normas gerais de caráter nacional para o desenvolvimento, implementação e uso responsável de sistemas de IA (Brasil, 2023, Projeto de lei nº 2338).

Os fundamentos centrais para a disciplina da IA no Brasil, segundo a proposta incluem a centralidade da pessoa humana, o respeito aos direitos humanos e valores democráticos, o livre desenvolvimento da personalidade, a proteção ao meio ambiente, a igualdade, a não discriminação, os direitos trabalhistas, o desenvolvimento tecnológico e a inovação, a livre iniciativa e concorrência, a defesa do consumidor, a privacidade, a proteção de dados, a autodeterminação informativa, e a promoção da pesquisa e acesso à informação e educação. O objetivo é conciliar a inovação e o desenvolvimento econômico com a proteção de direitos e liberdades fundamentais, a valorização do trabalho humano e a dignidade da pessoa humana (Brasil, 2023, Projeto de lei nº 2338).

A proposta no PL 2338 é baseada em riscos e a modelagem fundada em direitos. O desenvolvimento, implementação e uso de sistemas de IA devem observar a boa-fé e princípios orientadores como crescimento inclusivo, desenvolvimento sustentável, bem-estar, autodeterminação, liberdade de decisão, participação e supervisão humana, não discriminação, justiça, equidade, inclusão, transparência, explicabilidade, auditabilidade, confiabilidade, segurança da informação, devido processo legal, contestabilidade, rastreabilidade, prestação de contas, responsabilização, reparação integral de danos, prevenção, precaução, mitigação de riscos sistêmicos, não maleficência e proporcionalidade (Brasil, 2023, Projeto de lei nº 2338).

Na União Europeia (UE), a Lei da IA é a primeira regulamentação abrangente do mundo, focada em garantir melhores condições para o desenvolvimento e uso desta tecnologia inovadora, que pode trazer benefícios como melhores cuidados de saúde e transportes mais seguros (União Europeia, 2024). Os princípios e objetivos incluem: i) assegurar que a utilização da IA na UE seja regulada para proporcionar

melhores condições para o seu desenvolvimento e utilização; ii) garantir que os sistemas de IA utilizados na UE sejam seguros, transparentes, rastreáveis, não discriminatórios e respeitadores do ambiente; iii) assegurar que os sistemas de IA sejam supervisionados por pessoas, e não sejam totalmente automatizados, para evitar resultados prejudiciais; iv) estabelecer uma definição uniforme e tecnologicamente neutra para a IA, aplicável a futuros sistemas; v) apoiar a inovação e as *start-ups* de IA na Europa, permitindo que as empresas desenvolvam e testem modelos de IA de uso geral antes da divulgação pública. As autoridades nacionais devem oferecer testes de simulação de IA em condições próximas às do mundo real para ajudar as PMEs a competir (União Europeia, 2024)

Com base nas informações apresentadas, foi elaborado o quadro 1 que sistematiza os principais elementos com o propósito de oferecer uma compreensão clara e objetiva dos aspectos abordados.

Quadro 1: Princípios orientadores da regulamentação da IA: Brasil e União Europeia

Lei da UE sobre IA	Brasil (Projeto de Lei 2338)
Regulamentar para melhores condições de desenvolvimento e uso da IA.	Centralidade da pessoa humana e respeito aos direitos humanos.
Assegurar que sistemas de IA sejam seguros, transparentes, rastreáveis, não discriminatórios e respeitadores do ambiente.	Proteção de dados, privacidade e autodeterminação informativa.
Supervisão humana para evitar resultados prejudiciais.	Desenvolvimento tecnológico e inovação
Definição uniforme e tecnologicamente neutra para a IA.	Não discriminação, justiça, equidade e inclusão
Apoiar a inovação e startups, permitindo testes de modelos de uso geral em condições reais.	Transparência, explicabilidade, inteligibilidade e auditabilidade.
Não especificado diretamente, mas implícito no foco em segurança e direitos.	Participação e supervisão humana efetiva.
Não especificado diretamente, mas implícito no foco em segurança e direitos.	Confiabilidade, robustez e segurança da informação.
Não especificado diretamente, mas implícito na supervisão humana.	Rastreabilidade das decisões e prestação de contas.
Implícito na não discriminação.	Prevenção, precaução e mitigação de riscos sistêmicos.
A regulamentação é baseada em riscos.	Regulamentação baseada em riscos e modelagem fundada em direitos.

Fonte: Autor, 2025

Ambas as regulamentações buscam garantir que a IA seja desenvolvida e utilizada de forma segura e responsável. A Lei da IA da UE foca em segurança, transparência, rastreabilidade, não discriminação, respeito ambiental e supervisão humana, além de apoiar a inovação e o teste de modelos de IA. Já o PL 2338 brasileiro destaca a centralidade da pessoa humana, o respeito aos direitos humanos, a proteção de dados e privacidade, a não discriminação, a transparência, a participação e supervisão humana efetiva, a confiabilidade e a responsabilidade, adotando uma abordagem baseada em riscos e direitos.

ii) Mecanismos de responsabilização

As regulamentações brasileiras e da União Europeia para a Inteligência Artificial (IA) compartilham uma abordagem baseada em riscos, um elemento fundamental em seus mecanismos de responsabilização. Essa abordagem significa que quanto maior o risco de um sistema de IA, maiores são as obrigações e a intensidade dos controles e da responsabilização exigidos. No Brasil, por exemplo, o Projeto de Lei (PL) 2338/2023 classifica os riscos em "excessivo", "alto" e "não alto", tornando obrigatória uma avaliação preliminar de risco antes que qualquer sistema de IA seja introduzido no mercado ou utilizado.

Além disso, ambos os marcos legais impõem obrigações específicas tanto a fornecedores quanto a utilizadores/operadores de sistemas de IA, que também variam conforme a classificação de risco. Isso exige que os agentes de IA, independentemente de serem fornecedores ou operadores, estabeleçam estruturas de governança e processos internos robustos. O objetivo principal é garantir a segurança dos sistemas e a proteção dos direitos das pessoas afetadas pela IA. Outra semelhança importante é a exigência de avaliação durante todo o ciclo de vida dos sistemas de IA. Para aqueles classificados como de alto risco, tanto na proposta brasileira quanto na europeia, a avaliação não se limita à fase pré-mercado; ela se estende por todo o ciclo de vida do sistema, desde o desenvolvimento até a sua desativação.

A comunicação de incidentes graves também é um requisito comum em ambas as regulamentações. Em ambos os contextos, incidentes que possam comprometer a segurança ou os direitos dos usuários devem ser reportados às autoridades competentes. Na União Europeia, por exemplo, modelos de alto impacto, como o GPT-4, precisam comunicar incidentes graves à Comissão Europeia. No Brasil, a comunicação é exigida para incidentes graves de segurança, incluindo

aqueles que representam risco à vida ou que resultam em violações de direitos fundamentais. Por fim, a transparência atua como um pilar central de responsabilidade em ambas as regulamentações. Embora os documentos destaquem sua importância, o Projeto de Lei brasileiro detalha a transparência de forma mais específica para fins de responsabilização. Isso inclui medidas gerais de governança que exigem clareza no uso da IA e nas próprias medidas de supervisão implementadas.

Contudo, observam-se diferenças substantivas entre os dois documentos, notadamente em relação aos seguintes aspectos: i) responsabilidade civil; ii) sanções administrativas; iii) avaliação de impacto algorítmico (AIA); e iv) diretrizes de boas práticas e mecanismos de supervisão.

No que diz respeito à responsabilidade civil, o Projeto de Lei brasileiro é significativamente mais detalhado e explícito na responsabilidade civil. Ele estabelece que o fornecedor ou operador de sistema de IA que causar dano (patrimonial, moral, individual ou coletivo) é obrigado a repará-lo integralmente, independentemente do grau de autonomia do sistema. Para sistemas de alto risco ou risco excessivo, a responsabilidade é objetiva (independentemente de culpa), na medida da participação do agente. Para sistemas que não são de alto risco, a culpa do agente causador do dano será presumida, com inversão do ônus da prova em favor da vítima. Há também previsões para exclusão de responsabilidade em casos específicos, como culpa exclusiva da vítima ou de terceiros. Já a Lei da IA da UE impõe obrigações a fornecedores e utilizadores com base no nível de risco, mas não detalha explicitamente os tipos de responsabilidade civil (objetiva, presumida) ou o ônus da prova da mesma forma que o PL brasileiro nos excertos fornecidos. O foco está mais nas avaliações e conformidade para evitar danos, e no direito dos cidadãos de apresentar queixas.

Em relação às sanções administrativas, o PL brasileiro especifica um conjunto de sanções administrativas que a autoridade competente pode aplicar, incluindo advertência, multa (com limites financeiros claros, até R\$ 50.000.000,00 ou 2% do faturamento), publicização da infração, proibição/restrição em *sandboxes* regulatórios, suspensão (parcial, total, temporária ou definitiva) do desenvolvimento/fornecimento/operação do sistema de IA, além da proibição de tratamento de bases de dados. A aplicação dessas sanções considera diversos fatores como a gravidade, boa-fé e cooperação do infrator. Para sistemas de risco excessivo, há previsão de multa e, para pessoa jurídica, suspensão das atividades.

Embora o documento da União Europeia mencione que "novas regras impõem obrigações" e a avaliação exaustiva de modelos de alto impacto, e um prazo para a proibição de sistemas de risco inaceitável, ele não detalha as sanções administrativas com o mesmo nível de granularidade do PL brasileiro nos trechos fornecidos. O foco está mais na conformidade pré e pós-mercado e na supervisão.

No tocante à avaliação de Impacto Algorítmico (AIA) e boas práticas, o PL 2338 torna a Avaliação de Impacto Algorítmico (AIA) uma obrigação para sistemas de alto risco, devendo ser realizada por profissionais independentes e suas conclusões tornadas públicas (respeitando segredos). Além disso, o PL incentiva a formulação de códigos de boas práticas e governança por agentes de IA, que podem influenciar a aplicação de sanções. A Lei da IA da UE exige avaliações exaustivas para sistemas de alto risco e modelos de uso geral de alto impacto. Contudo, o documento não descreve um requisito explícito de AIA com a mesma metodologia detalhada e obrigatoriedade de publicidade como no PL brasileiro. As boas práticas são implícitas nos requisitos de conformidade, mas não são apresentadas como um mecanismo de responsabilização tão explicitamente quanto no PL.

No âmbito da Supervisão, o Parlamento Europeu criou um grupo de trabalho para supervisionar a implementação e aplicação do regulamento, cooperando com o Gabinete da Comissão Europeia para a IA na UE. Já o PL 2338 refere-se a uma "autoridade competente" que será responsável pela aplicação das sanções e reclassificação de riscos (BRASIL, 2023, Art.16)

O PL 2338 proíbe sistemas de risco excessivo, incluindo aqueles que utilizam técnicas subliminares ou exploram vulnerabilidades para induzir comportamento prejudicial, ou sistemas usados pelo poder público para avaliar pessoas de forma ilegítima. Para sistemas de alto risco, são exigidas medidas adicionais como documentação detalhada, registro automático da operação, testes de confiabilidade, gestão de dados para prevenir vieses discriminatórios (incluindo avaliação de dados e composição de equipe inclusiva) e medidas técnicas para viabilizar a explicabilidade dos resultados. A supervisão humana em sistemas de IA de alto risco é crucial, exigindo que os operadores compreendam as capacidades e limitações do sistema, controlem seu funcionamento e tenham a capacidade de intervir ou interromper seu uso. Embora a Lei da IA da União Europeia estabeleça uma estrutura de conformidade baseada em risco com foco em avaliação e supervisão contínuas, o Projeto de Lei 2338/2023 do Brasil se aprofunda nas particularidades da responsabilidade civil e das

sanções administrativas. O arcabouço brasileiro oferece um sistema mais detalhado para a reparação de danos e a punição de infrações, além de promover ativamente a governança e as boas práticas como parte integrante do regime de responsabilização.

Com base na análise, foram identificados pontos de convergência que revelam distintas abordagens no que diz respeito aos mecanismos de responsabilização. A fim de sistematizar as informações e oferecer uma visão comparativa, apresenta-se, a seguir o Quadro 2 que consolida os elementos centrais de cada marco regulatório.

Quadro 2: Síntese dos Mecanismos de Responsabilização da IA

Característica Principal	Lei Europeia	Projeto de Lei 2338
Abordagem Geral	Obrigações por nível de risco; avaliação contínua para alto risco.	Responsabilidades civil e sanções: foco em boas práticas.
Responsabilidade civil	Cidadãos podem queixar-se; modelos de alto impacto (ex: GPT-4) exigem avaliação exaustiva e comunicação de incidentes.	Reparação integral de danos (patrimonial/moral/coletivo) Responsabilidade objetiva para alto risco/excesso de risco; culpa presumida para outros.
Sanções administrativa	Aplicação de sanções não exclui reparação integral do dano (implícito na supervisão).	Advertência, multas altas (até R\$50 mil / 2% do faturamento), publicização, suspensão, multa e suspensão para risco excessivo.
Mecanismos específicos	Avaliação e supervisão para conformidade; Grupo de trabalho para implementação do regulamento.	AIA obrigatória e pública para alto risco; comunicação de incidentes graves (risco à vida/direitos); fomenta códigos de boas práticas e governança.

Fonte: Autor (2025)

iii) Critérios de risco associados à IA

A proposta europeia estabelece um sistema de classificação hierárquica que determina níveis de risco com base no potencial de dano que os sistemas podem representar aos usuários. A legislação define quatro categorias principais: i) risco inaceitável que abrange sistemas de IA proibidos por infringirem diretamente os direitos fundamentais. Incluem, por exemplo, técnicas de manipulação cognitivo-comportamental dirigidas a grupos vulneráveis (como crianças), pontuação social baseada em comportamento ou características pessoais, e reconhecimento biométrico remoto em tempo real em espaços públicos. Algumas exceções são previstas para o uso em contextos de segurança pública, desde que aprovadas

judicialmente. As proibições entraram em vigor em 2 de fevereiro de 2025; ii) risco elevado que compreende sistemas que podem afetar significativamente a segurança ou os direitos das pessoas. São divididos em duas subcategorias: a) sistemas integrados a produtos regulados por legislações europeias de segurança (como automóveis, brinquedos e dispositivos médicos) e b) sistemas aplicados em áreas sensíveis, como educação, emprego, segurança pública, migração e administração da justiça. Tais sistemas deverão ser registrados em uma base de dados da UE e passar por avaliações de conformidade contínuas ao longo de seu ciclo de vida. No campo da educação, sistemas aplicados à formação profissional e à avaliação de estudantes são enquadrados nesta categoria; iii) risco limitado: Inclui sistemas que requerem transparência no uso, como os que envolvem interação com usuários (por exemplo, *chatbots*), exigindo que estes sejam informados sobre a natureza automatizada da comunicação; iv) risco mínimo que refere-se a aplicações que não representam riscos relevantes à segurança ou aos direitos dos usuários, estando isentas de obrigações específicas, embora ainda possam requerer boas práticas de desenvolvimento.

No que tange à IA generativa e aos modelos, como o ChatGPT e o GPT-4, a legislação da União Europeia impõe requisitos de transparência mesmo que não sejam classificados como de risco elevado. Esses requisitos incluem: a obrigatoriedade de informar que o conteúdo foi gerado por IA, a concepção de sistemas que evitem gerar conteúdos ilegais, a publicação de resumos dos dados utilizados (especialmente quando protegidos por direitos autorais), e o monitoramento de incidentes relevantes. Conteúdos manipulados ou gerados por IA (como *deepfakes*) também devem ser rotulados de forma clara. Tais exigências entrarão em vigor 12 meses após a implementação da legislação.

A proposta brasileira adota estrutura normativa semelhante, baseada na avaliação de risco, mas introduz um processo preliminar de classificação obrigatória, a ser realizado pelo fornecedor antes da comercialização ou disponibilização do sistema. A autoridade reguladora pode, ainda, solicitar reclassificação ou a realização de uma AIA, especialmente nos casos de alto risco. A proposta define as seguintes categorias: i) risco excessivo: Inclui sistemas de IA cuja utilização é expressamente proibida. São exemplos: tecnologias que empreguem técnicas subliminares para manipulação prejudicial, exploração de vulnerabilidades de grupos específicos (por idade ou deficiência), e sistemas que promovam pontuação social para acesso a bens,

serviços ou políticas públicas. A identificação biométrica remota em espaços públicos, de forma contínua, somente será admitida mediante lei federal específica e autorização judicial e; ii) alto risco: Abrange sistemas com potencial de afetar direitos fundamentais, segurança ou dignidade humana, como os utilizados em: (i) educação e formação profissional (acesso, avaliação, monitoramento), (ii) recrutamento e gestão de trabalhadores, (iii) concessão de serviços públicos e privados, (iv) decisões na área da saúde e justiça, (v) segurança pública e controle migratório. Tais sistemas estarão sujeitos a rigorosos requisitos técnicos e regulatórios, além da exigência de AIA.

A AIA prevista no PL 2338/2023, deve ser conduzida por equipe com competências técnicas, jurídicas e científicas, resguardando a independência funcional. A metodologia deve abranger fases de preparação, identificação de riscos, mitigação e monitoramento. Deve também considerar benefícios, riscos previsíveis, lógica de funcionamento, medidas de precaução, e riscos residuais. Quando envolver grupos vulneráveis, é imperativo que o sistema seja compreensível por esses usuários e que contemple ações de conscientização e formação crítica.

No contexto educacional, destaca-se que tanto a proposta brasileira quanto a legislação europeia classificam os sistemas de IA aplicados à educação e formação profissional como de alto risco, o que demanda registro em base de dados, monitoramento contínuo, e avaliações de impacto ético e técnico, com vistas à proteção dos direitos dos estudantes e à promoção da equidade educacional.

Tendo em vista os critérios de risco na estrutura normativa, a análise possibilitou identificar diferentes formas de classificação, avaliação e tratamento nos documentos analisados. Para sistematizar essas distinções e facilitar a compreensão das abordagens de cada país, segue, portanto, o Quadro 3 que sintetiza os principais achados.

Quadro 3: Comparativo dos critérios de risco

Critério	Lei da UE sobre IA	PL 2338/2023 (Brasil)
Abordagem	Classificação em 4 níveis de risco.	Classificação baseada em avaliação preliminar.
Risco Inaceitável/ Excessivo	Proíbe pontuação social, manipulação e biometria em tempo real.	Proíbe manipulação subliminar, pontuação social e biometria sem autorização.
Risco Alto	Sistemas que afetam direitos (ex: educação, saúde, segurança).	Sistemas em áreas sensíveis (educação, saúde, justiça etc.).
Educação	Classificada como alto risco.	Classificada como alto risco.
Obrigatoriedade de Avaliação	Avaliação pré-mercado e durante o ciclo de vida.	Avaliação de Impacto Algorítmico (AIA) obrigatória.
Transparência	Exigida para IA generativa e sistemas interativos.	Exigida em decisões automatizadas que afetam direitos.
Proteção vulneráveis	Mecanismos específicos para crianças e minorias.	Exige linguagem acessível e direito à compreensão do sistema.

Fonte: Autor (2025)

Os dois marcos regulatórios estabelecem mecanismos claros para garantir a responsabilização no uso da IA. A Lei da IA da UE foca em obrigações diferenciadas por nível de risco, avaliação contínua de sistemas de alto risco, e o direito dos cidadãos de apresentar queixas. O PL 2338 do Brasil detalha a responsabilidade civil, diferenciando entre responsabilidade objetiva (para alto risco e risco excessivo) e culpa presumida (para outros riscos), e prevê um conjunto abrangente de sanções administrativas, além de fomentar a governança e as boas práticas.

iv) Aspectos diretamente relacionados à educação

O Plano Europeu estabelece um marco regulatório pioneiro para a IA em escala global, com o propósito de fomentar um ambiente propício ao desenvolvimento e à utilização segura e benéfica desta tecnologia. Um dos pilares desta regulamentação é a classificação dos sistemas de IA com base no risco inerente às suas aplicações. Consoante a este sistema, a área de "educação e formação profissional" é explicitamente designada como um domínio onde os sistemas de IA são categorizados como de "alto risco". Tal classificação implica que os sistemas de IA operantes neste setor deverão ser submetidos a registo numa base de dados da União Europeia e a avaliações rigorosas, tanto antes de serem disponibilizados no mercado quanto ao longo de todo o seu ciclo de vida. As obrigações relativas a estes

sistemas de alto risco terão aplicação plena 36 meses após a entrada em vigor da legislação.

Já a proposta legislativa brasileira visa instituir uma estrutura regulatória abrangente para a IA, balizada por princípios que procuram harmonizar a proteção de direitos com o incentivo à inovação tecnológica, conferindo previsibilidade e segurança jurídica ao desenvolvimento tecnológico. Um dos fundamentos basilares para o desenvolvimento, implementação e utilização de sistemas de inteligência artificial preceitua o acesso à informação e à educação, bem como a conscientização acerca dos sistemas de IA e suas aplicações.

No que tange à classificação de risco, o Projeto de Lei 2338/2023 alinha-se ao modelo europeu ao considerar de "alto risco" os sistemas de IA empregados na educação e formação profissional. Isso abrange, por exemplo, sistemas que determinam o acesso a instituições de ensino ou que procedem à avaliação e monitoramento de estudantes. A atribuição desta categoria de alto risco implica que tais sistemas serão submetidos a medidas rigorosas de governança e a Avaliações de Impacto Algorítmico (AIA), visando assegurar a salvaguarda dos direitos dos indivíduos por eles impactados.

Outro ponto relevante reside na determinação de que sistemas de IA direcionados a grupos vulneráveis, como crianças e adolescentes, devem ser concebidos de modo a facilitar a compreensão do seu funcionamento e dos direitos dos indivíduos perante os agentes de IA. A própria justificação do projeto de lei, ao buscar conciliar a proteção de direitos com a inovação, implicitamente aponta para a necessidade de educação e formação contínua para lidar com as novas tecnologias emergentes.

Enquanto a legislação europeia evidencia preocupação com a prática de direitos, a transparência algorítmica e os riscos associados à utilização da IA em contexto educativo, o projeto brasileiro apresenta lacunas sobre a temática. Para sistematizar os resultados, foi elaborado o quadro resumo 4 que destaca os aspectos relacionados à educação.

Quadro 4: Comparativo: IA na Educação - Lei da UE x. Projeto de Lei Brasileiro

Aspecto Chave	Lei da UE sobre a IA	Projeto de Lei Brasileiro sobre o Uso da IA
Classificação de risco	"Educação e formação profissional" é "alto risco".	Sistemas de IA na educação (acesso, avaliação, monitoramento) são "alto risco".
Implicações do risco	Sistemas de IA devem ser registrados na UE e passar por avaliações contínuas (antes de serem lançados no mercado e durante toda a sua utilização).	Sujeitos a rigorosas governança e Avaliação de Impacto Algorítmico (AIA) para proteger direitos.
Fundamentos	Melhora o desenvolvimento e uso da tecnologia, mas não foca na educação como princípio.	Acesso à informação e educação e conscientização são fundamentos para o desenvolvimento e uso da IA
Treinamento	Não detalha requisitos específicos, mas menciona supervisão humana e avaliações.	AIA deve registrar treinamento e ações de conscientização dos riscos.
Grupos vulneráveis	Proíbe manipulação cognitivo-comportamental (ex.: brinquedos ativados por voz em crianças).	Sistemas para crianças/adolescentes devem ser desenvolvidos para que estes entendam o funcionamento e seus direitos.
Outros pontos relevantes	Obrigações para sistemas de alto risco aplicáveis 36 meses após entrada em vigor.	Busca conciliar proteção de direitos com inovação tecnológica, visando previsibilidade e segurança jurídica.

Fonte: Autor (2025)

Triangulação com a literatura especializada

A triangulação de dados, ao articular a análise documental de marcos normativos e referenciais teóricos especializados, configura-se como um recurso metodológico para a construção de uma compreensão crítica acerca das dinâmicas regulatórias da IA, sobretudo no campo educacional. A análise comparativa entre o Projeto de Lei nº 2338/2023, em trâmite no Brasil, e o Regulamento Europeu de Inteligência Artificial (*AI Act*) evidencia convergências e divergências que refletem as especificidades sociopolíticas e culturais de cada contexto normativo. Ressalta-se, nesse cenário, que enquanto a proposta legislativa brasileira enfatiza mais a capacitação e a inclusão digital, a União Europeia prioriza a mitigação de riscos e a imposição de restrições no âmbito do *AI Act*. Essa análise evidencia não apenas estilos regulatórios diversos, mas também diferentes prioridades sociais e econômicas, enriquecendo o debate sobre os rumos possíveis da regulamentação da IA na educação.

O Projeto de Lei brasileiro sobre IA e a Lei da UE estão alinhados com as propostas de especialistas como Barroso e Mello (2024) e Santaella (2025), que defendem uma regulação baseada em valores éticos e interesse social. Os princípios comuns incluem: equidade (evitar vieses), confiabilidade e segurança, responsabilidade de desenvolvedores, transparência nas decisões automatizadas, análise de impacto social, supervisão humana e governança colaborativa (com participação do Estado, setor privado e sociedade). Essa convergência mostra um consenso: a IA deve ser regulada de forma proativa e centrada no ser humano, transformando princípios em práticas que equilibrem inovação e proteção de direitos. O desafio atual é criar mecanismos eficazes para implementar essas diretrizes.

Isso porque a literatura que trata da temática enfatiza que os algoritmos de IA podem perpetuar e amplificar vieses humanos existentes nos dados de treinamento, levando a decisões injustas e preconceituosas em áreas como contratação, concessão de crédito, saúde e justiça criminal. Há exemplos documentados de sistemas de IA que descartaram mulheres em processos de contratação ou criminalizaram homens negros, perpetuando desigualdades estruturais (Barroso; Mello, 2024). Além disso Santaella (2025) destaca que a Inteligência Artificial Generativa, ao produzir conteúdo original (imagens, textos, áudios) que se assemelham aos criados por humanos, o risco de disseminação de informações falsas ou prejudiciais, como os *deepfakes*, é substancialmente ampliado, demandando, por conseguinte, uma governança algorítmica ainda mais rigorosa.

Há um debate sobre a suficiência das leis existentes, como a Lei Geral de proteção de Dados (LGPD) e as leis de direitos autorais, para lidar com os desafios da IAG, havendo a percepção de que novas questões exigem leis específicas ou adaptações (Drullis, 2023). A velocidade das transformações tecnológicas supera o ritmo legislativo, e há preocupações com a regulação excessiva que possa inibir a inovação (Barroso; Melo, 2024). As propostas buscam defender os direitos fundamentais, proteger a democracia e promover a boa governança (Barroso; Mello, 2024).

Nessa óptica, os agentes de IA (fornecedores e operadores) devem estabelecer estruturas de governança e processos internos para garantir a segurança e a proteção dos direitos das pessoas afetadas. Medidas gerais incluem transparência no uso da IA e nas medidas de governança, gestão de dados para mitigar vieses, legitimidade do tratamento de dados (conforme LGPD), adoção de parâmetros para

dados de treinamento/teste/validação, e segurança da informação desde a concepção (BRASIL, 2023, Projeto de lei nº 2338).

Assim, para sistemas de alto risco, medidas adicionais são exigidas. Isso abrange documentação detalhada do funcionamento e decisões ao longo do ciclo de vida, registro automático da operação para avaliação de acurácia, robustez e vieses, testes de confiabilidade, medidas de gestão de dados para prevenir vieses discriminatórios (incluindo avaliação de dados e composição de equipe inclusiva), e medidas técnicas para viabilizar a explicabilidade dos resultados (Brasil, 2023, Projeto de lei nº 2338).

Estes direitos buscam garantir que, apesar da crescente automação e complexidade dos sistemas, a supervisão humana seja mantida e que a pessoa afetada não perca a capacidade de compreender e influenciar decisões que a impactam. Conforme Barroso e Mello (2023), a intervenção humana é indispensável em áreas que exigem inteligência emocional, valores éticos ou compreensão de nuances comportamentais, mesmo que a IA possa processar mais informações.

Para garantir a aplicação da lei e a proteção dos direitos, as propostas preveem a designação de uma autoridade competente pelo Poder Executivo, responsável por zelar pela implementação e fiscalização (Brasil, 2023, Projeto de lei nº 2338).

Essa autoridade terá diversas atribuições, incluindo proteger direitos, promover a estratégia brasileira de IA elaborar normas e estudos, articular com reguladores setoriais, fiscalizar, aplicar sanções administrativas (advertência, multa, publicização, proibição de *sandbox*, suspensão de atividades, proibição de tratamento de dados) e apreciar petições. As sanções serão aplicadas gradativamente com base na gravidade da infração, condição econômica do infrator, reincidência, grau do dano, cooperação, e adoção de boas práticas e mecanismos internos. A aplicação de sanções administrativas não exclui a obrigação de reparação integral do dano (BRASIL, 2023, Projeto de lei nº 2338).

Considerações finais

Em suma, as fontes retratam a IA como uma tecnologia de vasto potencial transformador, mas que exige um arcabouço regulatório robusto baseado em princípios éticos e direitos fundamentais, uma classificação de riscos para mitigar

efeitos adversos, e um sistema de supervisão e responsabilização claro. No âmbito educacional, a IA/IAG oferece oportunidades significativas para personalização e otimização, mas impõe desafios substanciais que demandam uma reestruturação pedagógica, o desenvolvimento de novas habilidades e uma reflexão crítica sobre o papel humano na era digital.

A comparação entre os contextos brasileiro e europeu revela a necessidade premente de uma regulamentação que transcenda diretrizes genéricas, incorporando elementos específicos para o setor educacional. No Brasil, a desigualdade digital — evidenciada pela pesquisa TIC Educação (CETIC, 2023) — contrasta com o alerta da UNESCO (2023) sobre o agravamento da "pobreza digital" no Sul Global, onde a concentração de dados e poder computacional nas mãos de grandes empresas do Norte Global perpetua assimetrias históricas. A dependência de modelos como o ChatGPT, treinados majoritariamente com dados ocidentais, pode marginalizar ainda mais realidades locais, reforçando visões de mundo hegemônicas e excluindo perspectivas regionais.

Além da infraestrutura precária, outros fatores agravam as disparidades no Brasil. Dados da pesquisa TIC Educação (CETIC, 2023) ilustram como a adoção não crítica de ferramentas digitais, sem políticas públicas robustas e formação docente adequada, pode ampliar desigualdades em vez de reduzi-las.

Enquanto o Plano de Ação para a Educação Digital da União Europeia pressupõe infraestrutura consolidada e formação docente avançada, o Brasil precisa de uma abordagem regulatória que considere suas particularidades. Incluindo infraestrutura como base, com o objetivo de garantir acesso universal a dispositivos e conectividade, condição mínima para equidade digital. Quanto à regulamentação ética e inclusiva, deve-se buscar a transparência algorítmica e proteção de dados, evitando vieses em sistemas de IA Generativa. Além disto a formação docente crítica, por meio de capacitações destes docentes, não apenas no uso técnico, mas na avaliação pedagógica e ética dessas ferramentas.

A regulação da IA na educação brasileira não pode ficar restrita a regulações abstratas, sem a previsão de dispositivos que viabilizam a efetividade das normativas. Pois, só assim será capaz de enfrentar desigualdades estruturais, assegurando que a tecnologia sirva à inclusão, à diversidade e ao desenvolvimento de capacidades humanas — conforme defendido pela UNESCO (2023). Caso contrário, corre-se o

risco de aprofundar exclusões, transformando a inovação em mais um vetor de marginalização.

Referências

ALVES, Lynn. Notas iniciais sobre Inteligência Artificial e Educação. In: ALVES, Lynn (Org.). **Inteligência Artificial e educação: refletindo sobre os desafios contemporâneos**. Salvador: Edufba, 2023.

AZAMBUJA, Celso Candido de; SILVA, Gabriel Ferreira. Novos desafios para a educação na era da Inteligência Artificial. **Filos. Unisinos** 25 (1) • 2024. Disponível em: <https://doi.org/10.4013/fsu.2024.251.07>. Acesso em: 15 jun. de 2025.

BARROSO, Luís Roberto; MELLO, Patrícia Perrone Campos. Inteligência Artificial: promessas, riscos e regulação. Algo de novo debaixo do sol. **Rev. Direito Práx.** 15 (04), 2024. Disponível em: <https://doi.org/10.1590/2179-8966/2024/84479>. Acesso em: 15 jun. 2025.

BRASIL. **Projeto de Lei nº 2338/2023**. Altera a Lei nº 12.865, de 9 de outubro de 2013, para dispor sobre a regulamentação dos serviços de ativos virtuais e das prestadoras de serviços de ativos virtuais. Brasília, DF: Câmara dos Deputados, 2023. Disponível em: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>. Acesso em: 15 de jun. de 2025.

CETIC. Pesquisa sobre o uso das tecnologias de informação e comunicação nas escolas brasileiras: **TIC Educação 2022**. São Paulo: Comitê Gestor da Internet no Brasil, 2023. Disponível em: https://www.cetic.br/media/docs/publicacoes/2/20231122125825/tic_educacao_2022_resumo_executivo.pdf. Acesso em: 12 abril de 2025.

COMISSÃO EUROPEIA. Digital Education Action Plan (2021-2027). Luxemburgo: União Europeia, 2021. Disponível em: <https://education.ec.europa.eu/focus-topics/digital-education/action-plan>. Acesso em: 15 jun. de 2025.

DRULLIS, Gustavo. **Inteligência Artificial Generativa: o que muda para a regulação?** Mobile Time. Disponível em: <https://www.mobiletime.com.br/noticias/19/05/2023/inteligencia-artificial-generativa-o-que-muda-para-a-regulacao/>. Acesso em: 15 jun. de 2025.

ESPÍRITO SANTO, Eniel et al. Um mosaico de ideias sobre a Inteligência Artificial Generativa no contexto da educação. In: ALVES, Lynn (Org.). **Inteligência Artificial e educação: refletindo sobre os desafios contemporâneos**. Salvador: Edufba, 2023.

FÁVERO, Altair Alberto; CENTENARO, Junior Bufon. (2019). A pesquisa documental nas investigações de políticas educacionais: potencialidades e limites. **Contrapontos**, 19(1), 170-184. Epub 22 de agosto de 2019. Disponível em: <https://doi.org/10.14210/contrapontos.v19n1>. Acesso em: 15 jun. de 2025.

FREIRE, Paulo. **Pedagogia da autonomia: saberes necessários à prática educativa**. São Paulo: Paz e Terra, 1996.

LUCKIN, R.; HOLMES, W. Griffiths; M; FORCIER, L. B. Intelligence Unleashed: An argument for AI. **Education**. London: Pearson, 2016. Disponível em: <https://www.researchgate.net/publication/29961597>. Acesso em: 15 jun. 2025.

LÜDKE, Menga. ANDRÉ, Marli E.D.A. **A Pesquisa em Educação: abordagens qualitativas**. 2 ed. Rio de Janeiro: E.P.U., 2013.

RIBEIRO, Renato Janine. Editorial: Regulação da Inteligência Artificial no Brasil e no Mundo. **Sociedade Brasileira para o Progresso da Ciência (SBPC)**. Abril de 2024. Disponível em: <https://portal.sbpcnet.org.br/noticias/editorial-regulacao-da-inteligencia-artificial-no-brasil-e-no-mundo/>. Acesso em: 15 jun. 2025.

SANTAELLA, Lucia. A inteligência artificial sob a égide da ética. **Cronos: Revista da Pós-Graduação em Ciências Sociais**, Natal, v. 26, n. 1, p. 7-19, jan./jun. 2025. Disponível em: <https://periodicos.ufrn.br/cronos/article/view/39309>. Acesso em: 14 de Junho de 2025.

UNESCO. **Guia para a IA generativa na educação e na pesquisa**. Paris: UNESCO, 2023. Disponível em: https://unesdoc.unesco.org/ark:/48223/pf0000385146_por. Acesso em: 12 de abril de 2025.

UNIÃO EUROPEIA. **Regulamento (UE) 2024/1689 do Parlamento Europeu e do Conselho, de 12 de julho de 2024, que estabelece regras harmonizadas em matéria de inteligência artificial (Lei da Inteligência Artificial) e altera determinados atos legislativos da União**. Jornal Oficial da União Europeia, L 2024/1689, 12 jul. 2024. Disponível em: <https://artificialintelligenceact.eu/the-act/>. Acesso em: 24 de jun. 2026.