



SERVIÇO PÚBLICO FEDERAL
MINISTÉRIO DA EDUCAÇÃO
FUNDAÇÃO UNIVERSIDADE FEDERAL DE MATO GROSSO DO SUL
Curso de Direito – Campus CPCX



GABRIELE REZENDE DIAS

**IMPLEMENTAÇÃO E USO RESPONSÁVEL DE SISTEMAS DE
INTELIGÊNCIA ARTIFICIAL NO BRASIL - ANÁLISE DA PL 2338/2023**

Coxim – MS

2025

Gabriele Rezende Dias

**IMPLEMENTAÇÃO E USO RESPONSÁVEL DE SISTEMAS DE INTELIGÊNCIA
ARTIFICIAL NO BRASIL: ANÁLISE DA PL 2338/2023**

Trabalho de Conclusão de Curso, na modalidade de Artigo Científico, apresentado à Coordenação do Curso de Direito da Universidade Federal do Mato Grosso do Sul - Campus CPCX, como exigência para obtenção do grau de Bacharel em Direito.

Orientador: Tiago Andreotti e Silva

Coxim – MS

2025

Gabriele Rezende Dias

**IMPLEMENTAÇÃO E USO RESPONSÁVEL DE SISTEMAS DE INTELIGÊNCIA
ARTIFICIAL NO BRASIL - ANÁLISE DA PL 2338/2023**

Trabalho de Conclusão de Curso, na modalidade de Artigo Científico, apresentado à Coordenação do Curso de Direito da Universidade Federal do Mato Grosso do Sul - Campus CPCX, como exigência para obtenção do grau de Bacharel em Direito.

Orientador: Tiago Andreotti e Silva

Coxim (MS) 15 de outubro de 2025

BANCA EXAMINADORA

Francisco Ilidio Ferreira Rocha

Guilherme Sampieri Santinho

RESUMO

O presente trabalho analisa a regulamentação da inteligência artificial (IA) sob a perspectiva jurídica, considerando os riscos decorrentes da ausência de um marco normativo específico no Brasil e examinando o Projeto de Lei nº 2338/2023 à luz das melhores práticas internacionais. Inicialmente, apresenta-se a discussão sobre os impactos da falta de regulação, destacando implicações jurídicas, sociais e econômicas, bem como a importância da supervisão humana em decisões automatizadas. Em seguida, são examinadas as principais disposições do PL 2338/2023, com ênfase na obrigatoriedade da avaliação de impacto algorítmico e na adoção de medidas técnicas para garantir a explicabilidade dos resultados. Posteriormente, realiza-se uma análise comparativa com o AI Act europeu, identificando convergências e divergências quanto à classificação de risco, obrigações de transparência, fiscalização e aplicação de sanções. A pesquisa evidencia que, embora o projeto brasileiro esteja alinhado a princípios internacionais, carece de maior detalhamento técnico para assegurar sua efetividade. Conclui-se que a adoção de um marco regulatório equilibrado, flexível e atualizado é essencial para garantir o uso ético, seguro e socialmente responsável da IA, de modo a proteger direitos fundamentais, promover a inovação e fortalecer a segurança jurídica no país.

Palavras-chave: Inteligência artificial. Regulamentação. Projeto de Lei 2338/2023. AI Act. Direito digital.

ABSTRACT

This study analyzes the regulation of artificial intelligence (AI) from a legal perspective, considering the risks arising from the absence of a specific regulatory framework in Brazil and examining Bill No. 2338/2023 in light of international best practices. It first discusses the impacts of regulatory gaps, highlighting legal, social, and economic implications, as well as the importance of human oversight in automated decision-making. Then, it examines the main provisions of Bill No. 2338/2023, emphasizing the mandatory algorithmic impact assessment and the adoption of technical measures to ensure the explainability of results. Subsequently, it conducts a comparative analysis with the European AI Act, identifying convergences and divergences in terms of risk classification, transparency obligations, enforcement, and sanctions. The research shows that, although the Brazilian bill aligns with international principles, it lacks sufficient technical detail to ensure its effectiveness. It concludes that adopting a balanced, flexible, and updated regulatory framework is essential to ensure the ethical, safe, and socially responsible use of AI, in order to protect fundamental rights, foster innovation, and strengthen legal certainty in the country.

Keywords: Artificial intelligence. Regulation. Bill No. 2338/2023. AI Act. Digital law.

SUMÁRIO

INTRODUÇÃO	7
1.O QUE É A INTELIGÊNCIA ARTIFICIAL	8
2.A IA E SEUS RISCOS	10
 <u>2.1 OS RISCOS INERENTES AO DESENVOLVIMENTO E USO DA IA</u>	10
 <u>2.2 RISCOS DA AUSÊNCIA DE REGULAMENTAÇÃO DA IA</u>	12
3. REGULAMENTO EUROPEU: OBJETIVOS E FUNDAMENTOS	14
4. O PROJETO DE LEI 2338/2023	18
 <u>4.1 MECANISMOS DE GOVERNANÇA</u>	18
 <u>4.2 MELHORES PRÁTICAS INTERNACIONAIS E A CONVERGÊNCIA COM O MODELO BRASILEIRO</u>	20
CONSIDERAÇÕES FINAIS	23
REFERÊNCIAS	25

INTRODUÇÃO

O avanço acelerado da tecnologia nas últimas décadas transformou profundamente setores econômicos, sociais e culturais. Entre as inovações mais impactantes, a inteligência artificial (IA) ocupa posição central, influenciando desde processos produtivos até a tomada de decisões estratégicas em áreas públicas e privadas. Essa tecnologia, capaz de processar grandes volumes de dados e executar tarefas complexas com mínima intervenção humana, já se faz presente em segmentos como segurança pública, saúde, educação, comunicação e finanças. Contudo, os mesmos atributos que conferem eficiência e alcance à IA também ampliam os riscos, especialmente quando não existe um marco regulatório claro e eficaz.

No Brasil, a ausência de regulamentação específica para a IA gera insegurança jurídica e dificulta a definição de responsabilidades. Embora a Lei Geral de Proteção de Dados (LGPD) represente um avanço na proteção de informações pessoais, ela não aborda de forma direta os desafios técnicos e operacionais que caracterizam os sistemas de inteligência artificial. Sem parâmetros adequados, abre-se espaço para práticas prejudiciais, como decisões automatizadas opacas, uso indevido de dados e implantação de sistemas sem supervisão adequada, potencialmente lesando direitos fundamentais.

Com o objetivo de enfrentar esse cenário, foi apresentado o Projeto de Lei nº 2338/2023, que busca estabelecer princípios e diretrizes para o desenvolvimento e uso responsável de IA no Brasil. O Projeto de Lei nº 2338/2023 representa uma tentativa de suprir essa lacuna, buscando estabelecer parâmetros para o desenvolvimento e uso responsável de sistemas de inteligência artificial.

A proposta pretende alinhar o país a padrões internacionais, definindo níveis de risco, exigências de transparência e mecanismos de prestação de contas. Ainda assim, seu conteúdo suscita debates quanto à clareza conceitual, à efetividade dos dispositivos e à capacidade real de prevenir abusos. Portanto, é necessário avaliar em que medida o texto proposto é capaz de atender aos desafios e riscos associados à tecnologia. Esse exame torna-se ainda mais relevante quando comparado a modelos internacionais já consolidados, que podem servir de referência para o aprimoramento do marco legal brasileiro.

A experiência internacional oferece referências importantes para essa análise. Na União Europeia, o debate formal sobre o tema ganhou força a partir da publicação, em 2020, do White

Paper on Artificial Intelligence, documento que delineou riscos, oportunidades e a necessidade de uma regulamentação robusta. Esse estudo inicial resultou, em 2024, na aprovação do AI Act, primeiro regulamento de inteligência artificial com força vinculante para todos os países-membros da UE. O texto europeu estabelece uma estrutura detalhada de classificação de riscos, requisitos técnicos, sanções e mecanismos de fiscalização, servindo como exemplo de abordagem normativa abrangente.

A análise dessa trajetória oferece elementos concretos para compreender como o debate sobre riscos evoluiu até a formulação de regras específicas, além de fornecer exemplos de boas práticas que podem ser adaptadas à realidade brasileira.

A partir dessa conjuntura, surge o seguinte problema de pesquisa: *em que medida o Projeto de Lei n.º 2.338/2023 é capaz de assegurar aplicabilidade prática e proteção de direitos fundamentais, em comparação ao modelo regulatório adotado pela União Europeia no AI Act?*

Diante desse panorama, a comparação entre o PL 2338/2023 e o modelo europeu torna-se fundamental para identificar pontos de convergência, lacunas e oportunidades de aprimoramento. A análise comparativa não se limita a avaliar textos legislativos, mas busca compreender como cada proposta enfrenta os riscos associados à IA e de que forma garante segurança jurídica sem sufocar a inovação.

Assim, este trabalho propõe-se a investigar os riscos da ausência de regulamentação específica para a inteligência artificial no Brasil, examinar os objetivos e diretrizes do Projeto de Lei nº 2338/2023 e confrontá-los com as melhores práticas internacionais, especialmente as expressas no White Paper e no AI Act europeu. Com isso, pretende-se contribuir para o debate legislativo e para a construção de um marco legal capaz de conciliar desenvolvimento tecnológico e proteção efetiva dos direitos fundamentais.

Para tanto, na seção 1 explico o que é a inteligência artificial. Na seção 2 apresento os seus riscos e os riscos de ausência de sua regulamentação. Na seção 3 explico os fundamentos do regulamento europeu de inteligência artificial. Na seção 4 faço uma análise do projeto de Lei 2338/2023 e de seus mecanismos de governança comparando os principais aspectos de ambos os instrumentos.

1. O QUE É A INTELIGÊNCIA ARTIFICIAL

Para iniciarmos a análise do projeto de lei que versa sobre a regulação da Inteligência artificial (IA) no Brasil, é necessário compreender o que é a Inteligência Artificial. Primeiramente, é fundamental destacar que, atualmente, não existe uma definição acadêmica consolidada, universalmente aceita ou padronizada sobre o que seria a inteligência artificial. Esta ausência de consenso reflete a complexidade e a natureza multidisciplinar desta área, que envolve ciência da computação, neurociência, filosofia, psicologia cognitiva, matemática e diversas outras disciplinas.

Outrossim, é interessante recorrer ao próprio ChatGPT, uma das plataformas de inteligência artificial mais utilizadas pelo cidadão comum:

A inteligência artificial (IA) é um campo da ciência da computação que busca criar sistemas ou máquinas capazes de executar tarefas que normalmente exigiriam inteligência humana. Essas tarefas incluem aprender com experiências, raciocinar, reconhecer padrões, tomar decisões, compreender linguagem natural, perceber o ambiente e até interagir de forma autônoma. De forma simples: é quando um computador ou programa é projetado para pensar e agir de maneira inteligente, imitando certos processos da mente humana. (OPENAI, 2025).

No entanto, existem dois conceitos fundamentais que precisamos compreender para entender melhor o funcionamento da inteligência artificial. Eles representam tecnologias essenciais para que um sistema seja capaz de simular o raciocínio lógico humano (BARBOSA; PORTES, 2023), são eles:

Machine Learning: termo que significa “aprendizado da máquina”. É a tecnologia que propicia aos sistemas a capacidade de aprenderem sozinhos e tomarem decisões autônomas, seguindo o processamento de dados e identificação de padrões. Sem o Machine Learning, a inteligência artificial da contemporaneidade; a que sai da ficção e está presente na vida da população, não seria possível. O termo que significa “aprendizado da máquina”. É a tecnologia que propicia aos sistemas a capacidade de aprenderem sozinhos e tomarem decisões autônomas, seguindo o processamento de dados e identificação de padrões. Deep Learning: parte do machinelearning, o “aprendizado profundo” diz respeito a uma maior capacidade de aprendizado do sistema, pois utiliza redes neurais complexas. Essas redes seguem a mesma lógica da ligação neurônio-cérebro humano. Um dos principais exemplos do Deep Learning é o reconhecimento facial e de voz. Para trabalhar com inteligência artificial é necessário conhecimentos básicos em informática, matemática e lógica de computadores, que é a fundação da maioria dos programas de inteligência artificial. Em termos mais simples, a Inteligência Artificial refere-se a sistemas ou máquinas que mimetizam a inteligência humana para executar tarefas e podem se aprimorar iterativamente com base nas informações que eles coletam. (BARBOSA; PORTES, 2023).

Diante do exposto, observa-se que a Inteligência Artificial constitui um campo vasto e interdisciplinar, cuja complexidade reflete a ausência de uma definição única e universalmente aceita. Ainda assim, é possível compreender seu núcleo conceitual a partir de tecnologias estruturantes, como o *Machine Learning* e o *Deep Learning*, que permitem aos sistemas

aprender, raciocinar e tomar decisões de forma autônoma, aproximando-se de processos característicos da cognição humana.

Além disso, a utilização cotidiana de ferramentas como o ChatGPT evidencia a presença concreta da IA na vida social contemporânea, extrapolando os limites da pesquisa científica e alcançando o cidadão comum.

Diante disso, é fundamental discutir os riscos que acompanham o avanço da inteligência artificial e os desafios que sua regulação necessita enfrentar.

2. A IA E SEUS RISCOS

A Inteligência Artificial (IA) representa um dos avanços tecnológicos mais transformadores dos últimos tempos, prometendo inovações em diversos setores, porém gera preocupações quanto aos seus riscos éticos, sociais, econômicos e regulatórios. No contexto brasileiro, onde a adoção acelerada de sistemas de IA ocorre sem um marco legal específico, emergindo desafios que vão desde a perda de autonomia humana e o desequilíbrio no mercado de trabalho até vulnerabilidades cibernéticas e a ampliação de desigualdades sociais. Esta seção explora esses perigos em duas perspectivas principais: os riscos inerentes ao desenvolvimento e uso da IA (2.1), incluindo dependência tecnológica, impactos socioeconômicos e falhas operacionais, e os agravantes decorrentes da ausência de regulamentação (2.2), que podem comprometer a transparência, a responsabilidade e a proteção de direitos fundamentais.

2.1 OS RISCOS INERENTES AO DESENVOLVIMENTO E USO DA IA

O desenvolvimento e a adoção da Inteligência Artificial (IA) têm suscitado diversos debates acerca de seus potenciais riscos e das incertezas éticas, sociais e econômicas que emergem dessa tecnologia. Entre essas discussões, destaca-se a preocupação com o impacto da IA sobre a autonomia humana e a forma como nos relacionamos com as ferramentas que criamos. Nesse sentido, Guimarães (2024) alerta para o risco de que a tecnologia, ao assumir funções antes realizadas por nós, transforme-se em uma presença que, ao mesmo tempo em que nos auxilia, também nos torna dependentes:

O desenvolvimento da IA traz ao cenário uma possibilidade angustiante: que de

vicários – agem em nosso lugar – passem a vampiros – suguem nossas energias. Todo artefato gera dependência (apropriamo-nos deles e passamos a os considerar como parte de nossos organismos). E dependência sempre gera vulnerabilidade. Quanto mais ‘agente’ a IA for em nosso lugar, menos ‘agentes’ seremos nós mesmos, restando passivos e obsequentes. (Guimarães, 2024, p. 371).

Além da questão da autonomia, outros riscos estão associados ao impacto da IA no mercado de trabalho. Segundo Carvalho (2021), diversos estudos indicam que profissões tradicionais tendem a desaparecer, enquanto novas ocupações, ainda desconhecidas, serão criadas. Esse processo remete à Revolução Industrial, quando a automação transformou profundamente a organização do trabalho, eliminando funções e, ao mesmo tempo, possibilitando novas oportunidades.

Nessa mesma linha, Kaufman (2022) aprofunda a análise ao demonstrar que o desequilíbrio atual decorre do fato de as novas funções não substituírem proporcionalmente aquelas eliminadas, o que se deve a múltiplos fatores:

O desequilíbrio do mercado de trabalho deve-se, em parte, ao fato de que as novas funções não estão substituindo proporcionalmente as funções eliminadas, basicamente por duas razões: primeiramente porque é maior o número de funções repetitivas e previsíveis, foco da substituição por máquinas inteligentes, em comparação com as novas funções ofertadas; o estímulo à transformação digital é reduzir custos e aumentar a eficiência, processos intensivos em tecnologia (e não em mão de obra). Em segundo lugar porque é um equívoco considerar as habilidades listadas pelos relatórios de consultorias, e enfatizadas por muitos “especialistas” – raciocínio lógico, empatia, pensamento crítico, compreensão de leitura, argumentação, comunicação clara e persuasiva, discernimento, bom senso, capacidade de tomar decisão, aprendizagem ativa, fluência de ideias, originalidade –, como inerentes aos seres humanos. Elas são potencialmente habilidades humanas, representam potencial vantagem comparativa dos trabalhadores humanos. Para florescerem, contudo, dependem de condições apropriadas, e essas condições apropriadas não estão disponíveis para a maioria da população nos países desenvolvidos e, particularmente, nos países em desenvolvimento. Elas são potencialmente habilidades humanas, representam potencial vantagem comparativa dos trabalhadores humanos. Para florescerem, contudo, dependem de condições apropriadas, e essas condições apropriadas não estão disponíveis para a maioria da população nos países desenvolvidos e, particularmente, nos países em desenvolvimento. Responda sinceramente: você possui todas essas “habilidades humanas”? Quantas pessoas você conhece que as possuem?. (KAUFMAN, 2022, p. 20).

Nesse contexto, Sichman (2021) complementa a análise ao destacar cinco classes de riscos relacionados ao uso da IA, conforme proposto por Dietterich e Horvitz (2015): falhas (bugs), segurança cibernética, aprendiz de feitiçeiro, autonomia compartilhada e impactos socioeconômicos:

Falhas (bugs): Quaisquer sistemas de software apresentam falhas. Vários sistemas de software convencionais foram desenvolvidos e validados para atingir altos níveis de garantia de qualidade; por exemplo, sistemas de piloto automático e de controle de espaçonaves são cuidadosamente testados e validados. Práticas semelhantes devem ser aplicadas aos sistemas de IA;

Segurança (cybersecurity): Os sistemas de IA são tão vulneráveis quanto qualquer outro software a ataques cibernéticos. Por exemplo, ao manipular dados de treinamento ou preferências e trade-offs codificados em modelos de utilidade, adversários podem alterar o comportamento desses sistemas; Aprendiz de feiticeiro (sorcerer's apprentice): Um aspecto importante de qualquer sistema de IA que interage com as pessoas é que ele deve raciocinar sobre o que estas pretendem, em vez de executar comandos literalmente. Um sistema de IA deve analisar e compreender se o comportamento que um ser humano está solicitando pode ser julgado como "normal" ou "razoável" pela maioria das pessoas; Autonomia compartilhada (Shared autonomy): Construir esses sistemas colaborativos levanta um quarto conjunto de riscos decorrentes de desafios sobre fluidez de engajamento e clareza sobre estados internos e objetivos dos envolvidos no sistema. Criar sistemas em tempo real onde o controle precisa mudar rapidamente entre as pessoas e os sistemas de IA é difícil; Impactos socioeconômicos: Precisamos entender as influências da IA na distribuição de empregos e na economia de forma mais ampla. Essas questões perpassam a ciência e engenharia da computação, chegando ao domínio das políticas e programas econômicos que podem garantir que os benefícios dos aumentos de produtividade baseados em IA sejam amplamente compartilhados. (Sichman, 2021).

Outro ponto crítico é a vulnerabilidade da IA a ataques cibernéticos. De acordo com Sichman (2021), assim como qualquer software, sistemas baseados em IA podem ser manipulados por adversários que exploram falhas em dados de treinamento ou nos modelos de decisão, gerando consequências potencialmente graves. Esse cenário revela que a IA, além de trazer novos riscos, também potencializa riscos já conhecidos no campo da computação e da segurança da informação.

Portanto, é evidente que, embora a IA seja capaz de gerar inúmeros benefícios, seu mau uso pode trazer impactos severos à sociedade. Eles abrangem desde a ameaça de decisões enviesadas e preconceituosas até impactos estruturais no mercado de trabalho e vulnerabilidades de segurança digital. Como destacam Carvalho (2021), Sichman (2021), Kaufman (2022) e Guimarães (2024), enfrentar esses riscos requer tanto regulação quanto o desenvolvimento de práticas de design responsáveis, capazes de tornar a IA mais justa, transparente e confiável.

2.2 RISCOS DA AUSÊNCIA DE REGULAMENTAÇÃO DA IA

A ausência de um marco legal específico para inteligência artificial (IA) no Brasil cria um cenário de vulnerabilidade regulatória, no qual tecnologias com potencial de impacto direto sobre direitos fundamentais são implementadas sem garantias mínimas de transparência, explicabilidade e supervisão. A Autoridade Nacional de Proteção de Dados – ANPD (2024)

ressalta que, sem diretrizes claras, há risco de “adoção indiscriminada de sistemas que processam grandes volumes de dados, influenciam decisões humanas e afetam a vida de indivíduos sem que estes tenham meios adequados de contestação”.

Além disso, a inexistência de um marco regulatório robusto para a IA no Brasil implica riscos concretos para a sociedade, o mercado e o próprio Estado. Em termos jurídicos, a falta de normas claras dificulta a responsabilização por danos decorrentes do uso de algoritmos, especialmente quando se trata de sistemas complexos cujas decisões não são facilmente auditáveis. Isso pode gerar um cenário de impunidade ou de litígios prolongados, nos quais vítimas de decisões automatizadas terão dificuldade em comprovar falhas ou vieses. Em suma, a ausência de parâmetros normativos claros fragiliza a responsabilização e compromete a efetividade do sistema jurídico diante dos impactos da IA.

No campo social, a adoção de sistemas de IA sem supervisão adequada pode amplificar desigualdades preexistentes. Algoritmos mal calibrados ou treinados com dados enviesados tendem a reproduzir e até acentuar discriminações, afetando negativamente grupos vulneráveis. Além disso, o uso de IA em decisões sensíveis, como concessão de benefícios, sentenças judiciais e diagnósticos médicos, exige transparência e possibilidade de revisão humana — condições que, sem regulação, não são garantidas. Portanto, a falta de regulação amplia desigualdades já existentes e ameaça a proteção de direitos fundamentais, sobretudo em setores sensíveis como justiça, saúde e assistência social.

O risco de desumanização de processos decisórios é particularmente preocupante no sistema de justiça. Conforme alertam Teigão e Fogaça (2024):

A IA pode servir como ferramenta poderosa de apoio, mas não substitui a escuta, a empatia, a prudência e a ponderação que caracterizam o julgamento verdadeiramente justo. Assim, os fundamentos ético-jurídicos analisados nesta seção orientam para uma incorporação prudente e responsável da IA, que reconheça a tecnologia como instrumento auxiliar, jamais como substituto da decisão judicial humana. (TEIGÃO; FOGAÇA, 2024, p. 141).

Essa advertência evidencia que, para além das preocupações técnicas, a regulamentação deve resguardar valores como dignidade, igualdade e devido processo legal evitando a desumanização das decisões. Sem salvaguardas, abre-se espaço para que decisões críticas sejam tomadas com base em parâmetros opacos, sem que as pessoas afetadas compreendam ou questionem o raciocínio por trás dos resultados.

No campo econômico, a ausência de regulação pode afastar investimentos. Grandes empresas e investidores tendem a preferir mercados com regras estáveis, previsíveis e alinhadas a padrões internacionais. A União Europeia, por exemplo, após a publicação do AI Act, consolidou-se como um espaço de referência para negócios em IA, pois sua legislação estabelece obrigações proporcionais ao risco, criando segurança para todos os atores do ecossistema. A experiência europeia demonstra que a regulamentação de IA não se limita a impor restrições, mas também atua como um instrumento de fomento à pesquisa e à inovação. O White Paper on Artificial Intelligence (COMISSÃO EUROPEIA, 2020) evidencia a importância de políticas de incentivo como parte integrante da governança tecnológica:

Uma das razões para a posição forte da Europa em termos de pesquisa é o programa de financiamento da União Europeia, que se mostrou fundamental para unir esforços, evitar duplicações e potencializar investimentos públicos e privados nos Estados-Membros. Nos últimos três anos, o financiamento da UE para pesquisa e inovação em IA aumentou para € 1,5 bilhão, ou seja, um aumento de 70% em comparação com o período anterior. (COMISSÃO EUROPEIA, 2020, p. 5).

Esse exemplo mostra que regulamentação e inovação não são conceitos antagônicos. Pelo contrário, quando bem articulados, geram um ambiente que une segurança jurídica e competitividade tecnológica. Assim, a experiência europeia mostra que a regulação, longe de ser um entrave, é também um motor de inovação e competitividade, gerando previsibilidade e atraindo investimentos.

O Brasil, se não avançar rapidamente nesse debate, corre o risco de ficar à margem dessa corrida tecnológica. No entanto, o Projeto de Lei nº 2338/2023 surge como tentativa de estabelecer princípios e diretrizes para o desenvolvimento e uso de IA buscando alinhar-se a práticas internacionais já consolidadas.

3. REGULAMENTO EUROPEU: OBJETIVOS E FUNDAMENTOS

O Regulamento Europeu sobre Inteligência Artificial (AI Act), aprovado em 2024, é o primeiro instrumento normativo de caráter vinculante no mundo a tratar de forma abrangente a IA. Seu objetivo central é criar um “ecossistema de excelência e confiança”, conciliando inovação tecnológica com proteção de direitos fundamentais. Nesse sentido, observa-se que o regulamento visa não apenas harmonizar o mercado interno, mas também assegurar que o desenvolvimento e a utilização da IA ocorram em conformidade com valores democráticos e centrada no ser humano:

O presente regulamento tem por objetivo a melhoria do funcionamento do mercado interno mediante a previsão de um regime jurídico uniforme, em particular para o desenvolvimento, a colocação no mercado, a colocação em serviço e a utilização de sistemas de inteligência artificial (sistemas de IA) na União, em conformidade com os valores da União, a fim de promover a adoção de uma inteligência artificial (IA) centrada no ser humano e de confiança, assegurando simultaneamente um elevado nível de proteção da saúde, da segurança, dos direitos fundamentais consagrados na Carta dos Direitos Fundamentais da União Europeia («Carta»), nomeadamente a democracia, o Estado de direito e a proteção do ambiente, a proteção contra os efeitos nocivos dos sistemas de IA na União, e de apoiar a inovação. O presente regulamento assegura a livre circulação transfronteiriça de produtos e serviços baseados em IA, evitando assim que os Estados-Membros imponham restrições ao desenvolvimento, à comercialização e à utilização dos sistemas de IA, salvo se explicitamente autorizado pelo presente regulamento. (REGULAMENTO (UE), 2024/1689, p. 1)

Segundo o Parlamento Europeu (2024), o AI Act estabelece condições proporcionais ao risco apresentado pelo sistema, classificando-o em categorias de risco mínimo, limitado, alto ou inaceitável. O regulamento define quatro condições principais que, quando atendidas, indicam que um sistema de IA não deve ser classificado como de risco elevado.

A primeira condição refere-se à natureza restrita e limitada da tarefa desempenhada pelo sistema. Exemplos de sistemas que transformam dados não estruturados em dados estruturados, classificam documentos em categorias específicas ou incluem duplicações em grandes volumes de dados. Essas tarefas processuais são consideradas de baixo risco, pois não ampliam os perigos mesmo quando aplicadas em contextos que, em outras circunstâncias, seriam de risco elevado.

A segunda condição contempla sistemas que atuam para melhorar resultados de atividades humanas já concluídas, relevantes para usos de risco elevado. Nesses casos, o sistema de IA funciona como uma camada adicional, aprimorando, por exemplo, a linguagem de documentos previamente redigidos, ajustando o tom, o estilo acadêmico ou o alinhamento com uma mensagem de marca. Tal abordagem reduz o risco, uma vez que o sistema não substitui a decisão humana, mas um complemento.

A terceira condição abrange sistemas destinados a detectar padrões ou desvios em decisões anteriores, sem substituir ou influenciar essas decisões sem revisão humana adequada. Um exemplo típico é um sistema que verifica ex post possíveis incoerências na atribuição de notas por um professor, sinalizando anomalias para análise posterior. Essa função de monitoramento contribui para a redução do risco, mantendo a decisão final sob controle humano.

Por fim, a quarta condição refere-se a sistemas que executam tarefas preparatórias para

avaliações posteriores relevantes para sistemas de risco elevado. Tais tarefas indexação, pesquisa, processamento de texto e voz, ligação de dados a outras fontes ou tradução inicial de documentos. O impacto desses sistemas diretos é limitado, o que justifica sua classificação como de risco limitado.

Para a garantia da transparência e a rastreabilidade, o Regulamento Europeu (2024), dispõe:

A fim de assegurar a rastreabilidade e a transparência, um prestador que considere que um sistema de IA não é de risco elevado com base nas condições referidas supra deverá elaborar a documentação da avaliação antes de esse sistema ser colocado no mercado ou colocado em serviço e deverá apresentar essa documentação às autoridades nacionais competentes mediante pedido. Tal prestador deverá ser obrigado a registar o sistema de IA na base de dados da UE criada ao abrigo do presente regulamento. Com vista a disponibilizar orientações adicionais para a aplicação prática das condições ao abrigo das quais os sistemas de IA referidos num anexo ao presente regulamento não são, a título excecional, de risco elevado, a Comissão deverá, após consulta do Comité, disponibilizar orientações que especifiquem essa aplicação prática, completadas por uma lista exaustiva de exemplos práticos de casos de utilização de sistemas de IA que sejam de risco elevado e de casos de utilização risco não elevado. (REGULAMENTO (UE), 2024/1689, p. 15)

Essa iniciativa visa facilitar a correta interpretação e aplicação do regulamento, promovendo a segurança jurídica e a proteção dos direitos fundamentais.

No que tange às sanções, o AI Act adota uma postura rigorosa ao remeter aos Estados membros a regulamentação de fiscalização e punição que asseguram a aplicação de suas disposições, já estabelecendo parâmetros claros para coimas e medidas de execução. Dessa forma, o texto legal dispõe:

O cumprimento do presente regulamento deverá ter força executória através da imposição de sanções e de outras medidas de execução. Os Estados-Membros deverão tomar todas as medidas necessárias para assegurar a aplicação das disposições do presente regulamento, inclusive estabelecendo sanções efetivas, proporcionadas e dissuasivas aplicáveis em caso de violação dessas disposições, e a respeito do princípio ne bis in idem. A fim de reforçar e harmonizar as sanções administrativas em caso de infração ao presente regulamento, deverão ser estabelecidos os limites máximos para a fixação de coimas para determinadas infrações específicas. Ao avaliar o montante das coimas, os Estados-Membros deverão, em cada caso individual, ter em conta todas as circunstâncias relevantes da situação específica, prestando a devida atenção à natureza, à gravidade e à duração da infração e às suas consequências, bem como à dimensão do prestador, em particular se o prestador for uma PME, incluindo uma empresa em fase de arranque. A Autoridade Europeia para a Proteção de Dados deverá ter competências para impor coimas às instituições, órgãos e organismos da União que se enquadram no âmbito do presente regulamento. (REGULAMENTO (UE), 2024/1689, p. 14 e 15).

O White Paper que deu origem à proposta alertava para a necessidade de superar a fragmentação das iniciativas de pesquisa no continente:

A Europa não pode se dar ao luxo de manter o atual cenário fragmentado de centros de competência, sem que nenhum alcance a escala necessária para competir com os principais institutos globais. É imperativo criar mais sinergias e redes entre os múltiplos centros de pesquisa europeus em IA e alinhar seus esforços para melhorar a excelência, reter e atrair os melhores pesquisadores e desenvolver a melhor tecnologia. A Europa precisa de um centro de referência em pesquisa, inovação e especialização que coordene esses esforços, seja uma referência mundial de excelência em IA e que possa atrair investimentos e os melhores talentos da área. (COMISSÃO EUROPEIA, 2020, p. 7).

O AI Act materializa essa visão ao instituir exigências técnicas e jurídicas para sistemas de IA de alto risco, prevendo mecanismos como auditorias obrigatórias, supervisão humana contínua, explicabilidade e certificações. Adicionalmente, proíbe práticas como vigilância biométrica em tempo real em espaços públicos, salvo em hipóteses restritas, e sistemas de pontuação social (social scoring) por autoridades públicas.

Um dos elementos mais inovadores do AI Act é a adoção de uma abordagem regulatória baseada em risco, calibrando o nível de exigência conforme o potencial de dano da aplicação. A Comissão Europeia (2024) afirma que essa estratégia permite proteger os cidadãos sem sufocar a inovação, assegurando que “aplicações de baixo risco continuem a ser desenvolvidas com flexibilidade, enquanto sistemas críticos sejam rigorosamente supervisionados”.

O White Paper destaca a importância de critérios objetivos para classificar sistemas de alto risco:

Uma abordagem baseada em risco é importante para ajudar a garantir que a intervenção regulatória seja proporcional. No entanto, ela exige critérios claros para diferenciar as diferentes aplicações de IA, em particular no que se refere à questão de saber se elas são ou não de ‘alto risco’. A determinação do que é uma aplicação de IA de alto risco deve ser clara e facilmente compreensível e aplicável para todas as partes interessadas. Ainda assim, mesmo que uma aplicação de IA não seja classificada como de alto risco, ela permanece totalmente sujeita às regras já existentes na União Europeia. (COMISSÃO EUROPEIA, 2020, p. 18).

Essa clareza é essencial para evitar interpretações divergentes e garantir previsibilidade regulatória. No Brasil, embora o PL 2338/2023 mencione a classificação por níveis de risco, o texto ainda não detalha os parâmetros técnicos que definem cada categoria, delegando essa tarefa a regulamentações posteriores — o que, segundo a ANPD (2024), pode gerar insegurança jurídica e atrasos na implementação prática das normas.

4. O PROJETO DE LEI 2338/2023

A inteligência artificial (IA) consolidou-se como um elemento transformador das estruturas sociais, econômicas e jurídicas contemporâneas. Seu potencial disruptivo é acompanhado por riscos significativos, especialmente quando utilizada em contextos críticos como saúde, segurança, crédito e justiça. A comparação entre o AI Act e o PL 2338/2023 revela convergências importantes, contudo, também evidencia diferenças substanciais na estrutura normativa e no detalhamento técnico.

O AI Act define de forma minuciosa obrigações específicas para cada categoria de risco, cria mecanismos de certificação e prevê sanções proporcionais à gravidade da infração. Já o PL 2338/2023, apesar de alinhado em valores, apresenta dispositivos mais genéricos, deixando lacunas que poderão comprometer a eficácia da lei no curto prazo. A Nota Técnica do Data Privacy Brasil (2024) enfatiza que, sem parâmetros claros, “o cumprimento da lei dependerá de regulamentações futuras, o que pode adiar a efetividade da proteção aos cidadãos”.

Assim, a experiência europeia oferece ao legislador brasileiro um modelo no qual a combinação de regras detalhadas, fiscalização efetiva e estímulos à pesquisa permitiu a criação de um ambiente equilibrado entre segurança e inovação. A adoção de elementos semelhantes no Brasil pode contribuir para reduzir riscos e promover um ecossistema de IA mais confiável e competitivo.

Diante desse cenário, propostas como o Projeto de Lei nº 2338/2023 ganham relevância ao tentar preencher a lacuna normativa, aproximando o Brasil de padrões já adotados em outros países e blocos econômicos, como a União Europeia. Este capítulo examina, de forma integrada, três eixos centrais para a construção de um marco regulatório efetivo: (4.1) as diretrizes e mecanismos previstos no PL 2338/2023 e (4.2) as melhores práticas internacionais, com destaque para o White Paper de 2020 e o AI Act de 2024.

4.1 MECANISMOS DE GOVERNANÇA

O PL 2338/2023 representa a proposta mais concreta de regulamentação de inteligência artificial atualmente em tramitação no Brasil. Seu texto estabelece princípios, direitos e deveres que buscam equilibrar a promoção da inovação com a proteção de direitos fundamentais.

Um de seus dispositivos mais significativos está no artigo 25, que determina:

Art. 25. A avaliação de impacto algorítmico de sistemas de IA é obrigação do desenvolvedor ou do aplicador que introduzir ou colocar sistema de IA em circulação no mercado, sempre que o sistema ou o seu uso forem de alto risco, considerando o papel e a participação do agente na cadeia. (BRASIL, 2023, p. 15).

Essa exigência estabelece a necessidade de um exame prévio detalhado sobre potenciais efeitos adversos de sistemas de alto risco antes de sua introdução no mercado. Ao vincular a obrigação tanto ao desenvolvedor quanto ao aplicador, o texto amplia o alcance da responsabilidade e impede que agentes envolvidos na cadeia de produção ou distribuição se eximam do dever de diligência.

Outro ponto central do projeto é a ênfase na explicabilidade, ou seja, na capacidade de tornar compreensíveis os resultados gerados pela IA. O inciso V do mesmo artigo dispõe:

V – adoção de medidas técnicas para viabilizar a explicabilidade dos resultados dos sistemas de inteligência artificial e de medidas para disponibilizar aos operadores e potenciais impactados informações gerais sobre o funcionamento do modelo de inteligência artificial empregado, explicitando a lógica e os critérios relevantes para a produção de resultados, bem como, mediante requisição do interessado, disponibilizar informações adequadas que permitam a interpretação dos resultados concretamente produzidos, respeitado o sigilo industrial e comercial. (BRASIL, 2023, p. 15).

Esse dispositivo responde a um dos principais problemas relacionados à IA: a chamada “caixa-preta algorítmica”. A opacidade nos critérios e processos de decisão de muitos sistemas automatizados impede que operadores e usuários compreendam como e por que determinado resultado foi produzido. A previsão legal de explicabilidade é, portanto, um avanço significativo, pois viabiliza tanto o controle social quanto a responsabilização jurídica em casos de falhas ou abusos.

Outrossim, o AI Act estabelece autoridades competentes nos Estados-membros e um comitê europeu para harmonização das práticas. O PL 2338/2023, por sua vez, ainda não define com clareza quais órgãos terão competência fiscalizatória, embora a ANPD figure como possível autoridade central.

Além desses dispositivos, o PL 2338/2023 traz princípios como transparência, não discriminação, supervisão humana e proteção de dados, alinhando-se a valores que também estruturam regulações internacionais. No entanto, diferentemente do AI Act europeu, o texto brasileiro ainda carece de detalhamento técnico para operacionalizar esses princípios, deixando lacunas que precisarão ser preenchidas por regulamentações infralegais.

4.2 MELHORES PRÁTICAS INTERNACIONAIS E A CONVERGÊNCIA COM O MODELO BRASILEIRO

A experiência internacional oferece subsídios importantes para o aperfeiçoamento do marco regulatório brasileiro. Um exemplo relevante é o White Paper on Artificial Intelligence, publicado pela Comissão Europeia em 2020, que serviu de base para a elaboração do AI Act. Entre as diretrizes propostas, destaca-se a necessidade de comunicação clara sobre o desempenho e as limitações dos sistemas:

Garantir que sejam fornecidas informações claras sobre as capacidades e limitações do sistema de IA, em particular sobre a finalidade para a qual os sistemas se destinam, as condições em que se pode esperar que funcionem conforme o previsto e o nível esperado de precisão para atingir a finalidade especificada. Essas informações são importantes especialmente para os implementadores dos sistemas, mas também podem ser relevantes para autoridades competentes e partes afetadas. (COMISSÃO EUROPEIA, 2020, p. 21).

A União Europeia avançou a partir dessas diretrizes para aprovar o AI Act em 2024, que institui um modelo regulatório baseado em risco, dividindo os sistemas de IA em quatro categorias: risco mínimo, risco limitado, alto risco e risco inaceitável. Para cada categoria, são previstas obrigações proporcionais, desde requisitos de transparência até a proibição total de determinadas práticas, como o uso de IA para avaliação social de cidadãos (social scoring).

Ao comparar o AI Act com o PL 2338/2023, percebe-se que ambos compartilham princípios como a obrigatoriedade de avaliação de impacto, a promoção da transparência e a supervisão humana. Contudo, o modelo europeu detalha minuciosamente os requisitos técnicos para cada categoria de risco, incluindo padrões para testes, auditorias e relatórios, enquanto o texto brasileiro delega parte significativa desses detalhes para regulamentações posteriores.

Essa diferença tem implicações práticas relevantes. O Brasil corre o risco de aprovar uma lei de princípios sem mecanismos claros de implementação, o que pode comprometer a efetividade da norma. Por outro lado, a abertura para regulamentações complementares também pode permitir maior flexibilidade para adaptar as exigências às especificidades nacionais, evitando engessamento excessivo.

A convergência entre o PL 2338/2023 e as práticas europeias é, portanto, promissora, mas exige um esforço adicional de detalhamento técnico e de fortalecimento institucional para que a lei seja aplicável na prática e capaz de inspirar confiança tanto na sociedade quanto no mercado.

Enquanto o AI Act estabelece um sistema detalhado e escalonado de requisitos, o PL 2338/2023 adota uma abordagem principiológica, deixando o detalhamento técnico para regulamentações futuras.

Essa diferença metodológica impacta diretamente a aplicabilidade prática das normas. De um lado, o AI Act garante clareza imediata para desenvolvedores, aplicadores e autoridades de fiscalização; de outro, o projeto brasileiro oferece maior flexibilidade, permitindo adaptações rápidas diante de inovações tecnológicas.

A seguir, apresenta-se uma síntese comparativa dos dois modelos, considerando seus principais aspectos estruturais e de implementação:

Tabela 1 – Comparativo entre o PL 2338/2023 (Brasil) e o AI Act (União Europeia)

Aspecto	PL 2338/2023 – Brasil	AI Act – União Europeia (2024)
Classificação por risco	Prevê identificação de sistemas de alto risco, com exigência de avaliação de impacto algorítmico (Art. 25). Não define categorias intermediárias.	Divide em quatro categorias: risco mínimo, risco limitado, alto risco e risco inaceitável, com requisitos proporcionais.
Avaliação de impacto	Obrigatória para sistemas de alto risco antes da introdução no mercado. Abrange desenvolvedores e aplicadores.	Obrigatória para sistemas de alto risco, com requisitos técnicos padronizados e documentação detalhada.
Explicabilidade e transparência	Exige medidas técnicas para explicabilidade, fornecendo informações sobre lógica e critérios utilizados, respeitando sigilo industrial (Art. 20, V).	Impõe requisitos de documentação técnica, rastreabilidade e comunicação clara ao usuário e autoridades competentes.
Supervisão	Prevê supervisão humana em	Define padrões para supervisão

humana	sistemas de IA que possam impactar direitos fundamentais, mas sem detalhar procedimentos.	humana eficaz, com responsabilidades específicas para operadores e desenvolvedores.
Sanções e fiscalização	Determina que regulamentações posteriores definam sanções e procedimentos de fiscalização; atualmente, não detalha valores ou prazos.	Prevê multas proporcionais à gravidade da infração, podendo alcançar até 6% do faturamento global da empresa.
Proteção de dados	Integra-se à LGPD, sem criar regras adicionais específicas para dados usados em IA.	Exige conformidade com o GDPR e estabelece controles adicionais para dados de treinamento de sistemas de alto risco.
Flexibilidade regulatória	Delega grande parte do detalhamento técnico para regulamentações infralegais, permitindo ajustes conforme evolução tecnológica.	Mais rígido, com detalhamento direto no texto legal e possibilidade de revisão periódica pelo Parlamento Europeu.

Fonte: Desenvolvido pela autora.

A análise desenvolvida ao longo deste capítulo evidencia que a regulamentação da inteligência artificial demanda equilíbrio entre segurança jurídica, proteção de direitos e estímulo à inovação. O PL 2338/2023 demonstra avanços importantes, como a previsão de avaliação de impacto e a exigência de explicabilidade, aproximando-se de práticas consagradas no cenário internacional. No entanto, sua efetividade dependerá da capacidade do Estado brasileiro de regulamentar, fiscalizar e atualizar continuamente essas normas.

Por sua vez, o AI Act europeu oferece um modelo de referência pela precisão técnica e pelo enfoque baseado em risco, capaz de inspirar ajustes no contexto brasileiro. Ainda que existam diferenças marcantes — especialmente quanto ao grau de detalhamento e à rigidez normativa —, há espaço para convergência e harmonização de princípios, de modo a inserir o

Brasil no debate global sobre governança de IA.

CONSIDERAÇÕES FINAIS

O avanço acelerado da inteligência artificial impõe ao direito o desafio de construir um arcabouço regulatório capaz de responder, de forma tempestiva e eficaz, aos riscos e oportunidades dessa tecnologia. Ao longo deste trabalho, constatou-se que a ausência de normas específicas no Brasil potencializa vulnerabilidades jurídicas, sociais e econômicas, especialmente em setores sensíveis onde decisões automatizadas podem afetar direitos fundamentais.

O Projeto de Lei nº 2338/2023 surge como uma tentativa concreta de suprir essa lacuna normativa, incorporando diretrizes que visam garantir transparência, explicabilidade, avaliação prévia de riscos e supervisão humana. A exigência de avaliação de impacto algorítmico e a previsão de medidas técnicas para explicabilidade representam avanços significativos, pois abordam dois dos principais desafios da governança em IA: a opacidade dos sistemas e a dificuldade de responsabilização em casos de falhas ou vieses.

A comparação entre o PL 2338/2023 e o AI Act europeu (2024) evidencia duas abordagens regulatórias distintas quanto ao uso da inteligência artificial. O projeto brasileiro adota um modelo mais aberto e flexível, priorizando a identificação de riscos e delegando o detalhamento técnico para regulamentações futuras, o que pode favorecer a adaptação à evolução tecnológica, mas também gera incerteza jurídica diante da ausência de definições claras e sanções específicas. Já o regulamento europeu se destaca por ser mais rígido e abrangente, com classificação detalhada de riscos, requisitos técnicos minuciosos, padrões claros de transparência e supervisão humana, além de sanções robustas que fortalecem sua efetividade.

Assim, enquanto o Brasil caminha para um modelo regulatório incipiente e adaptável, a União Europeia consolida um arcabouço estruturado e protetivo, que busca equilibrar inovação com segurança jurídica e tutela de direitos fundamentais. Essa diferença reflete a maturidade regulatória de cada sistema e levanta o debate sobre a necessidade de o Brasil encontrar um ponto de equilíbrio entre flexibilidade e efetividade normativa.

A análise comparativa com o AI Act europeu revelou que, embora ambos os modelos compartilhem princípios essenciais, como a centralidade dos direitos fundamentais e a

supervisão humana, há diferenças relevantes na forma de implementação. Enquanto o regulamento europeu adota um sistema de classificação por risco altamente detalhado e vinculante, o projeto brasileiro privilegia uma abordagem principiológica, delegando o detalhamento técnico a regulamentações posteriores. Essa distinção confere ao Brasil maior flexibilidade, mas também exige forte capacidade institucional para transformar princípios em práticas efetivas.

O diálogo com as melhores práticas internacionais, especialmente aquelas delineadas no White Paper on Artificial Intelligence e consolidadas no AI Act, mostrou-se fundamental para orientar o desenvolvimento de um marco regulatório nacional que seja compatível com padrões globais, sem perder de vista as especificidades sociais, econômicas e culturais do país. A convergência parcial entre os modelos sugere que o Brasil pode, a partir do PL 2338/2023, estabelecer um regime jurídico equilibrado, que estimule a inovação ao mesmo tempo em que protege cidadãos e garante segurança jurídica.

Por fim, reafirma-se que a regulamentação da inteligência artificial não deve ser concebida como um instrumento de mera contenção tecnológica, mas como um mecanismo de governança que assegure o uso ético, seguro e socialmente benéfico da IA. Para tanto, é imprescindível que o texto final da lei seja acompanhado por regulamentações técnicas claras, fiscalização eficiente e constante atualização normativa, de modo a acompanhar a rápida evolução dessa tecnologia. Somente assim será possível transformar o potencial disruptivo da IA em um vetor de desenvolvimento inclusivo e sustentável para o Brasil.

REFERÊNCIAS

ANPD – AUTORIDADE NACIONAL DE PROTEÇÃO DE DADOS. Análise preliminar do PL 2338/2023. Brasília, 2024. Disponível em: https://www.gov.br/anpd/pt-br/assuntos/noticias/analise-preliminar-do-pl-2338_2023-formatado-ascom.pdf. Acesso em: 15 ago. 2025.

BARBOSA, Lucia Martins; PORTES, Luiza Alves Ferreira. A inteligência artificial. Revista Tecnologia Educacional [online], Rio de Janeiro, n. 236, p. 16-27, 2023.

BOLETIM OFICIAL DA UNIÃO EUROPEIA. Regulamento (UE) 2024/1689 do Parlamento Europeu e do Conselho. Luxemburgo, 2024. Disponível em: <https://www.boe.es/buscar/doc.php?id=DOUE-L-2024-81079>. Acesso em: 15 ago. 2025.

BRASIL. Câmara dos Deputados. Projeto de Lei n. 2338/2023. Brasília, 2023. Disponível em: <https://www.camara.leg.br/proposicoesWeb/fichadetramitacao?idProposicao=2487262>. Acesso em: 15 ago. 2025.

BRASIL. Senado Federal. Projeto de Lei n. 2338/2023. Brasília, 2023. Disponível em: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>. Acesso em: 15 ago. 2025.

BRASIL. Senado Federal. Texto integral do Projeto de Lei n. 2338/2023. Brasília, 2023. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?disposition=inline&dm=9347622&ts=1733877727346>. Acesso em: 15 ago. 2025.

CARVALHO, A. C. P. de L. F. de. Inteligência artificial: riscos, benefícios e uso responsável. Estudos Avançados, v. 35, n. 101, p. 21–36, 2021. DOI: <https://doi.org/10.1590/s0103-4014.2021.35101.003>.

COMISSÃO EUROPEIA. Regulatory framework proposal for Artificial Intelligence. Bruxelas, 2024. Disponível em: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>. Acesso em: 15 ago. 2025.

COMISSÃO EUROPEIA. White Paper on Artificial Intelligence: a European approach to excellence and trust. Bruxelas, 2020. Disponível em: https://commission.europa.eu/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust_en. Acesso em: 15 ago. 2025.

DATAPRIVACY BRASIL. Nota técnica – Temas regulatórios em inteligência artificial. São Paulo, 2024. Disponível em: https://www.dataprivacybr.org/wp-content/uploads/2024/02/nota-tecnica-temas-regulatorios-ia_data.pdf. Acesso em: 15 ago. 2025.

DIGITAL.GOV.PT. AI Act: regulamentação europeia para inteligência artificial. Lisboa, 2024. Disponível em: <https://digital.gov.pt/regulamentacao/ai-act>. Acesso em: 15 ago. 2025.

DIREITOS NA REDE. PL 2338/2023: por uma regulação que defenda direitos! São Paulo, 2024. Disponível em: <https://direitosnarede.org.br/2024/06/18/projeto-de-lei-n-2338-2023/>. Acesso em: 15 ago. 2025.

GUIMARÕES, A. Possibilidades de autonomia da inteligência artificial. Revista Paranaense de Filosofia, v. 4, n. 1, p. 344–374, 2024. DOI: <https://doi.org/10.33871/27639657.2024.4.1.9192>.

INDEX LAW. Governança digital e PL 2338/2023. Revista de Direito Digital e Sistemas Sociais, São Paulo, 2024. Disponível em: <https://www.indexlaw.org/index.php/revistaddsus/article/view/10364>. Acesso em: 15 ago. 2025.

JUSBRASIL. A regulação da inteligência artificial no Brasil. São Paulo, 2024. Disponível em: <https://www.jusbrasil.com.br/artigos/a-regulacao-da-inteligencia-artificial-no-brasil/3475725911>. Acesso em: 15 ago. 2025.

KAUFMAN, Dora. Desmistificando a inteligência artificial. São Paulo: Autêntica Editora, 2022. E-book. ISBN 9786559281596. Disponível em: <https://app.minhabiblioteca.com.br/reader/books/9786559281596/>. Acesso em: 19 set. 2025.

MINDS.DIGITAL. Projeto de Lei 2338/2023 e o marco regulatório da inteligência artificial. São Paulo, 2024. Disponível em: <https://minds.digital/digital-transformation/projeto-de-lei-2338-2023/>. Acesso em: 15 ago. 2025.

OPENAI. ChatGPT [inteligência artificial]. São Francisco: OpenAI, 2025. Disponível em: <https://chat.openai.com/>. Acesso em: 22 ago. 2025.

PARLAMENTO EUROPEU. EU AI Act: first regulation on artificial intelligence. Bruxelas, 2024. Disponível em: <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>. Acesso em: 15 ago. 2025.

SICHMAN, J. S. Inteligência artificial e sociedade: avanços e riscos. Estudos Avançados, v. 35, n. 101, p. 37–50, 2021. DOI: <https://doi.org/10.1590/s0103-4014.2021.35101.004>.

TRIBUNAL DE JUSTIÇA DO PARANÁ. Uso ético e responsável da inteligência artificial no Judiciário brasileiro. Curitiba, 2024. Disponível em: <https://revista.tjpr.jus.br/gralhaazul/article/download/189/141/513>. Acesso em: 15 ago. 2025.