

ENTRE INOVAÇÃO E PROTEÇÃO: A REGULAMENTAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO BRASIL SOB A PERSPECTIVA DOS DIREITOS HUMANOS

BETWEEN INNOVATION AND PROTECTION: THE REGULATION OF ARTIFICIAL INTELLIGENCE IN BRAZIL UNDER THE PERSPECTIVE OF HUMAN RIGHTS

ENTRE LA INNOVACIÓN Y LA PROTECCIÓN: LA REGULACIÓN DE LA INTELIGENCIA ARTIFICIAL EN BRASIL DESDE LA PERSPECTIVA DE LOS DERECHOS HUMANOS

Vânia Ereni Lima Vieira¹, Edson Cosme Martins Filho², Renata Flávia Nobre Canela Dias³, Daniel Ferreira dos Santos⁴, Leonardo de Oliveira Lopes⁵, André Crisóstomo Fernandes⁶

DOI: 10.54899/dcs.v23i92.6262

Recibido: 01/06/2026 | Aceptado: 24/06/2026 | Publicación en línea: 01/07/2026.

RESUMO

A crescente incorporação da Inteligência Artificial (IA) em processos decisórios públicos e privados tem provocado transformações relevantes nas estruturas jurídicas e institucionais, suscitando novos desafios à proteção dos direitos humanos. Embora a inovação tecnológica represente avanços significativos, a utilização de sistemas algorítmicos opacos, baseados em grandes volumes de dados e decisões automatizadas, revela riscos éticos, jurídicos e sociais capazes de aprofundar desigualdades estruturais e fragilizar garantias fundamentais, como a igualdade, a não discriminação, a privacidade e o devido processo. Diante desse cenário, o presente artigo tem por objetivo analisar o processo de regulamentação da IA no Brasil sob a perspectiva da proteção dos direitos humanos, avaliando seus avanços e limitações à luz de experiências regulatórias internacionais. A pesquisa adota abordagem qualitativa, de natureza teórico-normativa, com procedimentos técnicos bibliográficos e documentais, fundamentando-se em literatura especializada nacional e internacional, bem como na análise de marcos normativos e institucionais brasileiros. Os resultados indicam que o Brasil dispõe de instrumentos relevantes para a governança da IA, como o Marco Civil da Internet, a Lei Geral de Proteção de Dados Pessoais (LGPD), a Estratégia Brasileira de Inteligência Artificial (EBIA) e o Projeto de Lei nº

¹ Doutoranda em Direito pela Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, Minas Gerais, Brasil.
E-mail: vaniaerenilimavieira@yahoo.com.br

² Mestrando em Direito pela Faculdade Milton Campos (FMC), Nova Lima, Minas Gerais, Brasil.
E-mail: edsoncosmemartins@gmail.com

³ Doutora em Educação pela Universidade de Uberaba (UNIUBE), Uberaba, Minas Gerais, Brasil.
E-mail: renatanobredias@gmail.com

⁴ Mestre em Proteção dos Direitos Fundamentais pela Universidade de Itaúna (UIT), Itaúna, Minas Gerais, Brasil.
E-mail: daniel.dir@hotmail.com

⁵ Doutorando em Direito pela Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, Minas Gerais, Brasil.
E-mail: leolopesadv@yahoo.com.br

⁶ Doutorando em Direito pela Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, Minas Gerais, Brasil.
E-mail: andregrisostomofernandes@gmail.com

2.338/2023, mas que tais iniciativas ainda apresentam fragilidades quanto à efetividade dos mecanismos de transparência, accountability, supervisão humana e controle democrático dos sistemas algorítmicos. Conclui-se que a proteção efetiva dos direitos humanos frente aos riscos da IA demanda o fortalecimento do sistema regulatório brasileiro, com a incorporação de salvaguardas jurídicas normativas e institucionais capazes de assegurar governança responsável, prevenção de danos e centralidade da pessoa humana no desenvolvimento e uso dessas tecnologias.

Palavras-chave: Inteligência Artificial. Direitos Humanos. Regulamentação da Inteligência Artificial no Brasil. Governança Algorítmica.

ABSTRACT

The growing incorporation of Artificial Intelligence (AI) into public and private decision-making processes has produced significant transformations in legal and institutional structures, raising new challenges for the protection of human rights. Although technological innovation offers substantial benefits, the use of opaque algorithmic systems based on large-scale data processing and automated decision-making reveals ethical, legal, and social risks capable of deepening structural inequalities and weakening fundamental guarantees, such as equality, non-discrimination, privacy, and due process. In this context, this article aims to analyze the process of AI regulation in Brazil from a human rights perspective, assessing its advances and limitations in light of international regulatory experiences. The research adopts a qualitative, theoretical-normative approach, using bibliographic and documentary methods grounded in national and international specialized literature, as well as in the analysis of Brazilian legal and institutional frameworks. The findings indicate that Brazil has relevant instruments for AI Governance such as the Brazilian Internet Bill of Rights, the General Data Protection Law (LGPD), the Brazilian Artificial Intelligence Strategy (EBIA), and Bill No. 2,338/2023, but that these initiatives still present weaknesses regarding the effectiveness of transparency, accountability, human oversight, and democratic control mechanisms over algorithmic systems. It is concluded that the effective protection of human rights in the face of AI-related risks requires the strengthening of the Brazilian regulatory system through the incorporation of normative and institutional legal safeguards capable of ensuring responsible governance, harm prevention, and the centrality of the human person in the development and use of these technologies.

Keywords: Artificial Intelligence. Human Rights. Regulation of Artificial Intelligence in Brazil. Algorithmic Governance.

RESUMEN

La creciente incorporación de la Inteligencia Artificial (IA) en los procesos de toma de decisiones, tanto públicos como privados, ha provocado transformaciones relevantes en las estructuras jurídicas e institucionales, lo que plantea nuevos retos para la protección de los derechos humanos. Aunque la innovación tecnológica supone avances significativos, el uso de sistemas algorítmicos opacos, basados en grandes volúmenes de datos y en decisiones automatizadas, pone de manifiesto riesgos éticos, jurídicos y sociales capaces de agravar las desigualdades estructurales y debilitar garantías fundamentales, como la igualdad, la no discriminación, la privacidad y el debido proceso. Ante este panorama, el presente artículo tiene por objeto analizar el proceso de regulación de la IA en Brasil desde la perspectiva de la protección de los derechos

humanos, evaluando sus avances y limitaciones a la luz de las experiencias reguladoras internacionales. La investigación adopta un enfoque cualitativo, de carácter teórico-normativo, con procedimientos técnicos bibliográficos y documentales, basándose en la literatura especializada nacional e internacional, así como en el análisis de los marcos normativos e institucionales brasileños. Los resultados indican que Brasil dispone de instrumentos relevantes para la gobernanza de la IA, como el Marco Civil de Internet, la Ley General de Protección de Datos Personales (LGPD), la Estrategia Brasileña de Inteligencia Artificial (EBIA) y el proyecto de ley n.º 2.338/2023, pero que dichas iniciativas aún presentan deficiencias en cuanto a la eficacia de los mecanismos de transparencia, rendición de cuentas, supervisión humana y control democrático de los sistemas algorítmicos. Se concluye que la protección efectiva de los derechos humanos frente a los riesgos de la IA exige el fortalecimiento del sistema regulador brasileño, con la incorporación de salvaguardias jurídicas, normativas e institucionales capaces de garantizar una gobernanza responsable, la prevención de daños y la centralidad de la persona humana en el desarrollo y el uso de estas tecnologías.

Palabras clave: Inteligencia Artificial. Derechos Humanos. Regulación de la Inteligencia Artificial en Brasil. Gobernanza Algorítmica.



Esta obra está bajo una [Licencia Creative Commons Atribución- NoComercial 4.0 Internacional](https://creativecommons.org/licenses/by-nc/4.0/)

INTRODUÇÃO

A IA tem sido progressivamente incorporada às dinâmicas sociais, econômicas e institucionais contemporâneas como tecnologia capaz de ampliar a eficiência decisória, impulsionar a inovação e oferecer soluções para problemas complexos em múltiplos setores. Sustentada pelo processamento massivo de dados, por modelos algorítmicos cada vez mais sofisticados e por sistemas de aprendizagem automática, a IA passou a ocupar posição central na organização da vida social, influenciando desde políticas públicas até relações de consumo, trabalho e acesso a serviços essenciais. Nesse cenário, consolidou-se a narrativa de que a automação inteligente poderia contribuir para o desenvolvimento social e para a racionalização das estruturas estatais e privadas.

Entretanto, a expansão acelerada do uso de sistemas de IA também revelou riscos relevantes para a proteção dos direitos humanos. A opacidade algorítmica, a reprodução de vieses estruturais, a assimetria informacional e a concentração de poder tecnológico em grandes atores econômicos evidenciam que a inovação tecnológica não é neutra, tampouco automaticamente compatível com valores democráticos. A utilização intensiva de dados pessoais e comportamentais como insumo central desses sistemas tensiona princípios fundamentais como a

dignidade da pessoa humana, a igualdade, a não discriminação, a privacidade e a autonomia individual, exigindo respostas jurídicas capazes de lidar com tecnologias que operam de forma autônoma, preditiva e em larga escala.

No contexto brasileiro, tais desafios assumem contornos específicos, embora o país disponha de uma estrutura normativa relevante para a proteção de direitos no ambiente digital, especialmente o Marco Civil da Internet (Lei nº 12.965/2014) e a Lei Geral de Proteção de Dados Pessoais (Lei nº 13.709/2018), a crescente adoção de sistemas de IA evidencia limites desses instrumentos diante das particularidades técnicas e dos riscos próprios da IA. A formulação da Estratégia Brasileira de IA e a tramitação do Projeto de Lei nº 2.338/2023 demonstram o reconhecimento institucional da necessidade de uma regulação específica, ao mesmo tempo em que revelam disputas conceituais, políticas e jurídicas sobre os modelos de governança mais adequados para compatibilizar inovação tecnológica, desenvolvimento econômico e tutela de direitos fundamentais.

É nesse cenário que se insere a presente pesquisa, orientada pela seguinte questão de pesquisa: o atual processo de regulamentação da IA no Brasil é capaz de assegurar a proteção efetiva dos direitos humanos diante dos riscos associados aos sistemas algorítmicos? Parte-se da hipótese de que, embora o Brasil tenha avançado na formulação de princípios e diretrizes gerais para o uso responsável da IA, o modelo regulatório em desenvolvimento ainda apresenta lacunas relevantes quanto à governança de riscos, à responsabilização e à incorporação sistemática de parâmetros de direitos humanos, o que compromete sua capacidade de oferecer proteção efetiva.

O objetivo geral da pesquisa consiste em analisar se o sistema regulatório brasileiro para a IA é adequado para assegurar a proteção efetiva dos direitos humanos diante dos riscos associados aos sistemas algorítmicos. Para alcançar esse propósito, estabelecem-se como objetivos específicos: (i) examinar os principais riscos éticos, jurídicos e sociais associados à IA e suas implicações para os direitos humanos e (ii) analisar a trajetória de construção de um sistema brasileiro de regulamentação da inteligência artificial, considerando seus avanços e limitações.

A relevância da pesquisa decorre da necessidade de reposicionar o debate jurídico sobre IA em um cenário marcado pela centralidade dos dados e dos sistemas algorítmicos nas dinâmicas sociais contemporâneas. Em um país atravessado por profundas desigualdades sociais e informacionais, a ausência de uma regulação sensível às vulnerabilidades estruturais pode aprofundar exclusões históricas e comprometer garantias constitucionais. Assim, a proteção dos

direitos humanos deve ser compreendida como eixo estruturante da governança da IA, e não como elemento acessório ou meramente declaratório.

REFERENCIAL TEÓRICO

IA, Riscos Estruturais e Proteção dos Direitos Humanos

A incorporação progressiva de sistemas de IA em processos decisórios socialmente relevantes tem suscitado preocupações quanto à capacidade do direito de responder adequadamente aos riscos associados a tecnologias que operam de forma autônoma, adaptativa e em larga escala. Scherer (2016) sustenta que a regulação da IA deve ser orientada por uma abordagem baseada em riscos, capaz de identificar desafios específicos relacionados à previsibilidade, à responsabilização e às competências institucionais necessárias para lidar com danos complexos. Segundo o autor, a ausência de estratégias regulatórias adequadas compromete o controle jurídico efetivo sobre sistemas que influenciam decisões com impacto direto sobre direitos e interesses individuais e coletivos.

No plano técnico-estrutural, os riscos da IA se intensificam com o desenvolvimento e a disseminação de modelos de larga escala. Bommasani *et al.* (2021) analisam os chamados foundation models e destacam que sua principal característica reside na reutilização do mesmo modelo em múltiplas tarefas e contextos. Essa característica amplia o alcance social dos sistemas de IA e transforma eventuais limitações presentes nos dados de treinamento ou na arquitetura do modelo em riscos estruturais, uma vez que falhas podem ser replicadas e amplificadas em diferentes aplicações, produzindo efeitos jurídicos e sociais de difícil contenção.

Um dos riscos mais recorrentes associados à IA diz respeito à opacidade algorítmica. Sartor (2022) observa que sistemas baseados em técnicas avançadas de aprendizagem automática frequentemente produzem resultados cujo processo decisório é de difícil compreensão, o que contrasta com as exigências de fundamentação e inteligibilidade próprias do direito. A introdução de decisões algorítmicas em contextos jurídicos ou administrativos, nesse sentido, desafia a racionalidade jurídica ao limitar a possibilidade de justificar decisões de forma transparente e de submetê-las ao escrutínio público e institucional (Sartor, 2022).

Essa dificuldade torna-se ainda mais evidente quando analisada à luz da teoria do raciocínio jurídico. Sartor (2009) desenvolve o conceito de defeasibility para demonstrar que a

aplicação do direito pressupõe a possibilidade de revisão de conclusões diante de novas razões, exceções ou circunstâncias específicas do caso concreto. Sistemas algorítmicos que operam a partir de correlações estatísticas rígidas tendem a reduzir essa flexibilidade interpretativa, o que pode comprometer garantias como o devido processo legal e o direito à contestação, especialmente quando decisões automatizadas produzem efeitos jurídicos relevantes.

Além da opacidade, a reprodução de desigualdades sociais por meio de vieses algorítmicos constitui risco significativo. Kaplan (2016) destaca que sistemas de IA aprendem padrões a partir de dados produzidos em contextos sociais marcados por assimetrias econômicas, raciais e de gênero. Quando esses dados são utilizados sem critérios normativos de correção, os algoritmos tendem a reproduzir e naturalizar discriminações existentes, sobretudo quando aplicados em áreas sensíveis como crédito, trabalho, segurança pública ou políticas sociais, afetando diretamente os princípios da igualdade e da não discriminação (Kaplan, 2016).

A centralidade dos direitos humanos na governança da IA é abordada de forma explícita por Cath *et al.* (2018), que relacionam o desenvolvimento da IA à ideia de uma “boa sociedade”. Os autores sustentam que sistemas de IA devem ser avaliados não apenas por critérios de eficiência ou inovação, mas por sua compatibilidade com valores democráticos e direitos fundamentais. Para Cath *et al.*, a ausência de um enquadramento normativo orientado por direitos humanos tende a deslocar escolhas políticas relevantes para esferas técnicas e privadas, enfraquecendo mecanismos de deliberação pública, participação social e accountability.

Essa perspectiva normativa é aprofundada no contexto europeu por Hildebrandt (2020), que analisa a relação entre IA, direitos fundamentais e Estado de Direito. A autora sustenta que a automação de decisões por sistemas algorítmicos compromete pilares centrais das democracias constitucionais, como a previsibilidade normativa, a possibilidade de contestação e o controle jurídico efetivo. Para Hildebrandt, o uso de sistemas de IA em processos decisórios públicos e privados afeta diretamente direitos fundamentais como privacidade, igualdade, não discriminação e devido processo legal, exigindo salvaguardas jurídicas robustas capazes de assegurar transparência, responsabilidade e supervisão humana.

No campo ético, Floridi (2018) propõe a noção de soft ethics como instrumento complementar à regulação jurídica, destinado a orientar o desenvolvimento e o uso de tecnologias digitais a partir de valores compartilhados. O autor reconhece que a ética pode oferecer diretrizes relevantes em contextos nos quais o direito positivo ainda é insuficiente, mas ressalta que essa abordagem não substitui a necessidade de normas juridicamente vinculantes, especialmente

quando estão em jogo direitos fundamentais e assimetrias estruturais de poder.

A discussão ética ganha contornos ainda mais complexos quando relacionada à soberania e à segurança. Timmers (2019) analisa a ética da IA em contextos nos quais a soberania estatal se encontra em risco, destacando que sistemas algorítmicos podem impactar infraestruturas críticas e decisões estratégicas. Para o autor, a proteção de direitos humanos no ambiente digital depende da existência de arranjos institucionais capazes de garantir controle, responsabilidade e segurança jurídica diante de tecnologias que operam em escala global e frequentemente fora das fronteiras regulatórias nacionais.

Em síntese, observa-se que os riscos éticos, jurídicos e sociais associados à IA não se limitam a falhas técnicas pontuais, mas assumem caráter estrutural quando sistemas operam de forma opaca, em larga escala e sobre bases de dados assimétricas. Esses riscos incidem diretamente sobre a proteção dos direitos humanos, ao comprometer princípios como dignidade, igualdade, não discriminação, devido processo legal e controle democrático. A centralidade dos direitos fundamentais como parâmetro de governança, conforme destacado pela literatura analisada, revela a necessidade de modelos regulatórios capazes de equilibrar inovação tecnológica e proteção jurídica, evitando que a IA se desenvolva à margem do Estado de Direito (Hildebrandt, 2020).

REGULAMENTAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO BRASIL: ANÁLISE JURÍDICO-INSTITUCIONAL DOS AVANÇOS E LIMITAÇÕES

A análise do modelo brasileiro de regulamentação da IA deve partir do reconhecimento de que o país construiu, ao longo da última década, um sistema normativo relevante para a proteção de direitos no ambiente digital, ainda que esse arcabouço não tenha sido concebido originalmente para lidar de forma específica com tecnologias baseadas em IA. O Marco Civil da Internet, instituído pela Lei nº 12.965/2014, estabelece princípios, garantias e direitos fundamentais para o uso da internet no Brasil, como a proteção da privacidade, a liberdade de expressão e a preservação da neutralidade de rede (Brasil, 2014). Embora não trate diretamente de sistemas de IA, o Marco Civil inaugura uma lógica normativa orientada por direitos fundamentais e pela responsabilização dos agentes que atuam no ambiente digital, constituindo uma base relevante para modelos de governança tecnológica comprometidos com valores democráticos.

Em perspectiva complementar, a Lei Geral de Proteção de Dados Pessoais, Lei nº 13.709/2018, representa um avanço significativo na tutela dos direitos fundamentais no contexto do tratamento de dados pessoais, elemento central para o funcionamento de sistemas de IA (Brasil, 2018). A LGPD estabelece princípios como finalidade, adequação, necessidade, transparência e não discriminação, além de assegurar direitos aos titulares de dados, como acesso, correção e a possibilidade de revisão de decisões automatizadas. Ao reconhecer o impacto do tratamento automatizado sobre a esfera jurídica dos indivíduos, a Lei aproxima-se de uma abordagem orientada por direitos humanos, ainda que enfrente desafios relevantes quanto à sua efetividade diante de sistemas algorítmicos complexos e opacos.

Nesse contexto normativo, a Estratégia Brasileira de Inteligência Artificial (EBIA), elaborada pelo Ministério da Ciência, Tecnologia e Inovações e publicada em 2021, constitui o primeiro documento oficial voltado especificamente à formulação de diretrizes nacionais para o desenvolvimento e o uso da IA no país (Brasil, 2021). A EBIA reconhece a importância de princípios éticos, da proteção de dados e da promoção de uma IA confiável e alinhada ao interesse público. Contudo, por se tratar de um instrumento de política pública sem força normativa vinculante, apresenta limitações quanto à imposição de deveres jurídicos concretos e à criação de mecanismos efetivos de responsabilização, o que fragiliza sua capacidade de assegurar, de forma robusta, a proteção de direitos humanos no uso de sistemas de IA.

No plano legislativo, o Projeto de Lei nº 2.338/2023 representa o esforço mais abrangente de regulamentação da IA no Brasil até o momento. O texto propõe um modelo regulatório baseado em riscos, prevendo categorias de sistemas de IA, deveres diferenciados para desenvolvedores e usuários e a incorporação de princípios como transparência, segurança e não discriminação (Brasil, 2023). Embora o projeto sinalize aproximação com modelos internacionais, especialmente europeus, ainda suscita debates quanto à suficiência de seus mecanismos de governança, à clareza das responsabilidades atribuídas aos diversos atores e à efetividade das salvaguardas voltadas à proteção dos direitos fundamentais e ao controle democrático.

A elaboração do Projeto de Lei nº 2.338/2023 foi precedida pelos trabalhos da Comissão de Juristas responsável por subsidiar a elaboração de substitutivo ao Projeto de Lei sobre Inteligência Artificial no Brasil (CJSUBIA). Os relatórios e debates da CJSUBIA evidenciam preocupação com a incorporação de princípios éticos, a mitigação de riscos e a compatibilização entre inovação tecnológica e proteção de direitos fundamentais (Brasil, 2022). Ao mesmo tempo,

revelam tensões entre interesses econômicos, estratégias de desenvolvimento tecnológico e a necessidade de salvaguardas jurídicas robustas, refletindo as dificuldades inerentes à construção de um marco regulatório em um país marcado por profundas desigualdades sociais e institucionais (Brasil, 2022).

Parentoni, Valentini e Alves (2020) destacam que, historicamente, a regulação da IA no Brasil caracterizou-se pela fragmentação normativa e pela ausência de um marco legal específico. Os autores apontam que iniciativas legislativas anteriores revelavam preocupações relevantes, mas careciam de sistematicidade e de uma abordagem integrada orientada pela proteção de direitos fundamentais, reforçando a necessidade de avaliação crítica das propostas mais recentes.

A perspectiva comparada oferece parâmetros adicionais para essa análise. Polido (2020a) examina a IA no contexto da corrida regulatória global, enfatizando que estratégias nacionais não podem ser compreendidas de forma isolada, mas devem ser situadas em dinâmicas transnacionais e disputas por poder normativo. Segundo o autor, a ausência de coordenação regulatória e de referenciais claros de proteção de direitos pode aprofundar assimetrias globais e comprometer a efetividade de modelos domésticos de governança, especialmente em países periféricos. Em complemento, Polido (2020b) destaca a necessidade de articular políticas domésticas e processos legais transnacionais, sustentando que a proteção de direitos fundamentais exige mecanismos que ultrapassem soluções exclusivamente nacionais e dialoguem com padrões internacionais.

A análise crítica da IA no contexto brasileiro também se estende às tecnologias generativas. Rosetti e Garcia (2023) examinam questões jurídicas e éticas relacionadas ao uso de sistemas de IA generativa, destacando riscos associados à produção automatizada de conteúdos, à desinformação e à violação de direitos autorais e da personalidade. As autoras apontam que a ausência de regras claras para esse tipo de tecnologia pode gerar impactos significativos sobre direitos fundamentais, reforçando a urgência de um marco regulatório que vá além de princípios genéricos e estabeleça deveres concretos de responsabilidade e controle.

Quando avaliados à luz de critérios normativos centrados em direitos humanos e Estado de Direito, como aqueles propostos por Cath *et al.* (2018) e Hildebrandt (2020), os avanços do modelo brasileiro mostram-se relevantes, mas ainda insuficientes. A incorporação de princípios e a preocupação declarada com a proteção de direitos representam passos importantes, porém persistem lacunas quanto à efetividade dos mecanismos de accountability, à clareza das responsabilidades e à capacidade institucional de fiscalização e controle. Nesse sentido, a regulamentação da IA no Brasil permanece em processo de consolidação, exigindo

aprofundamento normativo e institucional para assegurar que a inovação tecnológica não se desenvolva à margem da proteção dos direitos humanos e dos valores democráticos.

METODOLOGIA

Do ponto de vista metodológico, a pesquisa adota uma abordagem qualitativa, de natureza teórico-normativa, adequada à análise crítica de fenômenos jurídicos e institucionais complexos, como a regulamentação da IA e seus impactos sobre os direitos humanos. Conforme assinala Gil (2010), a pesquisa qualitativa é especialmente apropriada quando se busca compreender significados, fundamentos normativos e relações estruturais, em vez de mensurar variáveis ou produzir generalizações estatísticas. Trata-se, portanto, de um estudo voltado à interpretação e à problematização de marcos normativos, discursos institucionais e construções teóricas, com ênfase na dimensão jurídico-constitucional da governança da IA.

Os procedimentos técnicos utilizados consistem em pesquisa bibliográfica e documental. A pesquisa bibliográfica compreendeu o exame sistemático de obras acadêmicas nacionais e internacionais sobre IA, ética, governança algorítmica, direitos fundamentais e regulação tecnológica, selecionadas a partir de critérios de relevância teórica, reconhecimento acadêmico e aderência ao objeto de estudo, conforme orienta Gil (2010).

A pesquisa documental concentrou-se na análise de documentos normativos e institucionais relevantes para a regulamentação da IA, especialmente no contexto brasileiro. Foram examinadas leis, projeto de lei, estratégias nacionais, relatórios de comissões especializadas e documentos oficiais, considerados fontes primárias para a compreensão das opções regulatórias adotadas, de seus fundamentos e de suas limitações. De acordo com Gil (2010), esse tipo de fonte permite acessar diretamente os registros formais das decisões políticas e jurídicas, sendo essencial para estudos de caráter normativo e institucional.

O processo de análise seguiu uma leitura interpretativa e sistemática do material selecionado, com o objetivo de identificar categorias temáticas relacionadas aos riscos éticos, jurídicos e sociais da IA e à capacidade dos modelos regulatórios de proteger os direitos humanos. As informações foram organizadas em eixos analíticos interdependentes, voltados à compreensão (a) dos riscos estruturais da IA e suas implicações para direitos fundamentais; (b) das respostas normativas e institucionais construídas no ordenamento jurídico brasileiro; e (c) do diálogo comparado com modelos internacionais de governança da IA. Essa estratégia analítica favoreceu

uma leitura transversal do problema, articulando fundamentos teóricos, marcos normativos e critérios de avaliação orientados pelo Estado de Direito e pela proteção dos direitos humanos.

RESULTADOS E DISCUSSÕES

A análise desenvolvida ao longo deste estudo permite identificar, como resultado central, que o modelo brasileiro de regulamentação da IA apresenta avanços normativos, mas ainda não é plenamente capaz de oferecer proteção efetiva e sistemática aos direitos humanos diante dos riscos éticos, jurídicos e sociais associados à IA. Os instrumentos atualmente existentes revelam um esforço progressivo de incorporação de princípios orientados por direitos fundamentais, porém carecem de densidade normativa, institucional e operacional suficiente para enfrentar os desafios estruturais impostos por sistemas algorítmicos complexos e de larga escala.

No plano dos resultados normativos, observa-se que o Marco Civil da Internet e a LGPD constituem pilares importantes para a proteção de direitos no ambiente digital, especialmente no que se refere à privacidade, à liberdade de expressão e à proteção de dados pessoais (Brasil, 2014; Brasil, 2018). Esses diplomas oferecem fundamentos para a governança da IA, sobretudo ao reconhecerem limites ao tratamento automatizado de informações e ao estabelecerem deveres de transparência e responsabilização. Contudo, os resultados da análise indicam que tais normas foram concebidas para um ecossistema digital mais amplo e não respondem de forma específica às particularidades técnicas, decisórias e sistêmicas da IA.

A EBIA, por sua vez, evidencia um reconhecimento institucional da necessidade de orientar o desenvolvimento e o uso da IA a partir de princípios éticos e do interesse público (Brasil, 2021). Todavia, como resultado da análise crítica, constata-se que a EBIA possui caráter predominantemente programático, sem força normativa vinculante, o que limita sua capacidade de impor obrigações jurídicas concretas ou de estruturar mecanismos efetivos de accountability. Essa fragilidade compromete sua aptidão para garantir, de forma robusta, a proteção dos direitos humanos diante de riscos associados à automação decisória e à opacidade algorítmica.

O Projeto de Lei nº 2.338/2023 representa, nesse cenário, o principal esforço legislativo voltado especificamente à regulamentação da IA no Brasil. A análise do texto revela avanços importantes, como a adoção de uma abordagem baseada em riscos, a previsão de deveres diferenciados para os agentes envolvidos e incorporação de princípios como transparência, segurança e não discriminação (Brasil, 2023). Esses elementos aproximam o modelo brasileiro

de experiências regulatórias internacionais, especialmente do paradigma europeu. No entanto, os resultados indicam que o projeto ainda enfrenta desafios quanto à clareza das responsabilidades atribuídas aos diferentes atores, à definição dos mecanismos de fiscalização e à efetividade das sanções, o que pode comprometer sua capacidade de oferecer proteção consistente aos direitos humanos.

Os trabalhos da CJSUBIA responsável por subsidiar a elaboração de substitutivo sobre IA no Brasil revelam que o processo legislativo foi precedido por debates qualificados sobre ética, mitigação de riscos e proteção de direitos fundamentais (Brasil, 2022). Como resultado da análise desses documentos, observa-se uma preocupação explícita com a compatibilização entre inovação tecnológica e valores democráticos. Entretanto, também se evidenciam tensões entre interesses econômicos e a necessidade de salvaguardas jurídicas robustas, refletindo limites estruturais do modelo regulatório em construção (Brasil, 2022).

A literatura analisada contribui para a interpretação desses resultados, Parentoni, Valentini e Alves (2020) apontam que a regulação da IA no Brasil historicamente se caracterizou por fragmentação normativa e ausência de sistematicidade, o que ajuda a explicar as lacunas identificadas nos instrumentos atuais. Rosetti e Garcia (2023), ao analisarem a IA generativa, reforçam esse diagnóstico ao demonstrar que tecnologias emergentes ampliam riscos relacionados à desinformação, à violação de direitos da personalidade e à responsabilização difusa, exigindo respostas normativas mais precisas e específicas.

Quando os resultados obtidos são confrontados com experiências regulatórias internacionais, especialmente europeias, tornam-se ainda mais evidentes as limitações do modelo brasileiro. Polido (2020a) demonstra que a regulação da IA ocorre em um contexto de corrida regulatória global, no qual Estados disputam capacidade normativa e poder de influência. Segundo o autor, modelos domésticos que não dialogam de forma consistente com padrões internacionais tendem a aprofundar assimetrias e a comprometer a efetividade da proteção de direitos, sobretudo em países periféricos. Essa análise é reforçada por Polido (2020b), que sustenta a necessidade de articulação entre políticas domésticas e processos legais transnacionais para enfrentar a natureza global das tecnologias digitais.

À luz desses parâmetros comparativos, observa-se que o modelo brasileiro ainda carece de mecanismos equivalentes aos adotados na experiência europeia, especialmente no que se refere à exigência de avaliações de impacto, à supervisão institucional especializada e à garantia de contestabilidade das decisões algorítmicas. Cath *et al.*, (2018) destacam que a governança da

IA deve ser orientada por valores democráticos e direitos humanos, evitando a transferência de escolhas políticas para esferas técnicas e privadas. Os resultados da análise indicam que, no Brasil, esse deslocamento ainda não foi plenamente enfrentado pelo marco regulatório em construção.

Essa constatação é aprofundada quando se adota como critério normativo a perspectiva de Hildebrandt (2020). Sustenta que a proteção dos direitos fundamentais no contexto da IA exige salvaguardas jurídicas capazes de assegurar previsibilidade normativa, contestabilidade e controle efetivo das decisões automatizadas. Sob esse prisma, os resultados do estudo indicam que o modelo brasileiro ainda não dispõe de instrumentos suficientemente claros e operacionais para garantir tais requisitos, especialmente em contextos de uso da IA no setor público e em decisões que afetam diretamente a esfera jurídica dos indivíduos (Hildebrandt, 2020).

Do ponto de vista ético e institucional, Floridi (2018) e Timmers (2019) contribuem para a interpretação dos resultados ao evidenciarem que diretrizes éticas, embora relevantes, não são suficientes para proteger direitos humanos em contextos marcados por assimetrias de poder, soberania e segurança. A análise demonstra que, no Brasil, a ênfase em princípios e recomendações éticas ainda não foi acompanhada por uma arquitetura regulatória capaz de impor limites claros e mecanismos eficazes de responsabilização.

Diante desses resultados, a hipótese formulada na pesquisa é parcialmente confirmada. O modelo brasileiro de regulamentação da IA apresenta elementos normativos capazes de contribuir para a proteção dos direitos humanos, mas ainda não oferece garantias suficientes para enfrentar, de forma abrangente, os riscos éticos, jurídicos e sociais associados à IA. A ausência de mecanismos robustos de accountability, fiscalização e controle institucional limita a efetividade das normas existentes e das propostas em tramitação.

Como limitação do estudo, destaca-se o caráter teórico-normativo da análise, que não abrangeu a aplicação empírica dos instrumentos regulatórios em casos concretos. Pesquisas futuras podem aprofundar essa investigação a partir da análise de decisões administrativas e judiciais envolvendo sistemas de IA, bem como da comparação empírica entre a experiência brasileira e modelos internacionais já em implementação. Ainda assim, os resultados apresentados contribuem para o avanço do debate ao evidenciar que a efetiva proteção dos direitos humanos no contexto da IA depende não apenas da incorporação de princípios, mas da construção de uma governança jurídica sólida, alinhada ao Estado de Direito e capaz de equilibrar inovação tecnológica e justiça social.

CONCLUSÃO

A análise desenvolvida ao longo deste estudo demonstra que a regulamentação da IA no Brasil encontra-se em um estágio de construção incipiente, marcado por avanços normativos relevantes, mas ainda insuficientes para assegurar proteção efetiva e abrangente aos direitos humanos diante dos riscos éticos, jurídicos e sociais associados à IA. O exame dos principais marcos normativos e institucionais como o Marco Civil da Internet, a LGPD, a EBIA e o Projeto de Lei nº 2.338/2023, evidencia a incorporação gradual de princípios orientados por direitos fundamentais, como privacidade, transparência e não discriminação. Contudo, os resultados indicam que tais instrumentos ainda carecem de densidade normativa e de mecanismos institucionais robustos capazes de enfrentar a complexidade, a opacidade e o caráter sistêmico dos sistemas algorítmicos contemporâneos.

Sob a perspectiva comparada, especialmente à luz das experiências regulatórias europeias e dos referenciais teóricos centrados em direitos humanos e Estado de Direito, observa-se que o modelo brasileiro apresenta lacunas significativas quanto à efetividade de seus mecanismos de governança. A ausência de avaliações de impacto amplamente estruturadas, de instâncias especializadas de supervisão e de garantias claras de contestabilidade das decisões automatizadas limita a capacidade do ordenamento jurídico de exercer controle democrático sobre a IA. Assim, embora o discurso normativo brasileiro reconheça a centralidade dos direitos humanos, persiste um descompasso entre a formulação principiológica e a implementação de salvaguardas jurídicas capazes de prevenir violações e assegurar accountability em contextos decisórios mediados por algoritmos.

Conclui-se, portanto, que a proteção dos direitos humanos frente aos riscos da IA exige mais do que a incorporação de diretrizes éticas ou recomendações programáticas. Impõe-se a construção de uma governança jurídica sólida, articulada a padrões internacionais, dotada de mecanismos normativos vinculantes, fiscalização institucional eficaz e participação democrática. Somente a partir desse fortalecimento regulatório será possível equilibrar inovação tecnológica e justiça social, evitando que a IA se desenvolva à margem do Estado de Direito. O desafio colocado ao Brasil consiste em transformar princípios em garantias efetivas, assegurando que a IA não se torne instrumento para o agravamento das desigualdades e fragilização dos direitos fundamentais.

REFERÊNCIAS

BOMMASANI, Rishi *et al.* **On the opportunities and risks of foundation models.** Stanford: Center for Research on Foundation Models, 2021. Disponível em: <https://crfm.stanford.edu/report.html> . Acesso em: 11 dez. 2024.

BRASIL. **Lei nº 12.965**, de 23 de abril de 2014. Estabelece princípios, garantias, direitos e deveres para o uso da Internet no Brasil. Brasília, DF: Presidência da República, 2014. Disponível em: http://www.planalto.gov.br/ccivil_03/_ato2011-2014/2014/lei/112965.htm . Acesso em: 16 nov. 2024.

BRASIL. **Lei nº 13.709**, de 14 de agosto de 2018. Lei Geral de Proteção de Dados Pessoais. Brasília, DF: Presidência da República, 2018. Disponível em: http://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/L13709.htm . Acesso em: 16 nov. 2024.

BRASIL. **Ministério da Ciência, Tecnologia e Inovações. Estratégia Brasileira de IA.** Brasília, 2021. Disponível em: https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/transformacaodigital/arquivosinteligenciaartificial/ebia-documento_referencia_4-979_2021.pdf . Acesso em: 04 dez. 2025.

BRASIL. Projeto de **Lei nº 2.338**, de 2023. Dispõe sobre o uso da IA. Brasília, DF: Senado Federal, 2023. Disponível em: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233> . Acesso em: 02 jan. 2025.

BRASIL. **Senado Federal.** Comissão de Juristas responsável por subsidiar a elaboração de substitutivo sobre Inteligência Artificial no Brasil (CJSUBIA). Relatório Final. Brasília, DF, 2022. Disponível em: <https://legis.senado.leg.br/atividade/comissoes/comissao/2504/> . Acesso em: 16 nov. 2024.

CATH, Corinne *et al.* Artificial intelligence and the ‘Good Society’: the US, EU, and UK approach. **Science and Engineering Ethics**, v. 24, n. 2, p. 505–528, 2018. Disponível em: https://www.researchgate.net/publication/315705213_Artificial_Intelligence_and_the_'Good_Society'_the_US_EU_and_UK_approach . Acesso em: 16 nov. 2024.

FLORIDI, Luciano. Soft ethics and the governance of the digital. **Philosophy & Technology**, v. 31, p. 1–8, 2018. Disponível em: <https://link.springer.com/article/10.1007/s13347-018-0303-9> . Acesso em: 11 dez. 2024.

GIL, Antonio Carlos. **Como elaborar projetos de pesquisa.** 6. ed. São Paulo: Atlas, 2010.

HILDEBRANDT, Mireille. **Law for computer scientists and other folk.** Oxford: Oxford University Press, 2020. Disponível em: <https://academic.oup.com/book/40128> . Acesso em: 12 jan. 2025.

KAPLAN, Jerry. **Artificial intelligence: what everyone needs to know.** New York: Oxford University Press, 2016. Disponível em: <https://archive.org/details/artificialintell0000kapl> . Acesso em: 11 dez. 2024.

PARENTONI, Leonardo Netto; VALENTINI, Rômulo Soares; ALVES, Tárík César Oliveira e. Panorama da regulação da IA no Brasil: com ênfase no PLS n. 5.051/2019. **Revista Eletrônica do Curso de Direito da UFSM**, Santa Maria, v. 15, n. 2, e43730, 2020. Disponível em: <https://periodicos.ufsm.br/revistadireito/article/view/43730> . Acesso em: 07 jan. 2025.

POLIDO, Fabrício Bertini Pasquot. IA entre estratégias nacionais e a corrida regulatória global: rotas analíticas para uma releitura internacionalista e comparada. **Revista da Faculdade de Direito da UFMG**, Belo Horizonte, n. 76, p. 229–256, 2020a. Disponível em: <https://repositorio.ufmg.br/handle/1843/40727> . Acesso em: 02 jan. 2025.

POLIDO, Fabrício Bertini Pasquot. **Novas perspectivas para regulação da IA**: diálogos entre as políticas domésticas e os processos legais transnacionais. In: FRAZÃO, Ana; MULHOLLAND, Caitlin (orgs.). *IA e Direito: Ética, Regulação e Responsabilidade*. São Paulo: Thomson Reuters Brasil, 2020b.

ROSETTI, Regina; GARCIA, Kethly. **IA generativa**: questões jurídicas e éticas em torno do ChatGPT. *VirtuaJus*, Belo Horizonte, v. 8, n. 15, p. 253–264, 2023. Disponível em: <https://periodicos.pucminas.br/index.php/virtuajus/article/view/30769> . Acesso em: 03 jan. 2025.

SARTOR, Giovanni. **Defeasibility in legal reasoning**. European University Institute Working Papers, 2009. Disponível em: https://cadmus.eui.eu/bitstream/handle/1814/10768/LAW_2009_02.pdf . Acesso em: 02 jan. 2025.

SARTOR, Giovanni. **L'intelligenza artificiale e il diritto**. Bologna: Giappichelli, 2022. Disponível em: <https://www.giappichelli.it/media/catalog/product/openaccess/9788892177475.pdf> . Acesso em: 04 jan. 2025.

SCHERER, Matthew U. Regulating artificial intelligence systems: risks, challenges, competencies, and strategies. **Harvard Journal of Law & Technology**, v. 29, n. 2, p. 353–400, 2016. Disponível em: <https://jolt.law.harvard.edu/articles/pdf/v29/29HarvJLTech353.pdf> . Acesso em: 11 dez. 2024.

TIMMERS, Paul. Ethics of AI and cybersecurity when sovereignty is at stake. **Minds and Machines**, v. 29, n. 4, p. 635–645, 2019. Disponível em: <https://link.springer.com/article/10.1007/s11023-019-09508-4> . Acesso em: 15 out. 2023.