

**VI CONGRESSO INTERNACIONAL DE  
DIREITO E INTELIGÊNCIA  
ARTIFICIAL (VI CIDIA)**

**REGULAÇÃO DA INTELIGÊNCIA ARTIFICIAL I**

---

R344

Regulação da inteligência artificial I [Recurso eletrônico on-line] organização VI Congresso Internacional de Direito e Inteligência Artificial (VI CIDIA): Skema Business School – Belo Horizonte;

Coordenadores: Gabriel Oliveira de Aguiar Borges, Luiz Felipe de Freitas Cordeiro e Matheus Antes Schwede – Belo Horizonte: Skema Business School, 2025.

Inclui bibliografia

ISBN: 978-65-5274-357-2

Modo de acesso: [www.conpedi.org.br](http://www.conpedi.org.br) em publicações

Tema: Perspectivas globais para a regulação da inteligência artificial.

1. Compliance. 2. Ética. 3. Legislação. I. VI Congresso Internacional de Direito e Inteligência Artificial (1:2025 : Belo Horizonte, MG).

CDU: 34

---

**skema**  
BUSINESS SCHOOL

**LAW SCHOOL**  
FOR BUSINESS

# **VI CONGRESSO INTERNACIONAL DE DIREITO E INTELIGÊNCIA ARTIFICIAL (VI CIDIA)**

## **REGULAÇÃO DA INTELIGÊNCIA ARTIFICIAL I**

---

### **Apresentação**

A SKEMA Business School é uma organização francesa sem fins lucrativos, com presença em sete países diferentes ao redor do mundo (França, EUA, China, Brasil, Emirados Árabes Unidos, África do Sul e Canadá) e detentora de três prestigiadas creditações internacionais (AMBA, EQUIS e AACSB), refletindo seu compromisso com a pesquisa de alta qualidade na economia do conhecimento. A SKEMA reconhece que, em um mundo cada vez mais digital, é essencial adotar uma abordagem transdisciplinar.

Cumprindo esse propósito, o VI Congresso Internacional de Direito e Inteligência Artificial (VI CIDIA), realizado nos dias 18 e 19 de setembro de 2025, em formato híbrido, manteve-se como o principal evento acadêmico sediado no Brasil com o propósito de fomentar ricas discussões sobre as diversas interseções entre o direito e a inteligência artificial. O evento, que teve como tema central a "Regulação da Inteligência Artificial", contou com a presença de renomados especialistas nacionais e internacionais, que abordaram temas de relevância crescente no cenário jurídico contemporâneo.

Profissionais e estudantes dos cursos de Direito, Administração, Economia, Ciência de Dados, Ciência da Computação, entre outros, tiveram a oportunidade de se conectar e compartilhar conhecimentos, promovendo um ambiente de rica troca intelectual. O VI CIDIA contou com a participação de acadêmicos e profissionais provenientes de diversas regiões do Brasil e do exterior. Entre os estados brasileiros representados, estavam: Alagoas (AL), Bahia (BA), Ceará (CE), Goiás (GO), Maranhão (MA), Mato Grosso do Sul (MS), Minas Gerais (MG), Pará (PA), Paraíba (PB), Paraná (PR), Pernambuco (PE), Piauí (PI), Rio de Janeiro

Foram discutidos assuntos variados, desde a própria regulação da inteligência artificial, eixo central do evento, até as novas perspectivas de negócios e inovação, destacando como os algoritmos estão remodelando setores tradicionais e impulsionando a criação de empresas inovadoras. Com uma programação abrangente, o congresso proporcionou um espaço vital para discutir os desafios e oportunidades que emergem com o desenvolvimento algorítmico, reforçando a importância de uma abordagem jurídica e ética robusta nesse contexto em constante evolução.

A programação teve início às 13h, com o check-in dos participantes e o aquecimento do público presente. Às 13h30, a abertura oficial foi conduzida pela Prof.<sup>a</sup> Dr.<sup>a</sup> Geneviève Poulingue, que, em sua fala de boas-vindas, destacou a relevância do congresso para a agenda global de inovação e o papel da SKEMA Brasil como ponte entre a academia e o setor produtivo.

Em seguida, às 14h, ocorreu um dos momentos mais aguardados: a Keynote Lecture do Prof. Dr. Ryan Calo, renomado especialista internacional em direito e tecnologia e professor da University of Washington. Em uma conferência instigante, o professor explorou os desafios metodológicos da regulação da inteligência artificial, trazendo exemplos de sua atuação junto ao Senado dos Estados Unidos e ao Bundestag alemão.

A palestra foi seguida por uma sessão de comentários e análise crítica conduzida pelo Prof. Dr. José Luiz de Moura Faleiros Júnior, que contextualizou as reflexões de Calo para a realidade brasileira e fomentou o debate com o público. O primeiro dia foi encerrado às 14h50 com as considerações finais, deixando os participantes inspirados para as discussões do dia seguinte.

As atividades do segundo dia tiveram início cedo, com o check-in às 7h30. Às 8h20, a Prof.<sup>a</sup> Dr.<sup>a</sup> Margherita Pagani abriu a programação matinal com a conferência Unlocking Business

Após um breve e merecido coffee break às 9h40, os participantes retornaram para uma manhã de intensas reflexões. Às 10h30, o pesquisador Prof. Dr. Steve Ataky apresentou a conferência Regulatory Perspectives on AI, compartilhando avanços e desafios no campo da regulação técnica e ética da inteligência artificial a partir de uma perspectiva global.

Encerrando o ciclo de palestras, às 11h10, o Prof. Dr. Filipe Medon trouxe ao público uma análise profunda sobre o cenário brasileiro, com a palestra AI Regulation in Brazil. Sua exposição percorreu desde a criação do Marco Legal da Inteligência Artificial até os desafios atuais para sua implementação, envolvendo aspectos legislativos, econômicos e sociais.

Nas tardes dos dois dias, foram realizados grupos de trabalho que contaram com a apresentação de cerca de 60 trabalhos acadêmicos relacionados à temática do evento. Com isso, o evento foi encerrado, após intensas discussões e troca de ideias que estabeleceram um panorama abrangente das tendências e desafios da inteligência artificial em nível global.

Os GTs tiveram os seguintes eixos de discussão, sob coordenação de renomados especialistas nos respectivos campos de pesquisa:

a) Startups e Empreendedorismo de Base Tecnológica – Coordenado por Allan Fuezi de Moura Barbosa, Laurence Duarte Araújo Pereira, Cildo Giolo Júnior, Maria Cláudia Viana Hissa Dias do Vale Gangana e Yago Oliveira

b) Jurimetria Cibernética Jurídica e Ciência de Dados – Coordenado por Arthur Salles de Paula Moreira, Gabriel Ribeiro de Lima, Isabela Campos Vidigal Martins, João Victor Doreto e Tales Calaza

c) Decisões Automatizadas e Gestão Empresarial / Algoritmos, Modelos de Linguagem e Propriedade Intelectual – Coordenado por Alisson Jose Maia Melo, Guilherme Mucelin e

f) Regulação da Inteligência Artificial – III – Coordenado por Ana Júlia Silva Alves Guimarães, Erick Hitoshi Guimarães Makiya, Jessica Fernandes Rocha, João Alexandre Silva Alves Guimarães e Luiz Felipe Vieira de Siqueira

g) Inteligência Artificial, Mercados Globais e Contratos – Coordenado por Gustavo da Silva Melo, Rodrigo Gugliara e Vitor Ottoboni Pavan

h) Privacidade, Proteção de Dados Pessoais e Negócios Inovadores – I – Coordenado por Dineia Anziliero Dal Pizzol, Evaldo Osorio Hackmann, Gabriel Fraga Hamester, Guilherme Mucelin e Guilherme Spillari Costa

i) Privacidade, Proteção de Dados Pessoais e Negócios Inovadores – II – Coordenado por Alexandre Schmitt da Silva Mello, Lorenzo Antonini Itabaiana, Marcelo Fonseca Santos, Mariana de Moraes Palmeira e Pietra Daneluzzi Quinelato

j) Empresa, Tecnologia e Sustentabilidade – Coordenado por Marcia Andrea Bühring, Ana Cláudia Redecker, Jessica Mello Tahim e Maraluce Maria Custódio.

Cada GT proporcionou um espaço de diálogo e troca de experiências entre pesquisadores e profissionais, contribuindo para o avanço das discussões sobre a aplicação da inteligência artificial no direito e em outros campos relacionados.

Um sucesso desse porte não seria possível sem o apoio institucional do Conselho Nacional de Pesquisa e Pós-graduação em Direito - CONPEDI, que desde a primeira edição do evento provê uma parceria sólida e indispensável ao seu sucesso. A colaboração contínua do CONPEDI tem sido fundamental para a organização e realização deste congresso, assegurando a qualidade e a relevância dos debates promovidos.

Reitora – SKEMA Business School - Campus Belo Horizonte

Prof. Ms. Dorival Guimarães Pereira Júnior

Coordenador do Curso de Direito – SKEMA Law School

Prof. Dr. José Luiz de Moura Faleiros Júnior

Coordenador de Pesquisa – SKEMA Law School

## **VIESES ALGORÍTMICOS, DADOS E DIREITO: RUMO A UM MARCO LEGAL DA INTELIGÊNCIA ARTIFICIAL NO BRASIL**

### **ALGORITHMIC BIAS, DATA AND LAW: TOWARD AN ARTIFICIAL INTELLIGENCE LEGAL FRAMEWORK IN BRAZIL**

**Marina Silva Oliveira**

#### **Resumo**

O artigo analisa como os modelos tradicionais de responsabilidade civil enfrentam a crescente autonomia da inteligência artificial (IA). Após traçar uma breve linha do tempo tecnológica, examina se o art.186 e o art.927 do Código Civil brasileiro, contrapondo responsabilidade subjetiva e objetiva diante de sistemas autônomos. A pesquisa — fundamentada em revisão bibliográfica, estudo de casos e comparação legislativa — sustenta que a imputação de danos deve ser distribuída entre desenvolvedores, operadores e empresas, preservando a reparação integral e incentivando soluções éticas. Discute se ainda a proposta europeia de “personalidade eletrônica” para robôs, ponderando riscos de diluição da responsabilidade humana. O trabalho aprofunda a questão dos vieses algorítmicos, ilustrada por casos como Amazon e Cambridge Analytica, e descreve marcos regulatórios recentes, do GDPR ao PL2338/2023 no Brasil. Conclui-se que a adoção de responsabilidade objetiva solidária, alinhada a padrões de governança, é essencial para equilibrar inovação tecnológica, proteção de direitos fundamentais e segurança jurídica.

**Palavras-chave:** Inteligência artificial, Responsabilidade civil, Vieses algorítmicos, Regulação, Ética

#### **Abstract/Resumen/Résumé**

This article investigates how traditional civil liability models cope with the growing autonomy of artificial intelligence (AI). After outlining a technological timeline, it contrasts subjective and strict liability under Brazilian Civil Code articles 186 and 927 when autonomous systems are involved. Based on literature review, case studies and comparative



**Keywords/Palabras-claves/Mots-clés:** Artificialintelligence, Civilliability, Algorithmicbias, Regulation, Ethics

## 1. INTRODUÇÃO

Com o avanço acelerado da Inteligência Artificial e a sua súbita inserção em diversas áreas da vida humana, tais como saúde, segurança, mercado profissional, entretenimento, educação, dentre outras, questões complexas referentes à sua regulamentação e, mais especificamente, à atribuição da responsabilidade civil por eventuais ilegalidades cometidas por sistemas autônomos começam a ser levantadas.

Portanto, o presente artigo científico tem como objetivo realizar uma análise acerca da atual estrutura jurídica brasileira e de que forma esta se adequa à responsabilização em caso de ilicitudes cometidas por IA's autônomas. Busca, também, analisar questões éticas que conduzem as decisões tomadas pela IA e de que forma isso reflete na sociedade e no meio jurídico e apresentar casos práticos a respeito de ilegalidades relacionadas ao uso da IA.

Esta pesquisa irá explorar a seguinte questão: "Quem deve ser responsabilizado por ilegalidades resultantes do uso da Inteligência Artificial: os desenvolvedores, os operadores, as empresas que implementam a tecnologia, ou a própria IA e quais são as consequências éticas dessa responsabilidade?"

A metodologia utilizada é a pesquisa bibliográfica, com o levantamento de material literário que trata sobre os temas de responsabilidade civil, ética e regulamentação da Inteligência Artificial, trazendo as seguintes hipóteses: A responsabilidade civil deve ser distribuída entre os diferentes atores envolvidos no desenvolvimento e uso de IA e a responsabilização por danos causados por IA pode incentivar o desenvolvimento de sistemas mais éticos e seguros.

As pesquisas serão realizadas nas plataformas Google Acadêmico, Biblioteca PUC Minas e Canvas<sup>1</sup>, contando, também, com estudo de casos, a partir da investigação de casos reais de danos causados pelo uso da IA em contextos variados e, por fim, análise comparativa, com um paralelo entre legislações de diferentes países sobre a responsabilidade civil da IA.

Será dividida em três capítulos. O primeiro trará uma breve retrospectiva e alguns conceitos básicos atribuídos à trajetória da inteligência artificial até atingir a forma como é conhecida atualmente. O segundo capítulo trata sobre responsabilidade civil conforme o Código Civil brasileiro e aborda a atribuição de responsabilidade nos casos que envolvem sistemas e robôs autônomos. Por fim, o terceiro capítulo apresenta como ocorrem os danos resultantes de

---

<sup>1</sup> Portal do aluno do curso de pós-graduação em Direito 4.0: Estratégias com Inteligência Artificial da Pontifícia Universidade Católica de Minas Gerais.

vieses algorítmicos ou do mau uso de ferramentas tecnológicas e de que forma o sistema legislativo brasileiro e estrangeiro planejam mitigá-los.

## 2. A INTELIGÊNCIA ARTIFICIAL EM UMA BREVE LINHA DO TEMPO

Os seres humanos vêm desenvolvendo ferramentas destinadas ao auxílio no desempenho de suas atividades cotidianas desde o princípio do corpo social, como a criação de instrumentos de caça, ábaco ou, até mesmo, a descoberta do fogo. Com o passar do tempo e a constante evolução da sociedade, a tecnologia tem o seu marco principal datado das I e II Revoluções Industriais<sup>2</sup>, período em que o aprofundamento no estudo sobre a energia elétrica possibilitou a criação de diversos recursos a partir deste.

Em virtude do salto tecnológico, viabilizado pelas descobertas da época, a exploração de equipamentos digitais ganha destaque nos anos subsequentes, período em que nascem os primeiros computadores. A palavra computador tem origem no latim ‘*computare*’, que significa calcular ou adicionar e, embora a imagem desse aparelho seja associada às versões atuais, tem origem no ábaco, uma simples ferramenta de cálculo. Em 1946 surge a primeira máquina capaz de realizar cálculos matemáticos sem a necessidade de interação humana e, poucas décadas mais tarde, os supercomputadores.

Avançando ao presente, sabe-se que as últimas décadas foram marcadas, dentre diversos eventos, pela rápida evolução de tecnologias inovadoras em todo o mundo, especialmente durante o período da pandemia de Covid-19, que resultou no isolamento social e forçou a humanidade a buscar novos meios de conexão. Dentre os sistemas desenvolvidos nos últimos anos, a inteligência artificial popularizou-se, o que fez com que surgisse a crença de que a IA é uma descoberta recente. Todavia, ao contrário do que se pensa, seu desenvolvimento vem sendo objeto de estudos desde o século XX.

O termo surge, pela primeira vez, em 1955 durante a Conferência de Dartmouth<sup>3</sup>, evento no qual grandes nomes da ciência da computação reuniram-se e discutiram acerca da

---

<sup>2</sup> “Se as duas primeiras Revoluções Industriais marcaram os séculos XVIII e XIX – a primeira delas a partir de 1760 e a segunda por volta de 1848 –, foi o surgimento da eletrônica que propulsionou o desenvolvimento de novas tecnologias, já no século XX, que se eternizariam na história com a consolidação de uma Terceira Revolução Industrial.” FALEIROS JÚNIOR, José Luiz de Moura. *Evolução da inteligência artificial em breve retrospectiva*. In: BARBOSA, Mafalda Miranda et al. (org.). *Direito digital e inteligência artificial: diálogos entre Brasil e Europa*. Brasília: Editora Foco. p. 4.

<sup>3</sup> MCCARTHY, J.; MINSKY, M. L.; ROCHESTER, N.; SHANNON, C. E. A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. *AI Magazine*, [S. l.], v. 27, n. 4, p. 12, 2006. DOI:

possibilidade de máquinas desempenharem atividades humanas. Em consequência, foi proposto e desenvolvido, por John McCarthy, Marvin L. Minsky, Nathaniel Rochester e Claude E. Shannon, o projeto de pesquisa de verão, em Dartmouth College, New Hampshire, o qual visava estudar a possibilidade de se construir uma máquina com capacidade de produzir pensamentos autônomos.

Contudo, Alan Turing foi quem pôs à prova a capacidade de uma máquina de imitar o comportamento humano a ponto de não ser possível distingui-la de uma pessoa. O estudo foi apresentado em seu artigo “*Computing Machinery and Intelligence*”<sup>4</sup>, onde ele propõe o jogo de imitação, que fica conhecido como o Teste de Turing. Trata-se de um jogo entre três participantes: um avaliador humano, um participante humano e um robô. Durante este jogo, o avaliador faz perguntas aos dois participantes, a fim de identificar quem é o humano e quem é a máquina.

Entre os anos de 1956 a 1974 o campo de pesquisa acerca da IA obteve grande progresso. Com o passar do tempo, o estudo e a criação de máquinas capazes de desempenhar atividades diversas, atividades, estas, anteriormente performadas apenas por seres humanos, se popularizaram, dando início à exploração de um novo campo científico, o qual propôs a viabilidade de um robô com aptidão de executar funções sem supervisão direta.

Grande parte do desenvolvimento da IA se deu graças às pesquisas acadêmicas da época e financiamento governamental. A produção de artigos científicos e livros deram origem aos métodos que viriam a ser utilizados mais tarde. Na época, as duas principais teorias a respeito do desenvolvimento da inteligência artificial foram de Marvin Lee Minsky<sup>5</sup> e de Frank Rosenblatt<sup>6</sup>. Em 1957, Rosenblatt criou um programa de computador chamado Mark 1 Perceptron, que serviu como base para as redes neurais<sup>7</sup>.

---

10.1609/aimag.v27i4.1904.

Disponível

em:

<https://ojs.aaai.org/aimagazine/index.php/aimagazine/article/view/1904>. Acesso em: 25 jun. 2025.

<sup>4</sup> TURING, Alan M. *Computing machinery and intelligence*. *Creative Computing*, v. 6, n. 1, p. 44-53, 1980.

<sup>5</sup> “Uma teoria foi sugerida por Minsky, que disse que precisavam existir sistemas simbólicos. Isso significava que a IA devia basear-se na lógica tradicional do computador ou na pré-programação - ou seja, no uso de estruturas como if-then-else.” TAULLI, Tom. *Introdução à inteligência artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020. p. 26.

<sup>6</sup> “A segunda teoria, proposta por Frank Rosenblatt, acreditava que a IA precisava usar sistemas semelhantes ao cérebro como redes neurais (esse campo também era conhecido como conexionismo). Em vez de chamar as partes internas de neurônios, entretanto, ele as denominou "perceptrons". Um sistema seria capaz de aprender à medida que recebesse dados ao longo do tempo.” TAULLI, Tom. *Introdução à inteligência artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020. p. 26.

<sup>7</sup> “Uma rede neural é um processador maciçamente paralelamente distribuído constituído de unidades de processamento simples, que têm a propensão natural para armazenar conhecimento experimental e torná-lo disponível para o uso. Ela se assemelha ao cérebro em dois aspectos: 1. O conhecimento é adquirido pela rede a partir de seu ambiente através de um processo de aprendizagem. 2. Forças de conexão entre neurônios, conhecidas

De 1956 a 1974, o campo da IA foi um dos mais movimentados no mundo tecnológico. Um grande catalisador foi o rápido desenvolvimento na tecnologia dos computadores. Eles passaram de sistemas maciços – baseados em tubos de vácuo – para sistemas menores que funcionavam com circuitos integrados muito mais rápidos e dispunham de maior capacidade de armazenamento.<sup>8</sup>

Nos anos 70 houve um período que ficou conhecido como *AI Winter* (inverno da IA). A IA estava limitada ao meio acadêmico e passou por declínio em seu desenvolvimento, por conta de cortes nos financiamentos e descrença por parte dos pesquisadores. Alguns destes mantiveram seus projetos mesmo diante da desistência dos demais, que foi quando inovações como retropropagação e redes neurais recorrentes surgiram. Uma década mais tarde, com o fim do *AI Winter*, as ideias de Rosenblatt voltaram a ser discutidas e o aperfeiçoamento do que hoje é conhecido como *deep learning* se iniciou.

Desta forma, nos anos subsequentes, nascem os softwares desenvolvidos para executar uma alta gama de tarefas, das mais simples, às mais complexas. Estes sistemas podem ser divididos entre fraco e forte, em que a IA fraca, apesar de sua composição complexa, tem limitações em suas execuções. Ela apenas fará o que foi programada para fazer e, em geral, dependerá de intervenção ou supervisão humana. Como exemplo, pode-se apontar *chatbots*, reconhecimento facial, filtros ou assistentes virtuais. Esta é a IA mais utilizada, atualmente, no cotidiano da população mundial.

Já as máquinas autônomas, também conhecidas como sistemas de IA forte, são os softwares criados para que, em teoria, desenvolvam capacidades de compreensão e aprendizado, ou seja, o *machine learning*<sup>9</sup> e o *deep learning*<sup>10</sup>. Dessa forma, se utilizará de suas métricas para solucionar novos obstáculos, sem que haja, necessariamente, intervenção humana.<sup>11</sup> Com isso, inicia-se uma era de robôs independentes e capazes de tomar decisões, o

---

como pesos sinápticos, são utilizadas para armazenar o conhecimento adquirido.” HAYKIN, Simon. Redes neurais: princípios e prática; tradução Paulo Martins Engel. – 2. Ed. – Dados eletrônicos. Porto Alegre: Bookman, 2007. p. 28.

<sup>8</sup> TAULLI, Tom. *Introdução à inteligência artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020. p. 24.

<sup>9</sup> “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”. BARBOSA, Lucia Martins; PORTES, Luiza Alves Ferreira. A inteligência artificial. *Revista Tecnologia Educacional*, Rio de Janeiro, n. 236, p. 16-27, 2023. p. 19.

<sup>10</sup> “A tecnologia deep learning é uma subárea do machine learning. Esse tipo de sistema permite o processamento de enormes quantidades de dados para encontrar relacionamentos e padrões que os seres humanos são muitas vezes incapazes de detectar. A palavra “deep” (em português “profundo”) refere-se ao número de camadas ocultas na rede neural, as quais fornecem grande parte do poder de aprendizagem.” TAULLI, Tom. *Introdução à inteligência artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020. p. 98.

<sup>11</sup> “Em linhas gerais, tem-se que programas de computador capazes de aprender são, na verdade, capazes de extrair padrões estatísticos e correlações ocultas em seus (volumosos) dados de entrada.” NEGRI, Sergio M. C. A. ;

que levanta a seguinte questão: De que maneira estas máquinas autônomas podem afetar a sociedade?

Para que este conceito fique mais claro, cabe salientar que:

A Inteligência Artificial é um ramo das ciências da computação que busca construir mecanismos, físicos ou digitais, que simulem a capacidade humana de pensar e de tomar decisões. [...] é a capacidade de dispositivos eletrônicos de funcionar de maneira que lembre o pensamento humano. Isso implica em perceber variáveis, tomar decisões e resolver problemas. Enfim, operar em uma lógica que remete ao raciocínio.<sup>12</sup>

Neste caso, há dois exemplos a serem apresentados: O primeiro chama-se Baxter, um robô industrial equipado com dois braços e um algoritmo de aprendizado profundo pré-treinado em uma rede neural de duas camadas para que desenvolva a habilidade de identificar objetos, através de uma câmera, e determinar qual a melhor forma de apanhá-los e movimentá-los. Sua função se limitava a desempenhar tarefas repetitivas, o que permitia, por vezes, a intervenção humana.

O segundo exemplo é de um pequeno robô chamado Darwin, criado por Igor Mordatch. Este robô foi programado com um algoritmo de *deep learning* em camadas profundas de redes neurais que permitem que ele aprenda a partir dos próprios acertos ou erros e implemente este aprendizado em tempo real, tal qual uma criança humana.

Durante o período de experiências, o qual foi dividido em dois estágios, simulação e fase de testes, Darwin recebeu orientações básicas para que compreendesse o ambiente, sem demais intervenções. Desta forma, no decorrer da simulação, Darwin mapeou uma sequência de movimentos que lhe permitiram entender quais eram as posições necessárias para que ele pudesse permanecer em pé, girar o seu tronco ou locomover-se. Já na fase de testes, Darwin, de fato, empregou o conhecimento adquirido e transitou pela sala, sem o auxílio de seus criadores.

Essa capacidade de adquirir habilidades e executar tarefas para além do que foi projetado, ou seja, sem a intervenção humana, permite que robôs autônomos pratiquem atos imprevisíveis, até mesmo para os seus programadores, visto que algoritmos de *deep learning* podem gerar resultados inesperados, dada a complexidade de suas redes neurais. Isso fez com que o sistema jurídico de diversos países encontrasse dificuldade na atribuição de responsabilidade pelos danos causados por sistemas autônomos de IA.

---

LOPES, G. F. P. . DA PERSONALIDADE ELETRÔNICA À CLASSIFICAÇÃO DE RISCOS NA INTELIGÊNCIA ARTIFICIAL (IA). Teoria Jurídica Contemporânea , v. 6, p. 01-26, 2022. p. 5.

<sup>12</sup> BARBOSA, Lucia Martins; PORTES, Luiza Alves Ferreira. A inteligência artificial. *Revista Tecnologia Educacional*, Rio de Janeiro, n. 236, p. 16-27, 2023. p. 17.

### 3. PECULIARIDADES EM MATÉRIA DE RESPONSABILIDADE CIVIL

A ascensão da IA vem promovendo avanços em diversas áreas, tais como saúde, educação, direito, agronegócio, comunicação, setor financeiro, indústria, medicina, pesquisas científicas, dentre outras. Isso trouxe diversos benefícios, não apenas em escala global, mas, também, individual.

No entanto, essa transformação trouxe alguns desafios. À medida que algoritmos tomam decisões que impactam diretamente a sociedade, surge a necessidade de adaptar legislações existentes para que seja assegurada a proteção de cada indivíduo. Além dos obstáculos na garantia de segurança, capacitação, privacidade e parâmetros éticos, a atribuição de responsabilidade como consequência das ilegalidades cometidas com o uso da inteligência artificial ou causadas por tomadas de decisão não supervisionadas encontra lacunas.

Nos últimos anos a atribuição de responsabilidade por danos causados por tomadas de decisão autônomas vem sendo bastante discutida no meio jurídico, visto que a legislação atual carece de previsões expressas acerca do assunto. Este problema se agrava com a complexidade progressiva dos sistemas de IA desenvolvidos na atualidade, pois, além da imprevisibilidade, há dificuldade em identificar responsáveis e demonstrar a relação de causa e efeito.

[...] as implicações nos regimes de responsabilidade civil serão, no essencial, a expressão de uma tendência aparentemente contraditória nos seus termos: as novas tecnologias aproximam os utilizadores, mas diluem as identidades dos agentes responsáveis. A realidade obrigará à revisão dos pressupostos clássicos da responsabilidade civil.<sup>13</sup>

No ordenamento jurídico brasileiro, a responsabilidade civil vem como resultado da prática de ato ilícito que viola o direito de outrem. O art. 186 do Código Civil estabelece que “aquele que, por ação ou omissão voluntária, negligência ou imprudência, violar direito e causar dano a outrem, ainda que exclusivamente moral, comete ato ilícito”<sup>14</sup>. Assim, nasce, na forma do art. 927 do Código Civil<sup>15</sup>, o dever de reparar. Pela ótica do princípio da reparação integral,

<sup>13</sup> ANTUNES, Henrique Souza. Inteligência artificial e responsabilidade civil: enquadramento. *Revista de Direito da Responsabilidade*, Coimbra, v. 1, p. 139-154, 2019. p. 139.

<sup>14</sup> BRASIL. Lei n. 10.406, de 10 de janeiro de 2002. Institui o Código Civil. *Diário Oficial da União*: seção 1, Brasília, DF, 11 jan. 2002. Disponível em: [https://www.planalto.gov.br/ccivil\\_03/leis/2002/110406compilada.htm](https://www.planalto.gov.br/ccivil_03/leis/2002/110406compilada.htm). Acesso em: 5 maio 2025.

<sup>15</sup> “Art. 927. Aquele que, por ato ilícito (arts. 186 e 187), causar dano a outrem, fica obrigado a repará-lo. (Vide ADI nº 7055) (Vide ADI nº 6792) Parágrafo único. “Haverá obrigação de reparar o dano, independentemente de culpa, nos casos especificados em lei, ou quando a atividade normalmente desenvolvida pelo autor do dano implicar, por sua natureza, risco para os direitos de outrem.” BRASIL. Lei n. 10.406, de 10 de janeiro de 2002. Institui o Código Civil. *Diário Oficial da União*: seção 1, Brasília, DF, 11 jan. 2002. Disponível em: [https://www.planalto.gov.br/ccivil\\_03/leis/2002/110406compilada.htm](https://www.planalto.gov.br/ccivil_03/leis/2002/110406compilada.htm). Acesso em: 5 maio 2025.

busca-se atingir o estado em que se encontrava o lesado antes do dano. De acordo com Maria Helena Diniz:

Poder-se-á definir a responsabilidade civil como a aplicação de medidas que obriguem alguém a reparar dano moral ou patrimonial causado a terceiros em razão de ato do próprio imputado, de pessoa por quem ele responde, ou de fato de coisa ou animal sob sua guarda ou, ainda, de simples imposição legal. Definição esta que guarda, em sua estrutura, a ideia da culpa quando se cogita da existência de ilícito (responsabilidade subjetiva), e a do risco, ou seja, da responsabilidade sem culpa (responsabilidade objetiva).<sup>16</sup>

Sabe-se que a responsabilidade civil pode seguir duas linhas: subjetiva e objetiva. O principal elemento que as distingue é a culpa. De forma mais ampla, na responsabilidade civil subjetiva, o dever de reparação sucede, em conjunto, da comprovação de dano causado por conduta ilícita cometida pelo agente, da culpa e do nexo de causalidade entre a conduta e o resultado. Essa é a regra geral presente no ordenamento jurídico brasileiro.

Sobre a responsabilidade civil subjetiva, no contexto dos sistemas de inteligência artificial, comenta Luiza Loureiro Coutinho:

Cumpra, assim, pontuar sobre os elementos da culpa: a) a quebra do dever de cuidado que leva à conduta lesiva e a diligência a que o sujeito se obriga em um dado contexto, a qual se compara com a voluntariedade da conduta (na culpa *stricto sensu*, a conduta em si é praticada voluntariamente, porém o resultado não é desejado; no dolo, tanto a conduta como o resultado são intencionais e desejados); c) a previsibilidade (para se reputar culposa a quebra do dever de cuidado, o lesante precisa ter previsto, ao menos potencialmente, o dano)<sup>17</sup>.

Entretanto, nos casos em que a responsabilidade civil é objetiva, a culpa se torna um componente dispensável, sendo necessária, apenas, a comprovação da conduta, do dano e do nexo causal entre ambos. Trata-se de uma exceção aplicada aos cenários que envolvem risco. Isso pode ser observado em relações de consumo ou em atividades de risco. “Essa responsabilidade tem como fundamento a atividade exercida pelo agente, pelo perigo que pode causar dano à vida, à saúde ou a outros bens, criando risco de dano para terceiros.”<sup>18</sup>

Ao considerar a distinção fundamental entre os regimes de responsabilidade civil e o aumento da autonomia dos robôs<sup>19</sup>, a doutrina diverge acerca de questões de aplicabilidade

<sup>16</sup> DINIZ, Maria Helena. *Manual de Direito Civil*. 4. ed. São Paulo: Saraiva Jur, 2022. p. 285.

<sup>17</sup> COUTINHO, Luiza Loureiro. *Riscos em sistemas de inteligência artificial: definição, tipologias, correlações, principiologia, responsabilidade civil e regulação*. Salvador: Juspodivm, 2024. p. 372.

<sup>18</sup> DINIZ, Maria Helena. *Manual de Direito Civil*. 4. ed. São Paulo: Saraiva Jur, 2022. p. 288.

<sup>19</sup> “[...]a autonomia tem sido compreendida nos estudos sobre inteligência artificial como a capacidade de o sistema agir sem intervenção humana direta. No entanto, trata-se de característica de alguns sistemas ajustarem suas ações com base no feedback do ambiente, e, também, sobre sua própria situação atual. O aumento da autonomia é geralmente equiparado a uma maior adaptação ao ambiente e às vezes é apresentado como um aumento da “inteligência” para uma determinada tarefa.” FONSECA, Aline Klayse dos Santos. Interação ser humano-máquina: o padrão de “controle humano significativo” e seus impactos na imputação da responsabilidade civil por danos decorrentes de veículos autônomos. *Revista da Faculdade de Direito, Universidade de São Paulo, [S. l.]*, v.



da responsabilidade ao fato concreto. Nas palavras de Gustavo Melo, “a solução para os danos derivados da IA não vem de um regramento específico, mas sim de uma cláusula geral de responsabilidade civil”<sup>20</sup>.

Por um lado, acredita-se que o potencial lesivo dos novos robôs faz com que a responsabilidade civil objetiva se encaixe de melhor forma ao presente cenário<sup>21</sup>, fazendo com que a responsabilidade pelo produto danoso seja atribuída aos fabricantes, programadores e/ ou fornecedores.

Em alguns casos, a lei de responsabilidade civil impõe responsabilidade a réus (como fabricantes de um robô) que não são negligentes nem culpados de dolo intencional. Conhecida como responsabilidade objetiva, ou responsabilidade sem culpa, este ramo da responsabilidade civil busca regular aquelas atividades que são úteis e necessárias, mas que criam riscos anormalmente perigosos para a sociedade. A responsabilidade objetiva do produto se baseia na existência de um produto razoavelmente perigoso cujo uso previsível causou lesões. [...] Sob a doutrina da responsabilidade objetiva do produto, um fabricante deve garantir que seus bens sejam adequados para o uso pretendido quando são colocados no mercado para consumo público.<sup>22</sup>

Em contrapartida, alguns países já discutem sobre a criação de personalidade jurídica para robôs autônomos. Frente à imprevisibilidade e a habilidade de agir independentemente de instrução humana direta destes modelos de software<sup>23</sup>, o Parlamento

---

117, p. 357–376, 2022. Disponível em: <https://revistas.usp.br/rfdusp/article/view/218284>. Acesso em: 15 julho 2025.p. 369.

<sup>20</sup> MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas*: a reparação dos danos causados por sistemas de inteligência artificial. Londrina: Thoth, 2023, p. 78.

<sup>21</sup> “Mais recentemente, a Resolução do Parlamento Europeu, de outubro de 2020, contendo recomendações à Comissão sobre a responsabilidade civil por danos causados pela inteligência artificial<sup>37</sup>, acaba por nos oferecer, igualmente, pistas importantes em matéria de responsabilidade civil. Propõe-se, aí, a distinção entre duas vias de responsabilização: a responsabilidade objetiva, aplicável às hipóteses de sistemas de alto risco, e a responsabilidade por culpa, nas restantes situações” BARBOSA, Mafalda Miranda. Inteligência artificial, responsabilidade civil e causalidade: breves notas. *Revista de Direito da Responsabilidade*, Coimbra, ano 3, p. 605-625, 2021. Disponível em: <https://revistadireitoresponsabilidade.pt/2021/inteligencia-artificial-responsabilidade-civil-e-causalidade-breves-notas-mafalda-miranda-barbosa/>. Acesso em: 30 junho 2025. p. 622.

<sup>22</sup> BARFIELD, Woodrow. Liability for autonomous and artificially intelligent robots. *Paladyn, Journal of Behavioral Robotics*, v. 9, n. 1, p. 193-203, 2018. DOI: 10.1515/pjbr-2018-0018. Acesso em: 5 maio 2025. P. 197. Tradução livre. No original: “In some cases, tort law imposes liability on defendants (such as manufacturers of a robot) who are neither negligent nor guilty of intentional wrongdoing. Known as strict liability, or liability without fault, this branch of torts seeks to regulate those activities that are useful and necessary but that create abnormally dangerous risks to society. Strict products liability is predicated on the existence of an unreasonably dangerous product whose foreseeable use has caused injury. [...] And under the doctrine of strict products liability, a manufacturer must guarantee that its goods are suitable for their intended use when they are placed on the market for public consumption.”;

<sup>23</sup> “O que ocorre é que, na medida em que a influência do programador sobre a máquina diminui, a influência do ambiente operacional aumenta. Essencialmente, o programador transfere parte de seu controle sobre o produto para o ambiente, sobretudo em se tratando de máquinas que continuam aprendendo e se adaptando em seu ambiente operacional final. Como nessa situação elas têm que interagir com um número potencialmente grande de pessoas e situações, via de regra não será possível prever a influência do ambiente operacional.” NEGRI, Sergio M. C. A. ; LOPES, G. F. P. . DA PERSONALIDADE ELETRÔNICA À CLASSIFICAÇÃO DE RISCOS NA INTELIGÊNCIA ARTIFICIAL (IA). *Teoria Jurídica Contemporânea* , v. 6, p. 01-26, 2022. p.9.

Europeu, em 2017, propôs a criação de um novo status para os robôs com alto nível de independência: a personalidade eletrônica.

Ao conferir personalidade jurídica a um robô, este passa a ser reconhecido como sujeito titular de direitos e deveres perante a sociedade. Essa concessão implicaria na capacidade de ser responsabilizado diretamente por danos causados por este e, embora a proposta não se trate de lei vinculativa, há divergências a respeito da adaptação legislativa frente às atuais questões referentes aos robôs autônomos, bem como dúvidas acerca dos parâmetros éticos que determinarão estas mudanças.

No debate acerca da personificação da Inteligência Artificial, é comumente constatada a afirmação de que as normas legais existentes (a respeito, por exemplo, da responsabilidade civil em caso de danos causados por esses agentes) seriam incapazes de retratar e, consequentemente, disciplinar agentes artificiais autônomos.<sup>24</sup>

Em seu texto, o art. 59, alínea “f”, da Resolução do Parlamento Europeu (2015/2103) propõe o seguinte:

Criar um estatuto jurídico específico para os robôs a longo prazo, de modo a que, pelo menos, os robôs autônomos mais sofisticados possam ser determinados como detentores do estatuto de pessoas eletrônicas responsáveis por sanar quaisquer danos que possam causar e, eventualmente, aplicar a personalidade eletrônica a casos em que os robôs tomam decisões autônomas ou em que interagem por qualquer outro modo com terceiros de forma independente;<sup>25</sup>

Para Kurki<sup>26</sup>, no que diz respeito à personalidade jurídica para inteligências artificiais, há três contextos principais que devem ser analisados: Contexto do valor final, contexto da responsabilidade e contexto comercial. Estes têm ligação com os pilares da responsabilidade civil. No primeiro, o autor questiona o quão digna uma máquina, até mesmo as que mais se assemelham aos humanos, seria de receber as mesmas proteções legais que pessoas adquirem através dos direitos de personalidade.

Em sequência, no segundo contexto, analisa-se de que forma a IA poderia ter responsabilidade civil, ou responder criminalmente por ilegalidades. Além de tratar da possibilidade de antecipar estes danos e, o mais importante, a implicação de a IA poder ter propriedades para que arque com esses danos.

---

<sup>24</sup> NEGRI, Sergio M. C. A. ; LOPES, G. F. P. . DA PERSONALIDADE ELETRÔNICA À CLASSIFICAÇÃO DE RISCOS NA INTELIGÊNCIA ARTIFICIAL (IA). Teoria Jurídica Contemporânea , v. 6, p. 01-26, 2022. p.16.

<sup>25</sup> PARLAMENTO EUROPEU. Resolução do Parlamento Europeu, de 16 de fevereiro de 2017, que contém recomendações à Comissão sobre disposições de Direito Civil sobre Robótica (2015/2103(INL)). Disponível em: [https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051\\_PT.html](https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_PT.html). Acesso em: 5 maio 2025.

<sup>26</sup> KURKI, Visa A.J. *A theory of legal personhood*. Oxford: Oxford University Press, 2019. (Oxford Legal Philosophy). Disponível em: <https://doi.org/10.1093/oso/9780198844037.001.0001>. Acesso em: 5 maio 2025.

Por fim, no terceiro contexto, discorre sobre a IA como agente comercial, situação na qual estabelece vínculo de consumo através de compras e vendas. Neste contexto levanta-se questões acerca de direitos especiais, competência e propriedade.

Dessa forma, a concessão de direitos de personalidade às IA's seguiria uma linha de valores, em que a IA de valor poderia se tornar detentora de personalidade civil relativa, a mesma atribuída aos menores de 18 e maiores de 16 e aos pacientes em coma, por exemplo, passando a obter a personalidade civil absoluta na hipótese de provar ser capaz de agir forma independente. Já no caso de IA's de valor inestimável, estas poderiam adquirir direitos de personalidade apenas com base em sua capacidade de cumprirem deveres legais e/ou administrativos, assim como os demais seres humanos. Mas, nesse cenário, IA's e robôs teriam capacidade de compreender e lidar com as consequências de seus atos?

A concessão da personalidade jurídica plena aos agentes eletrônicos implicaria em adotar certas abordagens a fim de superar algumas deficiências atuais em relação à reparação de danos causador por robôs e IA's, como, por exemplo: a capacidade de estes agentes agirem, tanto em seu nome, quanto em nome de terceiros ao fazer declarações de vontade sobre os seus próprios atos, ou a possibilidade de gerarem e possuírem bens, recursos financeiros e, até mesmo, ter acesso à contas bancárias e gerir o seu patrimônio para que possam arcar com essa reparação.

Para Gunther Teubner<sup>27</sup>, a viabilidade da personalidade eletrônica apenas se demonstra em um cenário em que, ao ter direitos e deveres semelhantes aos dos seres humanos, o robô tenha meios de defender os seus próprios interesses, ou seja, autonomia o suficiente para que as suas decisões partam, única e exclusivamente, de suas vontades. No entanto, no presente momento, os robôs e IA's dotados de autonomia ainda não têm a capacidade de demonstrarem interesses próprios, eles apenas agem de acordo com os interesses de quem os programou ou de quem os utiliza. Sendo assim, a ausência de autonomia plena e de interesses próprios, não permite que estes agentes sejam classificados como sujeitos de direito.

A busca por maneiras mais eficazes de garantir a reparação de danos e a mitigação de riscos faz com que o legislador procure soluções que possam se adaptar à situação atual. Todavia, ainda que a ideia de concessão de personalidade civil às máquinas vise facilitar a atribuição direta de responsabilidade por infrações por elas cometidas, há obstáculos, não apenas legais, mas éticos e práticos, o que traz a necessidade de se considerar certos fatores.

---

<sup>27</sup> TEUBNER, Gunther. Digital personhood? The status of autonomous software agents in private law. Tradução de Jacob Watson. *Ancilla Iuris*, 2018. Disponível em <https://ssrn.com/abstract=3177096>. Acesso em 27 junho 2025. p. 114.

Ao permitir que robôs e sistemas algorítmicos passem a responder pelos seus próprios atos, desconsidera-se a atuação dos desenvolvedores, operadores e fornecedores, o que estabelece precedente para diluição da responsabilidade humana nos resultados negativos produzidos por estes softwares. Essa ótica acaba por ignorar o controle e toda a participação desses agentes sobre as ações executadas por suas máquinas e, por consequência, atrapalha as chances de reparação. Desta forma, a garantia de segurança jurídica aos usuários diretos e indiretos de ferramentas tecnológicas autônomas torna-se impraticável.<sup>28</sup>

#### **4. VIESES ALGORÍTMICOS, DADOS E O MARCO LEGAL DA INTELIGÊNCIA ARTIFICIAL: IMPLICAÇÕES ÉTICO-JURÍDICAS**

Ao tratar de questões éticas e legais envolvendo o uso da inteligência artificial, há que se considerar diversos pontos cruciais. Recentemente, debates acerca de temas como viés algorítmico, autonomia, transparência, explicabilidade, privacidade, segurança de dados, dentre outros, têm recebido maior relevância na esfera do direito digital. Nesse sentido, à medida que a utilização de inteligência artificial para automatizar processos de tomada de decisão vem demonstrando um avanço progressivo, a inquietação sobre justiça e equidade se intensifica.

Para que estes sistemas de tomada de decisão automatizada se desenvolvam, eles são abastecidos com uma quantidade massiva de dados, os quais servem de base de treinamento para os algoritmos da IA. A partir das informações geradas, o sistema passa a identificar padrões e absorver os aspectos desses dados.

Em geral, os dados coletados são utilizados de forma a beneficiar, tanto o usuário, para que este tenha acesso a uma experiência personalizada e otimizada, quanto o fornecedor ao auxiliar no desempenho de seu negócio. Atualmente, grande parte das plataformas obtêm esses dados através das informações passadas diretamente pelo usuário, como em registros, publicações e interações. Por outro lado, de forma indireta, os dados podem ser coletados e armazenados de acordo com os sites que são visitados, vídeos assistidos, cliques, curtidas e o tempo gasto em cada um destes.

---

<sup>28</sup> “Não há como se pretender transferir a responsabilidade inteiramente para o algoritmo, mesmo que este tenha autonomia considerável (ex.: IA do tipo forte), pois, em última instância, deve o administrador ser diligente, como determina a lei.” MEDON, Filipe. Inteligência artificial e responsabilidade civil: autonomia, riscos e solidariedade. São Paulo: JusPodivm, 2022. p. 404

Estes dados podem ser estruturados, não estruturados ou semiestruturados<sup>29</sup>. Assim que coletados, passam por uma fase de preparação, em que são feitas as filtragens e correções necessárias de forma a garantir que os algoritmos consigam interpretá-los de forma eficaz. Após essa fase inicia-se o treinamento de modelo algorítmico o qual é alimentado com os dados e adaptados de acordo com o objetivo final. Uma vez que o modelo apresente um resultado desejado, este estará pronto para a implementação.

Ainda, os dados precisam ser processados e trabalhados para que possam gerar valor: isso ocorre através do *data analytics*, em que há a possibilidade de extrair correlações, padrões e associações que possam ser consideradas informações. Essa análise de dados fornece incentivos ainda maiores para o uso de projeções no setor privado, na medida em que permitem que mais informações sejam processadas e que novas correlações entre dados e comportamentos futuros possam emergir.<sup>30</sup>

Contudo, mesmo após a limpeza, os dados utilizados para o treinamento do algoritmo podem carregar informações enviesadas, o que dá origem aos vieses algorítmicos. Os vieses se manifestam na implementação, momento em que a forma como as decisões são tomadas evidenciam sua presença, porém podem surgir a qualquer momento durante o ciclo de criação de um sistema de inteligência artificial.

Para que o conceito fique mais claro, vieses algorítmicos são erros presentes no sistema algorítmico de uma inteligência artificial que acabam por produzir resultados injustos. Sua fonte primária é composta por dados, os quais podem carregar informações que espelham desigualdades sociais, preconceitos e, até mesmo, ódio a grupos seletos. Caso treinado com dados como esses, o algoritmo passará a reproduzir padrões enviesados.

Para além dos dados, os vieses também podem surgir ao longo do desenvolvimento dos algoritmos, a partir de escolhas do desenvolvedor, mesmo que de forma não proposital.<sup>31</sup> Assim, há dois pontos que merecem destaque: vieses bons e vieses maus.

Ainda que os vieses sejam, em sua maioria, negativos, há alguns cenários em que o programador pode criar, de forma intencional, um viés capaz de garantir o melhor funcionamento do algoritmo que será utilizado, como, por exemplo, em um algoritmo hospitalar que priorize o atendimento de pacientes com maior risco de morte em relação a outros casos.

---

<sup>29</sup> “Dados estruturados são rotulados e formatados - e geralmente são armazenados em um banco de dados relacional ou em uma planilha. Dados não estruturados são informações que não têm formatação predefinida. Dados semiestruturados têm algumas marcas internas que ajudam na categorização.” TAULLI, Tom. *Introdução à inteligência artificial*. Uma abordagem não técnica. São Paulo: Novatec, 2020. p. 60.

<sup>30</sup> MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas: a reparação dos danos causados por sistemas de inteligência artificial*. Londrina: Thoth, 2023. p. 15.

<sup>31</sup> “Por meio do tratamento em massa de dados, as máquinas adquirem autoaprendizagem, podendo alcançar determinados resultados independentemente de qualquer supervisão humana” MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas: a reparação dos danos causados por sistemas de inteligência artificial*. Londrina: Thoth, 2023. p. 12.

Ou um sistema de contratação que promova maior equidade entre os colaboradores de uma empresa, garantindo a inclusão de diferentes etnias e gêneros naquele espaço.

Já os vieses maus são aqueles que produzem resultados discriminatórios ou injustos e perpetuam desigualdades sociais. Podem surgir através de dados que não refletem representatividade da população, quando o próprio sistema espelha o comportamento negativo de usuários, ou por dados que contenham informações discriminatórias.

Embora a utilização de algoritmos possa trazer maior eficiência, conforme visto no ponto anterior, o seu uso pode acarretar discriminações: são os chamados vieses algorítmicos. Nesse contexto, destaca-se que os algoritmos não estão imunes ao problema da discriminação, já que são programados por seres humanos, cujos valores estão embutidos no software. [...] Um sistema de IA será tão enviesado quanto forem os dados e os algoritmos utilizados por ele e que, por sua vez, foram selecionados por um humano. Portanto, seja pela má coleta, seja pelo mau processamento, corre-se o risco de que a decisão tomada tenha vieses.<sup>32</sup>

Para exemplificar, no ano de 2018 foi relatado incidente envolvendo a Amazon, a qual teria identificado uma falha em seu sistema de IA que fora desenvolvido para automatizar o processo de seleção curricular de candidatos a vagas de emprego dentro da empresa.

O objetivo era identificar e filtrar currículos que possuíam os atributos desejados de acordo com as vagas, porém, com o tempo, notou-se que o algoritmo demonstrava predisposição a discriminação em razão do gênero, evitando os currículos de candidatas do gênero feminino. Isso ocorreu em razão da base de dados utilizada para o treinamento do algoritmo, que foi composta por currículos antigos enviados à empresa, em grande parte provenientes de candidatos do gênero masculino.<sup>33</sup>

As decisões exclusivamente automatizadas se referem à capacidade de tomar decisões através de meios tecnológicos e sem intervenção humana. Embora o humano alimente o sistema com dados e interprete o resultado apresentado pelo software, o procedimento decisório, ainda assim, é automatizado. [...] Assim, o algoritmo objetiva auxiliar na tomada de decisões, utilizando-se de previsões probabilísticas pela análise dos dados fornecidos. Quanto maior a quantidade e qualidade dos dados disponibilizados ao algoritmo, maior a chance de o resultado estar próximo do real.<sup>34</sup>

Dessa forma, a busca pela imparcialidade algorítmica e equidade no tratamento de dados ainda encontra obstáculos, uma vez que a interferência humana pode trazer alteração nos

---

<sup>32</sup> MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas*: a reparação dos danos causados por sistemas de inteligência artificial. Londrina: Thoth, 2023. p. 17.

<sup>33</sup> “Consequentemente, o algoritmo “aprendeu” diante do viés estatístico a dar preferência a características mais comuns nos currículos masculinos, como certas palavras-chave e históricos de trabalho específicos, enquanto desvalorizava características frequentemente encontradas em currículos de candidatas do sexo feminino.” VIANA, Guilherme Manoel de Lima; MACEDO, Caio Sperandeo de. Inteligência artificial e a discriminação algorítmica: uma análise do caso Amazon. *Direito & TI*, v. 1, n. 19, p. 39-62, 2024. Disponível em: <https://direitoeti.com.br/direitoeti/article/view/212>. Acesso em: 5 maio 2025. p. 55.

<sup>34</sup> MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas*: a reparação dos danos causados por sistemas de inteligência artificial. Londrina: Thoth, 2023. p. 12.

resultados desejados, seja em virtude de dados enviesados, ou em razão de princípios, crenças e julgamentos particulares dos indivíduos envolvidos no processo de criação de um programa de inteligência artificial. Dessa forma, um olhar regulatório acerca de questões relacionadas a processamento de dados, decisões algorítmicas e equidade se faz necessário.

Essa demanda se torna ainda mais evidente em casos como o da Cambridge Analytica e Facebook que ocorreu em 2016. No período pré-eleição houve coleta não autorizada e uso indevido de dados pessoais de, aproximadamente, 87 milhões de usuários. Os dados foram obtidos através de um aplicativo destinado à pesquisa acadêmica, o qual traçava perfis dos usuários por meio de testes psicológicos e mais tarde os utilizou para influenciar decisões eleitorais.

Apesar de o consentimento dos usuários ter sido dado, exclusivamente, para a utilização de seus dados na pesquisa em questão, foram coletadas informações, não apenas sobre os indivíduos que haviam baixado o aplicativo, mas de seus amigos na rede social Facebook, a qual permitiu que o aplicativo Kogan tivesse acesso aos demais usuários.

Dessa forma, a Cambridge Analytica, uma empresa britânica de consultoria política, criou perfis psicológicos dos usuários, o que permitiu que os anúncios direcionados a cada um deles fossem personalizados de maneira a influenciá-los durante o período da campanha presidencial dos Estados Unidos.

Esse escândalo expôs, não apenas o quão frágil são as políticas de privacidade de dados aplicadas por grandes plataformas, mas a gravidade da falta de transparência e respeito aos parâmetros éticos.

Práticas como estas demonstram a premência da mitigação de vieses algorítmicos. Por mais que se trate de um desafio bastante complexo, é indispensável o desenvolvimento de estratégias capazes de monitorar e apontar erros existentes em todas as fases da criação de uma IA. Para além da promoção de melhorias internas na fase de coleta de dados, através de sistemas de verificação, utilização de dados mais representativos ou, até mesmo pelo treinamento de um algoritmo equipado com métricas de identificação de vieses, projetos externos que fomentam um desenvolvimento ético e a implementação de legislações que garantem direitos de transparência e, principalmente, de responsabilidade merecem destaque.

Nesse contexto, surgem legislações como o *GDPR*<sup>35</sup>, regulamento europeu de proteção de dados, e, anos mais tarde, o *EU AI Act*<sup>36</sup> como pioneiro na matéria de regulamentação da inteligência artificial. Ambos buscam garantir ao usuário o direito à transparência, uso ético da inteligência artificial, governança e tratamento justo de dados e a mitigação de práticas abusivas e vieses discriminatórios.

O Regulamento Geral de Proteção de Dados, que foi promulgado em 2018 na União Europeia, tem por objetivo estabelecer parâmetros éticos e legais para a proteção de dados pessoais. O regulamento determina os direitos de cada indivíduo sobre suas informações, bem como as obrigações atribuídas aos operadores e controladores. Em paralelo, no Brasil há a Lei Geral de Proteção de Dados, que vigora desde 2020 e tem estrutura semelhante ao modelo europeu, o que reforça a importância de se considerar a proteção de dados como direito fundamental dos usuários. Todavia, para além dos desafios referentes à proteção de dados, a inteligência artificial e questões ligadas a vieses, transparência, explicabilidade e possíveis consequências socioeconômicas se tornaram os propulsores para a criação de legislações que tratam sobre a matéria. Após a influência do AI Act, o Brasil avançou significativamente nos debates que estabelecem a criação de uma legislação própria.

Também há uma preocupação com a regulação da IA. Nesse sentido, a União Europeia criou algumas diretrizes e propostas para regular os sistemas de IA, em que se destaca a Proposta de Regulamento da Inteligência Artificial. No Brasil, recentemente, foi elaborado um anteprojeto do Marco Legal de Inteligência Artificial, após o Senado ter solicitado a uma comissão de juristas que debatessem a respeito do tema. Tal relatório foi entregue pelo Ministro do STJ Ricardo Cueva ao Presidente do Senado em 07/12/2022. Em maio de 2023, o referido anteprojeto virou o PL 2338/2023.<sup>37</sup>

O Projeto de Lei nº 2338/23 é o mais recente dentre as iniciativas de regulamentação da inteligência artificial no Brasil. O projeto foi protocolado em maio de 2023 pelo então presidente do Senado Federal, Rodrigo Pacheco, e visa estabelecer conjunto de regras e princípios a serem seguidos para que se garanta o bom desenvolvimento da inteligência

---

<sup>35</sup> UNIÃO EUROPEIA. Parlamento Europeu; Conselho da União Europeia. *Regulamento (UE) 2016/679, de 27 abr. 2016*. Relativo à proteção das pessoas singulares no que diz respeito ao tratamento de dados pessoais e à livre circulação desses dados, revogando a Diretiva 95/46/CE (Regulamento Geral sobre Proteção de Dados). *Jornal Oficial da União Europeia*, L 119, p. 1-88, 4 maio 2016. Disponível em: <https://eur-lex.europa.eu/eli/reg/2016/679/oj>. Acesso em: 5 maio 2025.

<sup>36</sup> UNIÃO EUROPEIA. Parlamento Europeu; Conselho da União Europeia. *Regulamento (UE) 2024/1689, de 13 jun. 2024*. Estabelece regras harmonizadas sobre inteligência artificial e altera os Regulamentos (CE) n.º 300/2008, (UE) n.º 167/2013, (UE) n.º 168/2013, (UE) 2018/858, (UE) 2018/1139 e (UE) 2019/2144, bem como as Diretivas 2014/90/UE, (UE) 2016/797 e (UE) 2020/1828 (Lei da Inteligência Artificial). *Jornal Oficial da União Europeia*, L 2024/1689, p. 1-144, 12 jul. 2024. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. Acesso em: 5 maio 2025.

<sup>37</sup> MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas: a reparação dos danos causados por sistemas de inteligência artificial*. Londrina: Thoth, 2023. p. 9.



artificial, bem como a segurança da população brasileira e a proteção dos seus direitos fundamentais. Atualmente o projeto tramita no Congresso Nacional e ainda pode passar por alterações.

Assim como o AI Act, o PL 2338/23 tem abordagem que se baseia em risco e seus níveis, sendo assim, divide diferentes categorias de inteligências artificiais entre risco limitado, que são os sistemas mais básicos, como *chatbots* ou IA's que geram texto, e risco excessivo, que são as IA's que apresentam uma ameaça inadmissível, como sistemas de manipulação comportamental que empreguem técnicas subliminares, por exemplo.<sup>38</sup> Dessa forma se mede a necessidade regulatória de cada sistema.

Além do sistema de classificação de riscos, há outros paralelos importantes, tais como um sistema sólido de governança de dados de treinamento a fim de se evitar vieses algorítmicos. Salvaguarda dos direitos das pessoas afetadas, garantindo acesso à informação, explicabilidade, transparência e não-discriminação. Há definição das obrigações dos agentes de IA, que estabelecem responsabilidades dos operadores e fornecedores de sistemas de IA incluindo a necessidade de supervisão humana, plano de gerenciamento de riscos. Abordam-se, também, pontos sobre incentivos à inovação de forma controlada.

Por fim, o projeto, em seu capítulo 'V', estabelece regras acerca de responsabilidade civil por danos causados por sistemas de IA. Nesta linha, propõe dois pontos importantes: responsabilidade civil objetiva e solidária entre fornecedor e operador. Logo no artigo nº 27 propõe-se a atribuição de responsabilidade aos operadores e fornecedores em caso de dano patrimonial, moral individual ou coletivo, sendo assim, a obrigação de reparar os danos causados existe independentemente de comprovação de culpa.

Essa abordagem tem como base a teoria do risco da atividade presente no Código de Defesa do Consumidor<sup>39</sup> e visa a proteção dos interesses do usuário, considerando a complexidade dos sistemas de IA em relação à parte vulnerável e os benefícios da parte mais forte.

De forma geral, o projeto busca, além de garantir o bom uso e desenvolvimento da inteligência artificial no Brasil, uma forma de se encaixar no atual sistema jurídico, abordando

---

<sup>38</sup> “Art. 13. Previamente a sua colocação no mercado ou utilização em serviço, todo sistema de inteligência artificial passará por avaliação preliminar realizada pelo fornecedor para classificação de seu grau de risco, cujo registro considerará os critérios previstos neste capítulo.” BRASIL. Projeto de Lei n. 2.338, de 2023. Regulamenta o uso da inteligência artificial no Brasil. Brasília, DF, 2023. Disponível em: <https://www.camara.leg.br/propostas-legislativas/2338-2023>. Acesso em: 5 maio 2025.

<sup>39</sup> BRASIL. Lei n. 8.078, de 11 de setembro de 1990. Dispõe sobre a proteção do consumidor e dá outras providências. *Diário Oficial da União*, Brasília, DF, 12 set. 1990. Disponível em: [https://www.planalto.gov.br/ccivil\\_03/leis/l8078.htm](https://www.planalto.gov.br/ccivil_03/leis/l8078.htm). Acesso em: 5 maio 2025.

os princípios presentes, principalmente, no Código Civil e Código de Defesa do Consumidor. Não apenas no Brasil, mas em todo o mundo, a busca por integração e segurança jurídica são resultados positivos do atual cenário trazido pela era tecnológica de maneira a traçar um marco regulatório que garanta às pessoas os benefícios da inteligência artificial, enquanto minimiza os seus riscos e protege os direitos fundamentais de todos.

## 5. CONCLUSÃO

Com o avanço crescente da autonomia e complexidade dos sistemas de inteligência artificial, a aplicação dos modelos tradicionais de responsabilidade civil tem sido desafiadora. Dessa forma, a criação de uma estrutura jurídica capaz de acompanhar as mudanças trazidas pela nova era tecnológica tornou-se indispensável.

Nessa linha, diversos países se mobilizam a fim de analisar abordagens regulatórias capazes de garantir a segurança de seu povo. Com destaque ao AI Act da União Europeia, que se tornou inspiração para demais legislações pelo mundo, incluindo o projeto de lei brasileira de regulamentação da IA, a discussão acerca da adoção de regimes de responsabilidade entrou em cena.

Apesar da tentativa de fazer com que a responsabilidade civil recaia sobre a própria IA, a falta de êxito e as atuais propostas e legislações apoiam dois pontos importantes. Primeiro, de que a responsabilidade deve, sim, ser distribuída entre os agentes e, segundo, que a responsabilidade deve ser objetiva, ou seja, independentemente da comprovação da culpa dos operadores e fornecedores.

O presente artigo teve como objetivo explorar o seguinte problema de pesquisa: "Quem deve ser responsabilizado por ilegalidades resultantes do uso da Inteligência Artificial: os desenvolvedores, os operadores, as empresas que implementam a tecnologia, ou a própria IA e quais são as consequências éticas dessa responsabilidade?"

A hipótese levantada foi de que: "A responsabilidade civil deve ser distribuída entre os diferentes atores envolvidos no desenvolvimento e uso de IA e a responsabilização por danos causados por IA pode incentivar o desenvolvimento de sistemas mais éticos e seguros."

Os resultados dessa análise corroboram a hipótese inicial de distribuição de responsabilidade entre os sujeitos envolvidos no desenvolvimento e uso da IA. Já o incentivo regulatório e ético advindo da responsabilização estimula, não apenas a criação de sistemas de IA mais seguros, justos e transparentes, mas a implementação de práticas de governança, maior controle e verificação de vieses e garantia de explicabilidade aos usuários.

A criação de um regime de responsabilidade civil eficaz é um passo crucial na garantia de segurança jurídica, diante do potencial inovador das novas tecnologias, para que a sociedade possa receber os benefícios que a IA tem a oferecer, com respeito aos seus direitos fundamentais, dignidade da pessoa humana e princípios éticos.

## REFERÊNCIAS

ANTUNES, Henrique Souza. A responsabilidade civil aplicável à inteligência artificial: primeiras notas críticas sobre a Resolução do Parlamento Europeu de 2020. *Revista de Direito da Responsabilidade*, Coimbra, ano 3, p. 139-154, 2021. Disponível em: <https://repositorio.ucp.pt/handle/10400.14/40695>. Acesso em: 5 maio 2025.

ANTUNES, Henrique Souza. Inteligência artificial e responsabilidade civil: enquadramento. *Revista de Direito da Responsabilidade*, Coimbra, v. 1, p. 139-154, 2019.

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

BARBOSA, Mafalda Miranda. Inteligência artificial, responsabilidade civil e causalidade: breves notas. *Revista de Direito da Responsabilidade*, Coimbra, ano 3, p. 605-625, 2021. Disponível em: <https://revistadireitoresponsabilidade.pt/2021/inteligencia-artificial-responsabilidade-civil-e-causalidade-breves-notas-mafalda-miranda-barbosa/>. Acesso em: 5 maio 2025.

BARFIELD, Woodrow. Liability for autonomous and artificially intelligent robots. *Paladyn, Journal of Behavioral Robotics*, v. 9, n. 1, p. 193-203, 2018. DOI: 10.1515/pjbr-2018-0018.

BOTELHO, Catarina Santos. The end of the deception? – Counteracting algorithmic discrimination in the digital age. 2023. Pré-print. Disponível em: <https://ssrn.com/abstract=4659451>. Acesso em: 5 maio 2025. (Forthcoming in: DE GREGORIO, Giovanni; POLLICINO, Oreste; VALCKE, Peggy (org.). *The Oxford handbook on digital constitutionalism*. Oxford: Oxford University Press, 2024).

BRASIL. Lei n. 10.406, de 10 de janeiro de 2002. Institui o Código Civil. *Diário Oficial da União*: seção 1, Brasília, DF, 11 jan. 2002. Disponível em: [https://www.planalto.gov.br/ccivil\\_03/leis/2002/110406compilada.htm](https://www.planalto.gov.br/ccivil_03/leis/2002/110406compilada.htm). Acesso em: 5 maio 2025.

BRASIL. Lei n. 8.078, de 11 de setembro de 1990. Dispõe sobre a proteção do consumidor e dá outras providências. *Diário Oficial da União*, Brasília, DF, 12 set. 1990. Disponível em: [https://www.planalto.gov.br/ccivil\\_03/leis/l8078.htm](https://www.planalto.gov.br/ccivil_03/leis/l8078.htm). Acesso em: 5 maio 2025.

BRASIL. Projeto de Lei n. 2.338, de 2023. Regulamenta o uso da inteligência artificial no Brasil. Brasília, DF, 2023. Disponível em: <https://www.camara.leg.br/propostas-legislativas/2338-2023>. Acesso em: 5 maio 2025.

COUTINHO, Luiza Loureiro. *Riscos em sistemas de inteligência artificial: definição, tipologias, correlações, principiologia, responsabilidade civil e regulação*. Salvador: Juspodivm, 2024.

DINIZ, Maria Helena. *Manual de direito civil*. 4. ed. São Paulo: Saraiva Jur, 2022.

FALEIROS JÚNIOR, José Luiz de Moura. Evolução da inteligência artificial em breve retrospectiva. In: BARBOSA, Mafalda Miranda et al. (org.). *Direito digital e inteligência artificial: diálogos entre Brasil e Europa*. Brasília: Editora Foco, 2021. p. 3-26.

FLORIDI, Luciano. The ethics of artificial intelligence: exacerbated problems, renewed problems, unprecedented problems – introduction to the special issue of the *American Philosophical Quarterly* dedicated to the ethics of AI. Abr. 2024. Pré-print. Disponível em: <https://ssrn.com/abstract=4801799>. Acesso em: 5 maio 2025.

FONSECA, Aline Klayse dos Santos. Interação ser humano-máquina: o padrão de “controle humano significativo” e seus impactos na imputação da responsabilidade civil por danos decorrentes de veículos autônomos. *Revista da Faculdade de Direito, Universidade de São Paulo, [S. l.]*, v. 117, p. 357–376, 2022. Disponível em: <https://revistas.usp.br/rfdusp/article/view/218284>. Acesso em: 15 julho 2025.

HAYKIN, Simon. *Redes neurais: princípios e prática*; tradução Paulo Martins Engel. – 2. Ed. – Dados eletrônicos. Porto Alegre: Bookman, 2007.

KURKI, Visa A. J. *A theory of legal personhood*. Oxford: Oxford University Press, 2019. (Oxford Legal Philosophy). DOI: 10.1093/oso/9780198844037.001.0001.

MCCARTHY, J.; MINSKY, M. L.; ROCHESTER, N.; SHANNON, C. E. A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. *AI Magazine, [S. l.]*, v. 27, n. 4, p. 12, 2006. DOI: 10.1609/aimag.v27i4.1904. Disponível em: <https://ojs.aaai.org/aimagazine/index.php/aimagazine/article/view/1904>. Acesso em: 25 jun. 2025.

MEDON, Filipe. *Inteligência artificial e responsabilidade civil: autonomia, riscos e solidariedade*. São Paulo: JusPodivm, 2022.

MEDON, Filipe; FALEIROS JÚNIOR, José Luiz de Moura. Discriminação algorítmica de preços, perfilização e responsabilidade civil nas relações de consumo. *Revista de Direito da Responsabilidade*, Coimbra, ano 3, p. 947-969, 2021. Disponível em: <https://revistadireitoresponsabilidade.pt/2021/discriminacao-algoritmica-de-precos-perfilizacao-e-responsabilidade-civil-nas-relacoes-de-consumo-jose-luiz-de-moura-faleiros-junior-filipe-medon/>. Acesso em: 5 maio 2025.

MELO, Gustavo da Silva. *Discriminação algorítmica na tomada de decisões automatizadas: a reparação dos danos causados por sistemas de inteligência artificial*. Londrina: Thoth, 2023.

NEGRI, Sergio M. C. A. ; LOPES, G. F. P. . DA PERSONALIDADE ELETRÔNICA À CLASSIFICAÇÃO DE RISCOS NA INTELIGÊNCIA ARTIFICIAL (IA). *Teoria Jurídica Contemporânea* , v. 6, p. 01-26, 2022.

PARÉ, Zaven. Real robots: estudo comparativo de robôs sociais. *Illuminuras*, Porto Alegre, v. 25, n. 69, p. 246-268, 2024. DOI: 10.22456/1984-1191.142279.

PARLAMENTO EUROPEU. Resolução do Parlamento Europeu, de 16 de fevereiro de 2017, que contém recomendações à Comissão sobre disposições de Direito Civil sobre Robótica (2015/2103(INL)). Disponível em: [https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051\\_PT.html](https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_PT.html). Acesso em: 5 maio 2025.

TAULLI, Tom. *Introdução à inteligência artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020.

TEFFÉ, Chiara Spadaccini de; MEDON, Filipe. Responsabilidade civil e regulação de novas tecnologias: questões acerca da utilização de inteligência artificial na tomada de decisões empresariais. *Revista Estudos Institucionais*, Rio de Janeiro, v. 6, n. 1, p. 301-333, 2020. DOI: 10.21783/rei.v6i1.383. Disponível em: <https://estudosinstitucionais.com/REI/article/view/383>. Acesso em: 5 maio 2025.

TEUBNER, Gunther. Digital personhood? The status of autonomous software agents in private law. Tradução de Jacob Watson. *Ancilla Iuris*, 2018. Disponível em <https://ssrn.com/abstract=3177096>. Acesso em: 27 junho 2025.

TURING, Alan M. Computing machinery and intelligence. *Creative Computing*, v. 6, n. 1, p. 44-53, 1980.

UNIÃO EUROPEIA. Parlamento Europeu; Conselho da União Europeia. *Regulamento (UE) 2016/679, de 27 abr. 2016*. Relativo à proteção das pessoas singulares no que diz respeito ao tratamento de dados pessoais e à livre circulação desses dados, revogando a Diretiva 95/46/CE (Regulamento Geral sobre Proteção de Dados). *Jornal Oficial da União Europeia*, L 119, p. 1-88, 4 maio 2016. Disponível em: <https://eur-lex.europa.eu/eli/reg/2016/679/oj>. Acesso em: 5 maio 2025.

UNIÃO EUROPEIA. Parlamento Europeu; Conselho da União Europeia. *Regulamento (UE) 2024/1689, de 13 jun. 2024*. Estabelece regras harmonizadas sobre inteligência artificial e altera os Regulamentos (CE) n.º 300/2008, (UE) n.º 167/2013, (UE) n.º 168/2013, (UE) 2018/858, (UE) 2018/1139 e (UE) 2019/2144, bem como as Diretivas 2014/90/UE, (UE) 2016/797 e (UE) 2020/1828 (Lei da Inteligência Artificial). *Jornal Oficial da União Europeia*, L 2024/1689, p. 1-144, 12 jul. 2024. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. Acesso em: 5 maio 2025

VIANA, Guilherme Manoel de Lima; MACEDO, Caio Sperandeo de. Inteligência artificial e a discriminação algorítmica: uma análise do caso Amazon. *Direito & TI*, v. 1, n. 19, p. 39-62, 2024. Disponível em: <https://direitoeti.com.br/direitoeti/article/view/212>. Acesso em: 5 maio 2025.

ZANATTA, Rafael. Perfilização, discriminação e direitos: do Código de Defesa do Consumidor à Lei Geral de Proteção de Dados Pessoais. [S.l.: s.n.], 2019. Disponível em: [https://www.researchgate.net/publication/331287708\\_Perfilizacao\\_Discriminacao\\_e\\_Direitos\\_do\\_Codigo\\_de\\_Defesa\\_do\\_Consumidor\\_a\\_Lei\\_Geral\\_de\\_Protecao\\_de\\_Dados\\_Pessoais](https://www.researchgate.net/publication/331287708_Perfilizacao_Discriminacao_e_Direitos_do_Codigo_de_Defesa_do_Consumidor_a_Lei_Geral_de_Protecao_de_Dados_Pessoais). Acesso em: 5 maio 2025.